

Answers

1. Answer: b

Explanation:

Understanding Analogies: Pipe and Moon Relationship

This question asks us to find the relationship between the first pair of words, 'Pipe' and 'Cylinder', and then apply that same relationship to the third word, 'Moon', to find the fourth word among the given options. This type of problem tests our ability to identify connections and patterns between different concepts.

Analyzing the First Pair: Pipe and Cylinder

Let's look closely at the relationship between 'Pipe' and 'Cylinder':

- A pipe is an object, usually hollow, used to convey fluids or gases.
- Geometrically, a pipe typically has the shape of a cylinder. Even though it's hollow, its outer and inner surfaces are cylindrical. The overall form and structure are based on a cylinder.
- So, the relationship is based on the fundamental geometric shape associated with the object. A pipe has a cylindrical shape.

Applying the Relationship to the Third Word: Moon

Now we need to apply the same type of relationship (object to its geometric shape) to the word 'Moon'.

- The Moon is a celestial body that orbits the Earth.
- What is the primary geometric shape of the Moon?
- Astronomically, large celestial bodies like the Moon are approximately spherical due to gravity.

Connecting Moon to its Shape

Following the pattern established by 'Pipe : Cylinder', where the object is related to its geometric shape, the 'Moon' should be related to its geometric shape. The shape of the Moon is a sphere.

Evaluating the Options

Let's examine the given options to see which one fits the identified relationship:

1. Starts: This refers to when something begins. Not a geometric shape.
2. Sphere: This is a three-dimensional geometric shape. The Moon is approximately spherical.
3. Night: The Moon is often visible at night, but 'night' is a time period, not a geometric shape.
4. Earth: This is another celestial body. The Moon orbits the Earth, but 'Earth' is not the geometric shape of the Moon itself.

Conclusion: Identifying the Related Term for Moon

Based on the analysis, the relationship is "Object is related to its geometric shape". A pipe is related to a cylinder because a pipe has a cylindrical shape. Similarly, the Moon is related to a sphere because the Moon has a spherical shape.

Therefore, the option that correctly completes the analogy 'Pipe : Cylinder :: Moon : ?' is 'Sphere'.

Revision Table: Key Relationships

First Term	Second Term	Relationship
Pipe	Cylinder	Geometric Shape of Object
Moon	Sphere	Geometric Shape of Object

Additional Information on Analogies and Shapes

Analogies questions help improve logical reasoning and vocabulary. They require identifying the precise connection between two items and finding another pair that shares the same connection.

Understanding basic geometric shapes is also useful in many contexts, including spatial reasoning and scientific descriptions. A cylinder is a 3D shape with two parallel circular bases and a curved surface connecting them. A sphere is a perfectly round 3D shape where every point on the surface is an equal distance from the center.

2. Answer: c

Explanation:

Solving the Given Linear Equation

We are asked to find the value of x in the linear equation:

$$\frac{5x}{3} - \frac{7}{2} \left(\frac{2x}{5} - \frac{1}{3} \right) = \frac{1}{3}$$

To solve this equation for x , we need to simplify it by removing the parentheses and combining like terms. Let's break down the process into several steps.

Step-by-Step Method to Find the Value of x

Here is how we can solve this linear equation:

1. Simplify the expression inside the parentheses:
2. First, distribute the $\frac{7}{2}$ to both terms inside the parentheses:
3. $\frac{7}{2} \left(\frac{2x}{5} - \frac{1}{3} \right) = \left(\frac{7}{2} \times \frac{2x}{5} \right) - \left(\frac{7}{2} \times \frac{1}{3} \right)$
4. $= \frac{14x}{10} - \frac{7}{6}$
5. Simplify the fraction $\frac{14x}{10}$:
6. $\frac{14x}{10} = \frac{7x}{5}$
7. So the simplified expression is $\frac{7x}{5} - \frac{7}{6}$.
8. Substitute the simplified expression back into the original equation:
9. The equation becomes:

10. $\frac{5x}{3} - \left(\frac{7x}{5} - \frac{7}{6}\right) = \frac{1}{3}$

11. Be careful with the minus sign before the parentheses. Distribute the minus sign:

12. $\frac{5x}{3} - \frac{7x}{5} + \frac{7}{6} = \frac{1}{3}$

13. **Isolate terms with x on one side and constant terms on the other:**

14. Subtract $\frac{7}{6}$ from both sides of the equation:

15. $\frac{5x}{3} - \frac{7x}{5} = \frac{1}{3} - \frac{7}{6}$

16. **Combine the terms with x :**

17. Find a common denominator for $\frac{5x}{3}$ and $\frac{7x}{5}$. The least common multiple (LCM) of 3 and 5 is 15.

18. $\frac{5x}{3} = \frac{5x \times 5}{3 \times 5} = \frac{25x}{15}$

19. $\frac{7x}{5} = \frac{7x \times 3}{5 \times 3} = \frac{21x}{15}$

20. Now combine them:

21. $\frac{25x}{15} - \frac{21x}{15} = \frac{25x-21x}{15} = \frac{4x}{15}$

22. So the left side of the equation is $\frac{4x}{15}$.

23. **Combine the constant terms on the right side:**

24. Find a common denominator for $\frac{1}{3}$ and $\frac{7}{6}$. The LCM of 3 and 6 is 6.

25. $\frac{1}{3} = \frac{1 \times 2}{3 \times 2} = \frac{2}{6}$

26. Now combine them:

27. $\frac{2}{6} - \frac{7}{6} = \frac{2-7}{6} = \frac{-5}{6}$

28. So the right side of the equation is $\frac{-5}{6}$.

29. **Equate the simplified sides and solve for x :**

30. The equation is now:

31. $\frac{4x}{15} = \frac{-5}{6}$

32. To isolate x , multiply both sides by 15:

33. $4x = \frac{-5}{6} \times 15$

34. $4x = \frac{-75}{6}$

35. Simplify the fraction $\frac{-75}{6}$ by dividing the numerator and denominator by their greatest common divisor, which is 3:

36. $\frac{-75 \div 3}{6 \div 3} = \frac{-25}{2}$

37. So, $4x = \frac{-25}{2}$.

38. Now, divide both sides by 4:

39. $x = \frac{-25}{2} \div 4$

40. $x = \frac{-25}{2} \times \frac{1}{4}$

41. $x = \frac{-25 \times 1}{2 \times 4}$

42. $x = \frac{-25}{8}$

Thus, the value of x that satisfies the given equation is $-\frac{25}{8}$.

Revision Table: Steps to Solve Linear Equations

Step	Description	Action in this problem
1	Simplify by removing parentheses (distribute)	Distribute $-\frac{7}{2}$ in $-\frac{7}{2}(\frac{2x}{5} - \frac{1}{3})$
2	Combine like terms on each side (if any)	No like terms on either side initially after distribution
3	Move terms with variable to one side, constants to the other	Move $\frac{7}{6}$ to the right side
4	Combine variable terms and constant terms	Combine $\frac{5x}{3} - \frac{7x}{5}$ and $\frac{1}{3} - \frac{7}{6}$ using common denominators
5	Solve for the variable	Isolate x by multiplying by the reciprocal of its coefficient

Additional Information on Solving Algebraic Equations

Solving linear equations involves isolating the variable using inverse operations. Key concepts include:

- **Distributive Property:** $a(b + c) = ab + ac$. This is used to remove parentheses.
- **Combining Like Terms:** Terms with the same variable raised to the same power (or constant terms) can be added or subtracted.
- **Properties of Equality:**
 - **Addition Property:** Adding the same number to both sides of an equation keeps the equation balanced.
 - **Subtraction Property:** Subtracting the same number from both sides keeps the equation balanced.
 - **Multiplication Property:** Multiplying both sides by the same non-zero number keeps the equation balanced.

- Division Property: Dividing both sides by the same non-zero number keeps the equation balanced.
- **Least Common Multiple (LCM):** Used to find a common denominator when adding or subtracting fractions. Multiplying the entire equation by the LCM of all denominators can clear the fractions, making the equation easier to solve. In step 6 above, instead of combining fractions first, we could have multiplied the entire equation $\frac{5x}{3} - \frac{7x}{5} + \frac{7}{6} = \frac{1}{3}$ by the LCM of 3, 5, and 6 (which is 30) to clear the denominators.

3. Answer: b

Explanation:

Understanding the Physics Problem: Kinetic Energy of a Combined System

This problem asks us to find the speed of a system consisting of a boy and a scooter, given their individual masses and the total kinetic energy of the system. The key concept here is kinetic energy, which is the energy of motion.

When the boy is riding the scooter, they move together at the same speed. This means we can treat them as a single system with a combined mass.

Given Information

- Mass of the boy (m_b): 50 kg
- Mass of the scooter (m_s): 100 kg
- Total kinetic energy of the boy and scooter (KE_{total}): 76.8 kJ
- We need to find the speed (v) in m/s.

Calculating the Combined Mass

Since the boy and the scooter are moving together, their masses add up to form the total mass of the system.

Combined mass (M) = Mass of boy (m_b) + Mass of scooter (m_s)

$$M = m_b + m_s$$

$$M = 50 \text{ kg} + 100 \text{ kg}$$

$$M = 150 \text{ kg}$$

Converting Kinetic Energy to Joules

The kinetic energy is given in kilojoules (kJ). To use it in standard physics formulas, we need to convert it to joules (J). $1 \text{ kJ} = 1000 \text{ J}$.

$$KE_{total} = 76.8 \text{ kJ}$$

$$KE_{total} = 76.8 \times 1000 \text{ J}$$

$$KE_{total} = 76800 \text{ J}$$

Applying the Kinetic Energy Formula

The formula for kinetic energy is:

$$KE = \frac{1}{2} M v^2$$

Where: **Your Personal Exams Guide**

- KE is the kinetic energy
- M is the mass of the object or system
- v is the speed of the object or system

We know the total kinetic energy (KE_{total}) and the combined mass (M). We need to solve for the speed (v).

Substitute the known values into the formula:

$$76800 \text{ J} = \frac{1}{2} \times 150 \text{ kg} \times v^2$$

$$76800 = 75 \times v^2$$

Solving for Speed (v)

Now, we rearrange the equation to find v^2 and then take the square root to find v .

Divide both sides by 75:

$$v^2 = \frac{76800}{75}$$

$$v^2 = 1024$$

Take the square root of both sides:

$$v = \sqrt{1024}$$

$$v = 32 \text{ m/s}$$

The speed of the scooter and the boy is 32 m/s.

Summary of Calculation Steps

Step	Description	Calculation	Result
1	Calculate Combined Mass	$50 \text{ kg} + 100 \text{ kg}$	150 kg
2	Convert KE to Joules	$76.8 \times 1000 \text{ J}$	76800 J
3	Set up KE equation	$76800 = \frac{1}{2} \times 150 \times v^2$	$76800 = 75v^2$
4	Solve for v^2	$v^2 = \frac{76800}{75}$	$v^2 = 1024$
5	Solve for v	$v = \sqrt{1024}$	32 m/s

The calculated speed matches option 2.

Revision Table: Key Concepts for Kinetic Energy and Mass

Concept	Definition/Formula	Units	Importance in Problem
Kinetic Energy (KE)	Energy due to motion, $KE = \frac{1}{2}mv^2$	Joules (J)	Given total energy, used to find speed
Mass (m, M)	Amount of matter in an object	Kilograms (kg)	Boy's mass + Scooter's mass = Combined mass
Speed (v)	Rate of change of position	meters per second (m/s)	The unknown quantity to be calculated
Combined Mass	Total mass of a system of objects moving together	Kilograms (kg)	Used as 'M' in the KE formula for the system

Additional Information: Energy and Motion

Kinetic energy is one of the fundamental forms of energy. It depends on both the mass of an object and its speed squared. This means that if you double the speed, the kinetic energy increases by a factor of four!

In this problem, we treated the boy and the scooter as a single rigid body because they move together. This is a common approach in physics problems involving systems of objects that are coupled and move with the same velocity.

Other forms of energy include potential energy (energy due to position or state, like gravitational potential energy or elastic potential energy), thermal energy (heat), chemical energy, nuclear energy, etc. The principle of conservation of energy states that energy cannot be created or destroyed, only transformed from one form to another in a closed system.

4. Answer: d

Explanation:

Finding the Third Proportional to 24 and 60

The question asks us to find the third proportional to the numbers 24 and 60. Understanding the concept of proportionality is key here.

When three numbers, say a , b , and c , are in continued proportion, it means that the ratio of the first to the second is equal to the ratio of the second to the third. This can be written as:

$$\frac{a}{b} = \frac{b}{c}$$

In this continued proportion, a is the first proportional, b is the mean proportional, and c is the third proportional.

In our problem, the two given numbers are 24 and 60. If 24, 60, and the unknown third proportional (c) are in continued proportion, then:

- The first number is $a = 24$.
- The second number (which is also the mean proportional in this context) is $b = 60$.
- The third proportional is the number we need to find, let's call it c .

So, we can set up the proportion as follows:

$$\frac{24}{60} = \frac{60}{c}$$

To find the value of c , we can use cross-multiplication. This involves multiplying the numerator of the first ratio by the denominator of the second ratio and setting it equal to the product of the denominator of the first ratio and the numerator of the second ratio.

Applying cross-multiplication to our equation:

$$24 \times c = 60 \times 60$$

Now, we simplify the equation:

$$24c = 3600$$

To isolate c , we divide both sides of the equation by 24:

$$c = \frac{3600}{24}$$

Let's perform the division:

$$c = \frac{3600}{24} = \frac{1800}{12} = \frac{900}{6} = \frac{450}{3} = 150$$

So, the third proportional to 24 and 60 is 150.

We can verify this by checking if 24, 60, and 150 are in continued proportion:

$$\frac{24}{60} = \frac{24 \div 12}{60 \div 12} = \frac{2}{5}$$

$$\frac{60}{150} = \frac{60 \div 30}{150 \div 30} = \frac{2}{5}$$

Since both ratios are equal to $\frac{2}{5}$, the numbers 24, 60, and 150 are indeed in continued proportion, and 150 is the correct third proportional.

Revision Table: Key Concepts in Proportion

Your Personal Exams Guide

Concept	Definition	Example
Ratio	A comparison of two quantities of the same kind by division.	The ratio of 10 to 5 is $\frac{10}{5} = 2$ or 2:1.
Proportion	An equality of two ratios. If $\frac{a}{b} = \frac{c}{d}$, then a, b, c, d are in proportion.	$\frac{2}{4} = \frac{3}{6}$, so 2, 4, 3, 6 are in proportion.
Continued Proportion	When the ratio of the first term to the second is equal to the ratio of the second term to the third, and so on. $\frac{a}{b} = \frac{b}{c} = \frac{c}{d} = \dots$	3, 6, 12 are in continued proportion because $\frac{3}{6} = \frac{6}{12} = \frac{1}{2}$.
Mean Proportional	For two numbers a and c , the mean proportional b satisfies $\frac{a}{b} = \frac{b}{c}$, which means $b^2 = ac$, so $b = \sqrt{ac}$.	The mean proportional between 4 and 9 is $\sqrt{4 \times 9} = \sqrt{36} = 6$.
Third Proportional	For two numbers a and b , the third proportional c satisfies $\frac{a}{b} = \frac{b}{c}$.	For 24 and 60, the third proportional c satisfies $\frac{24}{60} = \frac{60}{c}$.
Fourth Proportional	For three numbers a, b, c , the fourth proportional d satisfies $\frac{a}{b} = \frac{c}{d}$.	The fourth proportional to 1, 2, and 3 is d such that $\frac{1}{2} = \frac{3}{d}$, so $d = 6$.

Additional Information: Exploring Proportionality

Proportionality is a fundamental concept in mathematics that describes how two quantities are related. Beyond continued proportion, there are other important types of proportionality, such as direct proportion and inverse proportion.

- **Direct Proportion:** Two quantities are in direct proportion if they increase or decrease at the same rate. If y is directly proportional to x , we write $y \propto x$, which means $y = kx$ for some constant k . For example, the cost of purchasing apples is directly proportional to the number of apples bought.
- **Inverse Proportion:** Two quantities are in inverse proportion if an increase in one quantity leads to a decrease in the other, and vice versa, such that their product

is constant. If y is inversely proportional to x , we write $y \propto \frac{1}{x}$, which means $y = \frac{k}{x}$ or $xy = k$ for some constant k . For example, the time taken to complete a job is inversely proportional to the number of workers, assuming they work at the same rate.

Understanding these different types of proportionality helps in solving a wide range of problems in various fields, including physics, chemistry, economics, and everyday life.

5. Answer: b

Explanation:

Understanding Base and Derived Units in Physics

In the International System of Units (SI), physical quantities are measured using specific units. These units are broadly classified into two categories: base units and derived units.

What are SI Base Units?

SI base units are the fundamental units from which all other units can be derived. There are seven defined SI base units:

- Metre (m) for length
- Kilogram (kg) for mass
- Second (s) for time
- Ampere (A) for electric current
- Kelvin (K) for thermodynamic temperature
- Mole (mol) for amount of substance
- Candela (cd) for luminous intensity

These base units are considered independent of each other.

What are Derived Units?

Derived units are units of measurement that are obtained by combining the base units through multiplication or division. For example, the unit for speed is metres per second (m/s), which is derived from the base units of length (metre) and time (second).

Many physical quantities like force (Newton, N), energy (Joule, J), power (Watt, W), pressure (Pascal, Pa), etc., have derived units.

Analyzing the Given Options

We are asked to identify the option that is NOT a derived unit. Let's examine each option based on our understanding of base and derived units:

1. **Volt:** The Volt (V) is the SI derived unit of electric potential difference or electromotive force. It is defined as one joule per coulomb ($1 \text{ V} = 1 \text{ J/C}$). Both Joule (J) and Coulomb (C) are derived units ($1 \text{ J} = 1 \text{ kg} \cdot \text{m}^2/\text{s}^2$, $1 \text{ C} = 1 \text{ A} \cdot \text{s}$). Since Volt is defined in terms of other derived units (which are themselves derived from base units), Volt is a derived unit.
2. **Mole:** The Mole (mol) is the SI unit for the amount of substance. As listed above, the mole is one of the seven fundamental SI base units. It is not derived from other units; it is a foundational unit.
3. **Radian:** The Radian (rad) is the SI unit for measuring plane angles. It is defined as the angle subtended at the center of a circle by an arc equal in length to the radius. Dimensionally, it is the ratio of arc length to radius, i.e., m/m , making it a dimensionless unit. Historically, radian and steradian were considered supplementary units in SI, but in 1995, they were classified as derived units. While dimensionless, it is a unit derived from the ratio of two lengths.
4. **Lumen:** The Lumen (lm) is the SI derived unit of luminous flux. It is defined in terms of the candela (cd), which is a base unit, and the steradian (sr), which is the SI derived unit of solid angle. One lumen is equal to one candela times one steradian ($1 \text{ lm} = 1 \text{ cd} \cdot \text{sr}$). Since it is a combination of a base unit and a derived unit, Lumen is a derived unit.

Conclusion on Derived Units

Based on the analysis:

- Volt is a derived unit.
- Mole is a base unit.
- Radian is a derived unit (specifically, a dimensionless derived unit).
- Lumen is a derived unit.

The question asks which option is NOT a derived unit. The Mole is a base unit, and therefore it is not a derived unit.

Classification of Units

Unit	Symbol	Classification	Notes
Volt	V	Derived Unit	Unit of electric potential difference
Mole	mol	Base Unit	Unit of amount of substance
Radian	rad	Derived Unit	Unit of plane angle (dimensionless)
Lumen	lm	Derived Unit	Unit of luminous flux

Therefore, the unit that is NOT a derived unit among the given options is the Mole.

Revision Table: SI Units Classification

Key Differences: Base vs. Derived Units

Feature	Base Units	Derived Units
Definition	Fundamental, independent units	Combinations of base units
Number in SI	Seven	Infinite (as needed)
Examples	Metre, Kilogram, Second, Ampere, Kelvin, Mole, Candela	Newton, Joule, Watt, Volt, Pascal, Lumen, Radian, Steradian, etc.

Additional Information on Units

Understanding the distinction between base and derived units is crucial for dimensional analysis in physics and other sciences. Dimensional analysis helps check the consistency of equations and understand the relationships between different physical quantities.

While radian and steradian are dimensionless derived units, they are sometimes given special names and symbols (rad and sr) because they represent distinct physical quantities (plane angle and solid angle) that are important in many applications.

The definition of the mole is related to Avogadro's number (N_A), which is defined as $6.02214076 \times 10^{23}$ entities per mole. One mole of any substance contains exactly N_A elementary entities (like atoms, molecules, ions, etc.).

The correct option is the one that is a base unit, as base units are not derived units.

6. Answer: a

Explanation:

Calculating Cost Price with Loss Percentage

This problem requires us to find the original cost price of an item when the selling price and the percentage loss incurred are known. We are given that a vendor sells a coconut for Rs. 32 and incurs a 20% loss on this sale.

Understanding the Concepts

When a vendor incurs a loss, it means the selling price (SP) is less than the cost price (CP). The loss percentage is calculated based on the cost price.

- **Cost Price (CP):** The original price at which the vendor bought the item.
- **Selling Price (SP):** The price at which the vendor sells the item.

- **Loss:** The difference between the cost price and the selling price when $SP < CP$.
Loss = $CP - SP$.
- **Loss Percentage:** $(\text{Loss} / CP) * 100\%$.

Formula for Calculating Cost Price with Loss

We can relate the Selling Price, Cost Price, and Loss Percentage using the following formula:

If there is a loss of $L\%$, the Selling Price is calculated as:

$$SP = CP \times \left(1 - \frac{L}{100}\right)$$

To find the Cost Price when SP and $L\%$ are given, we can rearrange the formula:

$$CP = \frac{SP}{\left(1 - \frac{L}{100}\right)}$$

Step-by-Step Calculation

Given:

- Selling Price (SP) = Rs. 32
- Loss Percentage (L) = 20%

We need to find the Cost Price (CP).

Using the formula $CP = \frac{SP}{\left(1 - \frac{L}{100}\right)}$:

1. Substitute the given values into the formula: $CP = \frac{32}{\left(1 - \frac{20}{100}\right)}$
2. Simplify the fraction inside the parenthesis: $CP = \frac{32}{(1-0.20)}$
3. Perform the subtraction: $CP = \frac{32}{0.80}$
4. Calculate the final value of CP : $CP = \frac{32}{0.8} = \frac{320}{8}$ $CP = 40$

So, the cost price of the coconut is Rs. 40.

Verification

Let's check if selling at Rs. 32 results in a 20% loss on a CP of Rs. 40.

- Loss = CP – SP = Rs. 40 – Rs. 32 = Rs. 8
- Loss Percentage = $\left(\frac{\text{Loss}}{\text{CP}}\right) \times 100\% = \left(\frac{8}{40}\right) \times 100\%$
- Loss Percentage = $\left(\frac{1}{5}\right) \times 100\% = 20\%$

This matches the given information, confirming our calculated cost price is correct.

Therefore, the cost price of the coconut is Rs. 40.

Revision Table: Profit and Loss Key Formulas

Concept	Formula (in terms of CP & SP)	Formula (in terms of CP & % Change)
Profit	Profit = SP – CP (when SP > CP)	N/A
Loss	Loss = CP – SP (when CP > SP)	N/A
Profit %	Profit % = $\left(\frac{\text{Profit}}{\text{CP}}\right) \times 100$	N/A
Loss %	Loss % = $\left(\frac{\text{Loss}}{\text{CP}}\right) \times 100$	N/A
SP (with Profit P%)	N/A	$\text{SP} = \text{CP} \times \left(1 + \frac{P}{100}\right)$
SP (with Loss L%)	N/A	$\text{SP} = \text{CP} \times \left(1 - \frac{L}{100}\right)$
CP (with Profit P%)	N/A	$\text{CP} = \frac{\text{SP}}{\left(1 + \frac{P}{100}\right)}$
CP (with Loss L%)	N/A	$\text{CP} = \frac{\text{SP}}{\left(1 - \frac{L}{100}\right)}$

Additional Information: Profit and Loss Scenarios

Profit and loss calculations are fundamental in business and everyday transactions. They help determine the financial outcome of selling goods or services.

- **Profit Scenario:** If a vendor sells an item for more than its cost price, they make a profit. The profit percentage is calculated based on the cost price.
- **Break-even Point:** If a vendor sells an item at exactly its cost price ($SP = CP$), there is no profit or loss.
- **Calculating Percentage Change:** Both profit percentage and loss percentage are calculated relative to the cost price, not the selling price, unless explicitly stated otherwise. This is a common point of confusion.
- **Application:** These concepts are widely used in retail, wholesale, manufacturing, and financial analysis to evaluate performance and make pricing decisions.

7. Answer: a

Explanation:

Solving Cycling Speed and Time Problems

This problem involves the relationship between speed, time, and distance. The fundamental concept is that for a constant distance, speed and time are inversely proportional. This means if speed increases, time decreases, and vice versa, assuming the distance covered is the same.

The relationship is given by the formula:

$$\text{Distance} = \text{Speed} \times \text{Time}$$

Let's define the variables for Anwar's usual cycling journey to school:

- Usual Speed = S
- Usual Time = T minutes
- Distance to school = D

So, the distance to school is:

$$D = S \times T$$

Now, consider the journey when Anwar cycles at $\frac{7}{9}$ times his usual speed:

- New Speed = $\frac{7}{9}S$
- New Time = $T + 4$ minutes (since he reaches 4 minutes late)

The distance to school is still the same D . So, we can write the equation for the second journey:

$$D = \left(\frac{7}{9}S\right) \times (T + 4)$$

Since the distance D is the same in both cases, we can set the two expressions for D equal to each other:

$$S \times T = \frac{7}{9}S \times (T + 4)$$

We want to find the value of T , which is the usual time. We can solve this equation for T . Since S is the speed and must be greater than 0 (otherwise he wouldn't reach school), we can divide both sides of the equation by S :

$$T = \frac{7}{9} \times (T + 4)$$

Now, let's solve for T :

1. Multiply both sides by 9 to eliminate the fraction:

$$9 \times T = 9 \times \left(\frac{7}{9} \times (T + 4)\right)$$

$$9T = 7(T + 4)$$

2. Distribute the 7 on the right side:

$$9T = 7T + 7 \times 4$$

$$9T = 7T + 28$$

3. Subtract $7T$ from both sides to isolate the term with T :

$$9T - 7T = 28$$

$$2T = 28$$

4. Divide both sides by 2 to find T :

$$T = \frac{28}{2}$$

$$T = 14$$

So, Anwar takes 14 minutes to reach school at his usual cycling speed.

Revision Table: Speed, Time, and Distance Formulas

Concept	Formula	Relationship (Distance Constant)
Distance	Speed $\times$ Time	Speed $\propto \frac{1}{\text{Time}}$
Speed	$\frac{\text{Distance}}{\text{Time}}$	Speed ₁ $\times$ Time ₁ = Speed ₂ $\times$ Time ₂
Time	$\frac{\text{Distance}}{\text{Speed}}$	$\frac{\text{Speed}_1}{\text{Speed}_2} = \frac{\text{Time}_2}{\text{Time}_1}$

Additional Information: Solving Time Change Problems

Problems where speed changes and affects the time taken to cover a fixed distance are common. The key is to understand the inverse relationship between speed and time when distance is constant. If speed becomes $\frac{a}{b}$ times the original speed, the time taken will become $\frac{b}{a}$ times the original time (assuming $a \neq 0, b \neq 0$).

In this problem, the new speed is $\frac{7}{9}$ times the usual speed. Therefore, the new time taken will be $\frac{9}{7}$ times the usual time.

Let the usual time be T .

- New Time = $\frac{9}{7}T$
- The difference in time is given as 4 minutes (late), so:

$$\text{New Time} - \text{Usual Time} = 4 \text{ minutes}$$

$$\frac{9}{7}T - T = 4$$

- Find a common denominator to subtract:

$$\frac{9}{7}T - \frac{7}{7}T = 4$$

$$\left(\frac{9-7}{7}\right)T = 4$$

$$\frac{2}{7}T = 4$$

- Solve for T :

$$T = 4 \times \frac{7}{2}$$

$$T = \frac{28}{2}$$

$$T = 14$$

This method using the inverse proportion confirms the result obtained earlier. The usual cycling time is 14 minutes.

8. Answer: c

Explanation:

Understanding Chemicals that Damage the Ozone Layer

The ozone layer is a region of Earth's stratosphere that absorbs most of the Sun's ultraviolet (UV) radiation. It contains high concentrations of ozone (O_3) relative to other parts of the atmosphere, although still small relative to other gases in the stratosphere. This layer acts as a protective shield for life on Earth. However, certain human-made chemicals can reach the stratosphere and break down ozone molecules, leading to depletion of the ozone layer. We need to identify which class of chemicals is known for causing this damage.

Analyzing the Options for Ozone Layer Damage

Let's examine each option to determine its impact on the ozone layer:

- **Antimicrobials:** These are substances used to kill or inhibit the growth of microorganisms. While they are important in medicine, hygiene, and agriculture, common antimicrobials are not typically associated with significant depletion of the stratospheric ozone layer.
- **Aromatic compounds:** These are organic compounds that contain a benzene ring or similar ring structure. Examples include benzene, toluene, and xylene. While some volatile organic compounds can contribute to air pollution and ground-level ozone formation (which is harmful), they do not directly cause the depletion of the stratospheric ozone layer in the same way that certain halogenated hydrocarbons do.
- **Chlorofluorocarbons:** This class of chemicals, often abbreviated as CFCs, are organic compounds that contain carbon, chlorine, and fluorine atoms. They were widely used in refrigerants, aerosols, foam blowing agents, and solvents due to their stability and non-toxicity at ground level. However, their stability allows them to persist in the atmosphere and eventually reach the stratosphere. In the stratosphere, UV radiation breaks down CFCs, releasing chlorine atoms. These chlorine atoms then participate in catalytic cycles that efficiently destroy ozone molecules. This is the primary mechanism by which CFCs damage the ozone layer.
- **Phenols:** Phenols are organic compounds containing a hydroxyl group ($-OH$) bonded directly to an aromatic hydrocarbon group. Phenols are used in disinfectants, antiseptics, and in the production of resins and plastics. Like many other organic compounds, phenols are generally broken down in the lower atmosphere and are not known to cause significant stratospheric ozone depletion.

Identifying the Chemicals Damaging the Ozone Layer

Based on the analysis, Chlorofluorocarbons (CFCs) are the class of chemicals notorious for causing significant damage to the stratospheric ozone layer. Their release of chlorine atoms under stratospheric UV radiation triggers a chain reaction that destroys ozone molecules.

Chemical Class	Known Effect on Stratospheric Ozone Layer
Antimicrobials	Generally not known to cause significant damage
Aromatic compounds	Generally not known to cause significant damage (can contribute to ground-level ozone)
Chlorofluorocarbons (CFCs)	Known to cause significant depletion by releasing chlorine atoms
Phenols	Generally not known to cause significant damage

The Montreal Protocol, an international treaty, was established to phase out the production and consumption of ozone-depleting substances, including CFCs, because of their severe impact on the ozone layer.

Revision Table: Ozone Depleting Substances

Your Personal Exams Guide

Substance Class	Examples	Mechanism of Ozone Depletion
Chlorofluorocarbons (CFCs)	CCl_3F (CFC-11), CCl_2F_2 (CFC-12)	Release Cl atoms in stratosphere; $\text{Cl} + \text{O}_3 \rightarrow \text{ClO} + \text{O}_2$; $\text{ClO} + \text{O} \rightarrow \text{Cl} + \text{O}_2$. A single Cl atom can destroy many O_3 molecules.
Hydrochlorofluorocarbons (HCFCs)	CHClF_2 (HCFC-22)	Contain C-H bonds, less stable than CFCs, break down partially in lower atmosphere, but still release Cl in stratosphere. Lower ozone depletion potential than CFCs.
Halons	CBrF_3 (Halon-1301), CBrClF_2 (Halon-1211)	Contain bromine (Br), which is even more effective at destroying ozone than chlorine. Used in fire extinguishers.
Carbon Tetrachloride (CCl_4)		Releases Cl atoms in stratosphere. Used as a solvent.
Methyl Chloroform (CH_3CCl_3)		Releases Cl atoms in stratosphere. Used as a solvent.
Methyl Bromide (CH_3Br)		Releases Br atoms in stratosphere. Used as a fumigant.

Additional Information: Protecting the Ozone Layer

Efforts to protect the ozone layer have been largely successful due to global cooperation under the Montreal Protocol. Key aspects include:

- Phasing out production and consumption of Ozone Depleting Substances (ODS).
- Developing and using alternatives to CFCs and other ODS, such as Hydrofluorocarbons (HFCs). Note that while HFCs do not deplete ozone (they contain no chlorine or bromine), many are potent greenhouse gases.

Subsequent agreements like the Kigali Amendment address HFCs due to their climate impact.

- Monitoring the ozone layer to track its recovery. Scientists have observed signs of recovery in the ozone hole over Antarctica thanks to the reduction in ODS.

Understanding which chemicals damage the ozone layer is crucial for environmental protection and the study of atmospheric chemistry.

9. Answer: b

Explanation:

Solving the Steel Bar Cutting Problem

This problem asks us to find the decimal fraction of a steel bar that is left over after cutting multiple pieces of a specific length. We start with a total length of the steel bar and need to determine how many pieces can be cut and what length remains.

Step-by-Step Calculation for Leftover Steel Bar

Here's how we can solve this **steel bar cutting** problem:

1. **Calculate the maximum number of pieces:** We need to find out how many pieces of 5.25 m length can be cut from the 50 m long bar. This is done by dividing the total length by the length of each piece. The number of pieces must be a whole number, as you cannot cut a fraction of a piece.

Total length = 50 m

Length per piece = 5.25 m

$$\text{Number of pieces} = \frac{\text{Total length}}{\text{Length per piece}} = \frac{50}{5.25}$$

Let's perform the division:

$$\frac{50}{5.25} = \frac{50 \times 100}{5.25 \times 100} = \frac{5000}{525}$$

Now, simplify the fraction:

$$\frac{5000}{525} = \frac{1000 \times 5}{105 \times 5} = \frac{1000}{105}$$

$$\frac{1000}{105} = \frac{200 \times 5}{21 \times 5} = \frac{200}{21}$$

Dividing 200 by 21:

$$200 \div 21 \approx 9.52$$

Since only whole pieces can be cut, the workman can cut 9 pieces.

2. **Calculate the total length used:** Multiply the number of pieces cut by the length of each piece.

Number of pieces = 9

Length per piece = 5.25 m

Total length used = 9×5.25 m

$$9 \times 5.25 = 9 \times (5 + 0.25) = (9 \times 5) + (9 \times 0.25) = 45 + 2.25 = 47.25 \text{ m}$$

The total length used is 47.25 m.

3. **Calculate the length left:** Subtract the total length used from the original total length of the bar.

Original total length = 50 m

Total length used = 47.25 m

$$\text{Length left} = 50 \text{ m} - 47.25 \text{ m} = 2.75 \text{ m}$$

The length of the bar left is 2.75 m.

4. **Calculate the decimal fraction of the whole left:** Divide the length left by the original total length of the bar.

Length left = 2.75 m

Original total length = 50 m

$$\text{Decimal fraction left} = \frac{\text{Length left}}{\text{Original total length}} = \frac{2.75}{50}$$

To convert this fraction to a decimal, we can multiply the numerator and denominator by 2 to make the denominator 100:

$$\frac{2.75}{50} = \frac{2.75 \times 2}{50 \times 2} = \frac{5.5}{100} = 0.055$$

The decimal fraction of the whole bar left is 0.055.

Conclusion on the Decimal Fraction

After cutting as many 5.25 m pieces as possible from the 50 m **steel bar**, the length remaining is 2.75 m. When expressed as a **decimal fraction** of the original length, this leftover piece is 0.055 of the whole bar.

Revision Table: Steel Bar Problem

Description	Value	Unit
Original Bar Length	50	m
Length per Piece	5.25	m
Maximum Number of Pieces Cut	9	-
Total Length Used	47.25	m
Length Left	2.75	m
Decimal Fraction Left	0.055	-

Additional Information: Understanding Decimal Fractions

A **decimal fraction** is a fraction where the denominator is a power of ten (like 10, 100, 1000, etc.). However, the term is also commonly used to refer to any number written in decimal form, especially when it represents a part of a whole. In this problem, the **decimal fraction** represents the proportion of the original **steel bar** that remains after

cutting. Calculating this involves division and understanding how to represent quantities as parts of a whole.

When dealing with real-world problems like **steel bar cutting**, it's important to consider practical constraints, such as only being able to cut whole pieces. This often means using integer division (finding the quotient) to determine the number of items that can be made or cut, before calculating the remainder or leftover quantity.

10. Answer: b

Explanation:

The given pattern is as follows:

$a(\underline{x} \ y \ z)\underline{b} \ b(\underline{x} \ \underline{y} \ z)c \ \underline{c} \ c(\underline{x} \ \underline{y} \ z)$

If we divide the series into two parts, we obtain the following pattern:

xyz xyz xyz

a bb ccc

Therefore, the correct answer is 'xbcy'.

11. Answer: c

Explanation:

Understanding Resistors in Parallel Circuits

In this problem, we are given two resistors connected in parallel. One resistor has a resistance of $R \Omega$, and the other has a resistance of 15Ω . We are told that the combined or effective resistance of this parallel combination is 12Ω . Our goal is to find the value of the unknown resistance R .

Parallel Resistance Calculation

When resistors are connected in parallel, the reciprocal of the effective resistance (R_{eq}) is equal to the sum of the reciprocals of the individual resistances. The formula for two resistors, R_1 and R_2 , connected in parallel is:

$$\frac{1}{R_{eq}} = \frac{1}{R_1} + \frac{1}{R_2}$$

An alternative formula, often useful for just two resistors, is:

$$R_{eq} = \frac{R_1 \times R_2}{R_1 + R_2}$$

Step-by-Step Solution to Find R

Let's use the second formula with the given values:

- $R_1 = R$ (the unknown resistance)
- $R_2 = 15 \Omega$ (the known resistance)
- $R_{eq} = 12 \Omega$ (the effective resistance)

Substitute these values into the formula:

$$12 \Omega = \frac{R \times 15 \Omega}{R + 15 \Omega}$$

Now, we need to solve this equation for R :

1. Multiply both sides by $(R + 15)$ to remove the denominator:

$$12 \times (R + 15) = 15R$$

2. Distribute the 12 on the left side:

$$12R + (12 \times 15) = 15R$$

$$12R + 180 = 15R$$

3. Subtract $12R$ from both sides of the equation to isolate the terms with R on one side:

$$180 = 15R - 12R$$

$$180 = 3R$$

4. Divide both sides by 3 to find the value of R :

$$R = \frac{180}{3}$$

$$R = 60 \Omega$$

So, the value of the unknown resistor R is 60Ω .

Verifying the Result

Let's quickly check if a 60Ω resistor in parallel with a 15Ω resistor gives an effective resistance of 12Ω :

$$\frac{1}{R_{eq}} = \frac{1}{60} + \frac{1}{15}$$

Find a common denominator, which is 60:

$$\frac{1}{R_{eq}} = \frac{1}{60} + \frac{4}{60}$$

$$\frac{1}{R_{eq}} = \frac{1+4}{60}$$

$$\frac{1}{R_{eq}} = \frac{5}{60}$$

Simplify the fraction:

$$\frac{1}{R_{eq}} = \frac{1}{12}$$

Taking the reciprocal of both sides:

$$R_{eq} = 12 \Omega$$

This matches the given effective resistance, confirming our calculation is correct.

Revision Table: Key Concepts

Concept	Description	Formula (for two resistors)
Resistors in Parallel	Components are connected across the same two points, providing alternative paths for current.	$\frac{1}{R_{eq}} = \frac{1}{R_1} + \frac{1}{R_2}$ or $R_{eq} = \frac{R_1 R_2}{R_1 + R_2}$
Effective Resistance (R_{eq})	The total resistance of a circuit or part of a circuit. In parallel, R_{eq} is always less than the smallest individual resistance.	Calculated using the parallel resistance formula.

Additional Information: Parallel vs. Series Circuits

Understanding the difference between parallel and series connections is crucial in circuit analysis. Here's a brief comparison:

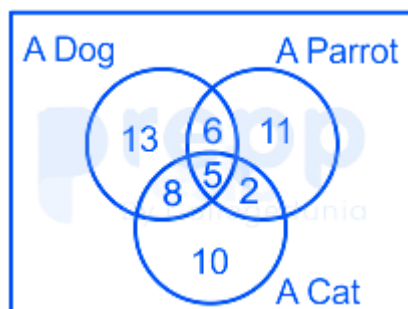
- **Series Connection:** Components are connected end-to-end, forming a single path for the current. The total resistance is the sum of individual resistances ($R_{eq} = R_1 + R_2 + R_3 + \dots$). The current is the same through all components.
- **Parallel Connection:** Components are connected across the same two points, providing multiple paths for the current. The reciprocal of the total resistance is the sum of the reciprocals of individual resistances. The voltage is the same across all components in parallel.

Parallel connections are commonly used in household wiring so that each appliance receives the full voltage and can be operated independently.

12. Answer: d

Explanation:

Houses with two pets is shown by the shaded region below:



Number of houses that have only two pets = $(8 + 2 + 6) = 16$

Hence, '16' is the correct answer.

13. Answer: a

Explanation:

Solving the Square Root of Decimals Problem

The problem asks us to calculate the sum of two square roots involving decimal numbers: $\sqrt{0.2025} + \sqrt{0.1225}$.

To solve this, we need to find the square root of each decimal number separately and then add the results.

Step 1: Find the Square Root of 0.2025

To find the square root of a decimal, we can first ignore the decimal point and find the square root of the integer part. The integer part is 2025.

Let's find the square root of 2025. We can try squaring numbers:

- $40 \times 40 = 1600$
- $50 \times 50 = 2500$

Since 2025 is between 1600 and 2500, its square root is between 40 and 50. Also, 2025 ends in 5, so its square root must end in 5.

Let's try 45:

- $45 \times 45 = 2025$

So, the square root of 2025 is 45.

Now, let's consider the decimal places. The number 0.2025 has 4 decimal places. When we take the square root of a number with decimal places, the result will have half the number of decimal places. So, the square root of 0.2025 will have $4 \div 2 = 2$ decimal places.

Therefore, $\sqrt{0.2025} = 0.45$.

Step 2: Find the Square Root of 0.1225

Again, ignore the decimal point and find the square root of the integer part, which is 1225.

Let's find the square root of 1225. We can try squaring numbers:

- $30 \times 30 = 900$
- $40 \times 40 = 1600$

Since 1225 is between 900 and 1600, its square root is between 30 and 40. The number 1225 ends in 5, so its square root must end in 5.

Let's try 35:

- $35 \times 35 = 1225$

So, the square root of 1225 is 35.

The number 0.1225 has 4 decimal places. Its square root will have $4 \div 2 = 2$ decimal places.

Therefore, $\sqrt{0.1225} = 0.35$.

Step 3: Add the Square Roots

Now we need to add the results from Step 1 and Step 2:

$$\sqrt{0.2025} + \sqrt{0.1225} = 0.45 + 0.35$$

Adding the two decimal numbers:

	Units	.	Tenths	Hundredths
	0	.	4	5
+	0	.	3	5
---	---	---	---	---
	0	.	8	0

$0.45 + 0.35 = 0.80$, which is equal to 0.8.

Conclusion

The sum of the square roots $\sqrt{0.2025} + \sqrt{0.1225}$ is 0.8.

Revision Table: Square Roots of Decimals

Concept	Explanation	Example
Finding Square Root of Decimal	1. Find the square root of the number ignoring the decimal point. 2. Count the number of decimal places in the original number (let it be n). 3. The square root will have $n/2$ decimal places.	$\sqrt{0.09}$. $\sqrt{9} = 3$. 0.09 has 2 decimal places. $\sqrt{0.09}$ has $2/2 = 1$ decimal place. So, $\sqrt{0.09} = 0.3$.
Adding Decimals	Align the decimal points vertically and add the numbers column by column, carrying over as needed.	$0.45 + 0.35 = 0.80$

Additional Information: Perfect Squares

Recognizing perfect squares can make finding square roots easier. A perfect square is an integer that is the square of another integer.

Examples of perfect squares:

- $1^2 = 1$
- $2^2 = 4$
- $3^2 = 9$
- ...
- $10^2 = 100$
- ...
- $30^2 = 900$
- $35^2 = 1225$
- ...
- $40^2 = 1600$
- $45^2 = 2025$
- ...

When dealing with decimals like 0.2025, recognizing that 2025 is a perfect square (45 squared) is the key to quickly finding its square root.

14. Answer: c

Explanation:

Understanding Equilibrium with a Uniform Meter Scale

This question is about the principle of moments, which is essential for understanding rotational equilibrium. A meter scale is a rigid body, and when forces act on it, it can rotate unless the moments are balanced.

For an object like a meter scale to be in rotational equilibrium, the sum of the clockwise moments about the pivot must be equal to the sum of the anticlockwise moments about the same pivot.

Analyzing Forces and Moments on the Meter Scale

A uniform meter scale has its weight acting at its geometrical center. Since it's a uniform meter scale, its weight acts at the 50 cm mark. The scale weighs 50 g.

The scale is pivoted at the 70 cm mark. This is the point around which rotation can occur, so we calculate moments about this point.

- **Force 1: Weight of the scale**
- Magnitude: 50 g
- Point of application: 50 cm mark
- Distance from pivot (70 cm mark): $|50 \text{ cm} - 70 \text{ cm}| = 20 \text{ cm}$
- Direction of moment: The 50 cm mark is to the left of the 70 cm pivot. This force will tend to rotate the scale anticlockwise.
- Anticlockwise Moment = Force $\times$ Distance = $50 \text{ g} \times 20 \text{ cm} = 1000 \text{ g} \cdot \text{cm}$

We need to place a 40 g mass on the scale to achieve equilibrium. Let's say this mass is placed at the mark 'x' cm.

- **Force 2: Weight of the 40 g mass**
- Magnitude: 40 g
- Point of application: x cm mark (unknown)
- Distance from pivot (70 cm mark): $|x \text{ cm} - 70 \text{ cm}|$

To balance the anticlockwise moment caused by the scale's weight (which is to the left of the pivot), the 40 g mass must create a clockwise moment. This means the 40 g mass must be placed to the right of the pivot (70 cm mark). Therefore, x must be greater than 70 cm, and the distance from the pivot is $(x - 70) \text{ cm}$.

- Direction of moment: Clockwise (since $x > 70$)
- Clockwise Moment = Force $\times$ Distance = $40 \text{ g} \times (x - 70) \text{ cm}$

Applying the Principle of Moments for Equilibrium

For the meter scale to be in equilibrium, the total clockwise moment must equal the total anticlockwise moment about the pivot:

$$\text{Anticlockwise Moment} = \text{Clockwise Moment}$$

$$1000 \text{ g} \cdot \text{cm} = 40 \text{ g} \times (x - 70) \text{ cm}$$

Solving for the Position of the 40g Mass

Now, we solve the equation for x:

$$\frac{1000}{40} = x - 70$$

$$25 = x - 70$$

$$x = 25 + 70$$

$$x = 95$$

So, the 40 g mass should be placed at the 95 cm mark.

Let's verify this:

- Anticlockwise moment due to scale's weight: $50 \text{ g} \times (70 \text{ cm} - 50 \text{ cm}) = 50 \text{ g} \times 20 \text{ cm} = 1000 \text{ g} \cdot \text{cm}$
- Clockwise moment due to 40 g mass at 95 cm: $40 \text{ g} \times (95 \text{ cm} - 70 \text{ cm}) = 40 \text{ g} \times 25 \text{ cm} = 1000 \text{ g} \cdot \text{cm}$

Since the anticlockwise moment equals the clockwise moment, the scale is in equilibrium when the 40 g mass is placed at the 95 cm mark.

The position where the 40 g mass should be placed is the 95 cm mark.

Component	Weight (Force)	Point of Action	Distance from Pivot (70 cm)	Moment Direction	Moment (Force × Distance)
Uniform Scale	50 g	50 cm	$ 50 - 70 = 20 \text{ cm}$	Anticlockwise	$50 \times 20 = 1000 \text{ g} \cdot \text{cm}$
40 g Mass	40 g	x cm (to be found)	$ x - 70 \text{ cm}$	Clockwise (for equilibrium)	$40 \times x - 70 \text{ g} \cdot \text{cm}$

Revision Table: Key Concepts in Rotational Equilibrium

Concept	Description	Importance in this Problem
Principle of Moments	For equilibrium, sum of clockwise moments = sum of anticlockwise moments about the pivot.	Foundation for setting up the equilibrium equation.
Moment of a Force	Moment = Force $\times$ Perpendicular distance from the pivot. Represents the turning effect of a force.	Used to calculate the moment created by the scale's weight and the added mass.
Center of Gravity	The point where the entire weight of an object is considered to act. For a uniform object, it's the geometrical center.	Determines the point of action of the scale's weight (50 cm mark).

Additional Information: Extending Your Understanding of Moments

The principle of moments is a specific condition for rotational equilibrium. For complete equilibrium (both translational and rotational), the net force acting on the object must also be zero. In this problem, the upward force at the pivot balances the downward weights of the scale and the added mass, ensuring translational equilibrium in the vertical direction.

Understanding moments is crucial for many applications, from simple levers and see-saws to complex structures and machinery. The concept helps engineers and physicists predict how objects will behave under various forces and ensure stability.

Always remember to choose a pivot point (usually the point of support or a point where an unknown force acts) and consistently define clockwise and anticlockwise moments when solving such problems.

15. Answer: b

Explanation:

Finding $a+b$ using Algebraic Identities

The problem asks us to find the value of $a + b$, given the values of $a^2 + b^2$ and ab . We are given:

- $a^2 + b^2 = 53$
- $ab = 14$

This problem can be solved using a fundamental algebraic identity that relates $(a + b)^2$, $a^2 + b^2$, and ab .

Applying the Algebraic Identity

The relevant algebraic identity is the square of a sum:

$$(a + b)^2 = a^2 + 2ab + b^2$$

We can rearrange this identity to group the terms we are given:

$$(a + b)^2 = (a^2 + b^2) + 2ab$$

Now, we can substitute the given values for $a^2 + b^2$ and ab into this rearranged identity.

Step-by-Step Calculation

Let's substitute the given values into the equation:

$$(a + b)^2 = (a^2 + b^2) + 2ab$$

Substitute $a^2 + b^2 = 53$ and $ab = 14$:

$$(a + b)^2 = 53 + 2(14)$$

Now, perform the multiplication:

$$(a + b)^2 = 53 + 28$$

Perform the addition:

$$(a + b)^2 = 81$$

We have found the value of $(a + b)^2$. To find the value of $a + b$, we need to take the square root of both sides of the equation.

$$\sqrt{(a + b)^2} = \sqrt{81}$$

$$a + b = \pm 9$$

The equation gives two possible values for $a + b$: $+9$ and -9 . Since the typical context for such problems and the nature of the options provided usually implies the positive root, we consider the positive value.

Thus, the value of $a + b$ is 9.

Summary of Steps

Step	Description	Calculation
1	Recall the identity for $(a + b)^2$	$(a + b)^2 = a^2 + 2ab + b^2$
2	Rearrange the identity	$(a + b)^2 = (a^2 + b^2) + 2ab$
3	Substitute given values	$(a + b)^2 = 53 + 2(14)$
4	Simplify	$(a + b)^2 = 53 + 28 = 81$
5	Take square root	$a + b = \sqrt{81}$
6	Find the value	$a + b = 9$ (considering the positive root)

Revision Table: Algebraic Identities

Understanding basic algebraic identities is crucial for solving many problems. Here's a quick revision of some important ones:

- $(a + b)^2 = a^2 + 2ab + b^2$

- $(a - b)^2 = a^2 - 2ab + b^2$
- $a^2 - b^2 = (a + b)(a - b)$
- $(a + b)^3 = a^3 + 3a^2b + 3ab^2 + b^3 = a^3 + b^3 + 3ab(a + b)$
- $(a - b)^3 = a^3 - 3a^2b + 3ab^2 - b^3 = a^3 - b^3 - 3ab(a - b)$
- $a^3 + b^3 = (a + b)(a^2 - ab + b^2)$
- $a^3 - b^3 = (a - b)(a^2 + ab + b^2)$

Additional Information: Solving for a and b Individually

While the problem only asked for $a + b$, it is sometimes possible to find the individual values of a and b from the given equations $a^2 + b^2 = 53$ and $ab = 14$. We can use the value of $a + b$ we just found ($a + b = 9$).

We have a system of two equations:

1. $a + b = 9$
2. $ab = 14$

From equation (1), we can express b as $b = 9 - a$. Substitute this into equation (2):

$$a(9 - a) = 14$$

$$9a - a^2 = 14$$

Rearrange into a quadratic equation:

$$a^2 - 9a + 14 = 0$$

This quadratic equation can be factored:

$$(a - 7)(a - 2) = 0$$

This gives two possible values for a : $a = 7$ or $a = 2$.

- If $a = 7$, then from $b = 9 - a$, $b = 9 - 7 = 2$.
- If $a = 2$, then from $b = 9 - a$, $b = 9 - 2 = 7$.

In both cases, the pair of values for (a, b) is either $(7, 2)$ or $(2, 7)$. Let's check if these pairs satisfy the original conditions:

- If $a = 7, b = 2$:
 - $a + b = 7 + 2 = 9$ (Matches our result for $a+b$)
 - $ab = 7 \times 2 = 14$ (Matches the given condition)
 - $a^2 + b^2 = 7^2 + 2^2 = 49 + 4 = 53$ (Matches the given condition)
- If $a = 2, b = 7$:
 - $a + b = 2 + 7 = 9$ (Matches our result for $a+b$)
 - $ab = 2 \times 7 = 14$ (Matches the given condition)
 - $a^2 + b^2 = 2^2 + 7^2 = 4 + 49 = 53$ (Matches the given condition)

Both pairs satisfy the given conditions, and for both pairs, $a + b$ is 9. This confirms our result.

16. Answer: d

Explanation:

Finding the Value of X Using the Median

Understanding the Median Concept

The median is a statistical measure representing the middle value of a dataset when it's arranged in ascending order. For a dataset with an even number of observations, the median is calculated as the average of the two middle values.

Analyzing the Given Data

We are given a set of 8 numbers: 9, 15, 1, 15, 14, 9, 4, and X. The total number of observations (n) is 8, which is an even number.

The median is given as 11.

Step-by-Step Solution

1. **Sort the known numbers:** First, let's arrange the known numbers (excluding X) in ascending order: 1, 4, 9, 9, 14, 15, 15

2. **Determine the median calculation formula:** Since there are 8 numbers in total (an even count), the median is the average of the 4th and 5th numbers in the final sorted list (including X). The formula is:

$$\text{Median} = \frac{(\text{4th term}) + (\text{5th term})}{2}$$

3. **Set up the equation:** We know the median is 11. So,

$$11 = \frac{(\text{4th term}) + (\text{5th term})}{2}$$

Multiplying both sides by 2, we get:

$$22 = (\text{4th term}) + (\text{5th term})$$

4. **Incorporate X and test possibilities:** We need to find the value of X that satisfies the condition $(\text{4th term}) + (\text{5th term}) = 22$ when X is included in the sorted list. Let's consider the sorted list of known numbers: 1, 4, 9, 9, 14, 15, 15. Now, we insert X into this sequence and check the median.

- If we place $X = 10$, the sorted list is: 1, 4, 9, 9, 10, 14, 15, 15. The 4th term is 9 and the 5th term is 10. Median = $(9 + 10)/2 = 9.5$. (Incorrect)
- If we place $X = 11$, the sorted list is: 1, 4, 9, 9, 11, 14, 15, 15. The 4th term is 9 and the 5th term is 11. Median = $(9 + 11)/2 = 10$. (Incorrect)
- If we place $X = 12$, the sorted list is: 1, 4, 9, 9, 12, 14, 15, 15. The 4th term is 9 and the 5th term is 12. Median = $(9 + 12)/2 = 10.5$. (Incorrect)
- If we place $X = 13$, the sorted list is: 1, 4, 9, 9, 13, 14, 15, 15. The 4th term is 9 and the 5th term is 13. Median = $(9 + 13)/2 = 22/2 = 11$. (Correct)

5. **Conclusion:** The value $X = 13$ makes the median of the dataset equal to 11. The sorted list becomes {1, 4, 9, 9, 13, 14, 15, 15}, and the median is calculated as the average of the 4th and 5th elements, which are 9 and 13.

$$\frac{9 + 13}{2} = \frac{22}{2} = 11$$

Final Answer Determination

Based on the calculations, the value of X that results in a median of 11 is 13.

17. Answer: a

Explanation:

Understanding Where Acceleration Due to Gravity is Highest

The acceleration due to gravity, often denoted by g , is the acceleration experienced by an object due to gravitational force. While we often use a standard value like 9.8 m/s^2 , the actual value of g varies slightly across the Earth's surface. This variation depends on factors like altitude, latitude, and the Earth's shape and rotation.

Factors Influencing Acceleration Due to Gravity (g)

Several factors determine the value of g at different locations:

- **Distance from the Earth's Center:** According to Newton's Law of Universal Gravitation, the force of gravity is inversely proportional to the square of the distance between the centers of two masses. The formula for acceleration due to gravity at the surface is given by:

$$g = G \frac{M}{r^2}$$

where G is the universal gravitational constant, M is the mass of the Earth, and r is the radius of the Earth.

- **Earth's Rotation:** The Earth rotates on its axis, creating a centrifugal effect that acts outwards, opposing the gravitational force. This effect is maximum at the equator and zero at the poles. The effective acceleration due to gravity is reduced by this centrifugal force. The approximate formula accounting for rotation is:

$$g_{eff} \approx g - \omega^2 r \cos^2(\lambda)$$

where ω is the angular velocity of Earth's rotation, r is the radius at the given latitude, and λ is the latitude.

- **Earth's Shape:** The Earth is not a perfect sphere; it's an oblate spheroid, meaning it bulges at the equator and is flattened at the poles. This results in the radius being larger at the equator than at the poles.

Analysis of Options

Let's analyze why the acceleration due to gravity is highest at the poles:

1. The Poles

- **Distance:** Locations at the poles are closer to the Earth's center because the Earth is flattened there. A smaller radius (r) leads to a higher value of g according to the formula $g \propto 1/r^2$.
- **Rotation:** The centrifugal effect due to Earth's rotation is zero at the poles ($\cos(90^\circ) = 0$). Therefore, there is no outward force counteracting gravity at the poles.
- **Conclusion:** Being closest to the center and experiencing no centrifugal effect results in the highest acceleration due to gravity at the poles.

2. At an Infinite Distance from the Earth

- As the distance from the Earth increases, the gravitational force decreases. At an infinite distance ($r \rightarrow \infty$), the acceleration due to gravity approaches zero.

3. The Equator

- **Distance:** The equator is farthest from the Earth's center due to the equatorial bulge. A larger radius (r) results in a lower value of g .
- **Rotation:** The centrifugal effect is maximum at the equator ($\cos(0^\circ) = 1$). This outward force significantly reduces the effective acceleration due to gravity.
- **Conclusion:** Due to the larger distance from the center and the maximum centrifugal effect, the acceleration due to gravity is lowest at the equator.

4. The Center of the Earth

- At the very center of the Earth, the gravitational forces from all parts of the Earth acting on an object cancel each other out. Effectively, there is no net

gravitational pull towards the center, and the acceleration due to gravity is zero.

Summary

Based on the analysis of the factors influencing gravity (distance from the center and the effect of Earth's rotation), the acceleration due to gravity is strongest at **the poles**.

18. Answer: a

Explanation:

Understanding the UN Security Council

The United Nations Security Council is one of the six principal organs of the United Nations (UN), charged with ensuring international peace and security, recommending the admission of new UN members to the General Assembly, and approving any changes to the UN Charter. Its powers include establishing peacekeeping operations, enacting international sanctions, and authorizing military action. Its powers are exercised through resolutions.

Permanent Members of the UN Security Council

The Security Council consists of fifteen members. These include five permanent members and ten non-permanent members. The five permanent members hold veto power over any substantive resolution. The composition of the permanent members is fixed and was decided at the founding of the United Nations.

The five permanent members are:

- China
- France
- Russian Federation (succeeded the Soviet Union)
- United Kingdom
- United States

Analysing the Options

The question asks which of the given nations is not a permanent member of the United Nations Security Council. Let's look at the options provided:

1. Germany
2. United Kingdom
3. France
4. China

Comparing this list to the list of the five permanent members, we can identify which country is not a permanent member.

Permanent Member	Listed in Options?
China	Yes
France	Yes
Russian Federation	No
United Kingdom	Yes
United States	No

From the options provided, United Kingdom, France, and China are all permanent members of the UN Security Council.

Identifying the Non-Permanent Member Option

The option that remains and is not listed as a permanent member is Germany.

Germany is a significant global power and a major contributor to the UN system, often serving as a non-permanent member of the Security Council. However, it is not one of the original five permanent members established in 1945.

Conclusion

Based on the established list of permanent members of the United Nations Security Council, Germany is the nation among the given options that is not a permanent member.

Revision Table: UN Security Council Facts

Feature	Description
Total Members	15
Permanent Members (P5)	5
Non-Permanent Members	10 (elected for two-year terms)
Veto Power	Held by Permanent Members
Primary Responsibility	Maintaining international peace and security

Additional Information: Security Council Dynamics

The composition of the Security Council's permanent members reflects the geopolitical landscape at the end of World War II. There have been discussions and proposals over the years to reform the Security Council, including potentially adding new permanent members to better reflect the current global power distribution. However, any changes to the composition of the permanent members require the agreement of the current permanent members, which has historically been difficult to achieve.

Non-permanent members are elected by the General Assembly for two-year terms. Geographical representation is a factor in the election of non-permanent members.

19. Answer: b

Explanation:

Understanding the Percentage Problem

This question asks us to find the percentage increase needed for one number to become equal to another, given that both numbers are a certain percentage less than a third number. Let's break down the problem step-by-step to understand the relationship between these three numbers.

Defining the Numbers

To make the calculation easy, let's assume the third number is a round figure, like 100. This makes calculating percentages straightforward.

- Let the third number be $T = 100$.

Calculating the First Number

The problem states that the first number is 10% less than the third number.

- Percentage less than the third number = 10%
- The first number is $100\% - 10\% = 90\%$ of the third number.
- So, the first number $F = 90\%$ of T
- $F = \frac{90}{100} \times 100 = 90$.

The first number is 90.

Calculating the Second Number

The problem states that the second number is 20% less than the third number.

- Percentage less than the third number = 20%
- The second number is $100\% - 20\% = 80\%$ of the third number.
- So, the second number $S = 80\%$ of T
- $S = \frac{80}{100} \times 100 = 80$.

The second number is 80.

Finding the Required Increase

We need to find by what percentage the second number (80) should be increased to make it equal to the first number (90). The difference between the first number and the second number is the amount of increase needed.

- Required increase amount = First number – Second number
- Required increase amount = $90 - 80 = 10$.

The second number (80) needs to be increased by 10 to become equal to the first number (90).

Calculating the Percentage Increase

The percentage increase is calculated with respect to the original value, which is the second number (80) in this case. The formula for percentage increase is:

$$\text{Percentage Increase} = \left(\frac{\text{Amount of Increase}}{\text{Original Value}} \right) \times 100\%$$

- Amount of Increase = 10
- Original Value (Second number) = 80
- Percentage Increase = $\left(\frac{10}{80} \right) \times 100\%$
- Percentage Increase = $\left(\frac{1}{8} \right) \times 100\%$
- Percentage Increase = $0.125 \times 100\%$
- Percentage Increase = 12.5%

Therefore, the second number should be increased by 12.5% to make it equal to the first number.

Summary of Steps

1. Assume a value for the third number (e.g., 100).
2. Calculate the first number based on its percentage less than the third number.
3. Calculate the second number based on its percentage less than the third number.
4. Find the difference between the first and second numbers.
5. Calculate the percentage increase needed for the second number using the difference and the second number as the base.

Revision Table: Key Concepts

Concept	Explanation	Formula/Example
Percentage Less Than	Value is a certain percentage reduced from a base value.	If value is 10% less than 100, it's $100 - (10\% \text{ of } 100) = 90$.
Percentage Increase	The relative increase between two values, expressed as a percentage of the original value.	$\left(\frac{\text{New Value} - \text{Original Value}}{\text{Original Value}} \right) \times 100\%$
Base for Percentage Increase	The starting value upon which the increase is calculated. In this problem, it's the second number.	If 80 increases to 90, the base is 80. Increase is 10. Percentage increase is $\left(\frac{10}{80} \right) \times 100\%$.

Additional Information: Working with Percentages

Understanding percentages is crucial for many quantitative problems. A percentage is simply a fraction out of 100. For example, 10% means $\frac{10}{100}$ or 0.10. When a value is described as being 'X% less than Y', it means the value is $Y - (X\% \text{ of } Y)$, which can also be written as $Y \times (1 - X/100)$.

In our problem:

- First number is 10% less than the third number: $\text{First} = \text{Third} \times (1 - 10/100) = \text{Third} \times 0.90$.
- Second number is 20% less than the third number: $\text{Second} = \text{Third} \times (1 - 20/100) = \text{Third} \times 0.80$.

To find the percentage increase from the second number to the first, we use the formula $\left(\frac{\text{First} - \text{Second}}{\text{Second}} \right) \times 100\%$. Substituting the expressions in terms of the third number:

$$\text{Percentage Increase} = \left(\frac{\text{Third} \times 0.90 - \text{Third} \times 0.80}{\text{Third} \times 0.80} \right) \times 100\%$$

$$\text{Percentage Increase} = \left(\frac{\text{Third} \times (0.90 - 0.80)}{\text{Third} \times 0.80} \right) \times 100\%$$

$$\text{Percentage Increase} = \left(\frac{\text{Third} \times 0.10}{\text{Third} \times 0.80} \right) \times 100\%$$

The 'Third' term cancels out:

$$\text{Percentage Increase} = \left(\frac{0.10}{0.80} \right) \times 100\%$$

$$\text{Percentage Increase} = \left(\frac{1}{8} \right) \times 100\% = 12.5\%.$$

This shows that the initial assumption of 100 for the third number was just a convenience; the result is independent of the actual value of the third number.

20. Answer: c

Explanation:

Calculating Speed in m/s: A Step-by-Step Guide

The question asks us to find the speed of a truck in meters per second (m/s), given the distance it travels and the time it takes. Speed is a measure of how quickly an object moves over a certain distance.

We are given:

- Distance traveled = 450 km
- Time taken = two and a half hours = 2.5 hours

The formula for speed is:

$$\text{Speed} = \frac{\text{Distance}}{\text{Time}}$$

However, the distance is given in kilometers (km) and the time in hours. We need to convert these units to meters (m) and seconds (s) respectively, because the required speed unit is m/s.

Converting Units for Speed Calculation

Let's convert the distance from kilometers to meters:

- 1 kilometer (km) = 1000 meters (m)
- So, 450 km = $450 \times 1000 \text{ m} = 450000 \text{ m}$

Next, let's convert the time from hours to seconds:

- 1 hour = 60 minutes
- 1 minute = 60 seconds
- So, 1 hour = $60 \times 60 \text{ seconds} = 3600 \text{ seconds}$
- Therefore, 2.5 hours = $2.5 \times 3600 \text{ seconds}$
- $2.5 \times 3600 = (2 + 0.5) \times 3600 = (2 \times 3600) + (0.5 \times 3600) = 7200 + 1800 = 9000 \text{ seconds}$

Now we have the distance in meters and the time in seconds:

- Distance = 450000 m
- Time = 9000 s

Calculating the Speed in m/s

Now we can use the speed formula with the converted units:

$$\text{Speed} = \frac{\text{Distance}}{\text{Time}}$$

$$\text{Speed} = \frac{450000 \text{ m}}{9000 \text{ s}}$$

We can simplify the fraction:

$$\text{Speed} = \frac{450000}{9000} \text{ m/s} = \frac{450}{9} \text{ m/s}$$

$$\text{Speed} = 50 \text{ m/s}$$

So, the speed of the truck is 50 m/s.

Revision Table: Units and Conversions for Speed

Quantity	Common Units	Standard Unit (SI)	Conversion Example (km to m)	Conversion Example (hours to s)
Distance	km, meters, miles	meter (m)	1 km = 1000 m	-
Time	hours, minutes, seconds	second (s)	-	1 hour = 3600 s 1 minute = 60 s
Speed	km/h, m/s, mph	meter per second (m/s)	-	-

Additional Information: Speed and Velocity

While often used interchangeably in everyday language, speed and velocity are slightly different in physics.

- **Speed:** This is a scalar quantity, meaning it only has magnitude (like 50 m/s). It tells you how fast something is moving.
- **Velocity:** This is a vector quantity, meaning it has both magnitude and direction (like 50 m/s East). It tells you how fast something is moving and in what direction.

In this problem, we were asked for speed, which is concerned only with the rate of distance covered over time, without specifying direction.

21. Answer: a

Explanation:

Calculation:

In 2015, No. of employees in all department = $20 + 40 + 50 + 70 + 80 + 30 = 290$

In 2018, No. of employees in all department = $30 + 70 + 80 + 90 + 60 + 30 = 360$

Difference = $360 - 290 = 70$

According to question,

No of employees added = 80

Fired employees from C = $80 - 70 = 10$

Strength of C in 2018 = 80 (Given)

If not fired, strength of C in 2018 = $80 + 10 = 90$

Strength of department C in 2018 would have been greater by = $(10/80) \times 100$

$\Rightarrow 12.5\%$

$\therefore$ Option 1 is the correct answer.

22. Answer: a

Explanation:

The question asks us to identify the physical quantity that is directly proportional to the current flowing through a resistor, assuming the temperature stays constant. This relationship is a fundamental concept in electricity, described by Ohm's Law.

Understanding Ohm's Law and Resistors

Ohm's Law describes the relationship between voltage (potential difference), current, and resistance in an electrical circuit. It states that the electric current through a metallic conductor is directly proportional to the voltage across its ends, provided the temperature and other physical conditions remain unchanged.

In simpler terms, if you increase the voltage across a resistor, the current through it will increase proportionally, as long as the resistor's temperature doesn't change

significantly.

Analyzing the Options

Let's look at the given options:

- **Potential difference:** This is the voltage across the ends of the resistor. Ohm's Law directly relates potential difference and current.
- **Resistivity:** This is an intrinsic property of the material the resistor is made of. It describes how strongly the material resists the flow of electric current. Resistivity is constant for a given material at a specific temperature and does not change with the current or potential difference across the resistor.
- **Charge:** This refers to the fundamental property of matter that experiences a force when placed in an electromagnetic field. While current is the rate of flow of charge, Ohm's Law describes the relationship between potential difference and the rate of charge flow (current), not charge itself.
- **Resistance:** This is the property of a resistor that opposes the flow of electric current. For an ohmic resistor (one that obeys Ohm's Law), resistance is constant under constant temperature and does not depend on the potential difference or current.

Applying Ohm's Law to the Question

The statement in the question is a direct rephrasing of Ohm's Law:

"The _____ across the ends of a resistor is directly proportional to the current through it, provided its temperature remains the same."

According to Ohm's Law, the potential difference (V) across a resistor is directly proportional to the current (I) flowing through it, provided the resistance (R) and temperature are constant. Mathematically, this is expressed as:

$$V \propto I$$

or

$$V = IR$$

Here:

- V is the potential difference across the resistor.
- I is the current flowing through the resistor.
- R is the resistance of the resistor (constant at a constant temperature).

Therefore, the quantity that is directly proportional to the current through the resistor, under constant temperature, is the potential difference across its ends.

Conclusion on Potential Difference and Current

Based on Ohm's Law, the potential difference across the ends of a resistor is directly proportional to the current flowing through it when the temperature is kept constant. This fits perfectly with the description provided in the question.

Quantity	Relation with Current (Constant Temperature)
Potential Difference	Directly Proportional ($V \propto I$)
Resistivity	Independent
Charge	Current is rate of flow of charge ($I = \frac{dQ}{dt}$), not direct proportionality
Resistance	Independent (for ohmic resistors)

Revision Table: Key Electrical Concepts

Term	Definition	Relevant Law/Formula
Potential Difference (Voltage)	The work done per unit charge to move it between two points in an electric field. Measured in Volts (V).	$V = IR$ (Ohm's Law)
Current	The rate of flow of electric charge. Measured in Amperes (A).	$I = \frac{V}{R}$ (Ohm's Law), $I = \frac{\Delta Q}{\Delta t}$
Resistance	Opposition to the flow of electric current. Measured in Ohms (Ω).	$R = \frac{V}{I}$ (Ohm's Law)
Resistivity	An intrinsic material property indicating resistance to current flow. Measured in Ohm-meters ($\Omega \cdot m$).	$R = \rho \frac{L}{A}$, where ρ is resistivity

Additional Information on Electrical Properties

While potential difference is directly proportional to current according to Ohm's Law, it's important to understand the other terms too. Resistivity is a material property, meaning it depends only on what the material is and its temperature, not its shape or the voltage/current applied. Resistance, on the other hand, depends on the material's resistivity, its length, and its cross-sectional area. For many materials, resistance stays constant at a constant temperature, making the potential difference-current proportionality valid.

23. Answer: a

Explanation:

Understanding Heat Absorption During Melting

When a substance melts, it absorbs energy without changing its temperature. This energy is used to break the bonds holding the particles in the solid state. The amount of energy required for a unit mass of a substance to melt at its melting point is called its specific latent heat of fusion.

Applying Specific Latent Heat of Fusion

The question provides the specific latent heat of fusion for ethyl alcohol and the mass of ethyl alcohol that melts. We need to find the total heat absorbed by the ethyl alcohol during this phase change (melting).

- Given: Specific latent heat of fusion (L_f) of ethyl alcohol = 100 Jg^{-1}
- Given: Mass (m) of ethyl alcohol = 2.25 g
- The melting point (-114°C) is given, but it's the temperature at which melting occurs, not needed for the calculation of heat absorbed during the phase change itself.

The formula to calculate the heat absorbed (Q) during melting is:

$$Q = m \times L_f$$

Step-by-Step Calculation of Heat Absorbed

Substitute the given values into the formula:

$$Q = 2.25 \text{ g} \times 100 \text{ Jg}^{-1}$$

$$Q = 225 \text{ J}$$

Therefore, the heat absorbed by 2.25 g of ethyl alcohol when it melts at its melting point is 225 Joules .

Result Summary

The calculation shows that 225 J of heat is absorbed by the ethyl alcohol during the melting process.

Quantity	Value	Unit
Specific Latent Heat of Fusion (L_f)	100	J/g
Mass (m)	2.25	g
Heat Absorbed (Q)	225	J

Revision Table: Key Concepts

Concept	Definition	Formula
Latent Heat	Heat energy absorbed or released during a phase transition (like melting, freezing, boiling, condensation) at constant temperature.	N/A
Specific Latent Heat of Fusion (L_f)	Heat energy absorbed per unit mass of a substance to change it from solid to liquid state at its melting point.	$L_f = Q/m$
Heat Absorbed During Melting (Q)	Total heat energy required to melt a given mass of a substance at its melting point.	$Q = m \times L_f$

Additional Information: Phase Changes and Energy

Phase changes, such as melting, freezing, boiling, condensation, sublimation, and deposition, involve the transfer of energy.

- **Melting (Solid to Liquid):** Energy is absorbed (endothermic). The substance absorbs heat equal to $m \times L_f$.
- **Freezing (Liquid to Solid):** Energy is released (exothermic). The substance releases heat equal to $m \times L_f$.
- **Boiling (Liquid to Gas):** Energy is absorbed (endothermic). This involves the specific latent heat of vaporization (L_v), and heat absorbed is $m \times L_v$.

- **Condensation (Gas to Liquid):** Energy is released (exothermic). The heat released is $m \times L_v$.

During a phase change, the temperature of the substance remains constant as the absorbed or released energy is used to change the state rather than increase or decrease kinetic energy of particles.

24. Answer: a

Explanation:

The question asks about the musical instrument associated with the renowned classical musician, Vilayat Khan. Vilayat Khan was a highly influential figure in Indian classical music.

Understanding Vilayat Khan's Instrument Association

Vilayat Khan (1928–2004) belonged to a lineage of musicians and is widely regarded as one of the greatest sitar players of his time. He was known for his unique style, known as the "gayaki ang" (singing style), which aimed to replicate the nuances of the human voice on the sitar.

Analyzing the Options

Let's look at the provided options:

- **Sitar:** This is a plucked string instrument used in Hindustani classical music. It is the instrument Vilayat Khan was famous for playing.
- **Flute:** A wind instrument. While important in classical music, it is not the instrument Vilayat Khan played.
- **Sarod:** A plucked string instrument without frets, common in Hindustani music. Not Vilayat Khan's primary instrument.
- **Santoor:** A hammered dulcimer, an ancient string instrument. Not associated with Vilayat Khan.

Based on his historical significance and musical legacy, Vilayat Khan is inextricably linked with the sitar.

Conclusion on Vilayat Khan's Instrument

Vilayat Khan dedicated his life to the sitar, evolving its playing techniques and popularizing it globally. His contribution to the sitar is immense and continues to influence musicians today.

Therefore, the musical instrument associated with classical musician Vilayat Khan is the Sitar.

Revision Table: Key Indian Classical Instruments

Instrument	Type	Associated with
Sitar	Plucked String	Ravi Shankar, Vilayat Khan, Nikhil Banerjee
Flute (Bansuri)	Wind	Hariprasad Chaurasia, Pannalal Ghosh
Sarod	Plucked String (Fretless)	Amjad Ali Khan, Ali Akbar Khan
Santoor	Hammered String	Shivkumar Sharma, Bhajan Sopori

Additional Information on Vilayat Khan and Sitar

Vilayat Khan came from the Imdadkhani gharana (school), a lineage of sitar and surbahar players known for their technical mastery and vocalistic style. He was a contemporary of other great sitar maestros like Ravi Shankar, and their contributions significantly shaped the sitar's place in classical music.

His approach to the sitar emphasized melodic expression and the ability to convey the emotional depth of the raga, similar to a vocalist. This "gayaki ang" became a

hallmark of his playing and a major contribution to the sitar tradition.

25. Answer: a

Explanation:

Decoding Blood Relations: P \$ Q * R # S

This question involves decoding a relationship expression using given symbol meanings. Let's break down the symbols and the expression step-by-step to determine the relationship between P and S.

Understanding the Symbols

- A \$ B means A is the sister of B. (A is female)
- A # B means A is the husband of B. (A is male, B is female)
- A * B means A is the daughter of B. (A is female)

Analyzing the Expression: P \$ Q * R # S

We need to work from right to left to determine the relationships sequentially:

1. **R # S:** Based on the definition, R is the husband of S. This tells us that R is male and S is female. They are a couple.
2. **Q * R:** Based on the definition, Q is the daughter of R. Since R is married to S, R and S are the parents of Q. Q is female.
3. **P \$ Q:** Based on the definition, P is the sister of Q. Since Q is the daughter of R and S, and P is Q's sister, P is also the daughter of R and S. P is female because she is Q's sister.

Establishing the Final Relationship

From the analysis:

- P is the daughter of R and S.

- Q is the daughter of R and S.
- P and Q are sisters.
- R is the husband of S.
- S is the wife of R and one of the parents of P and Q.

Therefore, S is the mother of P.

Evaluating the Options

Let's examine each option based on our findings:

- **Option 1:** S is the mother of P. This matches our conclusion.
- **Option 2:** S is the sister-in-law of P. This is incorrect. S is P's mother.
- **Option 3:** S is the grandmother of P. This is incorrect. S is P's mother.
- **Option 4:** S is the sister of P. This is incorrect. S is P's mother, and P is S's daughter.

Based on the decoding, the correct relationship is that S is the mother of P.

Revision Table: Symbol Meanings

Symbol	Meaning	Gender of A	Gender of B
\$	A is sister of B	Female	Undetermined
#	A is husband of B	Male	Female
*	A is daughter of B	Female	Undetermined

Additional Information: Tips for Solving Blood Relation Questions

Solving blood relation problems involves careful decoding and constructing a mental or physical diagram of the family tree. Here are some tips:

- Read the definitions of symbols carefully.
- Break down the expression into smaller segments (e.g., P \$ Q, Q * R, R # S).

- Work through the expression from right to left for clarity, establishing relationships sequentially.
- Determine the gender of individuals whenever possible based on the definitions (e.g., 'husband' implies male, 'daughter' implies female).
- Connect the segments to form a complete chain of relationships.
- Draw a family tree diagram if the relationships become complex.
- Finally, answer the specific question asked about the relationship between two individuals.

26. Answer: c

Explanation:

Calculating Force from Work and Distance

The question asks us to determine the force employed when a known amount of work is done over a specific distance. This involves the fundamental physics concept of work done by a constant force.

Understanding Work Done

Work is done when a force applied to an object causes displacement in the direction of the force. The amount of work done depends on the magnitude of the force and the distance over which it acts.

The relationship between work, force, and distance is given by the formula:

$$W = F \times d$$

Where:

- W is the work done
- F is the force applied
- d is the distance moved in the direction of the force

Applying the Formula to Find Force

We are given the following information:

- Work done (W) = 1,200 J (Joules)
- Distance moved (d) = 20 m (meters)

We need to find the force (F) in Newtons (N).

To find the force, we can rearrange the work formula:

$$F = \frac{W}{d}$$

Step-by-Step Calculation

Now, let's substitute the given values into the rearranged formula:

$$F = \frac{1200 \text{ J}}{20 \text{ m}}$$

Performing the division:

$$F = 60 \text{ N}$$

So, the force employed was 60 Newtons.

Comparing with Options

Let's check our calculated force against the given options:

- Option 1: 120 N
- Option 2: 90 N
- Option 3: 60 N
- Option 4: 30 N

Our calculated force of 60 N matches Option 3.

Revision Table: Work, Force, and Distance

Concept	Symbol	Unit (SI)	Formula
Work Done	W	Joule (J)	$W = F \times d$ (when force is constant and parallel to displacement)
Force	F	Newton (N)	$F = \frac{W}{d}$
Distance/Displacement	d	Meter (m)	$d = \frac{W}{F}$

Additional Information on Work and Force

It's important to remember a few key points about work:

- Work is a scalar quantity, meaning it only has magnitude, not direction.
- Work can be positive (force and displacement in the same direction), negative (force and displacement in opposite directions), or zero (force perpendicular to displacement, or no displacement).
- The formula $W = F \times d$ is a simplified version that applies when the force is constant and acts purely in the direction of the displacement. If the force and displacement are not parallel, or if the force is not constant, more advanced methods (like using dot products or integration) are needed to calculate the work done.
- In this specific problem, we assume the force is applied in the direction the trolley is pushed, making the calculation straightforward.

27. Answer: b

Explanation:

Henry Cavendish's Famous Discovery in 1766

The question asks about a significant discovery made by the English scientist Henry Cavendish in the year 1766. Cavendish was a brilliant natural philosopher known for his precise measurements and experiments.

In 1766, Henry Cavendish conducted a series of experiments studying "inflammable air," a gas produced when certain metals, like zinc or iron, react with acids. Although "inflammable air" had been observed before, Cavendish was the first to collect it, study its properties systematically, and recognize it as a distinct substance different from ordinary air or other known gases.

He noted that this "inflammable air" was much less dense than air and that it burned with a pale flame, producing water. This observation that burning "inflammable air" produced water was a crucial finding, although the full significance of this reaction involving oxygen (from air) was later explained by Antoine Lavoisier.

Cavendish is credited with identifying "inflammable air" as a unique element, which was later named "hydrogen" by Lavoisier, meaning "water-former." Therefore, his key discovery in 1766 was hydrogen.

Analyzing the Options and Discoveries

Let's look at the other options provided and their historical context regarding discovery dates and discoverers:

- **Oxygen:** Discovered independently by Carl Wilhelm Scheele around 1772 (published 1777) and Joseph Priestley in 1774.
- **Hydrogen:** Isolated and characterized as a distinct element by Henry Cavendish in 1766.
- **Helium:** First detected spectrographically in the Sun in 1868 by Pierre Janssen and Norman Lockyer, and isolated on Earth in 1895 by William Ramsay.
- **Chlorine:** Discovered by Carl Wilhelm Scheele in 1774.

Comparing the dates, only hydrogen was discovered by Henry Cavendish in 1766, matching the information in the question.

Conclusion on Cavendish's 1766 Discovery

Based on historical records of scientific discoveries, Henry Cavendish's major contribution in 1766 was his work on "inflammable air," which led to the identification of hydrogen as a unique chemical element.

Substance	Discoverer(s)	Year of Discovery/Identification
Oxygen	C.W. Scheele, J. Priestley	~1772, 1774
Hydrogen	Henry Cavendish	1766
Helium	P. Janssen, N. Lockyer (Sun); W. Ramsay (Earth)	1868 (Sun), 1895 (Earth)
Chlorine	C.W. Scheele	1774

Revision Table: Key Facts

Scientist	Key Discovery in 1766	Significance
Henry Cavendish	Hydrogen ("inflammable air")	Identified it as a distinct element, studied its properties (low density, combustion producing water).

Additional Information about Henry Cavendish and Hydrogen

Henry Cavendish was a reclusive but highly influential scientist. Besides his work on gases, he is also famous for the Cavendish experiment in 1798, where he measured the density of the Earth. This experiment effectively determined the gravitational constant 'G'.

Hydrogen (H_2) is the lightest element on the periodic table. It is the most abundant chemical substance in the universe, found primarily in stars and gas giant planets.

On Earth, it is rarely found in its pure elemental form and is usually combined with other elements, most notably in water (H_2O).

28. Answer: b

Explanation:

Calculating Heat Capacity of a Steel Vessel

The question asks us to find the heat capacity of a steel vessel given its mass, the specific heat capacity of steel, and a temperature change. While the temperature rise is given (10 degrees), it is actually not needed to calculate the heat capacity itself. The heat capacity is a property of the object (the steel vessel) and depends on its mass and the material's specific heat capacity.

Understanding Heat Capacity and Specific Heat Capacity

Let's first understand the key terms:

- **Specific Heat Capacity (c):** This is an intrinsic property of a substance. It is defined as the amount of heat energy required to raise the temperature of 1 kilogram of that substance by 1 degree Celsius or 1 Kelvin. Its standard unit is Joules per kilogram per Kelvin ($\text{Jkg}^{-1}\text{K}^{-1}$) or Joules per kilogram per degree Celsius ($\text{Jkg}^{-1}\text{C}^{-1}$). For steel, it's given as $500\text{Jkg}^{-1}\text{K}^{-1}$.
- **Heat Capacity (C):** This is a property of a specific object. It is defined as the amount of heat energy required to raise the temperature of the entire object by 1 degree Celsius or 1 Kelvin. Its standard unit is Joules per Kelvin (JK^{-1}) or Joules per degree Celsius (JC^{-1}).

The heat capacity (C) of an object is related to its mass (m) and the specific heat capacity (c) of the material it is made from by the formula:

$$C = m \times c$$

This formula tells us that a more massive object made of the same material will have a higher heat capacity, meaning it requires more energy to change its temperature by the same amount.

Applying the Formula to the Steel Vessel

We are given the following information for the steel vessel:

- Mass of the steel vessel, $m = 2.5 \text{ kg}$
- Specific heat capacity of steel, $c = 500 \text{ Jkg}^{-1}\text{K}^{-1}$

Using the formula $C = m \times c$, we can calculate the heat capacity of the steel vessel:

$$C = (2.5 \text{ kg}) \times (500 \text{ Jkg}^{-1}\text{K}^{-1})$$

Let's perform the multiplication:

$$C = 2.5 \times 500 \text{ JK}^{-1}$$

$$C = 1250 \text{ JK}^{-1}$$

So, the heat capacity of the steel vessel is 1250 JK^{-1} .

Analyzing the Options

Let's compare our calculated value with the given options:

- Option 1: $200 \text{ Jkg}^{-1}\text{K}^{-1}$ – This unit ($\text{Jkg}^{-1}\text{K}^{-1}$) is for specific heat capacity, not heat capacity (JK^{-1}). Also, the value is incorrect.
- Option 2: $1,250 \text{ JK}^{-1}$ – This matches our calculated value and has the correct unit for heat capacity.
- Option 3: 125 JK^{-1} – This value is incorrect. It seems like a calculation error might have occurred (e.g., $2.5 \times 50 = 125$).
- Option 4: $20 \text{ Jkg}^{-1}\text{K}^{-1}$ – This unit ($\text{Jkg}^{-1}\text{K}^{-1}$) is for specific heat capacity, not heat capacity (JK^{-1}). Also, the value is incorrect.

Based on our calculation, the correct heat capacity of the steel vessel is 1250 JK^{-1} .

Conclusion

The heat capacity of the steel vessel of mass 2.5 kg with a specific heat capacity of $500 \text{ Jkg}^{-1}\text{K}^{-1}$ is 1250JK^{-1} .

Revision Table: Heat Capacity Calculations

Concept	Symbol	Definition	Unit	Formula
Specific Heat Capacity	c	Heat per unit mass per degree temp change	$\text{Jkg}^{-1}\text{K}^{-1}$ or $\text{Jkg}^{-1}\text{C}^{-1}$	$c = \frac{Q}{m\Delta T}$
Heat Capacity	C	Heat per degree temp change for an object	JK^{-1} or JC^{-1}	$C = \frac{Q}{\Delta T}$ or $C = m \times c$
Heat Energy Transfer	Q	Amount of heat energy transferred	J (Joules)	$Q = mc\Delta T$ or $Q = C\Delta T$

Additional Information: Factors Affecting Heat Capacity

The heat capacity of an object is an extensive property, meaning it depends on the amount of substance present. Here are some factors that influence the heat capacity of an object:

- **Mass (m):** As shown in the formula $C = m \times c$, heat capacity is directly proportional to mass. A larger mass requires more heat to change its temperature by the same amount.
- **Material (Specific Heat Capacity, c):** Different materials have different specific heat capacities. For instance, water has a very high specific heat capacity compared to metals like steel. This means water can absorb a lot more heat energy than an equal mass of steel for the same temperature rise.
- **Phase of the Material:** The specific heat capacity of a substance can change depending on whether it is in a solid, liquid, or gaseous state. For example, the specific heat capacity of water is different from that of ice or steam.

Understanding heat capacity is crucial in various applications, from designing heating and cooling systems to studying climate patterns, where the high heat capacity of oceans plays a significant role in stabilizing Earth's temperature.

29. Answer: b

Explanation:

Given:

A can paint = in 50 days, B can paint = in 10 days, A + B + C can paint = in 6.25 days

Assuming total work is 100 units

Assessment:

A's efficiency = $100 \text{ units} / 50 \text{ days} = 2 \text{ units/day}$

B's efficiency = $100 \text{ units} / 10 \text{ days} = 10 \text{ units/day}$

A + B + C's efficiency = $100 \text{ units} / 6.25 \text{ days} = 16 \text{ units/day}$

If A + B work together,

then A + B's efficiency = $(2 + 10) \text{ units/day} = 12 \text{ units/day}$

Now, subtract A + B's efficiency from A + B + C's efficiency

$\Rightarrow 16 \text{ units/day} - 12 \text{ units/day} = 4 \text{ units/day}$

$\therefore$ C's efficiency = 4 units/day

C can paint the wall alone = $100 / 4 = 25 \text{ days}$

30. Answer: d

Explanation:

Understanding Paracetamol Use in First-Aid Kits

Paracetamol, also known as acetaminophen, is a very common medicine found in many household first-aid boxes. It is a popular choice because it helps with some common minor ailments that people experience.

The question asks when and why this drug should be taken. Let's look at the typical uses of Paracetamol to understand its role in a first-aid situation.

Why Take Paracetamol?

Paracetamol is primarily known for two main actions:

- **Pain Relief:** It helps to ease mild to moderate pain. This can include headaches, muscle aches, period pain, toothaches, and other common types of pain.
- **Fever Reduction:** It helps to lower a high body temperature, also known as fever. Fever is often a symptom of infections like the flu or a cold.

Because mild pain and fever are common symptoms that can occur unexpectedly, keeping Paracetamol in a first-aid kit allows for quick access to relief.

Analyzing the Options

Let's consider why other options are not the correct uses for Paracetamol:

- **To bring relief from asthma:** Asthma is a condition affecting the airways, requiring specific inhalers or other medications to open them up and ease breathing. Paracetamol does not treat asthma symptoms.
- **To ease the symptoms of hay fever and other allergies:** Allergies like hay fever are typically treated with antihistamines or other anti-allergy medications that block the body's reaction to allergens. Paracetamol does not have antihistamine properties and is not used for allergy symptoms.
- **To ease indigestion and heartburn:** Indigestion and heartburn are issues related to stomach acid and digestion. Medications like antacids or proton pump inhibitors are used to treat these conditions. Paracetamol does not affect stomach acid or digestion and is not used for indigestion or heartburn.
- **To ease mild pain and reduce high temperature (fever):** As discussed, this is the primary and correct use of Paracetamol. It effectively treats mild to

moderate pain and helps to lower fever, making it a valuable component of a first-aid kit for these common issues.

Mechanism of Action (Simple Terms)

While the exact way Paracetamol works is still being researched fully, it's believed to work mainly in the brain and spinal cord (central nervous system) to block the production of certain chemicals called prostaglandins, which are involved in causing pain and fever. It is different from NSAIDs (like ibuprofen or aspirin) which work more broadly throughout the body and also reduce inflammation.

Summary of Paracetamol Use

In summary, Paracetamol is a go-to medication in first-aid situations for managing mild to moderate pain and reducing fever. It addresses common symptoms encountered outside of severe medical emergencies.

Symptom	Is Paracetamol Used?	Alternative/Primary Treatment
Mild Pain (Headache, Muscle Ache, etc.)	Yes	NSAIDs (like Ibuprofen), Topical Pain Relievers
High Temperature (Fever)	Yes	NSAIDs, Hydration, Rest
Asthma Symptoms	No	Bronchodilators, Steroids
Allergy Symptoms (Hay Fever)	No	Antihistamines, Decongestants
Indigestion/Heartburn	No	Antacids, PPIs

Revision Table: Paracetamol in First-Aid

Here's a quick review of key points about Paracetamol's role in a first-aid kit:

- **What it is:** A common over-the-counter pain reliever and fever reducer.
- **Primary Uses:** Mild to moderate pain, high temperature (fever).

- **Where it's found:** Frequently in first-aid boxes and medicine cabinets.
- **Not for:** Asthma, allergies, indigestion, inflammation (unlike NSAIDs).

Additional Information on Paracetamol

It's important to always follow the dosage instructions on the packaging or provided by a healthcare professional. Taking too much Paracetamol can be harmful, especially to the liver. It's also important to check if other medications being taken already contain Paracetamol to avoid accidental overdose. If symptoms persist or worsen, medical advice should be sought.

31. Answer: b

Explanation:

This problem is a type of language puzzle where we need to decode the meaning of words in an artificial language by comparing given phrases and their translations. We are given three phrases:

- 'terake' means 'motherland'
- 'lorake' means 'mother tongue'
- 'loroli' means 'long tongue'

Our goal is to find the word that means 'long jump' in this artificial language.

Decoding Word Meanings in the Artificial Language

Let's compare the given phrases to find common parts and their corresponding meanings.

Compare 'terake' (motherland) and 'lorake' (mother tongue):

- Both phrases contain the word 'mother'.
- The artificial words are 'terake' and 'lorake'. They share the syllable 'ake'.
- This suggests that 'ake' means 'mother'.

Compare 'lorake' (mother tongue) and 'loroli' (long tongue):

- Both phrases contain the word 'tongue'.
- The artificial words are 'lorake' and 'loroli'. They share the syllable 'loro'.
- This suggests that 'loro' means 'tongue'.

Now we can use these findings to determine the meanings of other parts.

- In 'loroli' (long tongue), we found 'loro' means 'tongue'. The remaining part is 'oli', and the remaining English word is 'long'.
- This suggests that 'oli' means 'long'.

Let's summarise the words we have decoded so far:

Artificial Word Part	Meaning
ake	mother
loro	tongue
oli	long

We can also infer the meaning of 'tera' from 'terake' (motherland). If 'ake' means 'mother', then 'tera' must mean 'land'.

Determining the Structure and Finding 'long jump'

Let's look at the structure of the artificial words in relation to the English phrases:

- 'terake' = 'tera' + 'ake' which is potentially 'land' + 'mother' (motherland).
- 'lorake' = 'loro' + 'ake' which is potentially 'tongue' + 'mother' (mother tongue).
- 'loroli' = 'loro' + 'oli' which is potentially 'tongue' + 'long' (long tongue).

The pattern observed in 'loroli' (long tongue) is Noun + Adjective: 'tongue' (loro) followed by 'long' (oli). This suggests the artificial language might place the modifier (like an adjective) after the noun.

We need the word for 'long jump'. Applying the Noun + Adjective structure, this should be 'jump' + 'long'.

We know 'long' is 'oli'. So, the word should be [word for jump] + oli. This means the word for 'long jump' must end with 'oli'.

Evaluating the Options

Let's look at the given options:

1. akezit
2. neroli
3. hudter
4. dibnol

We are looking for a word that ends with 'oli'. Only option 2, 'neroli', fits this pattern.

Following the Noun + Adjective structure, if 'neroli' means 'long jump', then 'nero' must be the word for 'jump'.

Thus, 'neroli' is the word for 'long jump' in this artificial language.

Summary of Decoding

Your Personal Exams Guide

Artificial Word	Structure	Meaning
ake	-	mother
loro	-	tongue
oli	-	long
tera	-	land
terake	tera + ake	land mother (motherland)
lorake	loro + ake	tongue mother (mother tongue)
loroli	loro + oli	tongue long (long tongue)
neroli	nero + oli	jump long (long jump)

Revision Table: Artificial Language Words

Here's a table summarizing the derived meanings in this artificial language for quick review:

Artificial Word Part	Derived Meaning
ake	mother
loro	tongue
oli	long
tera	land
nero	jump

Additional Information: Solving Language Puzzles

Language decoding puzzles like this one test your logical reasoning and pattern recognition skills. Here are some tips for approaching them:

- **Look for common words:** Identify words that appear in multiple phrases. The corresponding artificial word parts are likely related to that common meaning.
- **Identify unique words:** Once common parts are found, the remaining unique parts in a phrase can be matched to the unique words in its translation.
- **Observe the structure:** Pay attention to the order of the artificial word parts. Does it follow the English word order (e.g., adjective + noun) or a different order (e.g., noun + adjective)? Apply this structure when translating new phrases.
- **Eliminate options:** Use the decoded meanings and the structure to eliminate options that do not fit the pattern or contain incorrect word parts.

These puzzles are similar to cracking simple codes or understanding grammatical rules of an unknown language.

32. Answer: c

Explanation:

Binary to Decimal Conversion of 101110110

Converting a binary number to a decimal number involves multiplying each digit by the corresponding power of 2 and summing the results. The powers of 2 start from 2^0 for the rightmost digit and increase by one for each position to the left.

Let's convert the binary number 101110110 to its decimal equivalent.

The binary number is 101110110. We can write out the digits and their positions, starting from position 0 on the right:

- The rightmost digit (0) is at position 0, corresponding to 2^0 .
- The next digit (1) is at position 1, corresponding to 2^1 .
- The next digit (1) is at position 2, corresponding to 2^2 .
- The next digit (0) is at position 3, corresponding to 2^3 .

- The next digit (1) is at position 4, corresponding to 2^4 .
- The next digit (1) is at position 5, corresponding to 2^5 .
- The next digit (1) is at position 6, corresponding to 2^6 .
- The next digit (0) is at position 7, corresponding to 2^7 .
- The leftmost digit (1) is at position 8, corresponding to 2^8 .

Now, we multiply each binary digit by the value of the power of 2 for its position:

Binary Digit	Position	Power of 2	Calculation	Result
1	8	$2^8 = 256$	1×256	256
0	7	$2^7 = 128$	0×128	0
1	6	$2^6 = 64$	1×64	64
1	5	$2^5 = 32$	1×32	32
1	4	$2^4 = 16$	1×16	16
0	3	$2^3 = 8$	0×8	0
1	2	$2^2 = 4$	1×4	4
1	1	$2^1 = 2$	1×2	2
0	0	$2^0 = 1$	0×1	0

Finally, we sum the results from the last column:

$$\text{Decimal value} = 256 + 0 + 64 + 32 + 16 + 0 + 4 + 2 + 0$$

$$\text{Decimal value} = 374$$

Therefore, the binary number 101110110 is equal to the decimal number 374.

Revision Table: Binary Number Conversion

Concept	Description	Example
Binary System	Base-2 number system using only digits 0 and 1.	101_2
Decimal System	Base-10 number system using digits 0 through 9.	5_{10}
Binary to Decimal Conversion	Sum of (digit * 2^{position}) for each digit.	$101_2 = 1 \times 2^2 + 0 \times 2^1 + 1 \times 2^0 = 4 + 0 + 1 = 5_{10}$

Additional Information on Number Systems

Number systems are ways of representing numbers. The most common system we use daily is the decimal system (base 10). Computers and digital electronics primarily use the binary system (base 2).

- **Base:** The base of a number system indicates the number of unique digits used (including zero) and is the value of the base raised to the power of the position. For example, the decimal system has a base of 10.
- **Place Value:** In any positional number system, the position of a digit determines its value. This value is the digit multiplied by the base raised to the power of the position index. For example, in 123_{10} , the '2' is in the tens place (10^1), so its value is $2 \times 10^1 = 20$. In 101_2 , the middle '0' is in the 2^1 place, so its value is $0 \times 2^1 = 0$.
- **Other Systems:** Besides binary and decimal, other number systems like Octal (base 8) and Hexadecimal (base 16) are also used, particularly in computing and digital electronics. Octal uses digits 0-7, while Hexadecimal uses digits 0-9 and letters A-F (representing 10-15).

33. Answer: d

Explanation:

Logical Reasoning: Analyzing Statements and Conclusions

This question requires us to analyze two given statements and determine which of the following conclusions logically follow from them. We must consider the statements as true, even if they contradict commonly known facts.

Understanding the Statements

Let's break down the two statements provided:

- **Statement 1:** No beakers are vases.
- **Statement 2:** All vases are jugs.

These statements describe relationships between three categories: beakers, vases, and jugs. Statement 1 tells us that the category of beakers and the category of vases are mutually exclusive; they have no items in common. Statement 2 tells us that the category of vases is a subset of the category of jugs; every item that is a vase is also a jug.

Visualizing the Relationships (Conceptual Venn Diagram)

We can imagine this relationship using sets or areas, similar to a Venn diagram:

- The set of vases is completely inside the set of jugs.
- The set of beakers is completely separate from the set of vases.

Based on this, we know that beakers are separate from the items within the 'Vases' area. However, the statements don't explicitly define the relationship between beakers and the entire 'Jugs' area (specifically, the part of 'Jugs' that is *not* 'Vases').

Evaluating the Conclusions

Now let's evaluate each conclusion based on the statements:

Conclusion I: Some jugs are beakers.

This conclusion suggests there is an overlap between the set of jugs and the set of beakers. From our statements, we know beakers and vases are separate, and all vases are jugs. This means beakers are separate from the part of jugs that consists of vases.

However, beakers could potentially overlap with the part of the jugs category that is **not** vases. For example, some items could be jugs but not vases, and some beakers could also be these types of jugs. Alternatively, beakers might be completely outside the entire set of jugs.

Since the conclusion "Some jugs are beakers" is not guaranteed to be true in all possible scenarios consistent with the statements (we can draw a diagram where beakers are outside jugs), this conclusion does not logically follow.

Conclusion II: Some jugs are vases.

This conclusion suggests there is an overlap between the set of jugs and the set of vases. Statement 2 explicitly says, "All vases are jugs."

If all vases are jugs, it means that the set of vases is contained within the set of jugs. Assuming there is at least one vase (standard assumption in syllogism unless stated otherwise), then that vase is also a jug. Therefore, there exists at least one item that is both a jug and a vase. This means "Some jugs are vases" is necessarily true if "All vases are jugs" is true and the set of vases is not empty.

This conclusion logically follows directly from Statement 2.

Summary of Conclusions

Based on our analysis:

- Conclusion I (Some jugs are beakers) does not necessarily follow from the statements.
- Conclusion II (Some jugs are vases) necessarily follows from the statements.

Statement/Conclusion	Relationship Described	Follows?
Statement 1: No beakers are vases.	$\text{Beakers} \cap \text{Vases} = \emptyset$	-
Statement 2: All vases are jugs.	$\text{Vases} \subseteq \text{Jugs}$	-
Conclusion I: Some jugs are beakers.	$\text{Jugs} \cap \text{Beakers} \neq \emptyset$	No
Conclusion II: Some jugs are vases.	$\text{Jugs} \cap \text{Vases} \neq \emptyset$ (implied by $\text{Vases} \subseteq \text{Jugs}$ if $\text{Vases} \neq \emptyset$)	Yes

Therefore, only conclusion II follows from the given statements.

Revision Table: Syllogism Basics

Type of Statement	Structure	Example	Implied Conclusions (if subject $\neq \emptyset$)
A (Universal Affirmative)	All S are P	All cats are mammals.	Some P are S
E (Universal Negative)	No S are P	No dogs are cats.	No P are S, Some S are not P, Some P are not S
I (Particular Affirmative)	Some S are P	Some students are athletes.	Some P are S
O (Particular Negative)	Some S are not P	Some birds are not eagles.	None directly imply reversal ($P \rightarrow S$)

In this question, Statement 2 ("All vases are jugs") is an A-type statement. It directly implies "Some jugs are vases," which is Conclusion II. Statement 1 ("No beakers are vases") is an E-type statement, describing mutual exclusion.

Additional Information: Logical Reasoning and Syllogisms

Logical reasoning questions, particularly those involving syllogisms, test your ability to deduce conclusions from given premises (statements). The key is to follow the rules of logic strictly and not bring in outside information or common sense that contradicts the statements.

Common methods for solving syllogisms include:

- **Venn Diagrams:** Drawing overlapping or separate circles to represent the sets described in the statements. This provides a visual way to see which areas must contain elements and which must be empty.
- **Rules of Inference:** Applying formal logical rules to derive conclusions directly from the statement structure.
- **Checking Counter-Examples:** Trying to construct a scenario where the statements are true but the conclusion is false. If you can construct such a scenario, the conclusion does not follow.

In this case, the relationship between Vases and Jugs (All Vases are Jugs) makes Conclusion II (Some jugs are vases) a direct implication, assuming vases exist. The relationship between Beakers and Vases (No Beakers are Vases), combined with Vases being inside Jugs, leaves the relationship between Beakers and the rest of Jugs undefined, which is why Conclusion I doesn't necessarily follow.

34. Answer: b

Explanation:

B	Brackets in order {}, {}, []	ब्रैकेट {}, {}, [] क्रम में
O	of	का
D	Division (÷)	विभाजन (÷)
M	Multiplication (×)	गुणा (×)
A	Addition (+)	जोड़ (+)
S	Subtraction (−)	घटाव (−)

Calculation:

Start with the innermost bracket

$$(14 \times 10 \div 35 + 12 \times 10 \div 15)$$

$$\Rightarrow \{(14 \times 2/7) + (12 \times 2/3)\}$$

12

Now,

$$\Rightarrow 16 - (25\% \text{ of } 12)$$

$$\Rightarrow 16 - \{(25/100) \times 12\}$$

$$\Rightarrow 16 - 3$$

∴ The value of the expression is 13.

35. Answer: a

Explanation:

Understanding Levers and Their Classes

Levers are simple machines that help us multiply force or change the direction of force. They consist of a rigid bar that pivots around a fixed point called the **fulcrum**. Three forces are involved in the operation of a lever:

- **Effort (E):** The force applied to the lever.
- **Load (L):** The resistance or weight that is being moved or lifted.
- **Fulcrum (F):** The pivot point around which the lever rotates.

Levers are classified into three types, or classes, based on the relative positions of the fulcrum, load, and effort.

Types of Levers: First, Second, and Third Class

Here's a breakdown of the three classes of levers:

- **First Class Lever:** In a first-class lever, the **fulcrum (F)** is located **between** the effort (E) and the load (L). Examples include a see-saw, a crowbar, and pliers. The mechanical advantage can be greater than, equal to, or less than 1, depending on the distances of the effort and load from the fulcrum.
- **Second Class Lever:** In a second-class lever, the **load (L)** is located **between** the fulcrum (F) and the effort (E). Examples include a wheelbarrow, a nutcracker, and a door (hinges are the fulcrum, the doorknob is where effort is applied, and the resistance from closing is the load). Second-class levers always provide a mechanical advantage greater than 1 ($MA > 1$), meaning the effort required is less than the load, but the effort must be moved over a greater distance.
- **Third Class Lever:** In a third-class lever, the **effort (E)** is located **between** the fulcrum (F) and the load (L). Examples include tweezers, ice tongs, a fishing rod, and the human arm. Third-class levers always have a mechanical advantage less than 1 ($MA < 1$). They are used to increase the speed or distance of movement at the cost of requiring greater effort.

Analyzing the Given Options for Second Class Lever

Let's examine each option to determine which one is an example of a second-class lever:

- **Wheelbarrow:** In a wheelbarrow, the wheel acts as the **fulcrum**. The load (contents in the bin) is positioned **between** the wheel (fulcrum) and where you lift the handles (effort). This arrangement perfectly matches the definition of a **second-class lever**.

- **Pliers:** Pliers have the pivot point (where the two handles cross) as the **fulcrum**. The effort is applied on the handles, and the load is applied on the jaws. The fulcrum is between the effort and the load, making pliers a **first-class lever**.
- **See-saw:** A see-saw pivots on a central support, which is the **fulcrum**. One person's weight acts as the load, and the other person's push/weight acts as the effort, with the fulcrum in between. This is a classic example of a **first-class lever**.
- **Ice tongs:** In ice tongs, the pivot point at the connected end is the **fulcrum**. The effort is applied in the middle of the arms, and the load (the ice) is held at the tips. The effort is applied between the fulcrum and the load, classifying ice tongs as a **third-class lever**.

Based on this analysis, the wheelbarrow is the only example among the options that fits the description of a second-class lever, where the load is situated between the fulcrum and the effort.

Option	Fulcrum (F)	Load (L)	Effort (E)	Arrangement	Lever Class
Wheelbarrow	Wheel	Contents in bin	Lifting handles	F - L - E	Second Class
Pliers	Pivot point	On jaws	On handles	E - F - L	First Class
See-saw	Center pivot	One person's weight	Other person's weight	L - F - E	First Class
Ice tongs	Connected end	On tips (ice)	Middle of arms	F - E - L	Third Class

Conclusion on Second Class Lever Example

Therefore, the wheelbarrow is a correct example of a second-class lever because its design places the load (material being carried) between the fulcrum (the wheel)

and the point where the effort is applied (lifting the handles). This configuration provides a mechanical advantage, making it easier to lift and move heavy loads.

Revision Table: Lever Classes

Lever Class	Arrangement (F=Fulcrum, L=Load, E=Effort)	Mechanical Advantage (MA)	Common Use	Examples
First Class	F is between L and E (L-F-E or E-F-L)	>1 , <1 , or $=1$	Changing direction or multiplying force/distance	See-saw, crowbar, pliers, scissors
Second Class	L is between F and E (F-L-E)	>1	Multiplying force (reducing effort)	Wheelbarrow, nutcracker, bottle opener, door
Third Class	E is between F and L (F-E-L)	<1	Increasing speed or distance of movement	Tweezers, ice tongs, fishing rod, human arm, broom

Your Personal Exams Guide

Additional Information: Simple Machines and Levers

Levers are fundamental simple machines. Simple machines are basic devices that change the direction or magnitude of a force. Besides levers, other simple machines include the wheel and axle, pulley, inclined plane, wedge, and screw. Understanding levers and their classes is important for understanding how many tools and everyday objects work, and how they provide mechanical advantage or facilitate movement.

The principle behind levers is the principle of moments or torques. The lever is in equilibrium when the sum of clockwise moments about the fulcrum equals the sum of anticlockwise moments about the fulcrum. Moment is calculated as force

multiplied by the perpendicular distance from the fulcrum to the line of action of the force.

The mechanical advantage of a lever is the ratio of the load to the effort ($MA = \frac{\text{Load}}{\text{Effort}}$). It can also be calculated based on the distances of the effort arm (d_E) and load arm (d_L) from the fulcrum: $MA = \frac{d_E}{d_L}$. For a second-class lever, the effort arm is always longer than the load arm ($d_E > d_L$), resulting in a mechanical advantage greater than 1.

36. Answer: c

Explanation:

Calculating Equilibrium Temperature When Mixing Water

This problem involves the mixing of two quantities of water at different temperatures. When substances at different temperatures are mixed in an insulated system (where no heat is lost or gained from the surroundings), heat energy transfers from the hotter substance to the colder substance until they reach a common final equilibrium temperature. The fundamental principle governing this process is that the heat lost by the hot substance equals the heat gained by the cold substance.

Understanding the Principle of Heat Exchange

The amount of heat energy (Q) transferred is given by the formula:

$$Q = mc\Delta T$$

Where:

- m is the mass of the substance.
- c is the specific heat capacity of the substance.

- ΔT is the change in temperature ($T_{final} - T_{initial}$). For heat loss, ΔT is often written as $T_{initial} - T_{final}$.

In this scenario, we are mixing hot water and cold water. Assuming the specific heat capacity (c) of water is constant over the given temperature range and the density (ρ) of water is also constant (approximately 1 kg/L), we can determine the mass of each quantity of water from its volume.

Given Information:

- Volume of hot water (V_1) = 2 litres
- Initial temperature of hot water (T_1) = 190°C
- Volume of cold water (V_2) = 6 litres
- Initial temperature of cold water (T_2) = 20°C
- No heat is lost to the surroundings (adiabatic mixing).

Assuming density of water $\rho \approx 1 \text{ kg/L}$:

- Mass of hot water (m_1) = $V_1 \times \rho = 2 \text{ L} \times 1 \text{ kg/L} = 2 \text{ kg}$
- Mass of cold water (m_2) = $V_2 \times \rho = 6 \text{ L} \times 1 \text{ kg/L} = 6 \text{ kg}$

Setting up the Heat Balance Equation

According to the principle of heat exchange in an isolated system:

Heat lost by hot water = Heat gained by cold water

$$m_1 c (T_1 - T_f) = m_2 c (T_f - T_2)$$

Where T_f is the final equilibrium temperature we need to find.

Since the specific heat capacity (c) is the same for both (water), we can cancel it from both sides of the equation:

$$m_1 (T_1 - T_f) = m_2 (T_f - T_2)$$

Solving for the Final Temperature

Now, we substitute the known values into the equation:

$$2 \text{ kg} \times (190^\circ\text{C} - T_f) = 6 \text{ kg} \times (T_f - 20^\circ\text{C})$$

Expand both sides of the equation:

$$2 \times 190 - 2 \times T_f = 6 \times T_f - 6 \times 20$$

$$380 - 2T_f = 6T_f - 120$$

Now, rearrange the terms to group T_f on one side and constants on the other:

$$380 + 120 = 6T_f + 2T_f$$

$$500 = 8T_f$$

Finally, solve for T_f :

$$T_f = \frac{500}{8}$$

$$T_f = 62.5^\circ\text{C}$$

Thus, the final equilibrium temperature after mixing the superheated water at 190°C and the cold water at 20°C is 62.5°C .

Summary of Calculation Steps

Step	Description	Equation/Calculation
1	Identify masses from volumes (assuming density 1 kg/L)	$m_1 = 2 \text{ kg}, m_2 = 6 \text{ kg}$
2	Set up heat balance equation (Heat lost = Heat gained)	$m_1(T_1 - T_f) = m_2(T_f - T_2)$ (after cancelling c)
3	Substitute given values	$2(190 - T_f) = 6(T_f - 20)$
4	Expand and rearrange equation	$380 - 2T_f = 6T_f - 120 \implies 500 = 8T_f$
5	Solve for T_f	$T_f = \frac{500}{8} = 62.5^\circ\text{C}$

Revision Table: Key Concepts in Mixing

Concept	Explanation
Heat Transfer	Energy transferred between systems due to temperature difference.
Specific Heat Capacity (c)	Amount of heat required to raise the temperature of 1 kg of a substance by 1°C. For water, it's approximately 4186 J/(kg·°C).
Heat Exchange Principle	In an isolated system, net heat transfer is zero: Heat lost by hot objects = Heat gained by cold objects.
Equilibrium Temperature	The final uniform temperature reached by all parts of a system after heat transfer ceases.
Adiabatic Process	A process where no heat is exchanged with the surroundings. The mixing here is assumed adiabatic.

Additional Information: Assumptions and Considerations

Several assumptions were made in solving this heat mixing problem:

- **Constant Specific Heat Capacity:** We assumed that the specific heat capacity of water remains constant over the temperature range from 20°C to 190°C. In reality, it varies slightly with temperature, but this assumption is common for simplified calculations.
- **Constant Density:** We assumed the density of water is constant at 1 kg/L. Density also varies slightly with temperature and pressure, but this is a reasonable approximation for this problem.
- **No Phase Change:** The term "superheated water at 190°C" indicates liquid water at a pressure above its saturation pressure at 190°C. The problem implicitly assumes that no boiling or phase change occurs during or after mixing, which would involve latent heat transfer and a more complex calculation. The final temperature (62.5°C) is well below the boiling point at standard pressure, confirming no boiling in the final state.
- **Perfect Mixing:** It is assumed that the two volumes of water mix completely and instantly to reach a uniform final temperature.

- **Isolated System:** The problem states "no heat is lost," meaning the system (the mixed water) is isolated from the surroundings, preventing any heat transfer to or from anything else.

37. Answer: a

Explanation:

Understanding the Unit Henry per Meter

The question asks to identify the physical quantity that has Henry per meter as its unit. Let's analyze the units of the given options to determine the correct answer.

What is Henry per Meter?

The Henry (H) is the unit of inductance. It is named after Joseph Henry. Inductance is the property of an electric conductor or circuit that causes an electromagnetic force to be generated in the conductor when the electric current changes.

The unit Henry per meter (H/m) suggests a property that is distributed over space, specifically per unit length or related to the material properties of a medium rather than a specific component.

Analyzing the Options and Their Units

Let's examine the units of each option provided:

1. permeability
2. watt per steradian
3. permittivity
4. electric conductance

1. Permeability (μ)

Permeability is a material property that describes the ability of a material to support the formation of a magnetic field within itself. It is defined as the ratio of magnetic flux density (**B**) to magnetic field strength (**H**).

The magnetic flux density **B** is measured in Tesla (T), and the magnetic field strength **H** is measured in Ampere per meter (A/m).

So, the unit of permeability (μ) is $\frac{\text{Tesla}}{\text{Ampere/meter}} = \frac{\text{T}}{\text{A/m}}$.

We also know that 1 Tesla is equal to 1 Weber per square meter ($\text{T} = \text{Wb/m}^2$), and 1 Henry is defined as 1 Weber per Ampere ($\text{H} = \text{Wb/A}$).

Let's express the unit of permeability using these relationships:

$$\text{Unit of } \mu = \frac{\text{T}}{\text{A/m}} = \frac{\text{Wb/m}^2}{\text{A/m}} = \frac{\text{Wb}}{\text{m}^2} \times \frac{\text{m}}{\text{A}} = \frac{\text{Wb}}{\text{Am}}.$$

Since $\text{H} = \text{Wb/A}$, we can substitute $\text{Wb} = \text{H} \times \text{A}$ into the unit expression:

$$\text{Unit of } \mu = \frac{\text{H} \times \text{A}}{\text{A} \times \text{m}} = \frac{\text{H}}{\text{m}}.$$

Thus, the unit of permeability is indeed Henry per meter (H/m).

2. Watt per Steradian (W/sr)

This is the unit of radiant intensity, which is the power emitted by a source per unit solid angle. Watt (W) is the unit of power, and steradian (sr) is the unit of solid angle. This unit is related to optics and radiation, not magnetic properties or inductance.

3. Permittivity (ϵ)

Permittivity is a material property that describes the ability of a material to store electrical energy in an electric field. It is defined as the ratio of electric displacement field (**D**) to electric field strength (**E**).

The electric displacement field **D** is measured in Coulomb per square meter (C/m^2), and the electric field strength **E** is measured in Volt per meter (V/m).

So, the unit of permittivity (ϵ) is $\frac{\text{C/m}^2}{\text{V/m}} = \frac{\text{C}}{\text{m}^2} \times \frac{\text{m}}{\text{V}} = \frac{\text{C}}{\text{Vm}}$.

We know that 1 Farad (F) is equal to 1 Coulomb per Volt ($F = C/V$).

Substituting this into the unit expression:

$$\text{Unit of } \epsilon = \frac{C}{V \cdot m} = \frac{F \times V}{V \times m} = \frac{F}{m}.$$

Thus, the unit of permittivity is Farad per meter (F/m), not Henry per meter.

4. Electric Conductance (G)

Electric conductance is a measure of how easily electric current flows through a material. It is the reciprocal of electrical resistance (R). The unit of resistance is Ohm (Ω), so the unit of conductance is the Siemens (S), which is equivalent to Ω^{-1} .

This unit is related to the flow of current, not magnetic properties.

Summary of Units

Quantity	Symbol	Unit	Derived Unit
Permeability	μ	Henry per meter	H/m
Radiant Intensity	I_e	Watt per steradian	W/sr
Permittivity	ϵ	Farad per meter	F/m
Electric Conductance	G	Siemens	S (Ω^{-1})

Based on the unit analysis, Henry per meter (H/m) is the unit of permeability.

Conclusion

The unit Henry per meter (H/m) corresponds to the physical quantity known as permeability.

The final answer is **permeability**.

Revision Table: Units in Electromagnetism

It's useful to remember the units of common electromagnetic quantities:

- Inductance (L): Henry (H)
- Capacitance (C): Farad (F)
- Resistance (R): Ohm (Ω)
- Magnetic Flux (Φ_B): Weber (Wb)
- Magnetic Flux Density (B): Tesla (T)
- Magnetic Field Strength (H): Ampere per meter (A/m)
- Electric Field Strength (E): Volt per meter (V/m)
- Electric Displacement Field (D): Coulomb per square meter (C/m²)
- Permeability (μ): Henry per meter (H/m)
- Permittivity (ϵ): Farad per meter (F/m)

Additional Information on Permeability

Permeability (μ) is a fundamental property of a material that affects the magnetic field lines within it. It determines how a material responds to an applied magnetic field.

- The permeability of free space (vacuum) is denoted by μ_0 and has a value of approximately $4\pi \times 10^{-7}$ H/m.
- The relative permeability (μ_r) of a material is the ratio of its absolute permeability (μ) to the permeability of free space (μ_0): $\mu_r = \mu/\mu_0$. Relative permeability is a dimensionless quantity.
- Materials are classified based on their relative permeability:
 - Diamagnetic materials: $\mu_r \lesssim 1$ (slightly less than 1)
 - Paramagnetic materials: $\mu_r \gtrsim 1$ (slightly greater than 1)
 - Ferromagnetic materials: $\mu_r \gg 1$ (much greater than 1)
- Permeability plays a crucial role in determining the inductance of coils and the behavior of magnetic fields in various media, important for designing transformers, inductors, and other magnetic components.

38. Answer: c

Explanation:

Statement: Height should be between 5 feet 6 inches and 6 feet — one of the conditions for recruiting fighter pilots.

Assumption I: Fighter pilots should be presentable as they are role models for the young. → **Not implicit.** A height range is a physical criterion; it says nothing about appearance or public image.

Assumption II: Only certain body types can fit in a fighter plane cockpit. → **Implicit.** A fighter cockpit is a highly confined space, so pilots must fall within specific physical dimensions — height being a key one — to operate controls and ejection systems safely.

Hence, the correct answer is “Only assumption II is implicit”.

39. Answer: b

Explanation:

Gear Train Analysis: Calculating Driver Turns

This question involves understanding the relationship between the number of teeth on gears in a simple gear train and their speeds (or turns). In a gear train, the speed ratio is inversely proportional to the number of teeth.

Understanding the Gear Ratio

In a simple gear train with a driver gear and a follower gear, the ratio of their speeds (angular velocities or turns) is given by the inverse ratio of their number of teeth. Mathematically, this can be expressed as:

$$\frac{\text{Speed of Follower}}{\text{Speed of Driver}} = \frac{\text{Number of teeth on Driver}}{\text{Number of teeth on Follower}}$$

Or, in terms of turns:

$$\frac{\text{Turns of Follower}}{\text{Turns of Driver}} = \frac{T_{\text{driver}}}{T_{\text{follower}}}$$

Where:

- N_{follower} = Number of turns of the follower gear
- N_{driver} = Number of turns of the driver gear
- T_{driver} = Number of teeth on the driver gear
- T_{follower} = Number of teeth on the follower gear

Applying the Formula to the Given Gear Train

From the question, we are given:

- Number of teeth on the driver (T_{driver}) = 24
- Number of teeth on the follower (T_{follower}) = 8
- Number of turns of the follower (N_{follower}) = 36

We need to find the number of turns of the driver (N_{driver}).

Using the gear ratio formula:

$$\frac{N_{\text{follower}}}{N_{\text{driver}}} = \frac{T_{\text{driver}}}{T_{\text{follower}}}$$

Substitute the given values into the equation:

$$\frac{36}{N_{\text{driver}}} = \frac{24}{8}$$

Solving for the Number of Driver Turns

First, simplify the ratio of teeth:

$$\frac{24}{8} = 3$$

So, the equation becomes:

$$\frac{36}{N_{\text{driver}}} = 3$$

Now, solve for N_{driver} :

$$36 = 3 \times N_{\text{driver}}$$

$$N_{driver} = \frac{36}{3}$$

$$N_{driver} = 12$$

Therefore, for every 12 turns of the driver gear, the follower gear turns 36 times in this gear train.

Summary of Calculation

Parameter	Value
Driver Teeth (T_{driver})	24
Follower Teeth ($T_{follower}$)	8
Follower Turns ($N_{follower}$)	36
Gear Ratio ($\frac{T_{driver}}{T_{follower}}$)	$\frac{24}{8} = 3$
Relationship ($\frac{N_{follower}}{N_{driver}} = \text{Ratio}$)	$\frac{36}{N_{driver}} = 3$
Driver Turns (N_{driver})	$\frac{36}{3} = 12$

The calculation shows that the driver must complete 12 turns for the follower to complete 36 turns.

Revision Table: Gear Train Fundamentals

Concept	Description
Gear Train	A system formed by meshing gears to transmit rotational motion and torque.
Driver Gear	The gear that initiates the motion.
Follower Gear (Driven Gear)	The gear that receives motion from the driver.
Gear Ratio (Velocity Ratio)	The ratio of the output speed to the input speed. For simple gears, it's the inverse ratio of teeth numbers.
Speed and Teeth Relationship	Speed is inversely proportional to the number of teeth: More teeth mean slower speed, fewer teeth mean higher speed.

Additional Information: Types of Gear Trains and Applications

Gear trains are fundamental components in many machines. There are different types beyond the simple gear train discussed here:

- **Simple Gear Train:** Gears are mounted on separate shafts, and each shaft carries only one gear. As in the question, motion is transmitted from one gear to the next in a sequence.
- **Compound Gear Train:** At least one shaft carries two gears. This allows for larger speed reductions or increases in a compact space.
- **Reverted Gear Train:** A compound gear train where the axis of the last driven gear is co-axial with the axis of the first driving gear. Found in clocks and automotive transmissions.
- **Epicyclic Gear Train (Planetary Gear Train):** Gears revolve around a central gear (sun gear). Used in automatic transmissions, power tools, and aerospace applications for high torque density and compact size.

Understanding gear ratios is crucial for designing systems that require specific speed and torque transformations, from simple mechanisms to complex

machinery.

40. Answer: a

Explanation:

Solving Pipe and Tank Problems: Finding Emptying Time

This problem involves understanding the concept of work rates for pipes filling and emptying a tank. We are given the filling time for one pipe (Pipe A), the combined filling time when Pipe A and an emptying pipe (Pipe B) work together, and we need to find the emptying time for Pipe B.

The key concept is to represent the work done by each pipe as a rate, which is the fraction of the tank filled or emptied per unit of time (in this case, per hour).

Calculating Individual and Combined Rates

- **Rate of Pipe A (filling):** Pipe A fills the tank in 12 hours. Its filling rate is the amount of tank filled per hour.
Rate of A = $\frac{1}{\text{Time taken by A}} = \frac{1}{12}$ tank per hour.
- **Rate of Pipe B (emptying):** Pipe B empties the tank in X hours. Its emptying rate is the amount of tank emptied per hour.
Rate of B = $\frac{1}{\text{Time taken by B}} = \frac{1}{X}$ tank per hour.
- **Combined Rate (filling):** When both Pipe A (filling) and Pipe B (emptying) are opened together, the tank is filled in 30 hours. The combined rate is the net amount of tank filled per hour.
Combined Rate = $\frac{1}{\text{Time taken when A and B work together}} = \frac{1}{30}$ tank per hour.

Setting Up the Equation

When a filling pipe and an emptying pipe work together, the net rate is the filling rate minus the emptying rate. Since the tank is being filled overall, the filling rate (Pipe A) must be greater than the emptying rate (Pipe B).

Combined Rate = Rate of A – Rate of B

Substituting the rates we found:

$$\frac{1}{30} = \frac{1}{12} - \frac{1}{X}$$

Solving for X (Emptying Time)

Our goal is to find the value of X. We need to rearrange the equation to isolate $\frac{1}{X}$.

Add $\frac{1}{X}$ to both sides and subtract $\frac{1}{30}$ from both sides:

$$\frac{1}{X} = \frac{1}{12} - \frac{1}{30}$$

To subtract the fractions on the right side, we need a common denominator. The least common multiple (LCM) of 12 and 30 is 60.

Convert the fractions to have a denominator of 60:

- $\frac{1}{12} = \frac{1 \times 5}{12 \times 5} = \frac{5}{60}$
- $\frac{1}{30} = \frac{1 \times 2}{30 \times 2} = \frac{2}{60}$

Now substitute these back into the equation:

$$\frac{1}{X} = \frac{5}{60} - \frac{2}{60}$$

Subtract the numerators:

$$\frac{1}{X} = \frac{5-2}{60}$$

$$\frac{1}{X} = \frac{3}{60}$$

Simplify the fraction $\frac{3}{60}$:

$$\frac{3}{60} = \frac{3 \div 3}{60 \div 3} = \frac{1}{20}$$

So, we have:

$$\frac{1}{X} = \frac{1}{20}$$

Taking the reciprocal of both sides gives us X:

$$X = 20$$

Therefore, Pipe B can empty the tank in 20 hours.

Revision Table: Pipe and Tank Concepts

Concept	Description	Formula
Work Rate	The fraction of total work done per unit of time.	$\text{Rate} = \frac{1}{\text{Time}}$
Filling Pipes Together	Rates add up to find the combined filling rate.	$\text{Rate}_{\text{Total}} = \text{Rate}_1 + \text{Rate}_2 + \dots$
Filling and Emptying Pipes Together	Filling rates are positive, emptying rates are negative. Net rate is the sum. If tank fills, net rate is positive.	$\text{Rate}_{\text{Net}} = \text{Rate}_{\text{Filling}} - \text{Rate}_{\text{Emptying}}$
Time Taken	The reciprocal of the work rate.	$\text{Time} = \frac{1}{\text{Rate}}$

Additional Information: Understanding LCM in Rate Problems

In pipe and tank or time and work problems, the Least Common Multiple (LCM) of the individual times taken is often used to assume a 'total work' unit. This can sometimes simplify calculations, especially when dealing with multiple pipes or workers.

For this problem, the LCM of 12 (Pipe A's time) and 30 (Combined time) is 60. We can think of the tank having a capacity of 60 units.

- Pipe A's filling rate: $\frac{60 \text{ units}}{12 \text{ hours}} = 5 \text{ units/hour}$ (filling)
- Combined rate: $\frac{60 \text{ units}}{30 \text{ hours}} = 2 \text{ units/hour}$ (filling)
- Let Pipe B's emptying rate be 'r' units/hour.

When A and B work together, the net filling rate is A's rate minus B's rate (since B is emptying):

Net Rate = Rate of A – Rate of B

$$2 = 5 - r$$

Solving for r:

$$r = 5 - 2 = 3 \text{ units/hour (emptying)}$$

If Pipe B empties 3 units per hour, the time taken to empty the whole tank (60 units) is:

$$\text{Time for B} = \frac{\text{Total Capacity}}{\text{Rate of B}} = \frac{60 \text{ units}}{3 \text{ units/hour}} = 20 \text{ hours.}$$

This method using LCM confirms the result obtained using fractions and provides an alternative way to think about the problem.

41. Answer: d

Explanation:

Understanding the Fraction Calculation Problem

This problem involves calculating a fraction of a number. Amit made a mistake by using the wrong fraction, which resulted in an answer that was greater than the correct one. We need to find the original number.

Setting Up the Equation to Find the Number

Let the unknown number be denoted by x .

According to the problem, Amit calculated $\frac{2}{5}$ of the number. This incorrect calculation is given by $\frac{2}{5}x$.

The correct calculation should have been $\frac{2}{15}$ of the number. This correct calculation is given by $\frac{2}{15}x$.

The problem states that Amit's answer ($\frac{2}{5}x$) was greater than the correct answer ($\frac{2}{15}x$) by 336.

We can write this relationship as an equation:

Incorrect calculation - Correct calculation = Difference

$$\frac{2}{5}x - \frac{2}{15}x = 336$$

Solving the Equation for the Unknown Number

To solve the equation $\frac{2}{5}x - \frac{2}{15}x = 336$, we first need to find a common denominator for the fractions $\frac{2}{5}$ and $\frac{2}{15}$. The least common multiple of 5 and 15 is 15.

Rewrite the fractions with the common denominator:

$$\frac{2}{5}x = \frac{2 \times 3}{5 \times 3}x = \frac{6}{15}x$$

Now substitute this back into the equation:

$$\frac{6}{15}x - \frac{2}{15}x = 336$$

Combine the terms on the left side:

$$\left(\frac{6}{15} - \frac{2}{15}\right)x = 336$$

$$\frac{6-2}{15}x = 336$$

$$\frac{4}{15}x = 336$$

To find x , we need to isolate it. Multiply both sides of the equation by the reciprocal of $\frac{4}{15}$, which is $\frac{15}{4}$:

$$x = 336 \times \frac{15}{4}$$

Now, perform the calculation. We can simplify by dividing 336 by 4:

$$336 \div 4 = 84$$

So, the equation becomes:

$$x = 84 \times 15$$

Calculate the final product:

$$84 \times 15$$

We can do this multiplication:

	8	4	
x	1	5	
—	—	—	
	4	2	0 (84 x 5)
8	4	0 (84 x 10)	
—	—	—	—
1	2	6	0

So, $x = 1260$.

Verifying the Answer

Let's check if the number 1260 satisfies the original condition.

- Correct calculation: $\frac{2}{15} \times 1260 = 2 \times \frac{1260}{15} = 2 \times 84 = 168$
- Incorrect calculation: $\frac{2}{5} \times 1260 = 2 \times \frac{1260}{5} = 2 \times 252 = 504$

The difference between the incorrect and correct calculations is $504 - 168$.

$$504 - 168 = 336$$

This matches the difference given in the problem. Therefore, the number is indeed 1260.

Comparing with Options

The number we found is 1260, which corresponds to one of the given options.

- Option 1: 1680
- Option 2: 1344

- Option 3: 1008
- Option 4: 1260

Our calculated number matches Option 4.

Revision Table: Fractions and Equations

Concept	Description	How it applies here
Fraction of a Number	Multiplying the fraction by the number (e.g., $\frac{a}{b}$ of x is $\frac{a}{b} \times x$).	Used to represent Amit's correct and incorrect calculations.
Solving Linear Equations	Finding the value of the unknown variable that satisfies the equation. Involves isolating the variable using inverse operations.	Used to find the unknown number x from the difference equation.
Combining Fractions	Adding or subtracting fractions requires a common denominator.	Used to simplify $\frac{2}{5}x - \frac{2}{15}x$.

Additional Information: Working with Fractions and Differences

When a problem states that one quantity is "greater than" another by a certain amount, it means the difference between the larger quantity and the smaller quantity is that amount. In this case, the incorrect answer was larger than the correct answer.

The fractions $\frac{2}{5}$ and $\frac{2}{15}$ represent parts of the same number. Since $\frac{2}{5}$ is a larger fraction than $\frac{2}{15}$ (because 5 is smaller than 15, or by finding a common denominator, $\frac{6}{15}$ vs $\frac{2}{15}$), calculating $\frac{2}{5}$ of a number will always yield a larger result than calculating $\frac{2}{15}$ of the same number (for positive numbers).

Setting up the correct equation based on the problem's wording is the crucial first step. The phrase "greater than... by" directly translates to a subtraction leading to the

given difference.

42. Answer: c

Explanation:

Understanding Sentence Sequencing

The question asks us to arrange the given sentences into a logical and coherent paragraph. This involves identifying the correct flow of ideas and events to create a smooth narrative or explanation. We need to consider how each sentence connects to the one before it and the one after it.

Analyzing the Sentences for Coherence

Let's look at the starting sentence and the labelled sentences provided:

- Starting Sentence: There were once two brothers who lived on the edge of a forest.
- A: One day, the elder brother went into the forest to find some firewood to sell in the market.
- B: As he went around chopping the branches of a tree after tree, he came upon a magical tree.
- C: The elder brother was very mean to his younger brother and ate up all the food and took all his good clothes.
- D: The tree said to him, 'Oh kind sir, please do NOT cut my branches.'

The starting sentence introduces the main characters: two brothers living by a forest. Now we need to see which of the labelled sentences logically follows this introduction.

Evaluating the Order of Sentences

Sentence C introduces a specific detail about one of the brothers – the elder brother's meanness. This seems like a good way to elaborate on one of the characters introduced in the first sentence before describing their actions.

After establishing the elder brother's character in C, sentence A describes an action he takes: going into the forest. This action is consistent with the setting mentioned in the starting sentence and provides a reason for him to be in the forest.

Sentence B describes something that happens while the elder brother is carrying out the action described in A (chopping branches): he finds a magical tree. Sentence B logically follows A as it details the events occurring during the trip to the forest.

Sentence D is a direct quote from the tree. It is a reaction from the magical tree found in sentence B. Therefore, D must follow B as it gives the tree's response to the elder brother's actions.

Putting this together, the logical sequence is: Starting Sentence → C → A → B → D.

Constructing the Coherent Paragraph

Let's form the paragraph using the sequence Starting Sentence → C → A → B → D:

There were once two brothers who lived on the edge of a forest. **C:** The elder brother was very mean to his younger brother and ate up all the food and took all his good clothes. **A:** One day, the elder brother went into the forest to find some firewood to sell in the market. **B:** As he went around chopping the branches of a tree after tree, he came upon a magical tree. **D:** The tree said to him, 'Oh kind sir, please do NOT cut my branches.'

This sequence creates a smooth and logical flow, starting with the characters and setting, introducing a key characteristic of the main acting character, describing his action, detailing an event during his action, and finally showing the immediate consequence or interaction related to that event.

Analyzing the Options

Let's compare our logical sequence (CABD) with the given options:

Option	Sequence	Analysis
1	ACBD	Starts with action (A) before character detail (C). C feels out of place after starting the forest trip.
2	ACDB	Starts with action (A) before character detail (C). Also, D (tree talking) comes before B (finding the tree), which is illogical.
3	CABD	Starts with character detail (C) after the introduction. Follows with action (A), then event during action (B), then reaction to event (D). This is the most logical flow.
4	CADB	Starts with character detail (C) and action (A). However, D (tree talking) comes before B (finding the tree), which is illogical.

Based on our analysis, the sequence CABD forms the most coherent paragraph.

Revision Table: Sentence Sequencing Skills

Your Personal Exams Guide

Concept	Key Point	Why it matters
Topic Sentence	Often introduces the main idea or topic. (Like the starting sentence here).	Provides context and sets the stage for the paragraph.
Logical Order	Arranging sentences by time, cause-effect, general-to-specific, etc.	Ensures the paragraph makes sense and is easy to follow.
Transitions	Words or phrases (like "One day", "As he went around") connecting ideas.	Help sentences flow smoothly from one to the next.
Pronoun Reference	Ensure pronouns (like "he") clearly refer to a previously mentioned noun.	Avoids confusion about who or what is being discussed.

Additional Information: Building Coherent Paragraphs

A coherent paragraph has sentences that are logically connected and flow smoothly. To build coherence, you can use several techniques:

- **Order of Events:** Arrange sentences chronologically, describing events as they happened in time. This story follows a chronological order.
- **Cause and Effect:** Show how one event leads to another.
- **Comparison and Contrast:** Group similar points together or highlight differences.
- **Using Transition Words:** Words and phrases like 'therefore', 'however', 'in addition', 'for example', 'meanwhile', 'as a result' help link ideas and show the relationship between sentences.
- **Repeating Keywords or Synonyms:** Repeating important terms or using synonyms can help maintain focus and connect ideas throughout the paragraph. Notice how "elder brother" is mentioned, and then "he" is used, clearly referring back.

Mastering sentence sequencing is crucial for effective writing and reading comprehension, as it helps you understand the intended message and structure of any text.

43. Answer: c

Explanation:

Only conclusion II follows.

44. Answer: a

Explanation:

Solving Blood Relation Puzzles: Analyzing the Statement

This question is a classic example of a blood relation puzzle where we need to decipher the relationship between two individuals, C and D, based on a statement made by C to D. The statement is: "You are my sister's husband's mother-in-law." Let's break down this statement part by part from the perspective of the speaker, C.

Step-by-Step Analysis of the Relationship

Let's analyze the statement segment by segment to determine how D is related to C:

1. "My sister": This refers to C's sister.
2. "My sister's husband": This person is the husband of C's sister. This individual is C's brother-in-law (assuming C is not the sister mentioned).
3. "My sister's husband's mother-in-law": This refers to the mother-in-law of C's brother-in-law.

Now let's consider the relationship from the perspective of "my sister's husband" (the brother-in-law):

- The mother-in-law of a person is the mother of that person's spouse.
- The spouse of C's brother-in-law is C's sister.
- Therefore, the mother-in-law of C's brother-in-law is the mother of C's sister.

The mother of C's sister is C's mother (assuming C and the sister share the same mother, which is standard in such puzzles). So, the statement "You are my sister's husband's mother-in-law" means "You are my mother".

Since C said this statement to D, it implies that D is C's mother.

Confirming the Relationship between C and D

Based on the breakdown of the statement, the person C is talking to (D) is C's mother. Therefore, D is related to C as her mother.

Evaluating the Options

Let's look at the given options based on our analysis:

- D is the mother of C.
- D is the daughter of C.
- D is the sister of C.
- D is the mother-in-law of C.

Our analysis concludes that D is the mother of C. This matches the first option provided.

Revision Table: Key Relationship Terms

Relationship	Description
Sister's husband	Brother-in-law
Husband's mother-in-law	Mother of the husband's wife (the wife's mother)
Mother of sister	Mother

Additional Information on Blood Relations

Blood relation questions test your ability to understand and decode family relationships. It's helpful to visualize a family tree or break down the statement step by step, working backward from the end of the phrase. Common relationships involved include:

- Parents: Father, Mother
- Siblings: Brother, Sister
- Children: Son, Daughter
- Grandparents: Grandfather, Grandmother
- Grandchildren: Grandson, Granddaughter
- Aunts and Uncles: Sister/Brother of parent
- Nieces and Nephews: Children of siblings
- In-laws: Relationships through marriage (e.g., mother-in-law, brother-in-law, daughter-in-law)

Practice with different types of statements helps improve speed and accuracy in solving these blood relation questions.

45. Answer: b

Explanation:

Understanding ATA Carnets and Issuing Bodies in India

The question asks us to identify the specific organization in India that is designated as the sole National Issuing & Guaranteeing Association for ATA Carnets. An ATA Carnet is often referred to as a 'Passport for Goods' because it simplifies customs procedures for the temporary importation of goods across international borders.

Let's look at the role of an ATA Carnet:

- It allows goods to be imported into a country temporarily without paying customs duties and taxes.
- It acts as a guarantee for the importing country that all duties and taxes will be paid if the goods are not re-exported within the allowed time frame.
- It covers a wide range of goods, including samples, professional equipment, and goods for exhibitions and trade fairs.

Each country that is part of the international ATA Carnet chain must have a National Issuing & Guaranteeing Association. This association is responsible for:

- Issuing ATA Carnets to their residents wishing to export goods temporarily.
- Guaranteeing to foreign customs authorities the payment of import duties and taxes if a Carnet holder defaults on their obligations.

Now, let's consider the options provided and determine which organization fulfills this specific role in India:

Your Personal Exams Guide

Organization	Role in ATA Carnet System (in the context of the question)
Confederation of Indian Industry (CII)	A major business association, but not the designated sole issuing/guaranteeing body for ATA Carnets.
Federation of Indian Chambers of Commerce & Industry (FICCI)	The officially recognized and sole National Issuing & Guaranteeing Association for ATA Carnets in India.
All India Management Association (AIMA)	Primarily focused on management education and development, not involved in issuing ATA Carnets.
International Resources for Fairer Trade (IRFT)	A non-profit organization focused on fair trade issues, not the body for ATA Carnets.

Based on the information regarding the ATA Carnet system and the roles of various organizations in India, the Federation of Indian Chambers of Commerce & Industry (FICCI) is the established and sole National Issuing & Guaranteeing Association for ATA Carnets in the country.

Therefore, FICCI is the correct answer to the question about which body is India's sole National Issuing & Guaranteeing Association for ATA Carnets.

Revision Table: Key Points on ATA Carnets in India

Concept	Description
ATA Carnet	International customs document ('Passport for Goods') for temporary import/export.
Purpose	Avoids duties/taxes on temporary imports, simplifies customs procedures.
National Association's Role	Issue Carnets, guarantee duties/taxes to foreign customs.
India's Sole Association	Federation of Indian Chambers of Commerce & Industry (FICCI).

Additional Information on FICCI and ATA Carnets

The ATA Carnet system is governed by international conventions administered by the World Customs Organization (WCO) and the International Chamber of Commerce (ICC). FICCI was appointed as the National Issuing & Guaranteeing Association in India by the Government of India and accepted into the international network of ATA Carnet associations. Their role is crucial in facilitating international trade and movement of goods for specific purposes like exhibitions, trade fairs, and professional assignments without the burden of extensive customs formalities and payments at each border.

Using an ATA Carnet through FICCI helps Indian businesses and professionals taking goods abroad temporarily, and similarly helps foreign entities bringing goods temporarily into India, ensuring compliance with international customs regulations.

46. Answer: a

Explanation:

Understanding Different Types of Water Bodies

This question asks us to identify the word that is different from the rest in the given list of words related to water bodies: Stream, Lake, Pool, and Pond.

Let's look at each word and its typical characteristics:

- **Stream:** A small, narrow river. A key characteristic of a stream is that the water is typically **flowing**.
- **Lake:** A large body of water surrounded by land. The water in a lake is generally **still**, although there might be currents or waves.
- **Pool:** A small area of still water. This can be a natural depression or an artificial structure (like a swimming pool). The water is usually **still**.
- **Pond:** A small body of still water. Ponds are smaller than lakes and contain **still** water.

Comparing the characteristics, we can see a clear difference:

Comparison of Water Body Types

Word	Typical Water Movement
Stream	Flowing
Lake	Generally Still (large body)
Pool	Still (small area)
Pond	Still (small body)

The main difference among these terms is the movement of the water. Stream implies flowing water, while Lake, Pool, and Pond primarily refer to bodies of still or non-flowing water.

Therefore, the word that is different from the rest is **Stream** because it is characterized by flowing water, unlike lakes, pools, and ponds which are bodies of still water.

Revision Table: Understanding Water Bodies

Here is a quick table summarizing the key difference discussed:

Revision Summary: Flowing vs Still Water

Category	Words
Flowing Water Body	Stream
Still Water Bodies	Lake, Pool, Pond

Additional Information on Water Bodies

Understanding different types of water bodies is important in geography and environmental studies. While the primary distinction here is flowing vs. still water, other factors also differentiate these bodies:

- **Size:** Lakes are generally larger than ponds. Pools can be natural or artificial and vary greatly in size.
- **Origin:** Lakes and ponds can be naturally formed (e.g., by glaciers, volcanic activity, or tectonic shifts) or artificial (reservoirs, ornamental ponds). Streams are natural channels formed by water flow.
- **Ecosystems:** The types of plants and animals found in flowing water (streams) are different from those found in still water (lakes, ponds, pools) due to differences in oxygen levels, temperature variation, and nutrient distribution.

This type of question helps in building vocabulary and understanding classifications based on key properties.

47. Answer: c

Explanation:

Correct answer: chloroform

Reviewing key aspects related to air contaminants and their health effects.

Regulatory bodies establish exposure limits for chemicals in the workplace and environment to protect human health. These limits are often expressed in ppm or

mg/m³ and specify concentrations considered safe for different durations (e.g., permissible exposure limit for an 8-hour workday, short-term exposure limit for 15 minutes). Acute toxicity data, like the symptoms described in the question at a certain concentration, are used to help determine these limits. Symptoms like fatigue, dizziness, and headache are often indicators that exposure levels are approaching or exceeding levels that can cause adverse effects, even if not immediately life-threatening.

48. Answer: d

Explanation:

Chennai Super Kings Captain in IPL 2018

The question asks about the captain of the Indian Premier League (IPL) team Chennai Super Kings (CSK) during the 2018 season. Understanding the team's leadership is key to following their performance in any given IPL season.

Let's look at the options provided:

- **Virat Kohli:** Known for captaining Royal Challengers Bangalore (RCB). He was the captain of RCB in 2018.
- **Rohit Sharma:** Known for captaining Mumbai Indians (MI). He was the captain of MI in 2018.
- **Gautam Gambhir:** Had previously captained Kolkata Knight Riders (KKR) to title victories. In 2018, he played for Delhi Daredevils (now Delhi Capitals) and briefly captained them before stepping down.
- **Mahendra Singh Dhoni:** Widely recognized as the long-standing captain of Chennai Super Kings (CSK). He led the team for many seasons.

In 2018, Chennai Super Kings made their return to the IPL after a two-year suspension. The team's leadership was a significant factor in their successful comeback season. **Mahendra Singh Dhoni**, who had been CSK's captain since the league's inception (except for the two seasons when CSK was suspended and he

played for Rising Pune Supergiant), resumed his role as the captain for the 2018 season.

Under Dhoni's captaincy, CSK had a remarkable season, eventually winning the IPL title in 2018. This victory marked their third IPL title.

Therefore, based on the team's history and the events of the 2018 IPL season, **Mahendra Singh Dhoni** was the captain of Chennai Super Kings.

IPL 2018 Team Captains Overview

IPL Team	Captain in 2018
Chennai Super Kings (CSK)	Mahendra Singh Dhoni
Delhi Daredevils (DD)	Gautam Gambhir (initially), Shreyas Iyer
Kings XI Punjab (KXIP)	Ravichandran Ashwin
Kolkata Knight Riders (KKR)	Dinesh Karthik
Mumbai Indians (MI)	Rohit Sharma
Rajasthan Royals (RR)	Ajinkya Rahane
Royal Challengers Bangalore (RCB)	Virat Kohli
Sunrisers Hyderabad (SRH)	Kane Williamson

The table above confirms that **Mahendra Singh Dhoni** was indeed the captain for Chennai Super Kings in the 2018 IPL season.

Revision Table: Key IPL Captains in 2018

Captain	Team in IPL 2018
Mahendra Singh Dhoni	Chennai Super Kings (CSK)
Virat Kohli	Royal Challengers Bangalore (RCB)
Rohit Sharma	Mumbai Indians (MI)
Gautam Gambhir	Delhi Daredevils (DD)

Additional Information about CSK in IPL 2018

The 2018 season was significant for Chennai Super Kings as it marked their return after a two-year ban. Under the leadership of **Mahendra Singh Dhoni**, the team successfully rebuilt and performed exceptionally well throughout the tournament. They finished second in the league stage and went on to win the final against Sunrisers Hyderabad. The team composition included experienced players like Suresh Raina, Shane Watson, Ambati Rayudu, and Dwayne Bravo, alongside younger talent, all guided by Dhoni's strategic captaincy. This season is often remembered as one of CSK's most memorable comeback stories.

Your Personal Exams Guide

49. Answer: b

Explanation:

Calculating the Angle Between Clock Hands at 3:15

Let's figure out the angle between the hour hand and the minute hand of a clock at exactly quarter past three, which is 3:15.

A standard analog clock face is a circle, representing 360 degrees. This circle is divided into 12 hours.

First, let's understand how fast each hand moves:

- **Minute Hand Speed:** The minute hand completes a full circle (360 degrees) in 60 minutes. So, its speed is:

$$\text{Minute hand speed} = \frac{360^\circ}{60 \text{ minutes}} = 6^\circ \text{ per minute}$$

- **Hour Hand Speed:** The hour hand completes a full circle (360 degrees) in 12 hours. In terms of minutes (12 hours * 60 minutes/hour = 720 minutes), its speed is:

$$\text{Hour hand speed} = \frac{360^\circ}{12 \text{ hours}} = \frac{360^\circ}{720 \text{ minutes}} = 0.5^\circ \text{ per minute}$$

Position of Hands at 3:15

We measure the angle of each hand clockwise from the 12 o'clock position.

- **Minute Hand Position:** At 3:15, the minute hand is exactly on the '3'. The '3' is 15 minutes past the '12'.

$$\text{Minute hand angle from 12} = 15 \text{ minutes} \times 6^\circ/\text{minute} = 90^\circ$$

- **Hour Hand Position:** At 3:15, the hour hand is past the '3' mark because it has moved for 15 minutes beyond 3:00. At 3:00 exactly, the hour hand is on the '3', which is 3 hours $\times$ $30^\circ/\text{hour} = 90^\circ$ from the 12. In the extra 15 minutes, the hour hand moves:

$$\text{Hour hand movement in 15 minutes} = 15 \text{ minutes} \times 0.5^\circ/\text{minute} = 7.5^\circ$$

So, the total angle of the hour hand from the 12 o'clock position is:

$$\text{Hour hand angle from 12} = \text{Angle at 3:00} + \text{Movement in 15 minutes}$$

$$\text{Hour hand angle from 12} = 90^\circ + 7.5^\circ = 97.5^\circ$$

Calculating the Angle Between the Hands

The angle between the two hands is the absolute difference between their angles from the 12 o'clock position:

$$\text{Angle between hands} = |\text{Hour hand angle} - \text{Minute hand angle}|$$

$$\text{Angle between hands} = |97.5^\circ - 90^\circ|$$

$$\text{Angle between hands} = |7.5^\circ|$$

$$\text{Angle between hands} = 7.5^\circ$$

The angle between the hour hand and the minute hand of a clock at quarter past three is 7.5 degrees.

Clock Angle Calculation Revision Table

Hand	Movement per hour	Movement per minute	Reference Point
Minute Hand	360°	6°	12 o'clock
Hour Hand	30°	0.5°	12 o'clock

Additional Information on Clock Angles

You can also use a general formula to find the angle between the hands at any given time, H hours and M minutes:

$$\text{Angle} = |30 \times H - 11/2 \times M|$$

Using this formula for 3:15 (H=3, M=15):

$$\text{Angle} = |30 \times 3 - 11/2 \times 15|$$

$$\text{Angle} = |90 - 165/2|$$

$$\text{Angle} = |90 - 82.5|$$

$$\text{Angle} = |7.5|$$

$$\text{Angle} = 7.5^\circ$$

This confirms the result obtained by calculating the individual hand positions. Remember that sometimes the angle calculated this way might be the larger reflex

angle ($>180^\circ$). If the question asks for the acute angle (the smaller one), subtract the result from 360° if it's greater than 180° . In this case, 7.5° is already the smaller angle.

50. Answer: b

Explanation:

Correct answer: Mortar

Mortar: Mortar is a paste used in construction to bind building blocks such as bricks or stones together. This substance is specifically used to join bricks.

Therefore, the analogy is: Paper : Glue :: Brick : Mortar.

51. Answer: b

Explanation:

Calculating Copper Calorimeter Mass

The question asks us to find the mass of a copper calorimeter given the amount of heat it absorbs, the resulting temperature rise, and the specific heat capacity of copper. This is a classic problem involving the concept of specific heat capacity and heat transfer.

When an object absorbs heat, its temperature changes. The amount of heat absorbed (Q) is related to the mass of the object (m), its specific heat capacity (c), and the change in temperature (ΔT) by the following formula:

$$Q = mc\Delta T$$

In this problem, we are given:

- Heat absorbed (Q) = 3.6 kJ

- Temperature rise (ΔT) = 45°C
- Specific heat capacity of copper (c) = $0.4 \text{ Jg}^{-1}\text{K}^{-1}$

We need to find the mass (m) of the copper calorimeter in grams.

First, let's make sure all our units are consistent. The specific heat capacity is given in $\text{Jg}^{-1}\text{K}^{-1}$. This means mass should be in grams (g), heat in Joules (J), and temperature change in Kelvin (K).

- The heat absorbed is given in kilojoules (kJ). We need to convert this to Joules (J).
 $1 \text{ kJ} = 1000 \text{ J}$
 So, $Q = 3.6 \text{ kJ} = 3.6 \times 1000 \text{ J} = 3600 \text{ J}$.
- The temperature rise is given in degrees Celsius ($^{\circ}\text{C}$). A change in temperature in Celsius is the same as a change in temperature in Kelvin.
 So, $(\Delta T = 45^{\circ}\text{C} = 45 \text{ K})$.
- The specific heat capacity is already in the correct units: $c = 0.4 \text{ Jg}^{-1}\text{K}^{-1}$.

Now we can rearrange the formula $Q = mc\Delta T$ to solve for mass (m):

$$m = \frac{Q}{c\Delta T}$$

Substitute the values we have (with consistent units) into the formula:

$$m = \frac{3600 \text{ J}}{(0.4 \text{ Jg}^{-1}\text{K}^{-1})(45 \text{ K})}$$

$$m = \frac{3600 \text{ J}}{(0.4 \times 45) \text{ Jg}^{-1}}$$

Calculate the product in the denominator:

$$0.4 \times 45 = 18$$

Substitute this back into the equation for mass:

$$m = \frac{3600 \text{ J}}{18 \text{ Jg}^{-1}}$$

Now, perform the division. The units cancel out to give grams ($\text{J} / (\text{J/g}) = \text{J} * \text{g} / \text{J} = \text{g}$).

$$m = \frac{3600}{18} \text{ g}$$

$$m = 200 \text{ g}$$

Therefore, the mass of the copper calorimeter is 200 grams.

Let's summarize the given information and the calculated result:

Quantity	Symbol	Value (with units)
Heat Absorbed	Q	3.6 kJ = 3600 J
Temperature Rise	ΔT	45°C = 45 K
Specific Heat Capacity (Copper)	c	0.4 Jg ⁻¹ K ⁻¹
Mass of Calorimeter	m	200 g

Key Concepts for Heat Transfer Calculations

Understanding the definitions of specific heat capacity, heat absorbed, and temperature change is crucial for solving problems like this.

- **Heat Absorbed (Q):** This is the amount of thermal energy transferred to an object, causing its temperature to change or its state to change. In this case, it's the energy absorbed by the copper calorimeter.
- **Temperature Change (ΔT):** This is the difference between the final and initial temperatures of the object. A positive ΔT means the temperature increased (heat was absorbed), and a negative ΔT means the temperature decreased (heat was released).
- **Specific Heat Capacity (c):** This is a physical property of a substance that represents the amount of heat required to raise the temperature of one unit of mass (like 1 gram or 1 kilogram) of the substance by one degree (like 1°C or 1 K). It's a measure of how much energy a substance can store as internal energy for a given change in temperature. Different materials have different specific heat capacities.

The formula $Q = mc\Delta T$ only applies when the heat transfer causes a change in temperature and not a change in the state of matter (like melting or boiling).

Revision Table: Thermodynamics Formulas

Here are some related formulas often used in thermodynamics and heat transfer problems:

Concept	Formula	Variables
Heat transfer causing temperature change	$Q = mc\Delta T$	Q : heat, m : mass, c : specific heat, ΔT : temp change
Heat transfer causing phase change (e.g., melting, boiling)	$Q = mL$	Q : heat, m : mass, L : latent heat of fusion/vaporization
Thermal Power	$P = \frac{Q}{t}$	P : power, Q : heat, t : time
Heat Conduction (through a material)	$Q = \frac{kA\Delta Tt}{d}$	k : thermal conductivity, A : area, ΔT : temp difference, t : time, d : thickness

Additional Information on Calorimetry

Calorimetry is the science of measuring heat. A calorimeter is a device used for this measurement. The copper calorimeter mentioned in the problem is a common type of calorimeter.

- Calorimeters are often insulated to minimize heat exchange with the surroundings, ensuring that the measured heat transfer occurs primarily within the calorimeter system.
- When an experiment is conducted inside a calorimeter, the heat absorbed or released by the substance being studied is equal to the heat released or absorbed by the calorimeter and the liquid (usually water) it contains (assuming ideal conditions).
- In this specific problem, we only considered the heat absorbed by the copper calorimeter itself. In a more complex calorimetry problem, you might need to account for the heat absorbed by a liquid inside the calorimeter as well, using the specific heat capacity of the liquid.

- The principle of conservation of energy is fundamental to calorimetry. In an isolated system, the total heat lost by some components equals the total heat gained by others.

52. Answer: d

Explanation:

Understanding Computer Interfaces and Device Connection

The question asks to identify the term for an interface on a computer where you can connect a device. Let's examine the options provided to understand which one fits this description of a computer interface.

Analyzing the Options for Computer Interfaces

- **amine:** An amine is a type of chemical compound derived from ammonia. This term has no relation to computer hardware or interfaces.
- **dongle:** A dongle is a small hardware device that plugs into a port on a computer or other electronic device. Examples include USB dongles for Wi-Fi, Bluetooth, or software licensing. While it connects to an interface, it is the device itself, not the interface.
- **array:** In computer science, an array is a data structure used to store a collection of elements, typically of the same data type, in contiguous memory locations. This is a programming concept and not a physical interface for connecting devices.
- **port:** In the context of computing, a port is a physical connection point or interface on a computer or other electronic device. These ports allow devices such as keyboards, mice, printers, external hard drives, monitors, and network cables to be plugged in and communicate with the computer. Examples include USB ports, HDMI ports, Ethernet ports, and audio ports.

Based on the analysis, the term that describes an interface on a computer to which you can connect a device is a **port**.

Detailed Explanation of Computer Ports

A computer port serves as a physical or logical interface point. Physical ports are visible connectors on the outside of the computer case, designed to match specific cables or plugs from external devices. They provide the means for power and data transfer between the computer and the connected peripheral.

Different types of computer ports are designed for specific purposes and types of devices:

- **USB Ports:** Used for connecting a wide range of devices like keyboards, mice, printers, cameras, external storage, and charging devices.
- **HDMI/DisplayPort:** Used for connecting monitors and other display devices.
- **Ethernet Port:** Used for connecting the computer to a wired network.
- **Audio Jacks:** Used for connecting headphones, microphones, and speakers.
- **Power Port:** Used for connecting the power supply or charger.

These ports are the essential connection points that enable a computer system to interact with external hardware, making them the interfaces where devices are connected.

Comparing the Terms

Let's quickly compare the meanings in the context of connecting devices:

- **Amine:** Chemical compound (Incorrect).
- **Dongle:** The device being connected (Incorrect).
- **Array:** Data structure (Incorrect).
- **Port:** The interface where a device connects (Correct).

Therefore, the correct term is **port**.

Summary of Terms

Term	Definition in Computing Context	Fits "Interface for Device Connection"?
Amine	N/A (Chemical)	No
Dongle	External hardware device	No (It connects *to* the interface)
Array	Data structure	No
Port	Connection point/interface on a computer	Yes

Revision Table: Computer Ports and Interfaces

Reviewing the key concept related to the question:

- **What is a Port?** A port is a physical connection point on a computer or device where external peripherals or network cables can be plugged in.
- **Purpose of Ports:** To allow communication and data transfer between the computer and connected devices.
- **Examples:** USB, HDMI, Ethernet, Audio Jack, Power.

Additional Information: Types of Computer Interfaces

While ports are physical interfaces, the term "interface" can be broader in computing. It can also refer to logical interfaces or user interfaces.

- **Physical Interface:** Hardware connection points like ports (USB, HDMI, etc.), slots (PCIe slots for expansion cards), or sockets (CPU sockets).
- **Logical Interface:** Software-based interfaces or protocols that allow communication between different software components or systems (e.g., API - Application Programming Interface).
- **User Interface (UI):** The means by which a user interacts with a computer system or software (e.g., graphical user interface with windows and icons,

command-line interface).

The question specifically refers to connecting a device, which points to a physical interface, and "port" is the most appropriate term among the options provided for this kind of physical connection point on a computer.

53. Answer: b

Explanation:

Finding the Missing Number in a Decimal Series

The problem asks us to find the missing number in the given series: -2.7, -2.1, -1.5, -0.9, -0.3, ?

To solve this type of series problem, we need to identify the pattern or relationship between the consecutive terms. Let's examine the differences between the terms:

- Difference between the second term and the first term:
 $-2.1 - (-2.7) = -2.1 + 2.7 = 0.6$
- Difference between the third term and the second term:
 $-1.5 - (-2.1) = -1.5 + 2.1 = 0.6$
- Difference between the fourth term and the third term:
 $-0.9 - (-1.5) = -0.9 + 1.5 = 0.6$
- Difference between the fifth term and the fourth term:
 $-0.3 - (-0.9) = -0.3 + 0.9 = 0.6$

We observe that the difference between any consecutive terms is a constant value, which is 0.6. This indicates that the given series is an arithmetic progression with a common difference (d) of 0.6.

Calculating the Next Term in the Arithmetic Series

In an arithmetic series, the next term is found by adding the common difference to the last term of the series. The last given term in the series is -0.3.

Next term = Last term + Common difference

Next term = $-0.3 + 0.6$

Next term = 0.3

So, the missing number in the series is 0.3.

Verifying the Pattern

Let's list the terms and see if the pattern holds:

Term Number	Term Value	Pattern (+0.6)
1	-2.7	-
2	-2.1	$-2.7 + 0.6 = -2.1$
3	-1.5	$-2.1 + 0.6 = -1.5$
4	-0.9	$-1.5 + 0.6 = -0.9$
5	-0.3	$-0.9 + 0.6 = -0.3$
6	0.3	$-0.3 + 0.6 = 0.3$

Your Personal Exams Guide

The pattern is consistent, and the next term is indeed 0.3.

Conclusion for the Missing Number

Based on our analysis of the pattern, the missing number in the series -2.7, -2.1, -1.5, -0.9, -0.3, ? is 0.3.

Revision Table: Key Concepts

Concept	Description	Relevance to Problem
Arithmetic Progression (AP)	A sequence of numbers where the difference between consecutive terms is constant.	The given series is an example of an AP.
Common Difference (d)	The constant difference between consecutive terms in an AP.	Identified as 0.6 in this series.
Finding Next Term	Add the common difference to the previous term.	Used to calculate the missing number.

Additional Information on Number Series

Number series problems are common in aptitude tests. They require you to find the pattern in a sequence of numbers. Patterns can be:

- Arithmetic (constant difference)
- Geometric (constant ratio)
- Differences of differences (second-order difference is constant)
- Squares, cubes, or other powers
- Alternating patterns
- Mixed patterns

Practicing different types of series helps in quickly identifying the underlying rule and solving for the missing term or the next term in the sequence. Pay attention to signs (positive/negative) and decimal points, as seen in this specific series problem involving negative decimal numbers.

54. Answer: b

Explanation:

Understanding the Question: Identifying the Author of 'Mrs. Funnybones'

The question asks to name the well-known personality who wrote the book titled 'Mrs. Funnybones'. This requires knowledge of famous authors, particularly those who might also be celebrities.

Analyzing the Book 'Mrs. Funnybones'

'Mrs. Funnybones' is a popular non-fiction book known for its humorous take on everyday life, family, and observations about society. The book gained significant attention upon its release due to its witty writing style and the author's celebrity status.

Evaluating the Options

Let's consider the provided options and see which one is associated with the book 'Mrs. Funnybones':

- **Bipasha Basu:** Known primarily as an actress, Bipasha Basu is not widely recognized as an author of a book like 'Mrs. Funnybones'.
- **Twinkle Khanna:** Twinkle Khanna is a former actress and interior designer who became a highly successful author and columnist. She is known for her humorous and insightful writing.
- **Shraddha Kapoor:** Known for her acting career, Shraddha Kapoor is not known to have authored a book titled 'Mrs. Funnybones'.
- **Sonakshi Sinha:** Primarily an actress, Sonakshi Sinha is not associated with the authorship of the book 'Mrs. Funnybones'.

Identifying the Author

Based on general knowledge about prominent Indian authors and celebrities, Twinkle Khanna is widely recognized as the author of the book 'Mrs. Funnybones'. She debuted as an author with this book, which became a bestseller.

Explanation of the Correct Choice

The book 'Mrs. Funnybones' was indeed authored by Twinkle Khanna. It marked her successful foray into writing, establishing her as a witty columnist and author. The book's title itself is derived from her popular column. Therefore, Twinkle Khanna is the correct answer.

Revision Table: Famous Celebrity Authors and Their Books

Celebrity Author	Notable Book(s)	Category
Twinkle Khanna	Mrs. Funnybones, The Legend of Lakshmi Prasad, Pyjamas Are Forgiving	Humor, Short Stories, Novel
Priyanka Chopra Jonas	Unfinished	Memoir
Naseeruddin Shah	And Then One Day: A Memoir	Memoir
Karan Johar	An Unsuitable Boy	Memoir

Your Personal Exams Guide

Additional Information on Twinkle Khanna's Writing

Twinkle Khanna started her writing career with a column for Daily News and Analysis (DNA). She then moved to The Times of India and writes a popular column under the pseudonym 'Mrs. Funnybones'. Her first book, 'Mrs. Funnybones', published in 2015, was a compilation and extension of her writings, focusing on her observations and experiences with humor. The book was a critical and commercial success, making her one of the highest-selling female writers in India that year. She has since published other successful books, including a collection of short stories and a novel.

55. Answer: a

Explanation:

Identifying Folk Dances of India: Himachal Pradesh

The question asks to identify the folk dance among the given options that originates from Himachal Pradesh. Folk dances are an integral part of India's rich cultural heritage, with each state and region having its unique dance forms reflecting its traditions, history, and social life.

Analyzing the Options

Let's examine each dance form provided in the options to determine its state of origin:

- **Nati:** The Nati is a traditional folk dance originating from the state of Himachal Pradesh and the Garhwal region of Uttarakhand. It is widely recognized and popular in Himachal Pradesh, often performed during festivals, weddings, and other social gatherings.
- **Lezim:** Lezim is a folk dance form popular in the state of Maharashtra. It uses a musical instrument also called Lezim, which is a small percussion instrument with jingling cymbals.
- **Bagurumba:** Bagurumba is a folk dance of the Bodo community in Assam and Northeast India. It is traditionally performed by women.
- **Giddha:** Giddha is a popular folk dance from the Punjab region, performed by women. It is often accompanied by rhythmic clapping and traditional folk songs.

Determining the Correct Answer

Based on the analysis of each option:

- Nati is associated with Himachal Pradesh and Uttarakhand.
- Lezim is associated with Maharashtra.
- Bagurumba is associated with Assam.

- Giddha is associated with Punjab.

Therefore, the folk dance from Himachal Pradesh among the given options is Nati.

Revision Table: Folk Dances and States

Folk Dance	State/Region
Nati	Himachal Pradesh / Uttarakhand
Lezim	Maharashtra
Bagurumba	Assam
Giddha	Punjab

Additional Information on Indian Folk Dances

India has a vast and diverse collection of folk dances, each telling a story of its people and place. Learning about these dances helps in understanding the cultural mosaic of the country.

- **Types of Folk Dances:** Folk dances can be categorized based on the occasion they are performed for (harvest, festivals, weddings), the community performing them, or the instruments/props used.
- **Importance:** These dances are not just entertainment; they preserve traditions, history, social values, and community identity across generations.
- **Examples from other states:**
 - Garba and Dandiya Raas from Gujarat.
 - Bhangra from Punjab (male counterpart of Giddha).
 - Kathakali and Mohiniyattam from Kerala.
 - Lavani from Maharashtra.
 - Bihu from Assam.
 - Chhau from Odisha, Jharkhand, and West Bengal.

Understanding the origin and characteristics of various folk dances is important for appreciating India's cultural diversity and is often tested in general knowledge

sections of exams.

56. Answer: d

Explanation:

Understanding Dimensions in Engineering Drawings

Engineering drawings are the universal language used to convey design information. Dimensions are critical components of these drawings, specifying the size, location, and geometric characteristics of features on a part. Each dimension provides specific information necessary for manufacturing and inspection.

Identifying Extra Information on Dimensions

Sometimes, a dimension might be included on a drawing for informational purposes only. It might be a calculated value or a dimension that is redundant because it can be derived from other dimensions. Such dimensions are included as supplementary information but are not required for manufacturing or inspecting the part based on the fundamental dimensions.

To distinguish these informational dimensions from critical manufacturing or inspection dimensions, special notation is used alongside the dimension value.

What Does 'REF' Mean on a Dimension?

In engineering drawings, the standard notation used to indicate that a dimension is for reference only is the abbreviation **REF**. This abbreviation is placed next to the dimension value.

- A dimension marked with **REF** means it is provided as extra information or for convenience.
- It is typically a dimension that is controlled by other dimensions on the drawing.

- Manufacturing and inspection should not rely on or check against a **REF** dimension directly; they should focus on the primary dimensions that control the part geometry.

Using **REF** clearly communicates to anyone reading the drawing that this particular dimension is not a driving dimension and does not need to be held to a specific tolerance for manufacturing or inspection purposes.

Analyzing the Given Options

Let's examine the provided options to determine which one correctly indicates extra information on a dimension in an engineering drawing:

- **EXT:** This abbreviation is not standard for indicating extra information or reference dimensions in this context. It might sometimes relate to external features or extensions depending on the specific drawing's conventions, but not to denote a reference dimension.
- **PER:** This abbreviation is not a standard notation for dimension modification in engineering drawings. It could potentially relate to perimeter in some calculations or notes, but not for marking a dimension as extra information.
- **NR:** While "NR" could conceptually mean "Not Required," it is not the standard notation used in engineering drawing standards (like ASME or ISO) to indicate a reference dimension or extra information. The universally accepted term for this is "REF".
- **REF:** As discussed, **REF** stands for "Reference" and is the standard and widely accepted notation in engineering drawings to indicate that a dimension is extra information, provided for reference only, and not required for manufacturing or inspection.

Therefore, the letters written on the dimension indicating it is extra information and NOT really required are **REF**.

Notation	Meaning in Dimensioning	Indicates Extra Information?
REF	Reference	Yes
EXT	Not a standard dimension modifier for this purpose	No
PER	Not a standard dimension modifier	No
NR	Not the standard notation for this purpose	No

Conclusion on Reference Dimensions

In summary, when you see the letters **REF** next to a dimension value on an engineering drawing, it signifies that the dimension is a reference dimension. It's supplementary information, often derived from other dimensions, and is not critical for the manufacturing or inspection processes which should instead focus on the primary, non-reference dimensions and their specified tolerances.

Revision Table: Engineering Drawing Dimensioning

Review key concepts related to dimensioning in engineering drawings:

- **Dimension Line:** A line used to indicate the extent of a dimension.
- **Extension Line:** Lines used to extend from the feature being dimensioned to the dimension line.
- **Leaders:** Lines pointing from a note or dimension to the feature it references.
- **Dimension Value:** The numerical value specifying the size or location.
- **Tolerances:** Permissible variations in size or location.
- **Reference Dimension (REF):** A dimension provided for information only, not for manufacturing or inspection verification.

Additional Information: Other Dimension Modifiers

Besides REF, engineering drawings use various symbols and notes to provide specific dimensioning instructions or information. Some common examples include:

- **TYP:** Typical (indicates a dimension or feature is typical for multiple locations).
- **MAX:** Maximum (indicates the maximum permissible size).
- **MIN:** Minimum (indicates the minimum permissible size).
- **R:** Radius (precedes a dimension value indicating a radius).
- **∅**: Diameter (precedes a dimension value indicating a diameter).
- **SQ:** Square (precedes a dimension value for a square feature).

These modifiers and symbols are crucial for accurately interpreting engineering drawings and ensuring parts are manufactured correctly.

57. Answer: c

Explanation:

Solving the Cab Speed Problem

This problem involves calculating the speed of a faster taxi given the distance, the speed of a slower taxi, and the difference in their travel times. We can use the relationship between speed, distance, and time to solve this.

Understanding the Given Information

We are given the following details:

- Distance between town C and town D (D) = 540 km
- Speed of the slower taxi (S_A) = 90 km/hr
- The slower taxi takes 1 hour more than the faster taxi.

Let S_B be the speed of the faster taxi in km/hr, T_A be the time taken by the slower taxi in hours, and T_B be the time taken by the faster taxi in hours.

From the problem, we know that $T_A = T_B + 1$.

Applying the Speed, Distance, Time Formula

The basic formula relating speed, distance, and time is:

$$\text{Speed} = \frac{\text{Distance}}{\text{Time}}$$

Rearranging this, we get the formula for time:

$$\text{Time} = \frac{\text{Distance}}{\text{Speed}}$$

Calculating the Time Taken by the Slower Taxi

Using the formula, we can calculate the time taken by the slower taxi (T_A):

$$T_A = \frac{D}{S_A}$$

Substituting the given values:

$$T_A = \frac{540 \text{ km}}{90 \text{ km/hr}}$$

$$T_A = 6 \text{ hours}$$

So, the slower taxi takes 6 hours to cover the 540 km distance.

Setting up the Equation for the Faster Taxi's Speed

We know that the slower taxi takes 1 hour more than the faster taxi. So, the time taken by the faster taxi (T_B) is:

$$T_B = T_A - 1$$

$$T_B = 6 \text{ hours} - 1 \text{ hour}$$

$$T_B = 5 \text{ hours}$$

Now, we can express the time taken by the faster taxi using its speed S_B and the distance D :

$$T_B = \frac{D}{S_B}$$

Substitute the known values into this equation:

$$5 \text{ hours} = \frac{540 \text{ km}}{S_B}$$

Solving for the Speed of the Faster Taxi

To find S_B , we can rearrange the equation:

$$S_B = \frac{540 \text{ km}}{5 \text{ hours}}$$

Performing the division:

$$S_B = 108 \text{ km/hr}$$

The speed of the faster taxi is 108 km/hr.

Analyzing the Options

Let's compare our calculated speed with the given options:

- Option 1: 117 km/hr
- Option 2: 126 km/hr
- Option 3: 108 km/hr
- Option 4: 99 km/hr

Our calculated speed of 108 km/hr matches Option 3.

Revision Table: Speed, Distance, Time Concepts

Here is a summary of the formulas used in speed, distance, and time problems:

Concept	Formula	Units (e.g., km, hr, km/hr)
Speed	$\text{Speed} = \frac{\text{Distance}}{\text{Time}}$	km/hr, m/s, miles/hr, etc.
Distance	$\text{Distance} = \text{Speed} \times \text{Time}$	km, meters, miles, etc.
Time	$\text{Time} = \frac{\text{Distance}}{\text{Speed}}$	hours, seconds, minutes, etc.

Additional Information: Solving Related Problems

Problems involving speed, distance, and time often appear in competitive exams. Understanding the basic formulas is crucial. Other variations include:

- Problems involving relative speed (when objects are moving towards or away from each other).
- Problems involving trains passing poles, platforms, or other trains.
- Problems involving boats and streams (considering the speed of the current).
- Problems where one part of the journey is covered at one speed and another part at a different speed.

Practice with different types of problems helps in mastering this topic.

58. Answer: b

Explanation:

Understanding Aerosols and Air Pollution

The question asks to identify the term that describes tiny particles suspended in the atmosphere, which are considered a subset of air pollution. Let's look at the options provided:

- **Loan:** A loan is a sum of money that is borrowed and is expected to be paid back, usually with interest. This term is completely unrelated to air pollution or atmospheric particles.
- **Aerosols:** Aerosols are defined as a suspension of tiny solid particles or liquid droplets in a gas, typically air. These particles can range in size from a few nanometers to tens of micrometers. They are indeed suspended everywhere in our atmosphere and are a significant component of air pollution. Examples include dust, soot, sea salt, pollen, and sulfates.

- **Genomes:** A genome is the complete set of genetic information in an organism. This relates to biology and genetics, not atmospheric science or air pollution.
- **Humus:** Humus is the dark, organic material that forms in soil when plant and animal matter decays. This relates to soil science and decomposition, not atmospheric particles.

Based on the definitions, **Aerosols** perfectly fit the description of tiny particles suspended in the atmosphere that are a subset of air pollution.

These atmospheric aerosols can have various origins, including both natural sources (like volcanic eruptions, dust storms, forest fires, and sea spray) and anthropogenic sources (like emissions from burning fossil fuels, industrial processes, and agriculture). Their presence in the atmosphere has significant impacts on air quality, visibility, cloud formation, precipitation, and climate.

Therefore, the term that correctly completes the sentence is Aerosols.

Statement Analysis

Let's analyze how the term 'Aerosols' fits the statement:

"_____ are a subset of air pollution that refers to the tiny particles suspended everywhere in our atmosphere."

Substituting 'Aerosols' into the blank:

"**Aerosols** are a subset of air pollution that refers to the tiny particles suspended everywhere in our atmosphere."

This statement is accurate. Aerosols are indeed tiny particles suspended in the atmosphere, and they are a major component contributing to air pollution.

Why Other Options Are Incorrect

- **Loan:** Irrelevant financial term.
- **Genomes:** Relevant to genetics and biology, not atmospheric science.
- **Humus:** Relevant to soil science, not atmospheric particles.

Revision Table: Key Terms in Air Pollution

Term	Definition	Relevance to Question
Aerosols	Suspension of tiny solid particles or liquid droplets in a gas (air).	Directly fits the description of tiny particles suspended in the atmosphere and a subset of air pollution.
Air Pollution	Contamination of the indoor or outdoor environment by any chemical, physical or biological agent that modifies the natural characteristics of the atmosphere.	Aerosols are a component of air pollution.
Particulate Matter (PM)	Also known as particle pollution, a mix of solid particles and liquid droplets found in the air. Includes aerosols. Often categorized by size (e.g., PM2.5, PM10).	A broader term that includes the types of particles described as aerosols. Aerosols are the suspension medium.

Additional Information: The Role of Aerosols

Aerosols play a crucial role in the Earth's atmosphere and climate system. Their effects are complex and can be both direct and indirect.

- **Direct Effects:** Aerosols can scatter and absorb solar radiation, affecting how much sunlight reaches the Earth's surface. This can lead to cooling or warming of the atmosphere, depending on the particle type.
- **Indirect Effects:** Aerosols serve as condensation nuclei for cloud droplets and ice crystals. Changes in aerosol concentration can influence cloud formation, brightness, and lifetime, thereby affecting precipitation patterns and the Earth's radiation balance.
- **Health Impacts:** Inhaling certain types of aerosols (particulate matter) can cause respiratory problems, cardiovascular issues, and other adverse health

effects, making them a significant public health concern and a key component of air pollution regulations.

- **Visibility:** High concentrations of aerosols can reduce visibility, creating haze.

Understanding aerosols is essential for studying air quality, climate change, and public health.

59. Answer: d

Explanation:

Syllogism Analysis: Statements and Conclusions

This question asks us to analyze two statements and determine which of the given conclusions logically follow from them. We must assume the statements are true, even if they contradict common knowledge. The statements involve the categories: Mansions, Palaces, and Castles.

Understanding the Given Statements

Let's break down the statements provided:

1. Statement 1: All mansions are palaces.
2. Statement 2: No castles are palaces.

These statements establish relationships between the three categories. Statement 1 tells us that the set of Mansions is a subset of the set of Palaces. Statement 2 tells us that the set of Castles has no members in common with the set of Palaces.

Evaluating Each Conclusion

We will now examine each conclusion based on the relationships established by the statements.

Conclusion I: Some palaces are mansions.

Statement 1 says "All mansions are palaces". If every single mansion is also a palace, then it logically follows that there must be some palaces that are mansions. This is a standard inference in logic: the converse of a universal affirmative ('All A are B') is a particular affirmative ('Some B are A'). As long as there is at least one mansion (which is usually assumed unless specified otherwise), this conclusion holds true.

Therefore, Conclusion I follows from the statements.

Conclusion II: Some castles are mansions.

Statement 2 says "No castles are palaces". This means the sets of Castles and Palaces are completely separate; they have no overlap. Statement 1 says "All mansions are palaces". This means that all mansions are located within the set of Palaces.

If all mansions are inside the set of Palaces, and the set of Castles is completely outside the set of Palaces, then the set of Castles must also be completely outside the set of Mansions. There can be no overlap between Castles and Mansions.

So, the statement "Some castles are mansions" is false based on the given statements.

Therefore, Conclusion II does not follow from the statements.

Conclusion III: No mansions are castles.

Let's revisit the relationships. We know from Statement 1 that all Mansions are within the Palaces set. We know from Statement 2 that the Castles set is entirely separate from the Palaces set.

If Mansions are inside Palaces, and Palaces have no common elements with Castles, then Mansions can also have no common elements with Castles. The statement "No mansions are castles" means that the set of Mansions and the set of Castles are completely separate. This is consistent with our understanding from the statements.

Therefore, Conclusion III follows from the statements.

Summary of Conclusions

Based on our analysis:

- Conclusion I: Some palaces are mansions. (Follows)
- Conclusion II: Some castles are mansions. (Does not follow)
- Conclusion III: No mansions are castles. (Follows)

The conclusions that follow are only Conclusion I and Conclusion III.

Statement/Conclusion	Relationship Implied	Follows?
Statement 1: All mansions are palaces.	Mansions $\subset$ Palaces	N/A
Statement 2: No castles are palaces.	Castles $\cap$ Palaces = $\emptyset$	N/A
Conclusion I: Some palaces are mansions.	Palaces $\cap$ Mansions $\neq \emptyset$	Yes (Converse of Statement 1)
Conclusion II: Some castles are mansions.	Castles $\cap$ Mansions $\neq \emptyset$	No (Contradicted by statements)
Conclusion III: No mansions are castles.	Mansions $\cap$ Castles = $\emptyset$	Yes (Derived from Statements 1 & 2)

Final Decision

Only Conclusion I and Conclusion III follow from the given statements.

Revision Table: Key Syllogism Rules

Statement Type	Representation	Converse
A (All X are Y)	$X \subset Y$	I (Some Y are X) – Simple Converse (valid if $X \neq \emptyset$;))
E (No X are Y)	$X \cap Y = \emptyset$;	E (No Y are X) – Simple Converse (always valid)
I (Some X are Y)	$X \cap Y \neq \emptyset$;	I (Some Y are X) – Simple Converse (always valid)
O (Some X are not Y)	$X \cap Y^c \neq \emptyset$;	None (No valid simple converse)

Additional Information: Syllogism Basics

Syllogisms are a form of logical argument that applies deductive reasoning to arrive at a conclusion based on two or more propositions that are asserted or assumed to be true. The basic structure involves:

- **Major Premise:** A general statement.
- **Minor Premise:** A specific statement related to the major premise.
- **Conclusion:** A statement that follows logically from the two premises.

In the given problem, Statement 1 is the minor premise (containing the subject of the conclusion, "mansions") and Statement 2 is the major premise (containing the predicate of the conclusion, "castles", when conclusion III is considered). The term "palaces" acts as the middle term, connecting the major and minor terms.

Visual tools like Venn diagrams are often helpful in solving syllogism problems. Drawing overlapping or separate circles representing the categories can make the relationships clearer and help verify which conclusions are valid.

60. Answer: c

Explanation:

300 g

The apparent mass is what the mass seems to be when the object is weighed in a fluid, and it's less than the real mass because of the buoyant force exerted by the fluid.

61. Answer: b

Explanation:

Finding the Value of Trigonometric Expressions

We are given that $\sin \theta = \frac{12}{13}$ and asked to find the value of the expression $2 \cot \theta + 13 \cos \theta$. To solve this, we first need to find the values of $\cos \theta$ and $\cot \theta$ using the given information about $\sin \theta$.

Using a Right-Angled Triangle

We can visualize this problem using a right-angled triangle. Recall that for a right-angled triangle, $\sin \theta = \frac{\text{Opposite side}}{\text{Hypotenuse}}$. Given $\sin \theta = \frac{12}{13}$, we can consider a right-angled triangle where:

- The length of the side opposite to angle θ is proportional to 12.
- The length of the hypotenuse is proportional to 13.

Let the opposite side be $12k$ and the hypotenuse be $13k$ for some positive constant k . We need to find the length of the adjacent side. Using the Pythagorean theorem:

$$(\text{Adjacent side})^2 + (\text{Opposite side})^2 = (\text{Hypotenuse})^2$$

$$(\text{Adjacent side})^2 + (12k)^2 = (13k)^2$$

$$(\text{Adjacent side})^2 + 144k^2 = 169k^2$$

$$(\text{Adjacent side})^2 = 169k^2 - 144k^2$$

$$(\text{Adjacent side})^2 = 25k^2$$

$$\text{Adjacent side} = \sqrt{25k^2} = 5k \text{ (Assuming } k > 0 \text{ and side length is positive).}$$

Calculating $\cos \theta$ and $\cot \theta$

Now that we have all three sides of the right-angled triangle, we can find $\cos \theta$ and $\cot \theta$.

Recall the definitions:

- $\cos \theta = \frac{\text{Adjacent side}}{\text{Hypotenuse}}$
- $\cot \theta = \frac{\text{Adjacent side}}{\text{Opposite side}}$

Substituting the values:

$$\cos \theta = \frac{5k}{13k} = \frac{5}{13}$$

$$\cot \theta = \frac{5k}{12k} = \frac{5}{12}$$

(Note: This assumes θ is in a quadrant where sine, cosine, and cotangent are positive, typically the first quadrant).

Evaluating the Expression $2\cot \theta + 13\cos \theta$

Now we substitute the calculated values of $\cot \theta$ and $\cos \theta$ into the given expression:

$$\text{Expression} = 2\cot \theta + 13\cos \theta$$

$$\text{Substitute } \cot \theta = \frac{5}{12} \text{ and } \cos \theta = \frac{5}{13}:$$

$$\text{Expression} = 2 \times \left(\frac{5}{12}\right) + 13 \times \left(\frac{5}{13}\right)$$

Simplify the terms:

$$\text{First term: } 2 \times \frac{5}{12} = \frac{10}{12} = \frac{5}{6}$$

$$\text{Second term: } 13 \times \frac{5}{13} = 5$$

Add the simplified terms:

$$\text{Expression} = \frac{5}{6} + 5$$

To add these, find a common denominator, which is 6:

$$5 = \frac{5 \times 6}{1 \times 6} = \frac{30}{6}$$

$$\text{Expression} = \frac{5}{6} + \frac{30}{6} = \frac{5+30}{6} = \frac{35}{6}$$

Thus, the value of $2 \cot \theta + 13 \cos \theta$ is $\frac{35}{6}$.

Step-by-Step Calculation Summary

Step	Action	Result
1	Identify known value from $\sin \theta$	Opposite = 12k, Hypotenuse = 13k
2	Calculate Adjacent side using Pythagoras theorem	Adjacent = 5k
3	Calculate $\cos \theta$	$\cos \theta = \frac{5}{13}$
4	Calculate $\cot \theta$	$\cot \theta = \frac{5}{12}$
5	Substitute values into $2 \cot \theta + 13 \cos \theta$	$2 \left(\frac{5}{12} \right) + 13 \left(\frac{5}{13} \right)$
6	Simplify the expression	$\frac{5}{6} + 5$
7	Add the fractions	$\frac{35}{6}$

Revision Table: Trigonometric Ratios in a Right Triangle

Ratio	Definition
$\sin \theta$	Opposite / Hypotenuse
$\cos \theta$	Adjacent / Hypotenuse
$\tan \theta$	Opposite / Adjacent
$\csc \theta$	Hypotenuse / Opposite ($= 1/\sin \theta$)
$\sec \theta$	Hypotenuse / Adjacent ($= 1/\cos \theta$)
$\cot \theta$	Adjacent / Opposite ($= 1/\tan \theta$ or $\cos \theta/\sin \theta$)

Additional Information on Trigonometric Identities

Besides using a right-angled triangle, trigonometric identities can also be used to find missing ratios.

- The fundamental identity: $\sin^2 \theta + \cos^2 \theta = 1$. Given $\sin \theta$, we can find $\cos \theta$. $\cos^2 \theta = 1 - \sin^2 \theta$ $\cos \theta = \pm \sqrt{1 - \sin^2 \theta}$ In our case, $\sin \theta = 12/13$. $\cos \theta = \pm \sqrt{1 - \left(\frac{12}{13}\right)^2} = \pm \sqrt{1 - \frac{144}{169}} = \pm \sqrt{\frac{169-144}{169}} = \pm \sqrt{\frac{25}{169}} = \pm \frac{5}{13}$. We chose $+\frac{5}{13}$ based on the assumption of θ being in the first quadrant where cosine is positive.
- Once $\sin \theta$ and $\cos \theta$ are known, $\cot \theta$ can be found using the identity $\cot \theta = \frac{\cos \theta}{\sin \theta}$. $\cot \theta = \frac{5/13}{12/13} = \frac{5}{13} \times \frac{13}{12} = \frac{5}{12}$.

Both the triangle method and the identity method give the same values for $\cos \theta$ and $\cot \theta$, leading to the same final result for the expression.

62. Answer: c

Explanation:

Understanding Kinetic Energy and Mass Calculation

The question asks us to find the mass of a car given the change in its kinetic energy and its initial and final speeds. Kinetic energy is the energy an object possesses due to its motion. When the speed of an object changes, its kinetic energy also changes. The relationship between kinetic energy (KE), mass (m), and speed (v) is given by the formula:

$$KE = \frac{1}{2}mv^2$$

The change in kinetic energy (ΔKE) is the difference between the final kinetic energy (KE_{final}) and the initial kinetic energy ($KE_{initial}$).

$$\Delta KE = KE_{final} - KE_{initial}$$

Since $KE = \frac{1}{2}mv^2$, we can write the change in kinetic energy as:

$$\Delta KE = \frac{1}{2}mv_{final}^2 - \frac{1}{2}mv_{initial}^2$$

We can factor out $\frac{1}{2}m$:

$$\Delta KE = \frac{1}{2}m(v_{final}^2 - v_{initial}^2)$$

Applying the Formula to the Car's Motion

We are given the following information about the car:

- Change in kinetic energy (ΔKE) = -200 kJ. Note that the car loses kinetic energy, so the change is negative.
- Initial speed ($v_{initial}$) = 25 m/s.
- Final speed (v_{final}) = 15 m/s.

First, let's convert the change in kinetic energy from kilojoules (kJ) to joules (J), which is the standard unit for energy in SI units. $1 \text{ kJ} = 1000 \text{ J}$.

$$\Delta KE = -200 \text{ kJ} = -200 \times 1000 \text{ J} = -200,000 \text{ J}$$

Step-by-Step Calculation of Mass

Now we can substitute the given values into the formula for the change in kinetic energy:

$$\Delta KE = \frac{1}{2}m(v_{final}^2 - v_{initial}^2)$$

$$-200,000 \text{ J} = \frac{1}{2}m((15 \text{ m/s})^2 - (25 \text{ m/s})^2)$$

Calculate the squares of the speeds:

$$15^2 = 15 \times 15 = 225$$

$$25^2 = 25 \times 25 = 625$$

Substitute these values back into the equation:

$$-200,000 = \frac{1}{2}m(225 - 625)$$

Calculate the difference in the squared speeds:

$$225 - 625 = -400$$

Substitute this difference back into the equation:

$$-200,000 = \frac{1}{2}m(-400)$$

Multiply $\frac{1}{2}$ by -400 :

$$\frac{1}{2} \times -400 = -200$$

So the equation becomes:

$$-200,000 = -200m$$

Now, solve for the mass m by dividing both sides by -200 :

$$m = \frac{-200,000}{-200}$$

$$m = 1000 \text{ kg}$$

Converting Mass to Tonnes

The question asks for the mass in tonnes. The conversion factor between kilograms (kg) and tonnes is:

$$1 \text{ tonne} = 1000 \text{ kg}$$

To convert the mass from kilograms to tonnes, we divide the mass in kilograms by 1000:

$$\text{Mass in tonnes} = \frac{\text{Mass in kg}}{1000}$$

$$\text{Mass in tonnes} = \frac{1000 \text{ kg}}{1000 \text{ kg/tonne}}$$

$$\text{Mass in tonnes} = 1 \text{ tonne}$$

Conclusion

The mass of the car is 1000 kg, which is equal to 1 tonne. Therefore, the mass of the car in tonnes is 1.

Summary of Calculation Steps

Step	Description	Calculation
1	Convert ΔKE to Joules	$-200 \text{ kJ} = -200,000 \text{ J}$
2	Write formula for ΔKE	$\Delta KE = \frac{1}{2}m(v_{final}^2 - v_{initial}^2)$
3	Substitute known values	$-200,000 = \frac{1}{2}m((15)^2 - (25)^2)$
4	Calculate squared speeds	$15^2 = 225, 25^2 = 625$
5	Substitute squared speeds	$-200,000 = \frac{1}{2}m(225 - 625)$
6	Calculate difference	$225 - 625 = -400$
7	Simplify equation	$-200,000 = \frac{1}{2}m(-400) = -200m$
8	Solve for mass (m) in kg	$m = \frac{-200,000}{-200} = 1000 \text{ kg}$
9	Convert mass to tonnes	$1000 \text{ kg} \times \frac{1 \text{ tonne}}{1000 \text{ kg}} = 1 \text{ tonne}$

Revision Table: Key Concepts

Physics Concepts for Car Kinetic Energy

Concept	Definition/Formula	Units
Kinetic Energy (KE)	Energy of motion, $KE = \frac{1}{2}mv^2$	Joules (J)
Mass (m)	Measure of inertia	Kilograms (kg), Tonnes
Speed (v)	Magnitude of velocity	Meters per second (m/s)
Change in KE (ΔKE)	$KE_{final} - KE_{initial}$	Joules (J)
Energy Loss	Indicated by negative ΔKE when speed decreases	Joules (J)
Unit Conversion	1 tonne = 1000 kg	-

Additional Information on Kinetic Energy Problems

Problems involving changes in kinetic energy often relate to the work-energy theorem, which states that the net work done on an object is equal to the change in its kinetic energy ($W_{net} = \Delta KE$). In this specific car problem, the loss of kinetic energy could be due to work done by friction, air resistance, or braking forces. However, the question provides the change in KE directly, simplifying the calculation needed to find the mass.

It's important to pay attention to units throughout the calculation. Using SI units (Joules for energy, kilograms for mass, meters per second for speed) consistently makes the calculation correct. The final answer might require conversion to a different unit, like tonnes in this case, as specified by the question.

When dealing with squares of speeds, remember that $(v)^2$ is always positive, but $(v_{final}^2 - v_{initial}^2)$ can be negative if v_{final} is less than $v_{initial}$, indicating a decrease in speed and a loss of kinetic energy.

63. Answer: c

Explanation:

Understanding Operator Substitution Problems

This type of question involves changing the meaning of standard mathematical operators according to a given set of rules. After substituting the operators based on the rules, we need to evaluate the new expression using the standard order of operations (BODMAS/PEMDAS).

Step 1: Identify the Operator Substitution Rules

The question provides the following substitution rules:

- '+' represents 'x'
- '÷' represents '+'
- '-' represents '÷'
- 'x' represents '-'

Step 2: Write Down the Original Expression

The expression given is:

$$8 + 4 \div 15 - 3 \times 1$$

Step 3: Substitute the Operators According to the Rules

Now, we will replace each original operator in the expression with its new meaning as per the rules:

- Replace '+' with 'x'
- Replace '÷' with '+'
- Replace '-' with '÷'
- Replace 'x' with '-'

Applying these rules to the original expression $8 + 4 \div 15 - 3 \times 1$:

- 8 (replace +) becomes $8 \times$
- 4 (replace $\div$) becomes $4 +$
- 15 (replace -) becomes $15 \div$
- 3 (replace $\times$) becomes $3 -$
- 1 remains 1

The new expression after substitution is:

$$8 \times 4 + 15 \div 3 - 1$$

Step 4: Evaluate the Substituted Expression Using BODMAS/PEMDAS

We evaluate the new expression $8 \times 4 + 15 \div 3 - 1$ following the order of operations:

BODMAS/PEMDAS Order:

- Brackets / Parentheses
- Orders / Exponents
- Division and Multiplication (from left to right)
- Addition and Subtraction (from left to right)

Let's evaluate step-by-step:

1. **Division:** First, calculate $15 \div 3$.

$$15 \div 3 = 5$$

The expression becomes: $8 \times 4 + 5 - 1$

2. **Multiplication:** Next, calculate 8×4 .

$$8 \times 4 = 32$$

The expression becomes: $32 + 5 - 1$

3. **Addition and Subtraction:** Finally, perform addition and subtraction from left to right.

$$32 + 5 = 37$$

Then, $37 - 1 = 36$

The final value of the expression is 36.

Conclusion on Operator Substitution

By correctly substituting the operators according to the given rules and then applying the standard order of operations (BODMAS/PEMDAS), we found the value of the expression $8 + 4 \div 15 - 3 \times 1$ with the new rules to be 36.

Revision Table: Operator Substitution Rules

Original Operator	Represents (New Operator)
+	$\times$
$\div$	+
-	$\div$
$\times$	-

Additional Information: Understanding BODMAS/PEMDAS Order of Operations

When solving mathematical expressions, it's crucial to follow a specific order to ensure consistency and accuracy. This order is commonly known as BODMAS or PEMDAS.

- **BODMAS:** Brackets, Orders (powers, square roots), Division, Multiplication, Addition, Subtraction.
- **PEMDAS:** Parentheses, Exponents, Multiplication, Division, Addition, Subtraction.

Both acronyms represent the same order of priority. Division and Multiplication have equal priority and are performed from left to right as they appear in the expression. Similarly, Addition and Subtraction have equal priority and are performed from left to right.

Failing to follow this order will likely result in an incorrect answer when evaluating expressions.

64. Answer: d

Explanation:

Understanding the Question: Finding the Different Letter Cluster

The question asks us to identify the letter cluster among the given options that follows a pattern different from the other clusters. This is a common type of verbal reasoning question where we need to find the underlying rule or sequence governing each group of letters.

Analyzing the Given Letter Clusters

Let's examine each letter cluster provided in the options and observe the relationship between the letters within each cluster. We will look at their positions in the standard English alphabet.

- **Option 1: RST**

These are three consecutive letters in the English alphabet: R is followed immediately by S, and S is followed immediately by T.

- **Option 2: KLM**

These are three consecutive letters in the English alphabet: K is followed immediately by L, and L is followed immediately by M.

- **Option 3: EFG**

These are three consecutive letters in the English alphabet: E is followed immediately by F, and F is followed immediately by G.

- Option 4: MOQ

Let's look at the sequence: M is followed by O. The letter skipped is N. So, it's $M \rightarrow (N) \rightarrow O$. Then, O is followed by Q. The letter skipped is P. So, it's $O \rightarrow (P) \rightarrow Q$.

Identifying the Pattern and the Different Cluster

From the analysis above, we can see a clear pattern in the first three options:

- RST: Each letter is immediately followed by the next letter in the alphabet (a difference of +1 position).
- KLM: Each letter is immediately followed by the next letter in the alphabet (a difference of +1 position).
- EFG: Each letter is immediately followed by the next letter in the alphabet (a difference of +1 position).

However, the fourth option, MOQ, does not follow this pattern:

- MOQ: M is followed by O (skipping N, a difference of +2 positions), and O is followed by Q (skipping P, a difference of +2 positions).

Therefore, the letter cluster **MOQ** is different from the rest because the letters within it are separated by one letter each, while the letters in the other clusters are consecutive.

Conclusion

The letter cluster that is different from the rest based on the pattern of consecutive letters is MOQ.

Revision Table: Comparing Letter Cluster Patterns

Letter Cluster	Sequence Pattern	Difference in Alphabetical Position	Follows "+1 Consecutive" Pattern?
RST	$R \rightarrow S \rightarrow T$	R to S: +1, S to T: +1	Yes
KLM	$K \rightarrow L \rightarrow M$	K to L: +1, L to M: +1	Yes
EFG	$E \rightarrow F \rightarrow G$	E to F: +1, F to G: +1	Yes
MOQ	$M \rightarrow O \rightarrow Q$	M to O: +2 (skips N), O to Q: +2 (skips P)	No

Additional Information: Types of Letter Series Patterns

Letter series and letter cluster questions in reasoning tests often involve patterns based on:

- **Consecutive letters:** Letters immediately following each other (like RST).
- **Skipping letters:** Letters are separated by a fixed number of skipped letters (like MOQ, skipping one letter each time).
- **Alphabetical position:** Patterns based on the numerical position of letters in the alphabet (A=1, B=2, etc.). For example, a sequence might add a constant number to the position of the previous letter.
- **Reverse alphabetical order:** Letters moving backward in the alphabet.
- **Combination of patterns:** Sometimes, a series might combine different rules, like increasing skips or alternating between skipping and consecutive letters.
- **Vowels and consonants:** Patterns might be based on whether letters are vowels or consonants.

Solving these questions requires carefully observing the relationship between letters and testing different possible rules.

65. Answer: c

Explanation:

Calculating the Mass of Kerosene in a Rectangular Tank

This problem requires us to find the mass of kerosene that fills a rectangular tank. We are given the dimensions of the tank and the density of the kerosene. To solve this, we first need to calculate the volume of the tank and then use the definition of density to find the mass.

Understanding the Concepts: Density, Mass, and Volume

Density is a physical property of a substance that relates its mass to its volume. It is defined as mass per unit volume. The formula connecting these quantities is:

$$\text{Density} = \frac{\text{Mass}}{\text{Volume}}$$

From this relationship, we can rearrange the formula to find the mass:

$$\text{Mass} = \text{Density} \times \text{Volume}$$

Calculating the Volume of the Tank

The tank is rectangular with dimensions 5 m × 2 m × 1 m. The volume of a rectangular prism (like this tank) is calculated by multiplying its length, width, and height.

- Length (l) = 5 m
- Width (w) = 2 m
- Height (h) = 1 m

$$\text{Volume (V)} = \text{Length} \times \text{Width} \times \text{Height}$$

$$V = 5 \text{ m} \times 2 \text{ m} \times 1 \text{ m}$$

$$V = 10 \text{ m}^3$$

So, the volume of the tank is 10 cubic meters.

Calculating the Mass of Kerosene

We are given the density of kerosene as 800 kg/m^3 . We have calculated the volume of the tank as 10 m^3 . Now we can use the mass formula:

$$\text{Mass} = \text{Density} \times \text{Volume}$$

- Density of kerosene = 800 kg/m^3
- Volume of tank = 10 m^3

$$\text{Mass} = 800 \text{ kg/m}^3 \times 10 \text{ m}^3$$

$$\text{Mass} = 8000 \text{ kg}$$

Therefore, the mass of kerosene filled up to the brim in the tank is 8000 kg.

Step-by-Step Solution

1. Identify the given values: Dimensions of the tank (length, width, height) and the density of kerosene.
2. Calculate the volume of the rectangular tank using the formula: $\text{Volume} = \text{length} \times \text{width} \times \text{height}$.
3. Use the density formula ($\text{Density} = \text{Mass}/\text{Volume}$) to find the mass by rearranging it to: $\text{Mass} = \text{Density} \times \text{Volume}$.
4. Substitute the calculated volume and given density into the mass formula and compute the result.

Revision Table: Key Quantities and Formulas

Quantity	Value	Formula/Concept
Tank Length	5 m	Given
Tank Width	2 m	Given
Tank Height	1 m	Given
Tank Volume	10 m^3	$\text{Length} \times \text{Width} \times \text{Height}$
Kerosene Density	800 kg/m^3	Given
Kerosene Mass	8000 kg	$\text{Density} \times \text{Volume}$

Additional Information on Density and Units

Density is an intensive property, meaning it does not depend on the amount of substance. The units of density are typically kg/m^3 or g/cm^3 . It is important to ensure that the units of mass, volume, and density are consistent when using the formula. In this problem, density is given in kg/m^3 and dimensions in meters, resulting in volume in m^3 . Therefore, the resulting mass is in kilograms (kg), which is the required unit.

Understanding density helps us compare how much 'stuff' is packed into a certain space for different materials. Kerosene is less dense than water (approx 1000 kg/m^3), which is why it floats on water.

66. Answer: a

Explanation:

Understanding the Banana Vendor Profit Problem

This problem involves calculating the percentage profit made by a vendor who buys bananas at a certain rate and sells them at a different rate. To find the percentage profit, we need to determine the cost price (CP) and the selling price (SP) for the same number of bananas.

Calculating Cost Price (CP) and Selling Price (SP)

The vendor buys bananas at the rate of 8 pieces for ₹5. This means the cost price of 8 bananas is ₹5.

The vendor sells bananas at the rate of 5 pieces for ₹8. This means the selling price of 5 bananas is ₹8.

To compare the buying and selling rates and calculate profit, it's easiest to find the CP and SP of a common number of bananas. A good number to choose is the Least Common Multiple (LCM) of the number of bananas in the buy and sell transactions, which are 8 and 5. The LCM of 8 and 5 is 40.

Cost Price (CP) for 40 Bananas:

Given: CP of 8 bananas = ₹5

CP of 1 banana = ₹ $\frac{5}{8}$

CP of 40 bananas = ₹ $(\frac{5}{8} \times 40) = ₹(5 \times 5) = ₹25$

Selling Price (SP) for 40 Bananas:

Given: SP of 5 bananas = ₹8

SP of 1 banana = ₹ $\frac{8}{5}$

SP of 40 bananas = ₹ $(\frac{8}{5} \times 40) = ₹(8 \times 8) = ₹64$

Calculating the Profit

Profit is calculated as the difference between the Selling Price and the Cost Price.

Profit = SP - CP

Profit = ₹64 - ₹25 = ₹39

The vendor makes a profit of ₹39 on 40 bananas.

Calculating the Percentage Profit

The percentage profit is calculated using the formula:

Percentage Profit = $(\frac{\text{Profit}}{\text{CP}}) \times 100\%$

Using the values we found:

Percentage Profit = $(\frac{39}{25}) \times 100\%$

$$\text{Percentage Profit} = 39 \times \left(\frac{100}{25}\right) \%$$

$$\text{Percentage Profit} = 39 \times 4\%$$

$$\text{Percentage Profit} = 156\%$$

Thus, the vendor's percentage profit is 156%.

Step-by-Step Solution Summary

1. Identify the buy rate: 8 bananas for ₹5.
2. Identify the sell rate: 5 bananas for ₹8.
3. Find the LCM of the number of items (8 and 5), which is 40.
4. Calculate the Cost Price (CP) for 40 bananas: $\left(\frac{5}{8}\right) \times 40 = ₹25$.
5. Calculate the Selling Price (SP) for 40 bananas: $\left(\frac{8}{5}\right) \times 40 = ₹64$.
6. Calculate the Profit: $SP - CP = ₹64 - ₹25 = ₹39$.
7. Calculate the Percentage Profit: $\left(\frac{\text{Profit}}{\text{CP}}\right) \times 100\% = \left(\frac{39}{25}\right) \times 100\% = 156\%$.

Item	Quantity	Cost/Selling Price
Bananas (Buy)	8	₹5 (CP)
Bananas (Sell)	5	₹8 (SP)
Bananas (Common Quantity: 40)	40	₹25 (Calculated CP)
Bananas (Common Quantity: 40)	40	₹64 (Calculated SP)

Revision Table: Profit and Loss Concepts

Concept	Formula/Explanation
Cost Price (CP)	The price at which an article is bought.
Selling Price (SP)	The price at which an article is sold.
Profit	$SP - CP$ (when $SP > CP$)
Loss	$CP - SP$ (when $CP > SP$)
Profit Percentage	$\left(\frac{\text{Profit}}{CP}\right) \times 100\%$
Loss Percentage	$\left(\frac{\text{Loss}}{CP}\right) \times 100\%$

Additional Information: Solving Buy/Sell Problems

Problems involving buying and selling items at different rates can often be simplified by finding the cost and selling price for the same number of items. The LCM of the quantities involved is a useful quantity to choose. This approach allows direct comparison of expenditure (CP) and revenue (SP) for an equal volume of trade.

Alternatively, you can calculate the CP and SP per single item and then use those values to find the profit percentage.

- CP per banana = ₹5 / 8
- SP per banana = ₹8 / 5
- Profit per banana = SP per banana - CP per banana = ₹8/5 - ₹5/8
- To subtract fractions, find a common denominator (LCM of 5 and 8 is 40):
- Profit per banana = ₹(64/40 - 25/40) = ₹39/40
- Percentage Profit = $\left(\frac{\text{Profit per banana}}{CP \text{ per banana}}\right) \times 100\%$
- Percentage Profit = $\left(\frac{39/40}{5/8}\right) \times 100\%$
- Percentage Profit = $\left(\frac{39}{40} \times \frac{8}{5}\right) \times 100\%$
- Percentage Profit = $\left(\frac{39}{5} \times \frac{1}{5}\right) \times 100\%$
- Percentage Profit = $\left(\frac{39}{25}\right) \times 100\% = 39 \times 4\% = 156\%$

Both methods yield the same percentage profit.

67. Answer: d

Explanation:

Understanding Minerals in Grains, Nuts, and Chocolates

Our daily diet includes various foods that provide essential minerals and nutrients. Grains, nuts, and chocolates are popular food items known for their energy content, but they also contribute significantly to our mineral intake.

Let's look at the question asking which element is mostly provided by grains, nuts, and chocolates in our diet and analyze the given options.

Analyzing the Element Contribution of Foods

Different foods are rich in different elements. We need to identify the element that is notably present in grains, nuts, and chocolates.

- **Grains:** Whole grains, in particular, are good sources of various minerals, including B vitamins, magnesium, iron, and certain trace minerals.
- **Nuts:** Nuts like almonds, walnuts, cashews, and peanuts are nutrient powerhouses, providing healthy fats, protein, fiber, vitamins, and minerals such as magnesium, zinc, selenium, and copper.
- **Chocolates:** Dark chocolate is especially known for its mineral content, including iron, magnesium, zinc, selenium, and copper.

Considering these common food sources, let's evaluate the options provided:

- **Chloride:** Chloride is an electrolyte, usually consumed as part of sodium chloride (table salt). While present in many foods, it is not typically highlighted as a primary mineral sourced *mostly* from grains, nuts, and chocolates compared to other sources.
- **Zinc:** Zinc is present in nuts and some grains (especially whole grains), as well as in chocolate. It is an important mineral found in these foods.
- **Iodine:** Iodine is essential for thyroid function and is primarily found in seafood, dairy products, and iodized salt. Grains, nuts, and chocolates are generally not

considered primary sources of iodine unless fortified or grown in iodine-rich soil.

- **Copper:** Copper is a trace mineral vital for various bodily functions. Grains (especially whole grains), nuts (like cashews, almonds, and walnuts), and dark chocolate are all recognized as good sources of copper. In fact, nuts and chocolate are particularly notable for their copper content.

Conclusion on Element Source

Based on the analysis of the mineral content of grains, nuts, and chocolates, **Copper** is the element that is significantly provided by all three food groups. While Zinc is also present, Copper is often highlighted as being particularly rich in nuts and dark chocolate, making it a strong candidate for the element mostly provided by this combination of foods.

Revision Table: Key Mineral Sources

Element	Common Food Sources	Found in Grains, Nuts, Chocolate?
Chloride	Table Salt, Processed Foods	Present, but not primary source
Zinc	Meat, Legumes, Seeds, Nuts, Whole Grains, Chocolate	Yes (Nuts, Whole Grains, Chocolate)
Iodine	Seafood, Dairy, Iodized Salt	Generally low or absent
Copper	Organ Meats, Seafood, Nuts, Seeds, Whole Grains, Chocolate	Yes (Nuts, Whole Grains, Chocolate)

Additional Information: Importance of Trace Minerals

Trace minerals, like copper and zinc, are needed by the body in smaller amounts compared to major minerals, but they are crucial for health. They play roles in

enzyme function, immune system support, energy metabolism, and antioxidant defense.

- **Copper:** Important for iron metabolism, nerve function, bone health, and collagen formation.
- **Zinc:** Essential for immune function, wound healing, DNA synthesis, and growth.

Ensuring a balanced diet that includes a variety of whole foods like grains, nuts, seeds, fruits, and vegetables helps in obtaining these essential trace minerals.

68. Answer: a

Explanation:

Understanding A0 Paper Size Area Calculation

The question asks about the area of A0 size paper. Paper sizes like A0, A1, A2, etc., are defined by the international standard ISO 216. This standard is widely used around the world.

ISO 216 Standard and A0 Definition

The ISO 216 standard is based on the German DIN 476 standard. The key principle behind this standard is that if you cut a sheet of paper in half along its longest side, the two resulting smaller sheets have the same aspect ratio as the original sheet. The aspect ratio is $1 : \sqrt{2}$.

The base size in this series is A0. A0 is defined as having an area of 1 square meter ($1 m^2$).

Converting Area from Square Meters to Square Centimeters

We need to find the area in square centimeters (cm^2). We know the relationship between meters and centimeters:

$$1\text{ m} = 100\text{ cm}$$

To convert square meters to square centimeters, we square the conversion factor:

$$1\text{ m}^2 = (1\text{ m}) \times (1\text{ m})$$

Substituting the equivalent in centimeters:

$$1\text{ m}^2 = (100\text{ cm}) \times (100\text{ cm})$$

$$1\text{ m}^2 = 10000\text{ cm}^2$$

Comparing Calculated Area to Options

So, the area of A0 size paper is 1 m^2 , which is equal to $10,000\text{ cm}^2$.

Let's look at the options provided:

- $10,000\text{ cm}^2$
- 100 cm^2
- $1,000\text{ cm}^2$
- 1 cm^2

Our calculated area of $10,000\text{ cm}^2$ matches the first option.

Conclusion on A0 Paper Area

The area of A0 size paper is indeed $10,000\text{ cm}^2$, based on its definition in the ISO 216 standard as having an area of 1 square meter.

Common A Series Paper Sizes
and Areas

Size	Area (m^2)	Area (cm^2)
A0	1.0	10,000
A1	0.5	5,000
A2	0.25	2,500
A3	0.125	1,250
A4	0.0625	625

Revision Table: Key Paper Area Facts

- **A0 Paper Definition:** Area is exactly 1 square meter.
- **Conversion Factor:** $1\text{ m} = 100\text{ cm}$.
- **Area Conversion:** $1\text{ m}^2 = 10000\text{ cm}^2$.
- **ISO 216 Aspect Ratio:** $1 : \sqrt{2}$.

Additional Information on Paper Sizes

The dimensions of A0 paper are approximately $841\text{ mm} \times 1189\text{ mm}$. The area calculated from these dimensions will be very close to $841 \times 1189 = 999949\text{ mm}^2$. Since $1\text{ cm} = 10\text{ mm}$, $1\text{ cm}^2 = (10\text{ mm}) \times (10\text{ mm}) = 100\text{ mm}^2$. So, $999949\text{ mm}^2 \approx 10000\text{ cm}^2$. The standard defines the area precisely as 1 m^2 or 10000 cm^2 , and the dimensions are derived to meet this area and the $1 : \sqrt{2}$ aspect ratio as closely as possible while being rounded to the nearest millimeter.

Each subsequent A series size (A1, A2, A3, etc.) has exactly half the area of the previous size.

- Area of A1 = $1/2$ Area of A0 = $0.5\text{ m}^2 = 5000\text{ cm}^2$.
- Area of A2 = $1/2$ Area of A1 = $0.25\text{ m}^2 = 2500\text{ cm}^2$.

- Area of A4, the most common paper size, is $\frac{1}{16}$ Area of A0 = $\frac{1}{16} m^2 = 0.0625 m^2 = 625 cm^2$.

69. Answer: d

Explanation:

Calculating Possible Total Weights from Box Combinations

We are given three boxes with specific weights: 3 kg, 8 kg, and 12 kg. The question asks which of the given options cannot be the total weight obtained by combining any selection of these boxes.

When we talk about a "combination of these boxes", it means we can choose to include one or more of the three distinct boxes in our total weight calculation. We find the total weight by adding up the weights of the boxes we choose.

Possible Combinations and Their Total Weights

Let's list all the ways we can combine these three boxes and calculate the resulting total weight. We can choose one box, two boxes, or all three boxes.

Number of Boxes	Boxes Combined	Calculation	Total Weight (kg)
1	The 3 kg box	3	3
1	The 8 kg box	8	8
1	The 12 kg box	12	12
2	The 3 kg and 8 kg boxes	$3 + 8$	11
2	The 3 kg and 12 kg boxes	$3 + 12$	15
2	The 8 kg and 12 kg boxes	$8 + 12$	20
3	The 3 kg, 8 kg, and 12 kg boxes	$3 + 8 + 12$	23

The possible total weights we can get by combining these boxes are 3 kg, 8 kg, 11 kg, 12 kg, 15 kg, 20 kg, and 23 kg.

Checking the Given Options

Now let's look at the given options for total weight and see if they appear in our list of possible total weights:

- Is 15 kg a possible total weight? Yes, it is obtained by combining the 3 kg and 12 kg boxes ($3 + 12 = 15$).
- Is 20 kg a possible total weight? Yes, it is obtained by combining the 8 kg and 12 kg boxes ($8 + 12 = 20$).
- Is 23 kg a possible total weight? Yes, it is obtained by combining all three boxes ($3 + 8 + 12 = 23$).
- Is 21 kg a possible total weight? Let's check our list: 3, 8, 11, 12, 15, 20, 23. The weight 21 kg is not in this list.

Conclusion on Possible and Impossible Total Weights

Based on the possible combinations of the 3 kg, 8 kg, and 12 kg boxes, the total weights that can be formed are 3 kg, 8 kg, 11 kg, 12 kg, 15 kg, 20 kg, and 23 kg.

The options provided are 15 kg, 20 kg, 23 kg, and 21 kg. We found that 15 kg, 20 kg, and 23 kg can all be formed.

The weight that **CANNOT** be the total weight of any combination of these boxes is 21 kg.

Revision Table: Box Combination Weights

Option (Total Weight in kg)	Can it be formed?	Combination (if possible)
15	Yes	3 kg + 12 kg
20	Yes	8 kg + 12 kg
23	Yes	3 kg + 8 kg + 12 kg
21	No	Cannot be formed from 3 kg, 8 kg, and 12 kg boxes

Additional Information on Weight Combinations

This problem involves finding sums from a given set of numbers. In more complex scenarios, this is related to the subset sum problem. Here, with only three distinct items, we can simply list all possible non-empty subsets and sum their elements.

Understanding how to find all possible sums from a set is useful in various mathematical and computer science problems.

For a set of n distinct items, there are $2^n - 1$ non-empty subsets (combinations) to consider. In this case, $n = 3$, so there are $2^3 - 1 = 8 - 1 = 7$ non-empty combinations.

The combinations are: $\{3\}$, $\{8\}$, $\{12\}$, $\{3, 8\}$, $\{3, 12\}$, $\{8, 12\}$, $\{3, 8, 12\}$.

Their sums are: 3, 8, 12, 11, 15, 20, 23.

Any weight that is not in this set of sums cannot be formed using these exact three boxes.

70. Answer: d

Explanation:

Finding Mechanical Advantage of a Pulley System

Let's determine the mechanical advantage of the pulley system using the given information: its efficiency, the distance the load is lifted, and the distance the rope is pulled by the effort.

Understanding Key Terms

- **Efficiency (η):** This is the ratio of the useful work done by the machine (output work) to the total work done on the machine (input work), usually expressed as a percentage. For a pulley system, it's also the ratio of the Mechanical Advantage (MA) to the Velocity Ratio (VR).
- **Velocity Ratio (VR):** This is the ratio of the distance moved by the effort to the distance moved by the load. It is determined purely by the geometry of the pulley system.
- **Mechanical Advantage (MA):** This is the ratio of the load lifted (output force) to the effort applied (input force). It tells us how much the machine multiplies the applied force.

Given Information

From the problem statement, we have:

- Efficiency (η) = 60% = 0.60 (as a decimal)
- Distance the load is lifted (d_L) = 3 m
- Distance the rope is pulled by the effort (d_E) = 12 m

Step 1: Calculate the Velocity Ratio (VR)

The Velocity Ratio is calculated using the distances moved by the effort and the load:

$$VR = \frac{\text{Distance moved by Effort}}{\text{Distance moved by Load}} = \frac{d_E}{d_L}$$

Substituting the given values:

$$VR = \frac{12\text{ m}}{3\text{ m}} = 4$$

So, the Velocity Ratio of the pulley system is 4.

Step 2: Use the Efficiency Formula to Find Mechanical Advantage (MA)

The efficiency of a machine relates Mechanical Advantage and Velocity Ratio:

$$\eta = \frac{MA}{VR} \times 100\%$$

When efficiency is given as a decimal (as we converted 60% to 0.60), the formula is:

$$\eta = \frac{MA}{VR}$$

We want to find MA, so we can rearrange the formula:

$$MA = \eta \times VR$$

Now, substitute the given efficiency (as a decimal) and the calculated Velocity Ratio:

$$MA = 0.60 \times 4$$

$$MA = 2.4$$

Conclusion

The mechanical advantage of the pulley system is 2.4.

Quantity	Symbol	Value
Efficiency	η	60% or 0.60
Load Distance	d_L	3 m
Effort Distance	d_E	12 m
Velocity Ratio	VR	4
Mechanical Advantage	MA	2.4

Revision Table: Pulley System Calculations

Concept	Formula	Description
Velocity Ratio (VR)	$VR = \frac{\text{Distance moved by Effort}}{\text{Distance moved by Load}}$	Depends on the pulley system's structure (number of ropes supporting the load).
Actual Mechanical Advantage (AMA)	$AMA = \frac{\text{Load}}{\text{Effort}}$	The real ratio of load to effort, considers friction and other losses.
Ideal Mechanical Advantage (IMA)	$IMA = VR$	The theoretical mechanical advantage assuming no friction or losses.
Efficiency (η)	$\eta = \frac{AMA}{IMA} \times 100\%$ or $\eta = \frac{\text{Output Work}}{\text{Input Work}} \times 100\%$	Ratio of AMA to IMA, or output work to input work. Always less than 100% for real machines.

Additional Information: Pulley System Efficiency

Efficiency is a crucial concept for understanding how real machines perform compared to ideal ones. For a pulley system, efficiency is typically less than 100% because of several factors:

- **Friction:** Friction exists in the pulley axles and bearings, opposing motion and requiring extra effort.
- **Weight of the Pulley(s):** The effort must not only lift the load but also the weight of the movable pulleys in the system.
- **Stretching of the Rope:** While often negligible, the rope can stretch slightly under load, affecting the distances moved.

A higher efficiency means that the actual mechanical advantage (AMA) is closer to the ideal mechanical advantage (IMA), meaning less effort is wasted on overcoming friction and lifting the pulleys themselves.

In ideal scenarios (no friction, massless pulleys), efficiency is 100%, and the mechanical advantage would equal the velocity ratio. However, in real-world applications, efficiency is always less than 100%, resulting in the actual mechanical advantage being less than the velocity ratio.

71. Answer: b

Explanation:

Calculating the Area of a Regular Hexagon

Let's find the area of a regular hexagon with a given side length. A regular hexagon is a six-sided polygon where all sides are equal in length and all interior angles are equal. A key property of a regular hexagon is that it can be divided into six identical equilateral triangles meeting at the center.

Understanding the Geometry

When a regular hexagon has a side length, say ' s ', each of the six equilateral triangles that form it also has a side length equal to ' s '. To find the total area of the hexagon, we can calculate the area of one of these equilateral triangles and multiply it by six.

Formula for Area of an Equilateral Triangle

The area of an equilateral triangle with side length 's' is given by the formula:

$$\text{Area of equilateral triangle} = \frac{\sqrt{3}}{4}s^2$$

Applying the Formula

In this problem, the side length of the regular hexagon is given as 10 cm. Since the hexagon is composed of six equilateral triangles, each with a side length of 10 cm, we can substitute $s = 10$ cm into the formula for the area of an equilateral triangle.

$$\text{Area of one equilateral triangle} = \frac{\sqrt{3}}{4}(10 \text{ cm})^2$$

$$\text{Area of one equilateral triangle} = \frac{\sqrt{3}}{4} \times 100 \text{ cm}^2$$

$$\text{Area of one equilateral triangle} = 25\sqrt{3} \text{ cm}^2$$

Calculating the Area of the Hexagon

Since the regular hexagon consists of six such equilateral triangles, the total area of the hexagon is 6 times the area of one triangle.

$$\text{Area of regular hexagon} = 6 \times (\text{Area of one equilateral triangle})$$

$$\text{Area of regular hexagon} = 6 \times 25\sqrt{3} \text{ cm}^2$$

$$\text{Area of regular hexagon} = 150\sqrt{3} \text{ cm}^2$$

Thus, the area of a regular hexagon with a side length of 10 cm is $150\sqrt{3} \text{ cm}^2$.

Summary of Steps

1. Recognize that a regular hexagon is made of 6 equilateral triangles.
2. Identify the side length of the equilateral triangles (same as the hexagon's side).
3. Use the formula for the area of an equilateral triangle: $\frac{\sqrt{3}}{4}s^2$.
4. Calculate the area of one equilateral triangle with the given side length.
5. Multiply the area of one triangle by 6 to get the total area of the hexagon.

Area Calculation Summary

Parameter	Value
Side length of hexagon (s)	10 cm
Number of equilateral triangles	6
Area of one equilateral triangle	$\frac{\sqrt{3}}{4}s^2 = \frac{\sqrt{3}}{4}(10)^2 = 25\sqrt{3} \text{ cm}^2$
Area of hexagon	$6 \times 25\sqrt{3} = 150\sqrt{3} \text{ cm}^2$

Revision Table: Area of Regular Hexagon

Key Formulas for Hexagon Area

Concept	Formula	Notes
Area of Equilateral Triangle (side s)	$\frac{\sqrt{3}}{4}s^2$	Building block for regular hexagon
Area of Regular Hexagon (side s)	$6 \times \frac{\sqrt{3}}{4}s^2 = \frac{3\sqrt{3}}{2}s^2$	Derived from 6 equilateral triangles

Your Personal Exams Guide

Additional Information: Properties of Regular Hexagons

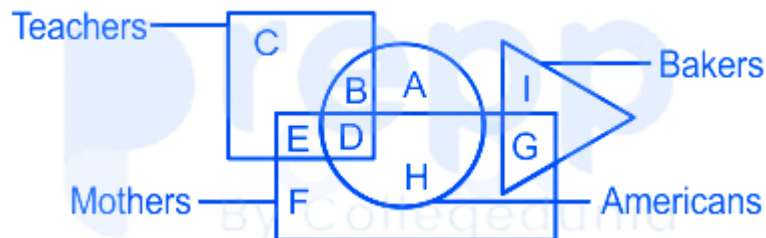
A regular hexagon is a fascinating shape with several interesting properties:

- It has 6 equal sides and 6 equal interior angles (120° each).
- The sum of interior angles is $(6 - 2) \times 180^\circ = 720^\circ$.
- It has 9 diagonals. The longest diagonals pass through the center and are twice the length of the side.
- It is a cyclic polygon (all vertices lie on a circle). The radius of the circumcircle is equal to the side length.
- It is a tangential polygon (all sides are tangent to a circle). The radius of the incircle (apothem) is $\frac{\sqrt{3}}{2}s$.

72. Answer: b

Explanation:

The letter or set of letters which represents the mothers who are neither Americans nor bakers, is shown below:



Hence, 'EF' is the correct answer.

73. Answer: c

Explanation:

Understanding Density Calculation

Density is a fundamental property of matter that tells us how much mass is packed into a certain volume. It is defined as the ratio of mass to volume.

The formula for density (ρ) is:

$$\rho = \frac{m}{V}$$

where:

- m is the mass of the object
- V is the volume of the object

In this problem, we are given the dimensions of a piece of wood and its weight. We need to find its density in kilograms per cubic meter (kg/m^3).

Calculating the Volume of the Wood

The piece of wood is rectangular, so its volume can be calculated by multiplying its length, width, and height. The dimensions are given in centimeters (cm), but we need the volume in cubic meters (m^3). We must first convert the dimensions from cm to meters (m).

Conversion factor: $1 \text{ cm} = 0.01 \text{ m}$

- Length (l) = $6 \text{ cm} = 6 \times 0.01 \text{ m} = 0.06 \text{ m}$
- Width (w) = $8 \text{ cm} = 8 \times 0.01 \text{ m} = 0.08 \text{ m}$
- Height (h) = $5 \text{ cm} = 5 \times 0.01 \text{ m} = 0.05 \text{ m}$

Now, calculate the volume (V):

$$\begin{aligned}
 V &= l \times w \times h \\
 V &= 0.06 \text{ m} \times 0.08 \text{ m} \times 0.05 \text{ m} \\
 V &= (6 \times 10^{-2}) \text{ m} \times (8 \times 10^{-2}) \text{ m} \times (5 \times 10^{-2}) \text{ m} \\
 V &= (6 \times 8 \times 5) \times 10^{(-2-2-2)} \text{ m}^3 \\
 V &= 240 \times 10^{-6} \text{ m}^3 \\
 V &= 2.4 \times 10^{-4} \text{ m}^3
 \end{aligned}$$

So, the volume of the piece of wood is $2.4 \times 10^{-4} \text{ m}^3$.

Calculating the Mass of the Wood from Weight

We are given the weight of the wood, which is 1.92 N . Weight is the force exerted on an object due to gravity. It is related to mass by the formula:

$$\text{Weight}(W) = \text{mass}(m) \times \text{acceleration due to gravity}(g)$$

We are given $W = 1.92 \text{ N}$ and $g = 10 \text{ m/s}^2$. We can rearrange the formula to find the mass (m):

$$m = \frac{W}{g}$$

$$m = \frac{1.92 \text{ N}}{10 \text{ m/s}^2}$$

$$m = 0.192 \text{ kg}$$

Thus, the mass of the piece of wood is 0.192 kg.

Calculating the Density of the Wood

Now that we have the mass ($m = 0.192 \text{ kg}$) and the volume ($V = 2.4 \times 10^{-4} \text{ m}^3$), we can calculate the density (ρ) using the density formula:

$$\rho = \frac{m}{V}$$

$$\rho = \frac{0.192 \text{ kg}}{2.4 \times 10^{-4} \text{ m}^3}$$

To make the calculation easier, we can write 0.192 as 192×10^{-3} and 2.4×10^{-4} as 24×10^{-5} (or keep it as is and simplify).

$$\rho = \frac{192 \times 10^{-3}}{2.4 \times 10^{-4}} \text{ kg/m}^3$$

Divide the numerical parts:

$$\frac{192}{2.4} = \frac{1920}{24} = 80$$

Combine the powers of 10:

$$10^{-3}/10^{-4} = 10^{-3-(-4)} = 10^{-3+4} = 10^1$$

So, the density is:

$$\rho = 80 \times 10^1 \text{ kg/m}^3$$

$$\rho = 800 \text{ kg/m}^3$$

The density of the piece of wood is 800 kg/m^3 .

Summary of Calculations for Density

Measurement	Value	Unit Conversion	Value in SI Units
Length (l)	6 cm	$6 \times 0.01 \text{ m}$	0.06 m
Width (w)	8 cm	$8 \times 0.01 \text{ m}$	0.08 m
Height (h)	5 cm	$5 \times 0.01 \text{ m}$	0.05 m
Weight (W)	1.92 N	Given	1.92 N
Gravity (g)	10 m/s ²	Given	10 m/s ²

Calculation	Formula	Values Used	Result
Volume (V)	$l \times w \times h$	0.06 m, 0.08 m, 0.05 m	$2.4 \times 10^{-4} \text{ m}^3$
Mass (m)	W/g	1.92 N, 10 m/s ²	0.192 kg
Density (ρ)	m/V	0.192 kg, $2.4 \times 10^{-4} \text{ m}^3$	800 kg/m ³

Final Answer for Wood Density

The calculated density of the piece of wood is 800 kg/m³.

Revision Table: Key Physics Concepts

Concept	Definition	Formula	Units (SI)
Volume	The amount of space an object occupies. For a cuboid, it's length $\times$ width $\times$ height.	$V = l \times w \times h$	m ³
Weight	The force of gravity on an object.	$W = m \times g$	Newtons (N)
Mass	The amount of matter in an object.	$m = W/g$	kilograms (kg)
Density	Mass per unit volume.	$\rho = m/V$	kg/m ³

Additional Information on Density and Measurement

Density is an intensive property, meaning it does not depend on the amount of substance. A small piece of wood from the same type of tree will have the same density as a large piece, assuming it's uniform.

Units are very important in physics calculations. Always ensure that all values are in consistent units (like SI units – meters, kilograms, seconds) before performing calculations. Conversion factors are necessary when units are mixed.

Weight is a force and is measured in Newtons (N). Mass is a measure of inertia and is measured in kilograms (kg). The relationship between them depends on the local acceleration due to gravity (g). On Earth, g is approximately 9.8 m/s², but for simpler calculations in problems like this, it is often given as 10 m/s².

Understanding how to calculate volume for different shapes (cubes, cylinders, spheres) and how to convert between different units of length, area, and volume is crucial for solving density problems.

74. Answer: a

Explanation:

Finding the Least Common Multiple (LCM) of 56 and 50

The Least Common Multiple (LCM) of two or more numbers is the smallest positive integer that is a multiple of all the given numbers. It's often used when working with fractions or in problems involving cycles or repetitions.

There are several ways to find the LCM, but the prime factorization method is very common and reliable. Let's find the prime factorization of 56 and 50.

Prime Factorization of 56

We break down 56 into its prime factors:

- 56 is divisible by 2: $56 \div 2 = 28$
- 28 is divisible by 2: $28 \div 2 = 14$
- 14 is divisible by 2: $14 \div 2 = 7$
- 7 is a prime number.

So, the prime factorization of 56 is $2 \times 2 \times 2 \times 7$, which can be written in exponential form as $2^3 \times 7^1$.

Prime Factorization of 50

Now, let's break down 50 into its prime factors:

- 50 is divisible by 2: $50 \div 2 = 25$
- 25 is divisible by 5: $25 \div 5 = 5$
- 5 is a prime number.

So, the prime factorization of 50 is $2 \times 5 \times 5$, which can be written in exponential form as $2^1 \times 5^2$.

Calculating the LCM using Prime Factors

To find the LCM of 56 and 50, we take the highest power of each prime factor that appears in either factorization. The prime factors involved are 2, 5, and 7.

- For the prime factor 2: The powers are 2^3 (from 56) and 2^1 (from 50). The highest power is 2^3 .
- For the prime factor 5: The powers are 5^0 (effectively, as 5 does not appear in 56) and 5^2 (from 50). The highest power is 5^2 .
- For the prime factor 7: The powers are 7^1 (from 56) and 7^0 (effectively, as 7 does not appear in 50). The highest power is 7^1 .

Now, we multiply these highest powers together to get the LCM:

$$\text{LCM}(56, 50) = 2^3 \times 5^2 \times 7^1 = 8 \times 25 \times 7$$

Let's calculate the product:

$$8 \times 25 = 200$$

$$200 \times 7 = 1400$$

So, the Least Common Multiple of 56 and 50 is 1400.

Prime Factorization and LCM Calculation

Number	Prime Factorization
56	$2^3 \times 7^1$
50	$2^1 \times 5^2$
$\text{LCM}(56, 50)$	$2^{\max(3,1)} \times 5^{\max(0,2)} \times 7^{\max(0,1)} = 2^3 \times 5^2 \times 7^1 = 8 \times 25 \times 7 = 1400$

Checking the options, we see that 1400 is one of the choices.

Revision Table: Key Concepts in LCM Calculation

Concept	Description	Relevance to LCM
Multiple	A number obtained by multiplying an integer by another integer.	LCM is the smallest common multiple .
Prime Number	A natural number greater than 1 that has no positive divisors other than 1 and itself.	Used in the prime factorization method.
Prime Factorization	Expressing a composite number as a product of its prime factors.	Essential step for calculating LCM using the prime factors method.
Highest Power	The largest exponent for a given prime factor in the factorizations of the numbers.	We use the highest power of each prime factor when calculating the LCM.

Additional Information on Finding LCM and HCF

The LCM is closely related to the Highest Common Factor (HCF) or Greatest Common Divisor (GCD). The HCF is the largest positive integer that divides two or more numbers without leaving a remainder.

There's a useful relationship between the LCM and HCF of two numbers, say 'a' and 'b':

$$\text{LCM}(a, b) \times \text{HCF}(a, b) = a \times b$$

Let's quickly find the HCF of 56 and 50 using their prime factorizations:

$$56 = 2^3 \times 7^1$$

$$50 = 2^1 \times 5^2$$

To find the HCF, we take the lowest power of each common prime factor. The only common prime factor is 2, and the lowest power is 2^1 .

$$\text{HCF}(56, 50) = 2^1 = 2$$

Now, let's check the relationship:

$$\text{LCM}(56, 50) \times \text{HCF}(56, 50) = 1400 \times 2 = 2800$$

$$a \times b = 56 \times 50$$

$$56 \times 50 = 56 \times 5 \times 10 = 280 \times 10 = 2800$$

The relationship holds true: $1400 \times 2 = 56 \times 50$, which is $2800 = 2800$. This confirms our calculation of the LCM is likely correct.

75. Answer: b

Explanation:

Given:

Total surface area = 236 cm^2 , Length = 8 cm, Height = 6 cm.

Formula:

Total surface area = $2(\text{Length} \times \text{Width} + \text{Length} \times \text{Height} + \text{Width} \times \text{Height})$

Evaluation:

$$236 = 2(8 \times B + 6 \times B + 8 \times 6)$$

$$\Rightarrow 118 = 8B + 6B + 48$$

$$\Rightarrow 70 = 14B$$

$$\Rightarrow B = 70/14 = 5 \text{ cm}$$

76. Answer: a

Explanation:

Understanding the Relationship Between Electrical Power, Resistance, Charge, Current, and Time

The question asks for the correct relation connecting electrical power (P), resistance (R), the total charge (Q) flowing through a wire, and the time (t) taken for this charge flow. The options provided also include electric current (I).

To find the correct relationship, we need to use fundamental electrical formulas that relate these quantities. The key concepts are power dissipated in a resistor, Ohm's Law, and the definition of electric current.

Fundamental Electrical Formulas

- Electric current (I) is defined as the rate of flow of electric charge. If charge Q flows through a cross-section of a wire in time t, the current I is given by:

$$I = \frac{Q}{t}$$

- Ohm's Law relates voltage (V), current (I), and resistance (R) in a circuit:

$$V = IR$$

- Electrical power (P) dissipated in a resistor can be expressed in several ways. One common formula related to current and resistance is:

$$P = I^2 R$$

We need to find a relation that involves P, R, Q, and t, potentially also including I as it appears in the options.

Deriving the Correct Relation

Let's start with the power formula that relates power, current, and resistance:

$$P = I^2 R$$

We need to introduce charge (Q) and time (t) into this equation. We know the relationship between current, charge, and time is $I = \frac{Q}{t}$. While we could substitute I directly, let's try multiplying the power equation by time 't':

$$P \times t = I^2 R \times t$$

$$Pt = I^2 Rt$$

Now, let's rearrange the terms on the right side: $I^2 Rt = I \times (I \times t) \times R$.

From the definition of current, we know that $I = \frac{Q}{t}$, which implies that $I \times t = Q$.

Substitute $I \times t = Q$ into the equation $Pt = I \times (I \times t) \times R$:

$$Pt = I \times Q \times R$$

Rearranging the terms on the right side, we get:

$$Pt = IRQ$$

This derived relation connects power (P), time (t), current (I), resistance (R), and charge (Q).

Checking the Options

Let's compare our derived relation $Pt = IRQ$ with the given options:

- Option 1: $Pt = IRQ$ - This matches our derived relation.
- Option 2: $PI = QRt$ - This does not match $Pt = IRQ$.
- Option 3: $PQ = IRt$ - This does not match $Pt = IRQ$.
- Option 4: $PR = QIt$ - This is equivalent to $PR = IQt$, which does not match $Pt = IRQ$.

Therefore, the correct relation is $Pt = IRQ$.

Let's quickly verify the units. Power (P) is in Watts (W), time (t) is in seconds (s). So Pt has units of Joule (J), which is energy. Current (I) is in Amperes (A), Resistance (R) is in Ohms (Ω), Charge (Q) is in Coulombs (C). The unit of IRQ would be $A \times \Omega \times C$. From $V=IR$, $A \times \Omega$ is Volts (V). So IRQ has units $V \times C$. Since Energy = Charge $\times$

Voltage ($E = QV$), the units $V \times C$ are Joules (J). The units match, adding confidence to the relation.

Quantity	Symbol	Standard Unit	Relationship to Others
Power	P	Watt (W)	$P = I^2R, P = VI$
Resistance	R	Ohm (Ω)	$R = V/I$
Charge	Q	Coulomb (C)	$Q = It$
Time	t	second (s)	—
Current	I	Ampere (A)	$I = Q/t, I = V/R$

The derivation shows how the relation $Pt = IRQ$ arises directly from fundamental principles of electricity relating power, current, resistance, charge, and time.

Revision Table: Key Electrical Formulas

Formula Name	Formula	Quantities Involved
Definition of Current	$I = \frac{Q}{t}$	Current (I), Charge (Q), Time (t)
Ohm's Law	$V = IR$	Voltage (V), Current (I), Resistance (R)
Power Dissipation (using I and R)	$P = I^2R$	Power (P), Current (I), Resistance (R)
Power Dissipation (using V and I)	$P = VI$	Power (P), Voltage (V), Current (I)
Power Dissipation (using V and R)	$P = \frac{V^2}{R}$	Power (P), Voltage (V), Resistance (R)
Work Done/Energy Dissipated	$W = Pt$	Work/Energy (W), Power (P), Time (t)

Additional Information: Concepts of Power, Charge, and Resistance

Electric Power (P): Power is the rate at which electrical energy is transferred or dissipated in a circuit. It is measured in Watts (W). In a resistor, electrical energy is converted into heat.

Electric Charge (Q): Charge is a fundamental property of matter. It is carried by particles like electrons (negative charge) and protons (positive charge). The unit of charge is the Coulomb (C).

Electric Current (I): Current is the flow of electric charge. It is typically the flow of electrons in metals. It is measured in Amperes (A). One Ampere is equal to one Coulomb of charge flowing per second ($1 \text{ A} = 1 \text{ C/s}$).

Resistance (R): Resistance is a measure of how much a material opposes the flow of electric current. It is measured in Ohms (Ω). Materials with low resistance conduct current easily, while materials with high resistance oppose current flow.

The relation $Pt = IRQ$ connects the energy dissipated (Pt , since Energy = Power $\times$ Time) with the circuit parameters (I , R) and the total charge flown (Q). While less common than $E = I^2 R t$, it is a valid relation derived from fundamental principles.

Your Personal Exams Guide

77. Answer: a

Explanation:

Let's break down this age word problem step by step to find the mother's present age using the given ratio and the age difference information.

Understanding the Age Problem and Ratio

The problem gives us two main pieces of information about Leela and her mother:

1. The ratio of their present ages is 3 : 8. This means for every 3 years Leela is old, her mother is 8 years old at present.
2. The mother was 25 years old when Leela was born. This tells us the constant age difference between the mother and Leela. The mother is always 25 years older than Leela.

Setting up the Equation for Present Ages

We can represent the present ages using the given ratio. Let the common ratio factor be x .

- Leela's present age = $3x$ years
- Mother's present age = $8x$ years

Now, we use the second piece of information: the mother is 25 years older than Leela. This age difference remains constant throughout their lives.

So, Mother's present age - Leela's present age = 25 years.

Substituting our expressions for their present ages:

$$8x - 3x = 25$$

Solving for the Unknown Factor x

Now we solve the equation for x :

$$5x = 25$$

To find x , divide both sides by 5:

$$x = \frac{25}{5}$$

$$x = 5$$

Calculating the Mother's Present Age

We found that the ratio factor x is 5. The mother's present age is represented by $8x$.

$$\text{Mother's present age} = 8 \times x = 8 \times 5 = 40$$

So, the mother's present age is 40 years.

Verifying the Solution

Let's check if the calculated ages satisfy both conditions:

- Leela's present age = $3x = 3 \times 5 = 15$ years.
- Mother's present age = $8x = 8 \times 5 = 40$ years.

Condition 1: Ratio of present ages = Leela : Mother = 15 : 40. Dividing both by 5, we get 3 : 8. This matches the given ratio.

Condition 2: Age difference = Mother's age - Leela's age = 40 - 15 = 25 years. This matches the given information that the mother was 25 when Leela was born.

Both conditions are satisfied, confirming our solution is correct.

Revision Table: Age Ratio Problem

Your Personal Exams Guide

Concept	Explanation	Application in Problem
Ratio of Ages	Expresses the proportional relationship between two ages at a specific point in time.	Present ages of Leela and Mother are in the ratio 3:8.
Constant Age Difference	The difference in age between two people remains the same throughout their lives.	Mother was 25 when Leela was born means Mother is always 25 years older than Leela.
Using a Variable for Ratio	Introducing a variable (x) allows us to represent the actual ages based on the ratio.	Leela's age = $3x$, Mother's age = $8x$.
Forming an Equation	Translating the problem's conditions (like age difference) into a mathematical equation using the variable.	$8x - 3x = 25$.

Additional Information: Solving Age Word Problems

Age word problems often involve ratios, differences, sums, or ages at different points in time (past, present, future). Here are some tips for solving them:

- Read the problem carefully to identify the unknowns and the given relationships.
- Define variables to represent the unknown ages, usually starting with the present age.
- Write down the relationships between the ages as equations. Remember that age difference is constant, but ratios and sums change over time.
- Solve the equation(s) to find the value of the variable(s).
- Use the variable's value to calculate the specific ages asked for in the question.
- Always check your answer by plugging the calculated ages back into the original problem statement to ensure all conditions are met.

- If the problem involves past or future ages, calculate those based on the present ages (e.g., age 5 years ago = present age - 5; age 10 years from now = present age + 10).

78. Answer: c

Explanation:

Locating the Famous Taj Lake Palace

The question asks about the city where the Taj Lake Palace hotel is located. This iconic hotel is known for its unique location in the middle of a lake. To answer this, we need to identify the lake and the city it is in.

Understanding the Taj Lake Palace Location

The Taj Lake Palace is situated in the middle of Lake Pichola. Lake Pichola is a historic artificial freshwater lake created in the year 1362 AD. It is one of the most famous lakes in a prominent city in Rajasthan, India.

Identifying the City

Lake Pichola is located in the city of Udaipur. Udaipur is often called the "City of Lakes" because of its many beautiful lakes. The Taj Lake Palace, originally known as Jag Niwas, was built between 1743 and 1746 as a royal summer palace and was later converted into a luxury hotel.

Let's look at the options provided:

- Jodhpur: Jodhpur is known as the "Blue City" and has forts and palaces but is not typically associated with Lake Pichola.
- Bikaner: Bikaner is located in the Thar Desert and is known for its forts and camels, not for Lake Pichola.
- Udaipur: Udaipur is famous for its lakes, including Lake Pichola, on which the Taj Lake Palace is located.

- Jaipur: Jaipur is the capital of Rajasthan and is known as the "Pink City," famous for its forts and palaces, but not Lake Pichola.

Based on the location of Lake Pichola, the Taj Lake Palace is located in Udaipur.

Revision Table: Famous Places in Rajasthan Cities

City	Key Attractions/Features	Associated with Lake Pichola?
Jodhpur	Mehrangarh Fort, Umaid Bhawan Palace	No
Bikaner	Junagarh Fort, Karni Mata Temple	No
Udaipur	Lake Pichola, City Palace, Jag Mandir, Taj Lake Palace	Yes
Jaipur	Hawa Mahal, Amer Fort, City Palace	No

Additional Information about Taj Lake Palace and Udaipur

The Taj Lake Palace is an exquisite example of Rajput architecture. Its location in the middle of Lake Pichola makes it appear to float on the water. This gives it a unique charm and makes it a major tourist attraction in Udaipur.

Udaipur itself is a historic city founded by Maharana Udai Singh II in 1559. Besides Lake Pichola and the Lake Palace, other significant sites include the City Palace, Jag Mandir island palace, and the Saheliyon Ki Bari garden. The city is a popular destination for its history, culture, scenic beauty, and luxurious heritage hotels.

The architecture often uses marble and intricate carvings, reflecting the royal heritage of the region. The palace offers stunning views of the Aravalli Hills and the city.

Understanding the geography and famous landmarks of major Indian cities, especially those known for tourism like those in Rajasthan, is important for general

knowledge questions.

79. Answer: d

Explanation:

Solving Work and Time Problems: M and L Working Together

This problem involves understanding the relationship between the efficiency of workers and the time taken to complete a task. Efficiency is inversely proportional to time; a more efficient worker takes less time to finish the same task.

Understanding Efficiency Relationships

The problem states that M is thrice as good a workman as L. This means that M's rate of work (efficiency) is three times the rate of work of L.

- Let the efficiency of L be E_L (amount of work L does in one day).
- The efficiency of M is E_M .
- According to the problem, $E_M = 3 \times E_L$.

Calculating Combined Efficiency

When M and L work together, their efficiencies add up to complete the task faster. The combined efficiency is the sum of their individual efficiencies.

- Combined efficiency, $E_{combined} = E_L + E_M$.
- Substituting $E_M = 3E_L$, we get $E_{combined} = E_L + 3E_L = 4E_L$.
- So, their combined efficiency is four times the efficiency of L alone.

Using the Given Time to Find Combined Efficiency

They finish the task together in 12 days. The combined efficiency is the reciprocal of the time taken when working together.

- Time taken together = 12 days.
- Combined efficiency $E_{combined} = \frac{1}{\text{Time taken together}} = \frac{1}{12}$ task per day.

Equating Combined Efficiencies to Find L's Efficiency

We have two expressions for combined efficiency: $4E_L$ and $\frac{1}{12}$. We can set them equal to find the value of E_L .

$$4E_L = \frac{1}{12}$$

To find E_L , divide both sides by 4:

$$E_L = \frac{1}{12 \times 4} = \frac{1}{48} \text{ task per day.}$$

L's efficiency is $\frac{1}{48}$ task per day. This means L can complete $\frac{1}{48}$ of the task in one day.

Calculating Time for L Alone

The time taken by L alone to finish the task is the reciprocal of L's efficiency.

$$\text{Time for L alone} = \frac{1}{E_L} = \frac{1}{1/48} = 48 \text{ days.}$$

Therefore, L alone will finish the same task in 48 days.

Summary of Steps

1. Define efficiencies: $E_M = 3E_L$.
2. Calculate combined efficiency: $E_{combined} = E_L + E_M = 4E_L$.
3. Calculate combined efficiency from time taken: $E_{combined} = \frac{1}{12}$.
4. Equate and solve for E_L : $4E_L = \frac{1}{12} \implies E_L = \frac{1}{48}$.
5. Calculate time for L alone: Time for L alone = $\frac{1}{E_L} = 48$ days.

Work and Time Summary

Worker(s)	Efficiency (Task/Day)	Time Taken (Days)
L	$E_L = \frac{1}{48}$	$\frac{1}{E_L} = 48$
M	$E_M = 3E_L = \frac{3}{48} = \frac{1}{16}$	$\frac{1}{E_M} = 16$
M & L Together	$E_{combined} = 4E_L = \frac{4}{48} = \frac{1}{12}$	$\frac{1}{E_{combined}} = 12$

Revision Table: Work and Time Concepts

Understanding these basic concepts is crucial for solving work and time problems.

Key Work and Time Concepts

Concept	Description	Formula
Work	The total task to be completed (often considered as 1 unit).	Total Work = Efficiency × Time
Efficiency (Rate of Work)	Amount of work done by a person or machine in unit time.	Efficiency = $\frac{\text{Total Work}}{\text{Time}}$
Time	Duration taken to complete the work.	Time = $\frac{\text{Total Work}}{\text{Efficiency}}$
Combined Efficiency	Sum of individual efficiencies when multiple workers work together.	$E_{combined} = E_1 + E_2 + \dots$

Additional Information: Solving Work Problems

Work and time problems often involve scenarios where individuals or groups work together or separately to complete a task. The key is to convert the information about time and relative efficiencies into rates of work (efficiency).

- **Inverse Relationship:** Efficiency and time are inversely proportional. If efficiency doubles, time taken halves.

- **Assuming Total Work:** Sometimes, it's helpful to assume the total work is a specific value, often the Least Common Multiple (LCM) of the times given, to work with integers instead of fractions. In this problem, assuming total work = 48 units would also yield the same result. L's efficiency = 1 unit/day, M's efficiency = 3 units/day. Combined efficiency = 4 units/day. Time taken together = $48/4 = 12$ days. Time for L alone = $48/1 = 48$ days.
- **Consistent Units:** Ensure that time units (days, hours, minutes) are consistent throughout the problem.

By focusing on the rate of work (efficiency) per unit of time, complex work and time problems can be broken down into simpler steps.

80. Answer: d

Explanation:

Solving Sentence Rearrangement Questions

Sentence rearrangement questions test your ability to understand the logical flow and grammatical structure of a sentence. You are given a main part of a sentence and several jumbled segments. Your task is to arrange the segments in the correct order to form a complete, coherent sentence.

Analyzing the Given Sentence Parts

Let's look at the parts we have:

- Unjumbled part: I'd just remembered that our
- Segment X: every once in a while, so our changing
- Segment Y: rooms were movable, on wheels
- Segment Z: owner liked to change the decor of the boutique

Step-by-Step Rearrangement Process

To find the correct sequence, we can try connecting the segments to the unjumbled part and to each other, looking for grammatical and logical connections.

1. **Start with the Unjumbled Part:** The sentence begins with "I'd just remembered that our...". We need a segment that logically follows "our".

Let's look at the start of each segment:

- X starts with "every once in a while..." - Doesn't fit directly after "our".
- Y starts with "rooms were movable..." - Could potentially fit ("our rooms"), but let's check others.
- Z starts with "owner liked to change..." - This fits perfectly after "our" ("our owner").

So, Segment Z is the most likely part to follow the unjumbled beginning.

Partial sentence: I'd just remembered that our owner liked to change the decor of the boutique...

2. **Connect the Next Segment:** Now we need to see which segment follows Z. Z ends with "...of the boutique". Let's look at the starts of X and Y.

- X starts with "every once in a while..." - This phrase often describes the frequency of an action. Changing decor is an action mentioned in Z. This seems like a good fit. The phrase could follow "...of the boutique every once in a while".
- Y starts with "rooms were movable..." - This doesn't logically follow "...of the boutique".

So, Segment X seems to follow Z.

Partial sentence: I'd just remembered that our owner liked to change the decor of the boutique every once in a while, so our changing...

3. **Complete the Sentence:** We have the unjumbled part + Z + X. The partial sentence ends with "...so our changing...". Now, the only remaining segment is Y, which starts with "rooms were movable, on wheels".

Let's see if Y follows X:

- o X ends with "...so our changing". Y starts with "rooms were movable...". This fits well: "so our changing rooms".

So, Segment Y follows X.

The Complete Rearranged Sentence

Putting all parts together in the order Unjumbled + Z + X + Y:

I'd just remembered that our **owner liked to change the decor of the boutique (Z)** **every once in a while, so our changing (X) rooms were movable, on wheels (Y).**

The complete sentence is: "I'd just remembered that our owner liked to change the decor of the boutique every once in a while, so our changing rooms were movable, on wheels."

This sentence makes grammatical and logical sense.

Identifying the Correct Order

The correct order of the jumbled segments X, Y, and Z is ZXY.

Segment	Position
Z	1st after unjumbled part
X	2nd after unjumbled part (after Z)
Y	3rd after unjumbled part (after X)

Therefore, the segments should be arranged as ZXY.

Revision Table: Key Takeaways

Concept	Explanation
Sentence Rearrangement	Ordering jumbled parts to form a meaningful sentence.
Identifying Clues	Look for words that connect logically (e.g., "our" followed by "owner", "changing" followed by "rooms").
Checking Flow	Read the sentence with the segments in the proposed order to ensure it makes sense grammatically and logically.

Additional Information: Tips for Jumbled Sentences

- Always read the unjumbled part carefully to understand the beginning of the sentence.
- Look for connecting words or phrases between segments (e.g., pronouns, conjunctions like "so", time phrases like "every once in a while").
- Identify the subject and verb of the main clause if possible.
- Try to find the segment that seems to complete an idea started by another segment.
- Once you have a potential order, read the entire sentence aloud to check if it flows naturally.

Your Personal Exams Guide

81. Answer: b

Explanation:

Understanding the Pilgrim's Journey

This problem involves calculating the final position of a pilgrim relative to their starting point after a series of movements in different directions and distances. To solve this, we can track the pilgrim's net movement in the east-west direction and the north-south direction.

Step-by-Step Calculation of Pilgrim's Position

Let's consider the starting position as the origin $(0, 0)$ on a 2D plane, where East is the positive x-direction, West is the negative x-direction, North is the positive y-direction, and South is the negative y-direction.

1. **First movement:** The pilgrim walks 13 km towards the east.

Change in position: +13 km East, 0 km North/South.

Current position relative to start: $(13, 0)$.

2. **Second movement:** Turns towards the north and walks 11 km.

Change in position: 0 km East/West, +11 km North.

Current position relative to start: $(13, 0) + (0, 11) = (13, 11)$.

3. **Third movement:** Turns towards the west and walks 7 km.

Change in position: -7 km West, 0 km North/South.

Current position relative to start: $(13, 11) + (-7, 0) = (6, 11)$.

4. **Fourth movement:** Turns towards the south and walks 6 km.

Change in position: 0 km East/West, -6 km South.

Current position relative to start: $(6, 11) + (0, -6) = (6, 5)$.

5. **Fifth movement:** Turns to her right and walks 6 km.

When facing South, turning right means turning towards the West.

Change in position: -6 km West, 0 km North/South.

Current position relative to start: $(6, 5) + (-6, 0) = (0, 5)$.

Summarizing Pilgrim's Displacement

Let's look at the net displacement in the East-West and North-South directions:

Movement	Direction	Distance (km)	East/West Change (km)	North/South Change (km)
Start	-	0	0	0
1	East	13	+13	0
2	North	11	0	+11
3	West	7	-7	0
4	South	6	0	-6
5	Right (West)	6	-6	0

Net change in East-West direction = $13 - 7 - 6 = 13 - 13 = 0$ km.

Net change in North-South direction = $11 - 6 = 5$ km.

Final Position Relative to Starting Point

The net change in the East-West direction is 0 km, meaning the pilgrim is neither East nor West of the starting point along the East-West axis.

The net change in the North-South direction is +5 km. Since North is the positive y-direction, this means the pilgrim is 5 km North of the starting point along the North-South axis.

Therefore, the final position is 5 km North with respect to the starting position.

Revision Table: Pilgrim's Journey Steps

Step	Action	Direction	Distance (km)	Position Update
1	Walks	East	13	13 km East
2	Turns & Walks	North	11	13 km East, 11 km North
3	Turns & Walks	West	7	$(13 - 7) = 6$ km East, 11 km North
4	Turns & Walks	South	6	6 km East, $(11 - 6) = 5$ km North
5	Turns Right & Walks	West (from South)	6	$(6 - 6) = 0$ km East, 5 km North

Additional Information on Displacement and Distance

In physics and vector mathematics, displacement is the shortest distance from the initial to the final position of a point. It is a vector quantity, having both magnitude and direction. In this problem, we calculated the displacement of the pilgrim from the starting point.

- **Distance:** The total path length covered by the pilgrim is $13 + 11 + 7 + 6 + 6 = 43$ km. Distance is a scalar quantity.
- **Displacement:** The straight-line distance and direction from the start to the end point. In this case, the displacement is 5 km North. It is a vector from (0,0) to (0,5). The magnitude of the displacement is $\sqrt{(0 - 0)^2 + (5 - 0)^2} = \sqrt{0 + 25} = \sqrt{25} = 5$ km. The direction is along the positive y-axis, which is North.

Understanding the difference between distance and displacement is crucial in problems involving movement.

82. Answer: b

Explanation:

This question asks us to find the value of $8 \# 12$ by identifying the pattern used in the given examples. We are given three equations relating two numbers with a '#' symbol between them to a result. Let's look closely at these examples to understand the operation that '#' represents.

Understanding the Logical Pattern

The given relationships are:

- $14 \# 2 = 80$
- $10 \# 4 = 70$
- $20 \# 6 = 130$

We need to figure out what mathematical operation or combination of operations the '#' symbol stands for. Let's examine the numbers in each case and their result:

- In the first case, we have 14 and 2 giving 80.
- In the second case, we have 10 and 4 giving 70.
- In the third case, we have 20 and 6 giving 130.

Let's try combining the numbers 14 and 2 in simple ways, like adding or subtracting, and see if there's a relationship with 80. The sum of 14 and 2 is $14 + 2 = 16$. How can we get 80 from 16? We can multiply 16 by 5, since $16 \times 5 = 80$. This looks promising.

Let's test this potential pattern with the other examples:

- For $10 \# 4 = 70$: The sum of 10 and 4 is $10 + 4 = 14$. If we multiply 14 by 5, we get $14 \times 5 = 70$. This matches the given result!
- For $20 \# 6 = 130$: The sum of 20 and 6 is $20 + 6 = 26$. If we multiply 26 by 5, we get $26 \times 5 = 130$. This also matches the given result!

So, the pattern is consistently: add the two numbers together and then multiply the sum by 5. If the two numbers are 'a' and 'b', the operation $a \# b$ is equivalent to $(a + b) \times 5$.

We can write the pattern using mathematical notation as:

$$a \# b = (a + b) \times 5$$

Applying the Pattern to Find 8 # 12

Now that we have discovered the logical pattern, we can use it to find the value of 8 # 12. Here, 'a' is 8 and 'b' is 12.

Using the pattern $a \# b = (a + b) \times 5$, we substitute $a = 8$ and $b = 12$:

$$8 \# 12 = (8 + 12) \times 5$$

First, calculate the sum inside the parentheses:

$$8 + 12 = 20$$

Now, multiply the sum by 5:

$$20 \times 5 = 100$$

Therefore, the value of 8 # 12 is 100.

Revision Table: Symbol Operations

Example	First Number (a)	Second Number (b)	Operation $a + b$	Operation $(a + b) \times 5$	Result	Given Result
1	14	2	$14 + 2 = 16$	$16 \times 5 = 80$	80	80
2	10	4	$10 + 4 = 14$	$14 \times 5 = 70$	70	70
3	20	6	$20 + 6 = 26$	$26 \times 5 = 130$	130	130
Question	8	12	$8 + 12 = 20$	$20 \times 5 = 100$	100	?

Additional Information: Logic Puzzles and Pattern Recognition

This question is a type of logic puzzle that requires pattern recognition. Logic puzzles often involve finding a hidden rule or relationship between given elements. These types of questions are common in reasoning and quantitative aptitude tests because they assess your ability to analyze information, identify patterns, and apply logical thinking.

Tips for solving pattern recognition problems:

- Carefully examine the given examples. Look for relationships between the numbers or elements.
- Try simple mathematical operations (addition, subtraction, multiplication, division) or combinations of operations.
- Test your potential pattern with all the given examples to ensure it works consistently.
- Once the pattern is confirmed, apply it to the specific case asked in the question.
- Don't overcomplicate things initially; sometimes the simplest pattern is the correct one.

Solving logic puzzles helps improve critical thinking and problem-solving skills, which are valuable in many areas.

The final answer is 100.

Your Personal Exams Guide

83. Answer: a

Explanation:

Understanding Fahrenheit to Celsius Temperature Conversion

This question asks us to find the equivalent temperature in degrees Celsius for a given temperature in degrees Fahrenheit. We need to perform a standard temperature conversion.

Fahrenheit to Celsius Conversion Formula

The relationship between the Fahrenheit scale ($^{\circ}\text{F}$) and the Celsius scale ($^{\circ}\text{C}$) is defined by a specific formula. To convert a temperature from Fahrenheit to Celsius, we use the following equation:

$$C = \frac{5}{9}(F - 32)$$

Where:

- C represents the temperature in degrees Celsius.
- F represents the temperature in degrees Fahrenheit.

This formula involves subtracting 32 from the Fahrenheit temperature and then multiplying the result by $\frac{5}{9}$.

Step-by-Step Calculation for 167 Fahrenheit

We are given the temperature in Fahrenheit as 167. Let's apply the conversion formula step-by-step:

1. **Start with the Fahrenheit temperature:** The given temperature is $F = 167^{\circ}\text{Fahrenheit}$.

2. **Subtract 32:** First, we subtract 32 from the Fahrenheit value:

$$167 - 32 = 135$$

3. **Multiply by $\frac{5}{9}$:** Next, we multiply the result (135) by $\frac{5}{9}$:

$$C = \frac{5}{9} \times 135$$

To calculate this, we can first divide 135 by 9:

$$\frac{135}{9} = 15$$

Now, multiply this result by 5:

$$C = 5 \times 15$$

$$C = 75$$

Conclusion of Conversion

Following the calculation steps using the standard conversion formula, we find that 167 degrees Fahrenheit is equivalent to 75 degrees Celsius.

Therefore, the blank in the statement "_____ °Celsius = 167 Fahrenheit" should be filled with 75.

84. Answer: b

Explanation:

Solving for Voltage using Work and Charge

The question asks us to find the voltage (V) when a certain amount of work is done to move a specific amount of charge. This problem involves the relationship between work, charge, and potential difference (voltage) in an electrical circuit.

The fundamental formula connecting these quantities is:

$$\text{Work Done (W)} = \text{Charge (Q)} \times \text{Voltage (V)}$$

This formula can be rearranged to solve for Voltage:

$$\text{Voltage (V)} = \frac{\text{Work Done (W)}}{\text{Charge (Q)}}$$

In this specific problem, we are given the following values:

- Work Done (W) = 90 Joules (J)
- Charge (Q) = 2,000 Coulombs (C)

Now, we can substitute these values into the rearranged formula to find the voltage:

$$V = \frac{W}{Q}$$

$$V = \frac{90 \text{ J}}{2000 \text{ C}}$$

$$V = \frac{90}{2000} \text{ Volts}$$

To simplify the fraction, we can divide both the numerator and the denominator by 10:

$$V = \frac{9}{200} \text{ Volts}$$

Now, perform the division:

$$V = 0.045 \text{ Volts}$$

Therefore, the voltage across which the charge is moved is 0.045 volts.

Calculation Summary: Work, Charge, and Voltage

Quantity	Symbol	Value	Unit
Work Done	W	90	Joules (J)
Charge	Q	2000	Coulombs (C)
Voltage (to find)	V	?	Volts (V)

Formula used:

$$V = \frac{W}{Q}$$

Calculation:

$$V = \frac{90}{2000} = 0.045 \text{ V}$$

Revision Table: Key Concepts

Term	Definition	Unit	Formula Relationship
Work	Energy transferred when a force moves an object over a distance; in electrical terms, energy transferred to move charge.	Joule (J)	$W = QV$
Charge	A fundamental property of matter that causes it to experience a force when placed in an electromagnetic field.	Coulomb (C)	$Q = W/V$
Voltage (Potential Difference)	The difference in electrical potential energy per unit charge between two points in a circuit. It is the energy required per unit charge to move charge between the points.	Volt (V)	$V = W/Q$

Additional Information: Understanding Electrical Work and Voltage

Work done in moving a charge is the energy transferred to move that charge against an electric field. Voltage, or potential difference, represents how much energy is needed per unit of charge to move it from one point to another. A higher voltage means more energy is available to push each unit of charge through a circuit.

The relationship $W = QV$ is a core concept in electricity. It tells us that the total work done is proportional to both the amount of charge moved and the potential difference it moved across. If you double the charge or double the voltage, you double the work done.

Units are very important here:

- Work is measured in Joules (J).
- Charge is measured in Coulombs (C).
- Voltage is measured in Volts (V).

The formula $V = W/Q$ shows that 1 Volt is equivalent to 1 Joule per Coulomb ($1 \text{ V} = 1 \text{ J/C}$). This unit relationship is consistent with our calculation.

85. Answer: c

Explanation:

Calculating Satellite Weight on Jupiter

This question asks us to find the weight of a satellite on Jupiter, given its mass, the acceleration due to gravity on Earth, and the relationship between gravity on Jupiter and Earth.

Understanding Mass and Weight

- **Mass:** This is a measure of the amount of matter in an object. It is a fundamental property and remains constant regardless of where the object is (on Earth, Jupiter, or in space). The satellite's mass is given as 250 kg.
- **Weight:** This is the force of gravity acting on an object's mass. It depends on the mass of the object and the acceleration due to gravity at its location. Weight is a force and is measured in Newtons (N). The formula for weight is:

$$\text{Weight}(W) = \text{Mass}(m) \times \text{Acceleration due to gravity}(g)$$

Given Information

- Mass of the satellite (m) = 250 kg
- Acceleration due to gravity on Earth (g_{earth}) = 10 m/s^2
- Acceleration due to gravity on Jupiter (g_{jupiter}) is two and a half times that on Earth, meaning $g_{\text{jupiter}} = 2.5 \times g_{\text{earth}}$.

Step-by-Step Calculation

Step 1: Find the acceleration due to gravity on Jupiter.

We know that $g_{jupiter} = 2.5 \times g_{earth}$. Substitute the given value for g_{earth} .

$$g_{jupiter} = 2.5 \times 10 \text{ m/s}^2$$

$$g_{jupiter} = 25 \text{ m/s}^2$$

So, the acceleration due to gravity on Jupiter is 25 m/s^2 .

Step 2: Calculate the weight of the satellite on Jupiter.

Use the formula for weight: $W_{jupiter} = m \times g_{jupiter}$. Substitute the mass of the satellite and the calculated gravity on Jupiter.

$$W_{jupiter} = 250 \text{ kg} \times 25 \text{ m/s}^2$$

$$W_{jupiter} = 6250 \text{ N}$$

The weight of the 250 kg satellite on Jupiter would be 6,250 Newtons.

Comparing with Options

Let's compare our calculated weight with the given options:

Option	Weight (N)
1	10
2	625
3	6,250
4	100

Our calculated value, 6,250 N, matches Option 3.

Revision Table: Satellite Weight on Jupiter Calculation

Concept	Formula/Value	Notes
Mass (m)	250 kg	Constant value
Gravity on Earth (g_{earth})	10 m/s ²	Given value
Gravity on Jupiter ($g_{jupiter}$)	$2.5 \times g_{earth}$	Given relationship
Calculated $g_{jupiter}$	$2.5 \times 10 = 25 \text{ m/s}^2$	Step 1 result
Weight on Jupiter ($W_{jupiter}$)	$m \times g_{jupiter}$	Formula for weight
Calculated $W_{jupiter}$	$250 \times 25 = 6250 \text{ N}$	Step 2 result

Additional Information on Gravity and Weight

- Acceleration due to gravity varies from planet to planet because it depends on the planet's mass and radius. Larger, denser planets generally have stronger gravity.
- Weight is a force vector, pointing towards the center of the gravitational source. Mass is a scalar quantity.
- On the Moon, where gravity is much weaker than on Earth (about 1/6th), the same 250 kg satellite would weigh significantly less than on Earth or Jupiter.
- While gravity is often approximated as 9.8 m/s² on Earth, using 10 m/s² simplifies calculations in many problems like this one and is explicitly stated in the question.

86. Answer: b

Explanation:

Understanding Binary Addition

Binary addition is a fundamental operation in digital electronics and computer systems. Unlike decimal addition which uses base 10, binary addition uses base 2. The rules for binary addition are simple:

- $0 + 0 = 0$ (with no carry)
- $0 + 1 = 1$ (with no carry)
- $1 + 0 = 1$ (with no carry)
- $1 + 1 = 0$ (with a carry of 1 to the next higher position)
- $1 + 1 + 1 = 1$ (with a carry of 1 to the next higher position)

When adding two binary numbers, we align them by their rightmost digit and add column by column, propagating any carries to the left, just like in decimal addition.

Let's find the sum of the given binary numbers: 1100100 and 101110.

Step-by-Step Binary Addition of 1100100 and 101110

To add 1100100 and 101110, we can align the numbers and add column by column from right to left.

Position (from right)	7	6	5	4	3	2	1
First Number: 1100100	1	1	0	0	1	0	0
Second Number: 101110 (padded)	0	1	0	1	1	1	0
Carry from right		1	0	1	1	0	0
Sum (current column)		0	0	1	0	1	0
Carry to left	1	0	0	1	1	0	0

Let's perform the addition column by column:

- **Rightmost column (Position 1):** $0 + 0 = 0$. Carry = 0. Result digit = 0.
- **Position 2:** $0 + 1 = 1$. Carry = 0. Result digit = 1.
- **Position 3:** $1 + 1 = 0$. Carry = 1. Result digit = 0.

- **Position 4:** $0 + 1 + \text{carry } (1) = 1 + 1 = 0$. Carry = 1. Result digit = 0.
- **Position 5:** $0 + 0 + \text{carry } (1) = 0 + 1 = 1$. Carry = 0. Result digit = 1.
- **Position 6:** $1 + 1 + \text{carry } (0) = 1 + 1 = 0$. Carry = 1. Result digit = 0.
- **Position 7:** $1 + 0 + \text{carry } (1) = 1 + 1 = 0$. Carry = 1. Result digit = 0.
- **Final Carry:** The carry from Position 7 is 1, which becomes the leftmost digit of the sum. Result digit = 1.

Reading the resulting digits from left to right (including the final carry), we get 10010010.

Final Sum of Binary Numbers

The sum of the binary numbers 1100100 and 101110 is 10010010.

This matches one of the provided options.

Revision Table: Key Binary Addition Rules

Operation	Sum	Carry
$0 + 0$	0	0
$0 + 1$	1	0
$1 + 0$	1	0
$1 + 1$	0	1
$1 + 1 + 1$	1	1

Additional Information on Binary Numbers

Binary numbers are the foundation of digital computing. A binary digit (bit) can only be 0 or 1. Larger binary numbers represent values using powers of 2, similar to how decimal numbers use powers of 10. Understanding binary arithmetic, including binary addition, subtraction, multiplication, and division, is crucial for studying computer architecture, digital logic design, and low-level programming.

For instance, converting binary numbers to decimal can help verify the addition:

- $1100100_2 = 1 \times 2^6 + 1 \times 2^5 + 0 \times 2^4 + 0 \times 2^3 + 1 \times 2^2 + 0 \times 2^1 + 0 \times 2^0$
- $1100100_2 = 64 + 32 + 0 + 0 + 4 + 0 + 0 = 100_{10}$
- $101110_2 = 1 \times 2^5 + 0 \times 2^4 + 1 \times 2^3 + 1 \times 2^2 + 1 \times 2^1 + 0 \times 2^0$
- $101110_2 = 32 + 0 + 8 + 4 + 2 + 0 = 46_{10}$

Their decimal sum is $100 + 46 = 146_{10}$.

Let's convert the binary sum 10010010 to decimal:

- $10010010_2 = 1 \times 2^7 + 0 \times 2^6 + 0 \times 2^5 + 1 \times 2^4 + 0 \times 2^3 + 0 \times 2^2 + 1 \times 2^1 + 0 \times 2^0$
- $10010010_2 = 128 + 0 + 0 + 16 + 0 + 0 + 2 + 0 = 146_{10}$

The decimal sum matches the decimal conversion of the binary sum, confirming the correctness of the binary addition.

87. Answer: d

Explanation:

Understanding Pump Efficiency Calculation

Efficiency is a measure of how well a device converts input energy or power into useful output energy or power. For a pump, the useful output is the work done in lifting water, and the input is the electrical power consumed by the motor.

The efficiency (η) is calculated using the formula:

$$\eta = \left(\frac{\text{Useful Output Power}}{\text{Input Power}} \right) \times 100\%$$

In this problem, we are given the input power of the pump and the task it performs (lifting water). We need to calculate the useful output power based on the work done and the time taken.

Calculating Useful Work Done by the Pump

The pump lifts 500 kg of water by a height of 30 m. The useful work done against gravity is the increase in potential energy of the water. The formula for work done against gravity (or potential energy gain) is:

$$\text{Work done} = \text{mass} \times \text{acceleration due to gravity} \times \text{height}$$

Given:

- Mass (m) = 500 kg
- Acceleration due to gravity (g) = 10 m/s²
- Height (h) = 30 m

Calculation of work done:

$$\text{Work done} = 500 \text{ kg} \times 10 \text{ m/s}^2 \times 30 \text{ m}$$

$$\text{Work done} = 5000 \text{ N} \times 30 \text{ m}$$

$$\text{Work done} = 150,000 \text{ Joules}$$

Determining Useful Output Power

Power is the rate at which work is done. The work of 150,000 Joules is done in 10 minutes. First, we convert the time from minutes to seconds:

$$\text{Time} = 10 \text{ minutes} \times 60 \text{ seconds/minute}$$

$$\text{Time} = 600 \text{ seconds}$$

Now, we calculate the useful output power using the formula:

$$\text{Useful Output Power} = \frac{\text{Work done}}{\text{Time taken}}$$

$$\text{Useful Output Power} = \frac{150,000 \text{ J}}{600 \text{ s}}$$

$$\text{Useful Output Power} = \frac{1500}{6} \text{ W}$$

$$\text{Useful Output Power} = 250 \text{ W}$$

Final Calculation of Pump Efficiency

We have the input power and the useful output power:

- Input Power = 400 W
- Useful Output Power = 250 W

Now, we use the efficiency formula:

$$\eta = \left(\frac{\text{Useful Output Power}}{\text{Input Power}} \right) \times 100\%$$

$$\eta = \left(\frac{250 \text{ W}}{400 \text{ W}} \right) \times 100\%$$

$$\eta = \left(\frac{25}{40} \right) \times 100\%$$

$$\eta = \left(\frac{5}{8} \right) \times 100\%$$

$$\eta = 0.625 \times 100\%$$

$$\eta = 62.5\%$$

Resulting Pump Efficiency

The calculated efficiency of the pump is 62.5%.

Revision Table: Pump Efficiency Formulas

Concept	Formula	Units
Work Done (Potential Energy)	$W = mgh$	Joules (J)
Power	$P = \frac{W}{t}$	Watts (W)
Efficiency	$\eta = \left(\frac{P_{\text{out}}}{P_{\text{in}}} \right) \times 100\%$	Percentage (%)

Additional Information: Energy and Power Concepts

Work: In physics, work is done when a force causes displacement. Lifting water against gravity is a form of work. The unit of work is the Joule (J), which is equal to one Newton-meter (Nm).

Potential Energy: This is the energy an object possesses due to its position or state. When a pump lifts water, it increases the water's gravitational potential energy, which is the useful work done.

Power: Power is the rate at which work is done or energy is transferred. The unit of power is the Watt (W), which is equal to one Joule per second (J/s). The input power is what the pump consumes (e.g., electrical power), and the output power is the useful power delivered (e.g., power used to lift water).

Efficiency: Efficiency is always a ratio (or percentage) of useful output to total input. It indicates how much of the input energy or power is converted into the desired form of output, with the rest typically lost as heat or sound.

88. Answer: b

Explanation:

Understanding the Code Logic Puzzle

This question is a type of coding-decoding problem where letters are assigned specific digits based on a hidden rule. We are given two examples of words coded into numbers: EXACT and MILK. Our goal is to use these examples to figure out the rule (the mapping of letters to digits) and then apply that rule to find the code for the word MELT.

Decoding the Given Words: EXACT and MILK

We are told that:

- EXACT is coded as 91685
- MILK is coded as 7243

Let's examine the letters in each word and their corresponding digits in the code. We assume a direct mapping where each letter consistently represents the same digit.

From the code for EXACT (91685):

- E corresponds to the 1st digit, 9
- X corresponds to the 2nd digit, 1
- A corresponds to the 3rd digit, 6
- C corresponds to the 4th digit, 8
- T corresponds to the 5th digit, 5

From the code for MILK (7243):

- M corresponds to the 1st digit, 7
- I corresponds to the 2nd digit, 2
- L corresponds to the 3rd digit, 4
- K corresponds to the 4th digit, 3

Establishing the Letter-to-Digit Mapping

By comparing the letters and their positions in the original words with the digits in the coded numbers, we can establish the following consistent mappings:

Letter	Coded Digit
E	9
X	1
A	6
C	8
T	5
M	7
I	2
L	4
K	3

This table shows the code for each individual letter derived from the examples.

Applying the Code to MELT

Now that we have the mapping for each letter, we can find the code for the word MELT. We just need to replace each letter in MELT with its corresponding digit from our established mapping.

- The first letter is M. From our mapping, M is coded as 7.
- The second letter is E. From our mapping, E is coded as 9.
- The third letter is L. From our mapping, L is coded as 4.
- The fourth letter is T. From our mapping, T is coded as 5.

Combining these digits in order, the code for MELT is 7945.

Step-by-Step Solution to Find the Code for MELT

Here is a summary of the steps:

1. Analyze the first example: EXACT is coded as 91685. This gives the codes for E, X, A, C, and T.
2. Analyze the second example: MILK is coded as 7243. This gives the codes for M, I, L, and K.
3. Compile a list of letter-to-digit mappings from both examples.
4. Take the word MELT.
5. Replace each letter in MELT with its corresponding digit based on the compiled mapping.
6. $M \rightarrow 7$
7. $E \rightarrow 9$
8. $L \rightarrow 4$
9. $T \rightarrow 5$
10. Combine the digits: 7945.

Thus, the code for MELT is 7945.

Revision Table: Key Code Mappings

Word	Letters	Code	Letter Mappings
EXACT	E, X, A, C, T	91685	E=9, X=1, A=6, C=8, T=5
MILK	M, I, L, K	7243	M=7, I=2, L=4, K=3
MELT	M, E, L, T	?	Using mappings: M=7, E=9, L=4, T=5 $\rightarrow$ 7945

Additional Information on Coding-Decoding Questions

Coding and decoding questions test your ability to identify patterns and rules in how information is transformed. Common types include:

- **Letter Coding:** Letters are replaced by other letters (e.g., shifting positions in the alphabet).
- **Number Coding:** Words or letters are coded into numbers, as in this example. The mapping can be direct, based on alphabetical position, or follow other arithmetic rules.

- **Symbol Coding:** Letters or words are replaced by symbols.
- **Mixed Coding:** A combination of letters, numbers, or symbols is used in the code.

To solve these problems, always look for consistent patterns in the provided examples. Identify how each element (letter, number, or symbol) in the original relates to its counterpart in the code. Once the rule is found, apply it systematically to the word or message you need to decode or code.

89. Answer: d

Explanation:

In option 1, 2 and 3, all the letters are same, while in option '4', H is missing and 'I' is present which is not there in other three options.

Hence, 'option 4' is the odd one out.

90. Answer: b

Explanation:

The correct answer is 100.



Key Points

- The resistance of the wire is 100 k Ω .
- The formula for calculating resistance is: $R = V/I$
- Here, we are given the potential difference (V) and current (I), so we can calculate the resistance as follows:

$R = 500 / 5 \times 10^{-3}$ (since the current is given in milliamperes and 'milli' refers to 10^{-3})

$R = 100 \times 10^3 \Omega$

$R = 100 \text{ k}\Omega$ (since 10^3 refers to 'kilo').

91. Answer: c

Explanation:

Understanding System Diagrams

The question asks for the name of a specific type of schematic drawing that uses simple shapes to illustrate the relationships between entities within a system. Let's look at the options provided to identify the correct diagram type.

Analyzing the Diagram Types

Different types of diagrams serve different purposes in representing information visually. Understanding what each option represents will help us identify the one that fits the description.

- **Venn Diagram:** This diagram uses overlapping circles to show the logical relationships between sets. It's primarily used in set theory, logic, and statistics to represent commonalities and differences between groups.
- **Circuit Diagram:** This is a visual representation of an electrical circuit. It uses standardized symbols to show the components (like resistors, capacitors, transistors) and their connections. It's specific to electrical and electronic systems.
- **Block Diagram:** This diagram uses simple shapes, typically rectangles (blocks), connected by lines or arrows, to represent the major parts or functions of a system and show their relationships or signal flow. It provides a high-level view of the system structure or process.
- **Scatter Diagram:** Also known as a scatter plot, this diagram plots data points on a Cartesian plane to show the relationship between two variables. It's used

in statistics to observe patterns or correlations.

Identifying the Correct Diagram

The description mentions a "schematic drawing that shows the relationships in a system, using simple shapes for the entities." Comparing this description with the definitions above:

- A Venn diagram focuses on set relationships, not system entities in this context.
- A circuit diagram uses specific electrical symbols, not just simple shapes for general entities.
- A scatter diagram shows the relationship between variables, not components or functions of a system.
- A block diagram uses simple shapes (blocks) to represent entities (parts, components, functions) and lines to show their relationships or interactions within a system. This perfectly matches the description.

Therefore, the diagram type that fits the description is a block diagram.

Example: A simple block diagram for a radio might show blocks for the antenna, tuner, amplifier, and speaker, with lines indicating the flow of the audio signal.

Conclusion

Based on the analysis of the options and the definition provided in the question, a block diagram is the schematic drawing that shows the relationships in a system using simple shapes for the entities.

Diagram Type	Primary Use	Shape Representation
Venn Diagram	Set relationships	Overlapping circles
Circuit Diagram	Electrical circuits	Standardized electrical symbols
Block Diagram	System structure/relationships	Simple shapes (blocks) for entities
Scatter Diagram	Relationship between variables	Points on a graph

Revision Table: Key Diagram Definitions

Let's quickly review the key characteristics of the diagram types discussed:

- **Block Diagram:** Represents a system's main components as blocks and shows connections/relationships with lines. Used for overall system view.
- **Venn Diagram:** Shows relationships between sets using overlapping circles.
- **Circuit Diagram:** Represents electrical circuits using specific component symbols.
- **Scatter Diagram:** Plots data points to show the relationship between two variables.

Additional Information: Uses of Block Diagrams

Block diagrams are widely used in various fields, including engineering, system design, and process mapping, because they provide a clear and simplified overview of a system's structure or process. They help in:

- Understanding complex systems at a high level.
- Identifying the main components and their interactions.
- Planning and designing systems before detailed design begins.
- Troubleshooting by following the flow between blocks.
- Communicating system architecture to non-experts.

They are particularly useful in control systems engineering to represent feedback loops and system dynamics.

92. Answer: b

Explanation:

Understanding Compound Interest for the Second Year

The question asks us to find the compound interest earned in the second year, given the compound interest for the first year and the rate of interest.

In compound interest, the interest for each period is calculated on the principal amount plus any interest accumulated from previous periods. This is different from simple interest, where interest is only calculated on the initial principal.

Calculating the Principal Amount

For the first year, compound interest is the same as simple interest because there is no previously accumulated interest. We are given that the compound interest for the first year is Rs. 1,440 at a 10% rate.

Let the principal amount be P . The interest for the first year (I_1) is given by the formula:

$$I_1 = P \times \text{Rate}$$

We have $I_1 = 1440$ and Rate = 10% or 0.10.

So, we can write the equation:

$$1440 = P \times 0.10$$

To find the principal P , we rearrange the equation:

$$P = \frac{1440}{0.10}$$

$$P = 14400$$

The principal amount is Rs. 14,400.

Calculating Compound Interest for the Second Year

The compound interest for the second year is calculated on the amount accumulated at the end of the first year. The amount at the end of the first year is the initial principal plus the interest earned in the first year.

- Principal (P) = Rs. 14,400
- Interest for the first year (I_1) = Rs. 1,440

$$\text{Amount at the end of the first year} = P + I_1$$

$$\text{Amount at the end of the first year} = 14400 + 1440 = 15840$$

Now, the interest for the second year (I_2) is calculated on this amount (Rs. 15,840) at the same rate of 10%.

$$I_2 = (\text{Amount at end of 1st year}) \times \text{Rate}$$

$$I_2 = 15840 \times 0.10$$

$$I_2 = 1584$$

The compound interest for the second year is Rs. 1,584.

Summary of Calculations

Given:	Compound Interest for 1st year	Rs. 1,440
Given:	Rate of Interest	10% per annum
Step 1:	Calculate Principal (P)	$P = \frac{I_1}{\text{Rate}} = \frac{1440}{0.10} = 14400$
Step 2:	Calculate Amount at end of 1st year	$P + I_1 = 14400 + 1440 = 15840$
Step 3:	Calculate Interest for 2nd year (I_2)	$I_2 = (\text{Amount}) \times \text{Rate} = 15840 \times 0.10 = 1584$

Thus, the compound interest for the second year will be Rs. 1,584.

Revision Table: Key Compound Interest Concepts

Concept	Description	Formula (A = Amount, P = Principal, r = rate, n = time in years)
Compound Interest	Interest calculated on the initial principal and also on the accumulated interest from previous periods.	$A = P(1 + r)^n$ $CI = A - P$
Simple Interest	Interest calculated only on the initial principal amount.	$SI = \frac{P \times r \times n}{100}$ (if rate is in %)
Interest in 1st Year (Compound)	Same as simple interest on the principal.	$I_1 = P \times r$
Interest in 2nd Year (Compound)	Calculated on the amount at the end of the 1st year ($P + I_1$).	$I_2 = (P + I_1) \times r$

Additional Information on Compound Interest Calculation

Compound interest is a powerful concept often used in finance for investments and loans. The frequency of compounding (annually, semi-annually, quarterly, etc.)

affects the total interest earned or paid. In this problem, compounding is annual as we are given yearly rates and interest for specific years.

Understanding the difference between interest earned in a specific period (like the second year) and the total compound interest accumulated over multiple years is important. The question specifically asked for the interest in the **second year**, not the total interest over two years.

The general formula for the amount at the end of 'n' years with annual compounding at rate 'r' is $A = P(1 + r)^n$. The compound interest for the nth year can be found by taking the difference between the amount at the end of n years and the amount at the end of (n-1) years.

For example, Amount at end of 2 years = $P(1 + r)^2$. Amount at end of 1 year = $P(1 + r)^1$. Interest in the 2nd year = Amount at end of 2 years - Amount at end of 1 year = $P(1 + r)^2 - P(1 + r) = P(1 + r)((1 + r) - 1) = P(1 + r)r$. Notice that $P(1 + r)$ is the amount at the end of year 1, which is $P + I_1$, so the formula $I_2 = (P + I_1)r$ used in the solution is consistent with the general formula.

93. Answer: a

Explanation:

Understanding the Slope of a Line

The slope of a line is a measure of its steepness. It tells us how much the vertical position (y-coordinate) changes for every unit change in the horizontal position (x-coordinate). The slope is often denoted by the letter 'm'.

Formula for Calculating Slope

To find the slope of a straight line passing through two distinct points, say $P_1(x_1, y_1)$ and $P_2(x_2, y_2)$, we use the following formula:

$$m = \frac{\text{Change in y}}{\text{Change in x}} = \frac{y_2 - y_1}{x_2 - x_1}$$

Applying the Slope Formula to the Given Points

We are asked to find the slope of the line joining the points $(3, -4)$ and $(5, 2)$.

Let's identify our points and their coordinates:

Point	x-coordinate	y-coordinate
Point 1 (P_1)	$x_1 = 3$	$y_1 = -4$
Point 2 (P_2)	$x_2 = 5$	$y_2 = 2$

Now, we substitute these coordinates into the slope formula:

$$m = \frac{y_2 - y_1}{x_2 - x_1}$$

$$m = \frac{2 - (-4)}{5 - 3}$$

First, calculate the difference in the y-coordinates ($y_2 - y_1$):

$$2 - (-4) = 2 + 4 = 6$$

Next, calculate the difference in the x-coordinates ($x_2 - x_1$):

$$5 - 3 = 2$$

Now, divide the change in y by the change in x:

$$m = \frac{6}{2}$$

$$m = 3$$

Resulting Slope Calculation

The slope of the line joining the points $(3, -4)$ and $(5, 2)$ is 3.

Revision Table: Slope Calculation Key Points

Concept	Description
Slope Definition	Measure of line steepness; $\frac{\text{Change in } y}{\text{Change in } x}$
Slope Formula	$m = \frac{y_2 - y_1}{x_2 - x_1}$ for points (x_1, y_1) and (x_2, y_2)
Given Points	$(3, -4)$ and $(5, 2)$
Calculated Slope	3

Additional Information on Linear Equations

Understanding the slope is fundamental in coordinate geometry and linear equations. Here are some related concepts:

- **Types of Slopes:** A line can have a positive slope (rises from left to right), a negative slope (falls from left to right), a zero slope (horizontal line), or an undefined slope (vertical line).
- **Equation of a Line:** The slope-intercept form of a linear equation is $y = mx + c$, where m is the slope and c is the y-intercept (the point where the line crosses the y-axis).
- **Parallel Lines:** Two distinct non-vertical lines are parallel if and only if they have the same slope.
- **Perpendicular Lines:** Two non-vertical lines are perpendicular if and only if the product of their slopes is -1 . A vertical line and a horizontal line are also perpendicular.

94. Answer: a

Explanation:

Understanding Moore's Law and Transistor Doubling

Moore's Law is a famous observation made by Gordon Moore, one of the co-founders of Intel. It's not a physical law of nature but rather an empirical rule of thumb that has accurately predicted the trend in the semiconductor industry for decades.

What is Moore's Law?

The core idea of Moore's Law is about the incredible pace of improvement in microchip technology. Specifically, it states that the number of transistors that can be placed on an integrated circuit, or microchip, tends to double periodically.

The Time Frame for Transistor Doubling

The question asks about the specific time period over which this doubling occurs, according to Gordon Moore's initial statement. Let's examine the options provided:

- 18 months
- 30 months
- 12 months
- 24 months

According to Gordon Moore's original observation in 1965, he predicted that the number of components (initially resistors, later transistors) on an integrated circuit would double approximately every year. However, this was later refined. The commonly cited version of Moore's Law, which has held true for a long time, suggests the doubling occurs approximately every 18 to 24 months, with 18 months being the most frequently quoted figure associated with the original prediction's refinement.

Given the options and the standard understanding of Moore's Law, the period for the number of transistors on a chip to double is typically cited as 18 months.

Significance of Moore's Law

Moore's Law has been a driving force behind the rapid advancements in computing and electronics over the past several decades. The ability to pack more and more transistors onto a smaller chip has led to:

- Increased processing power and speed
- Reduced cost per transistor
- Smaller and more energy-efficient devices

This continuous improvement has made possible everything from personal computers and smartphones to advanced AI systems and global communication networks.

Conclusion on Moore's Law Doubling Time

Based on the historical context and common understanding of Moore's Law, the number of transistors on a chip doubles approximately every 18 months.

Revision Table: Key Concepts of Moore's Law

Concept	Description
What it is	An observation/prediction about the semiconductor industry.
Stated by	Gordon Moore (Intel co-founder).
Core Idea	Number of transistors on a microchip doubles periodically.
Time Frame	Approximately every 18 months (commonly cited).
Impact	Drives innovation in computing, leads to more powerful, cheaper, smaller devices.

Additional Information on Moore's Law and Semiconductors

While often stated as exactly 18 months, it's worth noting that the specific period has varied over time and depends on various factors including technological breakthroughs and economic considerations. Some interpretations suggest the

doubling of performance or cost-efficiency every 18 months, while others focus strictly on transistor count.

Moore's Law has faced challenges in recent years as transistors approach atomic limits. While the rate of progress may be slowing down or changing, the semiconductor industry continues to find innovative ways to improve chip capabilities, although perhaps not strictly following the original 18-month doubling curve for transistor density alone.

The concept remains highly relevant as a benchmark and driver for research and development in microelectronics and computing power.

95. Answer: c

Explanation:

Understanding the Directive Principles of State Policy (DPSP)

The Directive Principles of State Policy, commonly known as DPSP, are a set of guidelines given in Part IV of the Constitution of India. These principles are not enforceable by any court, meaning you cannot go to court if the government doesn't follow them. However, they are considered fundamental in the governance of the country and it is the duty of the State to apply these principles in making laws.

The main idea behind DPSP is to guide the State towards establishing a welfare state where there is social and economic democracy. They aim to create a just society and ensure the well-being of all citizens.

Origin of Directive Principles in the Indian Constitution

When the Constitution of India was being drafted, the makers looked at different constitutions from around the world to borrow ideas that would be suitable for India.

The concept of Directive Principles of State Policy is one such idea that was adopted from another country's constitution.

The framers of the Indian Constitution were particularly inspired by the socio-economic objectives listed in the constitution of a specific nation. This inspiration led to the inclusion of DPSP in Part IV, Articles 36 to 51.

Analyzing the Options for DPSP Source

Let's look at the options provided and see which country's constitution was the source for the Directive Principles of State Policy:

- **Germany:** The Weimar Constitution of Germany influenced the provisions regarding the Emergency powers of the Union in the Indian Constitution. It is not the source for DPSP.
- **Canada:** The Canadian Constitution influenced the Indian federation system, particularly the idea of a strong Centre, residuary powers vesting with the Centre, and the appointment of state governors by the Centre. It is not the source for DPSP.
- **Ireland:** The Constitution of Ireland (1937) contains a set of 'Directive Principles of Social Policy'. The makers of the Indian Constitution found these principles similar in spirit and purpose to what they envisioned for India's governance and borrowed this concept.
- **Britain:** The British Constitution is an unwritten constitution based on conventions and parliamentary sovereignty. India borrowed the parliamentary system of government, rule of law, legislative procedure, single citizenship, cabinet system, etc., from Britain. It is not the source for DPSP.

Based on the historical facts of the Constitution drafting process, the Directive Principles of State Policy were directly inspired by and borrowed from the Constitution of Ireland.

Conclusion

The concept of Directive Principles of State Policy in the Indian Constitution is a significant feature aimed at achieving socio-economic goals. This concept was

adopted from the Constitution of Ireland.

Revision Table: Sources of Indian Constitution

Feature Borrowed	Source Country/Act
Directive Principles of State Policy	Ireland
Parliamentary government, Rule of Law, Legislative procedure, Single citizenship, Cabinet system, Prerogative writs	Britain
Federal Scheme, Office of Governor, Judiciary, Public Service Commissions, Emergency provisions, Administrative details	Government of India Act, 1935
Fundamental rights, Independent judiciary, Judicial review, Impeachment of the president, Removal of Supreme Court and High Court judges, Post of vice-president	USA
Federation with a strong Centre, Vesting of residuary powers in the Centre, Appointment of state governors by the Centre, Advisory jurisdiction of the Supreme Court	Canada
Emergency provisions (suspension of Fundamental Rights during Emergency)	Germany (Weimar Constitution)

Additional Information: Types and Significance of DPSP

The Directive Principles cover a wide range of objectives and can be broadly classified into three categories, although the Constitution does not formally classify them:

- **Socialist Principles:** Aim at providing social and economic justice and setting the path towards the welfare state. Examples: securing a social order for the promotion of the welfare of the people (Article 38), adequate means of livelihood (Article 39).

- **Gandhian Principles:** Based on the principles of Gandhi, representing the program of reconstruction advocated by Gandhi during the national movement. Examples: organisation of village panchayats (Article 40), promotion of cottage industries (Article 43).
- **Liberal-Intellectual Principles:** Reflect the ideology of liberalism. Examples: uniform civil code (Article 44), separation of judiciary from executive (Article 50).

The significance of DPSP lies in their role as moral principles for the government, acting as a yardstick to judge the performance of the government, and helping courts examine the constitutional validity of a law. They are crucial for the socio-economic development envisioned by the Constitution makers.

96. Answer: c

Explanation:

Understanding Oblique Projections and Their Types

Technical drawing and drafting use various methods to represent three-dimensional objects on a two-dimensional surface. These methods are called projections. One category of projection is oblique projection.

What is Oblique Projection?

In an oblique projection, the object is positioned with one face parallel to the viewing plane (or projection plane). Parallel lines in the object that are perpendicular to this face are projected at an angle to the projection plane, rather than perpendicular as in orthographic projection. The key characteristic of oblique projection is that the lines perpendicular to the projection plane are drawn at an angle, typically 30° , 45° , or 60° , and can be scaled to represent the depth.

There are two main types of oblique projections, distinguished by how the depth is represented:

- **Cavalier Projection:** In a cavalier projection, the lines representing the depth are drawn at full size (true length). This means if an object is 10 units deep, the lines showing that depth in the drawing are also 10 units long (before scaling the overall drawing size). The angle for the depth lines is often 45 degrees, but other angles are used. Cavalier projections can sometimes make objects appear longer than they are because the depth is shown at full scale.
- **Cabinet Projection:** In a cabinet projection, the lines representing the depth are foreshortened, typically drawn at half size (half of the true length). This helps to make the drawing look more realistic than a cavalier projection, reducing the appearance of elongation. The angle for the depth lines is also typically 45 degrees, but can vary.

Analyzing the Options

Let's look at the provided options in the context of projection types:

1. **perspective:** Perspective projection mimics how the human eye sees objects. Parallel lines appear to converge at vanishing points on the horizon line. This is different from oblique projection, where parallel lines remain parallel.
2. **fisheye:** A fisheye projection is a very wide-angle perspective projection that creates a distorted, curved image, similar to looking through a fisheye lens. This is not a standard type of technical drawing projection.
3. **cavalier:** As discussed above, a cavalier projection is a type of oblique projection where the depth is shown in full size. This directly matches the description in the question.
4. **orthographic:** Orthographic projection shows an object from different views (like top, front, side) where lines of sight are parallel and perpendicular to the projection plane. It's used for detailed technical drawings, but it represents views rather than a single pictorial representation with depth shown at an angle. While important, it's not an oblique projection.

Based on the definitions, the type of oblique projection in which the depth of the object is shown in full size is the cavalier projection.

Conclusion

The question asks for an oblique projection type where the depth is shown in full size. The analysis of projection types confirms that the cavalier projection fits this description perfectly.

Projection Type	Description	Depth Scaling
Cavalier Projection	Oblique projection with face parallel to view plane, depth lines at angle.	Full size (true length)
Cabinet Projection	Oblique projection with face parallel to view plane, depth lines at angle.	Half size (foreshortened)
Orthographic Projection	Views projected perpendicularly onto planes (e.g., Front, Top, Side).	Not applicable (shows views)
Perspective Projection	Lines converge to vanishing points, mimics human vision.	Varies based on distance

Revision Table: Oblique Projection Key Points

Term	Key Characteristic
Oblique Projection	One face parallel to viewing plane; depth lines at angle (e.g., 30°, 45°, 60°).
Cavalier Projection	Oblique projection where depth is full size.
Cabinet Projection	Oblique projection where depth is half size.
Projection Plane	The 2D surface onto which the 3D object is projected.

Additional Information: Related Projection Concepts

While oblique projection is useful for quick pictorial representations, other projection types are also widely used in technical drawing:

- **Orthographic Projection:** Provides multiple 2D views showing the true size and shape of features, essential for manufacturing and detailed design. Examples include front, top, and side views.
- **Axometric Projection:** A type of parallel projection where the object is rotated so that none of its faces are parallel to the projection plane. All parallel lines in the object remain parallel in the drawing.
 - **Isometric Projection:** A common type where the three axes are equally foreshortened, and the angles between them are 120 degrees. Lines along the axes are often drawn at their true length (isometric lines).
 - **Dimetric Projection:** Two of the three axes are equally foreshortened.
 - **Trimetric Projection:** All three axes are foreshortened by different amounts.
- **Perspective Projection:** Used in architectural drawings and illustrations to create a realistic view by simulating how objects appear smaller as they recede into the distance.

97. Answer: a

Explanation:

Understanding the Technical Notation: $2 \times \varnothing 6$

The notation $2 \times \varnothing 6$ is commonly found in technical drawings, engineering plans, or product specifications. It's a concise way to convey specific geometric information about a part or feature.

Let's break down the components of this notation:

- $\varnothing$: This symbol is the standard engineering and technical drawing symbol for **diameter**. It indicates that the dimension following it is the diameter of a circular feature.

- **6:** This number represents the value of the diameter, usually in specified units (e.g., millimeters, inches, units as defined elsewhere in the drawing). In this context, it means the diameter is 6 units.
- **2 ×:** This part indicates multiplicity. It means that the feature described by the subsequent notation ($\varnothing 6$) appears 2 times.

Putting it all together, $2 \times \varnothing 6$ means that there are two instances of a feature that has a diameter of 6 units. Given the options involve circles, this notation most accurately describes two circles, each having a diameter of 6 units.

Analyzing the Given Options

Let's examine each option in light of our understanding of the $2 \times \varnothing 6$ notation:

- **Option 1: 2 circles of diameter 6 units** – This directly matches our interpretation. The '2 ×' signifies two items, and ' $\varnothing 6$ ' specifies that each item is a circle with a diameter of 6 units.
- **Option 2: Scale the circle of diameter 6 units by a factor of 2** – This option suggests scaling, which means changing the size of a single circle. The notation $2 \times \varnothing 6$ does not imply scaling; it implies having multiple identical features. Scaling a diameter 6 circle by a factor of 2 would result in a diameter 12 circle.
- **Option 3: 2 circles of radius 6 units** – The notation uses the diameter symbol ($\varnothing$), not the radius symbol (R). If the diameter is 6 units, the radius is half of that, which is 3 units. This option is incorrect because it uses radius instead of diameter, and the value is incorrect for the specified diameter.
- **Option 4: Scale the circle of radius 6 units by a factor of 2** – This option is incorrect because it mentions radius 6 (instead of diameter 6) and implies scaling by a factor of 2 (instead of indicating two separate circles). A circle with radius 6 units has a diameter of 12 units. Scaling that by a factor of 2 would result in a circle with radius 12 or diameter 24.

Why "2 circles of diameter 6 units" is the Correct Interpretation

Based on the standard conventions of technical drawing notations, the prefix number followed by the multiplication symbol ($\times$) indicates the quantity of the feature being specified. The $\varnothing$ symbol always denotes diameter. Therefore, $2 \times \varnothing 6$ unambiguously indicates that there are two separate circular features, and each of these features has a diameter measuring 6 units.

Part of Notation	Meaning
$2 \times$	Quantity: There are 2 of the specified features.
$\varnothing$	Geometric shape/property: Diameter.
6	Size: The value is 6 units.

Revision Table: Key Concepts in Geometric Notation

Term	Symbol	Definition	Relationship
Diameter	$\varnothing$ (or d)	A straight line passing from side to side through the center of a circle or sphere.	Diameter = $2 \times$ Radius ($d = 2r$)
Radius	R (or r)	A straight line from the center to the circumference of a circle or sphere.	Radius = Diameter / 2 ($r = d/2$)
Circle	—	A round plane figure whose boundary (the circumference) consists of points equidistant from a fixed center point.	Defined by its center and either its radius or diameter.

Additional Information on Technical Drawing Notations

Technical drawings use various symbols and conventions to convey precise information about dimensions, tolerances, and features. Understanding these notations is crucial for correctly interpreting engineering specifications.

- Notations like $2 \times \text{Ø}6$ are part of dimensioning, which specifies the size and location of features.
- Other common symbols include R for radius, $\square$ for square, $\bigcirc$ for spherical diameter, etc.
- The position of the dimension and notation on the drawing also provides context, indicating which feature is being described.

98. Answer: c

Explanation:

Understanding Speed, Distance, and Time

This problem involves the relationship between speed, distance, and time. When an object travels at a constant speed, the distance covered is the product of its speed and the time taken.

The fundamental formula connecting these quantities is:

$$\text{Distance} = \text{Speed} \times \text{Time}$$

We can rearrange this formula to find Speed or Time:

- $\text{Speed} = \text{Distance} / \text{Time}$
- $\text{Time} = \text{Distance} / \text{Speed}$

In this problem, the distance traveled is the same in both scenarios. We need to find this distance first, and then use it to calculate the new speed required and the increase in speed.

Step 1: Calculate the Distance Traveled

First, let's identify the given information for the initial journey:

- Original Speed = 64 km/h
- Original Time = 50 minutes

Since the speed is given in km/h, we must convert the time from minutes to hours to maintain consistent units. There are 60 minutes in 1 hour.

$$\text{Original Time in hours} = \frac{50 \text{ minutes}}{60 \text{ minutes/hour}} = \frac{50}{60} \text{ hours} = \frac{5}{6} \text{ hours}$$

Now, we can calculate the distance traveled using the formula Distance = Speed × Time:

$$\text{Distance} = 64 \text{ km/h} \times \frac{5}{6} \text{ hours}$$

$$\text{Distance} = \frac{64 \times 5}{6} \text{ km} = \frac{320}{6} \text{ km} = \frac{160}{3} \text{ km}$$

So, the distance traveled is $\frac{160}{3}$ km.

Step 2: Calculate the New Speed Required

The train needs to travel the same distance in a shorter time. The new time is given as 40 minutes.

- Distance = $\frac{160}{3}$ km (from Step 1)
- New Time = 40 minutes

Again, convert the new time from minutes to hours:

$$\text{New Time in hours} = \frac{40 \text{ minutes}}{60 \text{ minutes/hour}} = \frac{40}{60} \text{ hours} = \frac{2}{3} \text{ hours}$$

Now, we calculate the new speed required using the formula Speed = Distance / Time:

$$\text{New Speed} = \frac{\text{Distance}}{\text{New Time}}$$

$$\text{New Speed} = \frac{\frac{160}{3} \text{ km}}{\frac{2}{3} \text{ hours}}$$

$$\text{New Speed} = \frac{160}{3} \times \frac{3}{2} \text{ km/h}$$

$$\text{New Speed} = \frac{160}{2} \text{ km/h} = 80 \text{ km/h}$$

The train needs to travel at a speed of 80 km/h to cover the same distance in 40 minutes.

Step 3: Calculate the Increase in Speed

The question asks for the increase in speed, which is the difference between the new speed and the original speed.

- Original Speed = 64 km/h
- New Speed = 80 km/h

$$\text{Increase in Speed} = \text{New Speed} - \text{Original Speed}$$

$$\text{Increase in Speed} = 80 \text{ km/h} - 64 \text{ km/h}$$

$$\text{Increase in Speed} = 16 \text{ km/h}$$

The train should increase its speed by 16 km/h.

Summary of Train Speed Calculation

To travel the same distance in less time, the train's speed must increase. We first calculated the distance and then the required speed for the shorter time, finding the difference.

Quantity	Initial Scenario	New Scenario
Speed	64 km/h	80 km/h (Calculated)
Time	50 min ($\frac{5}{6}$ h)	40 min ($\frac{2}{3}$ h)
Distance	$64 \times \frac{5}{6} = \frac{160}{3}$ km (Calculated)	$80 \times \frac{2}{3} = \frac{160}{3}$ km (Same Distance)

Revision Table: Train Speed Problem

Item	Value
Original Speed	64 km/h
Original Time	50 minutes ($\frac{5}{6}$ hours)
Calculated Distance	$\frac{160}{3}$ km
New Time	40 minutes ($\frac{2}{3}$ hours)
Calculated New Speed	80 km/h
Increase in Speed	16 km/h

Additional Information: Speed, Distance, and Time

Understanding the relationship between speed, distance, and time is crucial for solving problems like this. Here are a few key points:

- **Inverse Relationship:** When the distance is kept constant, speed and time are inversely proportional. This means if you want to cover the same distance in less time, you must increase your speed. Similarly, if you decrease your speed, it will take more time to cover the same distance.
- **Unit Consistency:** Always ensure that the units for speed, distance, and time are consistent when using the formulas. If speed is in km/h, distance should be in km and time in hours. If speed is in m/s, distance should be in meters and time in seconds. Unit conversion (like minutes to hours) is a common step in these problems.
- **Constant Speed Assumption:** Problems like this usually assume the speed is constant throughout the journey. In reality, speeds can vary.

99. Answer: c

Explanation:

Understanding the Question: Dadasaheb Phalke Award and India's First Colour Film

The question asks to identify a recipient of the prestigious Dadasaheb Phalke Award who holds the distinction of producing and directing India's first colour film, titled 'Sairandhri'. This requires knowledge about prominent figures in early Indian cinema and their contributions, specifically concerning the transition to colour films.

Identifying the Pioneer of India's First Colour Film

India's journey into colour cinema began with the film 'Sairandhri'. This historical film was known for its use of the Agfacolor process. The individual responsible for bringing this cinematic milestone to life as both its producer and director was a significant figure in Indian film history, who was later honoured with the Dadasaheb Phalke Award.

Among the options provided, we need to find the person who fits this description:

- Sivaji Ganesan: A legendary actor, primarily known for his work in Tamil cinema. While a highly respected figure, he was not known for directing or producing 'Sairandhri'.
- LV Prasad: A prominent actor, director, producer, and businessman. He played a crucial role in establishing film production houses in South India, but he is not associated with directing 'Sairandhri'.
- V Shantaram: A renowned filmmaker, actor, and producer known for his socially relevant films and technical innovations. He was indeed associated with early colour experiments in Indian cinema and later received the Dadasaheb Phalke Award.
- Birendranath Sircar: A key figure in Bengali cinema, founder of New Theatres. He was a pioneering producer but not associated with directing 'Sairandhri'.

Historical records indicate that the filmmaker V. Shantaram produced and directed 'Sairandhri'. Although the film was shot in colour using Agfacolor, it was mostly shown in monochrome in India due to technical and logistical constraints at the time. However, it is recognised as the first film where significant portions were shot and processed in colour.

Conclusion: V Shantaram's Contribution

Based on the historical facts, V. Shantaram is the filmmaker who produced and directed 'Sairandhri', the first film where significant colour filming was attempted in India, and he is also a recipient of the Dadasaheb Phalke Award. Therefore, he is the correct answer.

Revision Table: Key Information about Sairandhri

Aspect	Detail
Film Title	Sairandhri
Claim to Fame	Recognised as India's first film with significant colour shooting (using Agfacolor).
Producer	V. Shantaram
Director	V. Shantaram
Year (Release)	1933

Additional Information on Dadasaheb Phalke Award and Early Indian Cinema

The Dadasaheb Phalke Award is India's highest award in cinema, presented annually by the Directorate of Film Festivals to individuals for their outstanding contribution to the growth and development of Indian cinema. It was established in 1969.

V. Shantaram was a prolific filmmaker known for films like 'Duniya Na Mane', 'Dr. Kotnis Ki Amar Kahani', 'Do Aankhen Barah Haath', and others. His work often explored social themes and pushed technical boundaries. His involvement with 'Sairandhri' highlights the early efforts to introduce colour technology into Indian filmmaking, paving the way for the vibrant colour films we see today.

The transition from silent films to talkies, and then to colour, marked significant evolutionary steps in Indian cinema, each driven by visionary filmmakers like V. Shantaram.

100. Answer: b

Explanation:

Given:

Total marks = 1800

Calculation:

$$\Rightarrow \text{Total degree of marks} = 360^\circ$$

$$\Rightarrow \text{Marks scored in Science} = \frac{82}{360} \times 1800 = 82 \times 5 = 410$$

$$\Rightarrow \text{Marks scored in Mathematics} = \frac{84}{360} \times 1800 = 84 \times 5 = 420$$

$$\Rightarrow \text{Difference} = 420 - 410 = 10$$

Therefore, the difference between the marks scored in Mathematics and Science is 10.

101. Answer: b

Explanation:

Understanding Phase Difference in Three-Phase Induction Motors

A three-phase induction motor is a type of AC electric motor in which the electric current in the rotor needed to produce torque is obtained by electromagnetic

induction from the magnetic field of the stator winding. The stator of a three-phase induction motor typically has three separate windings. These windings are physically placed around the stator core.

Phase Windings and Their Arrangement

In a three-phase system, the AC voltage supplied has three distinct phases. For a three-phase induction motor to operate correctly and produce a rotating magnetic field, the three stator windings are designed to interact with these three phases of the power supply. The windings are typically distributed spatially around the stator bore.

Angular Phase Difference Explained

The term "angular phase difference" refers to the difference in the phase angle of the voltage or current waveforms between the different phases in an AC system. In a balanced three-phase system, the three phase voltages or currents are displaced from each other by a specific angle.

For a standard three-phase system, the phase voltages (or currents) are displaced by an angular difference of 120 degrees electrically. This means that if phase A is at its peak, phase B will reach its peak 120 electrical degrees later, and phase C will reach its peak 120 electrical degrees after phase B (or 240 electrical degrees after phase A).

To effectively utilize this three-phase power and create a smoothly rotating magnetic field, the three windings in the stator of a three-phase induction motor are typically placed and connected in such a way that they are also electrically spaced 120 degrees apart. When the three-phase voltages are applied to these windings, they produce magnetic fields that are also displaced by 120 electrical degrees in time and space, resulting in a rotating magnetic field.

Analyzing the Options

Let's look at the given options for the angular phase difference:

- **360°:** This would mean the phases are in sync, which is not a characteristic of a multi-phase system like three-phase.
- **120°:** This is the standard electrical phase difference between phases in a balanced three-phase system. The windings are designed to match this.
- **180°:** This represents an anti-phase relationship, typical in split-phase systems or between opposite ends of a single winding, but not the difference between distinct phases in a standard three-phase supply or motor winding arrangement.
- **90°:** This is the phase difference found in two-phase systems (like some types of single-phase motors using auxiliary windings), not in a three-phase system.

The angular phase difference between each phase winding of a three-phase induction motor is designed to align with the electrical phase difference of the supply, which is 120 degrees.

Summary Table

Concept	Angular Phase Difference	Relevance to Three-Phase Motor
Three-phase supply voltage/current	120° electrical	Input to the motor windings
Three-phase motor windings	Designed for 120° electrical displacement	Arranged to produce a rotating magnetic field matching the supply phase difference

Therefore, the angular phase difference between each phase winding of a three-phase induction motor is 120°.

Revision Table: Three-Phase System Basics

Feature	Description
Number of Phases	Three
Phase Displacement	120° electrical degrees
Voltage/Current Waveforms	Sinusoidal, offset by 120°
Application in Motors	Creates a rotating magnetic field

Additional Information: Electrical vs. Mechanical Degrees

It's important to distinguish between electrical degrees and mechanical degrees. Electrical degrees relate to the phase difference in time of the waveforms, while mechanical degrees relate to physical angle in space (like shaft rotation). In a motor, these are related by the number of poles:

$$\text{Electrical degrees} = \text{Mechanical degrees} \times (\text{Number of poles} / 2)$$

The 120° phase difference discussed for the windings is typically referred to in electrical degrees, matching the electrical phase difference of the power supply.

Your Personal Exams Guide

102. Answer: b

Explanation:

Understanding Inductor Function in Filter Circuits

Filter circuits are essential components in electronics used to select or reject specific frequencies from a signal. Different components like resistors, capacitors, and inductors behave differently towards direct current (DC) and alternating current (AC), making them suitable for filtering applications. This question asks specifically about the role of an inductor in such a circuit.

How Inductors Behave with AC and DC

An inductor is a passive electrical component that stores energy in a magnetic field when electric current flows through it. Its key characteristic is its opposition to changes in current flow. This property affects its behaviour towards DC and AC signals differently.

- **DC Signals:** A DC signal is a constant current. Once a steady state is reached, there is no change in current. Therefore, the inductor offers very little opposition to DC. Ideally, a pure inductor acts like a short circuit (zero resistance) for DC. Its reactance to DC is $X_L = 0$ because the frequency f is 0 for DC ($X_L = 2\pi fL$).
- **AC Signals:** An AC signal is a continually changing current. An inductor actively opposes these changes. The opposition an inductor offers to AC is called inductive reactance, denoted by X_L . The value of inductive reactance depends on the frequency of the AC signal and the inductance value (L) of the inductor, given by the formula:

$$X_L = 2\pi fL$$

This formula shows that inductive reactance is directly proportional to the frequency (f). Higher frequencies encounter greater opposition from the inductor.

Inductor's Role in Filter Circuits: Blocking AC, Passing DC

Based on the behavior described above, an inductor offers low opposition (ideally zero) to DC and high opposition to high-frequency AC. This property makes inductors useful in filter circuits, particularly in applications where you need to allow DC to pass while attenuating or blocking AC components.

For example, in a DC power supply, after rectification, the output contains a DC component along with residual AC components called ripple. An inductor placed in series with the load will have low impedance to the DC component, allowing it to pass through to the load easily. However, it will have high impedance to the AC ripple component, significantly reducing its amplitude reaching the load. This effectively 'blocks' AC ripple while 'passing' the DC voltage.

Let's summarize the behavior:

Signal Type	Inductor Opposition (Reactance X_L)	Effect on Signal
DC (Frequency $f = 0$)	Very Low ($X_L \approx 0$)	Passes Easily
AC (Frequency $f > 0$)	High and Increases with Frequency ($X_L = 2\pi fL$)	Opposes Strongly, Blocks Effectively (especially high frequencies)

Analyzing the Options

- **Option 1: block DC and pass AC** – This is incorrect. An inductor passes DC easily (low opposition) and blocks AC (high opposition).
- **Option 2: block AC and pass DC** – This is correct. An inductor's high reactance to AC and low reactance to DC allows it to block AC signals while letting DC signals pass through relatively unimpeded.
- **Option 3: block both AC and DC** – This is incorrect. An inductor does not block DC; it offers very little opposition to it.
- **Option 4: pass both AC and DC** – This is incorrect. While it passes DC, it significantly opposes and blocks AC, especially at higher frequencies.

Therefore, the primary function of an inductor in filter circuits, leveraging its fundamental electrical properties, is to block AC and pass DC.

Revision Table: Inductors and Filters

Component	Behavior with DC	Behavior with AC	Filter Application
Inductor	Low Opposition (Passes DC)	High Opposition (Blocks AC, increases with frequency)	Used to pass DC and block AC (e.g., in series for filtering DC power supply ripple)
Capacitor	High Opposition (Blocks DC)	Low Opposition (Passes AC, decreases with frequency)	Used to block DC and pass AC (e.g., for coupling AC signals) or shunt AC to ground (e.g., in parallel for filtering ripple)
Resistor	Constant Opposition (Resists both AC and DC)	Constant Opposition (Resistance is generally independent of frequency)	Used for voltage division or current limiting, not typically for frequency-selective blocking/passing based on its inherent resistance property alone in simple filters.

Additional Information on Filter Circuits

Filter circuits are broadly classified based on the frequency range they allow to pass:

- **Low-Pass Filter:** Allows low frequencies (including DC) to pass and attenuates high frequencies. An inductor in series or a capacitor in parallel can form a basic low-pass filter.
- **High-Pass Filter:** Allows high frequencies to pass and attenuates low frequencies (including DC). A capacitor in series or an inductor in parallel can form a basic high-pass filter.
- **Band-Pass Filter:** Allows a specific range of frequencies to pass and attenuates frequencies outside this band. Often uses a combination of inductors and capacitors.

- **Band-Stop (Notch) Filter:** Attenuates a specific range of frequencies and allows frequencies outside this band to pass. Also typically uses a combination of inductors and capacitors.

The question specifically highlights the inductor's role which is fundamental to its use in low-pass filtering applications where DC and low frequencies need to be passed, and higher frequencies need to be blocked.

103. Answer: c

Explanation:

Understanding Constant Voltage Chargers and Fast Charging

A constant voltage charger is a type of battery charger that maintains a constant voltage output across the battery terminals while charging. This method is commonly used in certain phases of charging for various battery types, especially lithium-ion batteries, and is particularly associated with achieving faster charging times under specific conditions compared to older or simpler charging methods.

How Constant Voltage Charging Works

When a battery is charged using the constant voltage method, the charger keeps the output voltage at a set level. Initially, when the battery's voltage is low, it draws a high current. As the battery's voltage rises and gets closer to the charger's set voltage, the charging current automatically decreases. This process continues until the current drops to a very low level, at which point the battery is considered fully charged or the charging transitions to a different stage (like trickle charging for some battery types, although less common for lithium-ion).

Comparing Charging Methods

Let's look at the options provided and how they relate to constant voltage charging:

- **Trickle charging:** This is a very low rate of charging, used to maintain a battery at full charge or to charge very slowly over a long period. It's the opposite of fast charging and not typically the primary function associated with the constant voltage *phase* which initially delivers high current.
- **Slow charging:** This involves charging at a relatively low current rate. While a constant voltage charger will transition to lower current as the battery fills, its initial phase often involves higher current than a dedicated slow charger. Slow charging is generally safer and better for battery health over the long term but takes much longer.
- **Fast charging:** This aims to charge the battery quickly by delivering a higher current initially. Constant voltage charging, especially when combined with an initial constant current phase (as in CC/CV charging for lithium-ion), is a key technique used to achieve fast charging. The constant current phase rapidly increases the voltage, and then the constant voltage phase brings the battery to full charge while the current tapers off.
- **Boost voltage:** While a charger does output a voltage higher than the battery's current voltage to push current into it, "boost voltage" isn't a standard term for a charging *method*. It might refer to the voltage level used, but it doesn't describe the overall charging strategy itself. Constant voltage charging sets a specific target voltage for the charging process.

Why Constant Voltage Supports Fast Charging

Modern fast charging systems, particularly for technologies like lithium-ion batteries, often use a combination of constant current (CC) and constant voltage (CV). The constant current phase quickly raises the battery voltage, and then the constant voltage phase brings it to full capacity efficiently without overcharging. The constant voltage stage allows the charger to deliver decreasing current as the battery voltage nears the target, managing the final stage of charge safely and relatively quickly compared to purely slow methods.

Therefore, among the given options, the constant voltage charger is most directly associated with enabling or being part of a system designed for fast charging.

Charging Method	Description	Relation to Constant Voltage	Speed
Constant Voltage (CV)	Maintains a fixed voltage; current decreases as battery charges. Often follows Constant Current (CC) phase.	Core method or phase in modern charging.	Part of Fast charging (especially combined with CC).
Trickle Charging	Very low current charging to maintain charge.	Distinct from CV; low current phase.	Very Slow
Slow Charging	Charging at a low, steady current.	Different method; lower current than initial CV/CC.	Slow
Fast Charging	Charging at high current/power to reduce time.	Often uses CC/CV method.	Fast
Boost Voltage	Not a standard charging method term; refers to voltage level used.	CV sets the target "boost" voltage.	N/A (not a method)

Conclusion on Constant Voltage Chargers

Based on the analysis of charging methods, the constant voltage method, particularly in modern battery charging schemes like CC/CV, plays a crucial role in achieving fast charging. It allows for efficient and relatively quick completion of the charging cycle after the initial high-current phase.

Revision Table: Constant Voltage Charger Benefits

Concept	Explanation	Relation to Constant Voltage Charger
Fast Charging	Reducing the time required to charge a battery significantly.	Constant voltage is a key phase (often following constant current) in modern fast charging profiles for batteries like Li-ion.
Battery Health	Maintaining the lifespan and performance of the battery.	CV phase safely tapers current to prevent overcharging towards the end, contributing to health.
Efficiency	Converting electrical energy into stored chemical energy effectively.	The CV phase ensures efficient topping off of the battery charge.

Additional Information: Constant Current vs. Constant Voltage

Most advanced battery charging systems, especially for lithium-ion batteries found in phones, laptops, and electric vehicles, don't use just one method but employ a combination. The most common is the Constant Current/Constant Voltage (CC/CV) method:

- 1. Constant Current (CC) Phase:** The charger delivers a constant high current to the battery. The battery voltage increases steadily during this phase. This is the fastest part of the charging process.
- 2. Constant Voltage (CV) Phase:** Once the battery voltage reaches a specific level (the target voltage, e.g., 4.2V for a standard Li-ion cell), the charger switches to maintaining this constant voltage. The current gradually decreases as the battery approaches full charge. This phase ensures the battery is fully charged without overcharging.

A constant voltage charger directly implements the CV phase of this combined method, which is essential for completing the fast charging cycle safely and

efficiently.

104. Answer: a

Explanation:

Understanding Electrical Safety for Household Wiring

Electrical safety is paramount in any wiring system, especially in residential settings and small units. Devices are installed to protect against faults like overcurrents (overload) and short circuits, which can cause damage to appliances, wiring, and even lead to fires or electric shocks.

Safety Measures in Household Wiring

Various safety devices are used in electrical circuits. The question asks about the appropriate safety measure for household wiring and small units among the given options, which are types of circuit breakers.

- **Circuit Breaker:** A circuit breaker is an automatic electrical switch designed to protect an electrical circuit from damage caused by overcurrent or short circuit. Unlike a fuse, which operates once and then must be replaced, a circuit breaker can be reset (either manually or automatically) to resume normal operation.

Analyzing the Options

Let's look at the options provided:

1. **MCB (Miniature Circuit Breaker):** These are electromechanical devices that protect an electric circuit from overcurrents, which can result from both overload and short circuits. They are designed for low voltage electrical networks, typically used in residential, commercial, and light industrial applications. Their current ratings usually range from 0.5 A to 63 A.

2. **OCB (Oil Circuit Breaker):** These circuit breakers use oil as a dielectric or insulating medium and also for arc extinction. They are typically used in high voltage applications (like transmission and distribution substations) and are not suitable for household use.
3. **ACB (Air Circuit Breaker):** These circuit breakers use air as the arc quenching medium at atmospheric pressure. They are generally used in low voltage installations but for much higher current ratings (typically from 800 A up to 10,000 A) compared to MCBs, often found in main distribution panels in large buildings or industrial settings.
4. **MCCB (Moulded Case Circuit Breaker):** These are also used for protection against overcurrents and short circuits but are designed for higher currents (typically from 100 A to 2500 A) than MCBs. They offer adjustable trip settings in some models and are commonly used in commercial and industrial applications as main breakers or for protecting large equipment.

Why MCB is Suitable for Household Wiring

Household wiring and small units involve relatively low current requirements compared to large commercial or industrial setups. Fault currents are also within the range that MCBs can safely interrupt. MCBs are compact, easy to install in distribution boards, reliable, and cost-effective for residential use. Their current ratings match the typical loads found in homes (lights, fans, appliances). The quick response of an MCB to an overload or short circuit helps protect the relatively thin wires used in household circuits from overheating and potential fire hazards.

Comparing the options, OCBs, ACBs, and MCCBs are designed for much higher power applications and are physically larger and more expensive. Using them in a typical household circuit would be impractical and unnecessary overkill.

Therefore, the Miniature Circuit Breaker (MCB) is the appropriate and commonly used safety measure for protection against overcurrents and short circuits in household wiring and small units.

Summary of Circuit Breaker Types

Type	Medium	Typical Application	Current Range (Approx.)
MCB	Air	Household, Small Commercial	0.5 A – 63 A
OCB	Oil	High Voltage Substations	High
ACB	Air	Large Commercial, Industrial Mains	800 A – 10,000 A
MCCB	Air	Commercial, Industrial Distribution	100 A – 2500 A

Revision Table: Electrical Safety Devices

Device	Function	Typical Use Case	Key Feature
MCB	Overload and Short Circuit Protection	Household Wiring, Small Units	Resettable, Low Current Ratings
Fuse	Overcurrent Protection (One-time use)	Older Installations, Specific Appliances	Melts and Breaks Circuit
RCCB/ELCB	Earth Leakage Protection	Protection against Electric Shock	Detects Current Imbalance

Additional Information: Household Electrical Safety Concepts

Overload vs. Short Circuit

- **Overload:** Occurs when too many appliances are connected to a circuit, drawing more current than the wiring and protective device are designed to handle. This causes overheating over time.
- **Short Circuit:** Occurs when there is an unintended, low-resistance path between two points in a circuit (e.g., live wire touching neutral or earth wire). This results in a very large current flowing instantaneously.

Both overload and short circuits cause overcurrent, but a short circuit typically involves a much higher and faster rise in current, which can be far more destructive.

MCB Trip Characteristics

MCBs have different tripping characteristics (e.g., Type B, Type C, Type D) based on how quickly they react to different levels of overcurrent. Type B and Type C are commonly used in residential settings depending on the type of loads (resistive vs. inductive).

Importance of Proper Rating

Choosing the correct current rating for an MCB is crucial. An undersized MCB will trip frequently under normal load, while an oversized MCB will fail to protect the circuit adequately during a fault, potentially leading to wiring damage or fire.

105. Answer: b

Explanation:

Correct answer: Galvanized rigid conduit pipe

Galvanized rigid conduit pipe (GRC): Made of steel and coated with a layer of zinc. This galvanization process provides good protection against rust and corrosion. GRC offers high mechanical strength, making it very durable and resistant to physical damage. It is suitable for both indoor and outdoor industrial wiring installations and can withstand harsh conditions.

Based on this analysis, Galvanized Rigid Conduit pipe stands out as a highly suitable and commonly used type of conduit pipe for industrial wiring applications due to its robust nature.

106. Answer: c

Explanation:

Understanding Energy Meter Torque Production

A single-phase induction type energy meter works based on the principle of electromagnetic induction. It has two main electromagnets: the shunt magnet and the series magnet. These magnets produce magnetic fluxes that interact with an aluminum disc, causing it to rotate. The speed of rotation of the disc is proportional to the power being consumed, and the total number of rotations over time indicates the total energy consumed.

Key Components and Fluxes

- **Shunt Magnet:** Connected across the supply voltage. It has a voltage coil with a large number of turns and is highly inductive. This magnet produces a magnetic flux (shunt flux, Φ_s) proportional to the supply voltage.
- **Series Magnet:** Connected in series with the load. It has a current coil with a few turns. This magnet produces a magnetic flux (series flux, Φ_{se}) proportional to the load current.
- **Aluminum Disc:** Placed between the air gaps of the shunt and series magnets. Eddy currents are induced in this disc due to the changing magnetic fluxes.

How Torque is Generated in an Induction Energy Meter

The driving torque that rotates the aluminum disc is produced by the interaction between the magnetic fluxes (Φ_s and Φ_{se}) and the eddy currents induced in the disc by the other flux. Specifically, the torque is primarily due to:

- Interaction of Φ_s with the eddy currents induced by Φ_{se} .
- Interaction of Φ_{se} with the eddy currents induced by Φ_s .

The total driving torque (T_d) is proportional to the product of the two fluxes and the sine of the phase angle between them. However, for accurate energy

measurement, the torque needs to be proportional to the instantaneous power ($VI \cos \phi$).

Phase Relationship for Maximum Torque and Correct Operation

In an ideal induction type energy meter, the shunt flux (Φ_s) produced by the voltage coil should lag the supply voltage (V) by exactly 90 degrees. The series flux (Φ_{se}) produced by the current coil is nearly in phase with the load current (I).

Let's consider the ideal phase relationships:

- Supply Voltage: V at angle 0° (reference).
- Supply Current: I at angle $-\phi$ (assuming lagging power factor, ϕ is the power factor angle).
- Ideal Shunt Flux: Φ_s lags V by 90° , so Φ_s is at angle -90° .
- Series Flux: Φ_{se} is in phase with I , so Φ_{se} is at angle $-\phi$.

The phase angle between the shunt flux (Φ_s) and the series flux (Φ_{se}) is $\theta = (-\phi) - (-90^\circ) = 90^\circ - \phi$.

The driving torque is proportional to $\Phi_s \Phi_{se} \sin(\text{angle between fluxes})$. For correct energy measurement, the torque should be proportional to $VI \cos \phi$. The torque generated is actually proportional to $\Phi_s \Phi_{se} \sin(90^\circ - \phi)$, which is proportional to $\Phi_s \Phi_{se} \cos \phi$. Since $\Phi_s \propto V$ and $\Phi_{se} \propto I$, the torque is proportional to $VI \cos \phi$, which is the instantaneous power.

To achieve the ideal 90° lag of the shunt flux behind the supply voltage, the voltage coil is designed to be highly inductive. Additionally, a compensating coil or an inductive shunt (known as the 'lag adjustment') is used on the shunt magnet to ensure this precise phase relationship.

While maximum torque between *any* two interacting fluxes for a given magnitude occurs when they are 90 degrees apart, in the context of an energy meter measuring power ($VI \cos \phi$), the crucial requirement for accurate registration is that the shunt flux lags the supply voltage by 90 degrees. This specific phase shift ensures the resulting torque is proportional to the true power.

Analyzing the Options

- **Is in phase with the supply voltage:** If Φ_s is in phase with V , the torque relationship will not result in proportionality to $VI \cos \phi$.
- **Lags the supply voltage by 45 degrees:** This would also result in incorrect torque proportionality.
- **Lags the supply voltage by 90 degrees:** This is the ideal condition for an induction type energy meter to produce driving torque that is proportional to the power being measured ($VI \cos \phi$), leading to accurate energy registration.
- **Leads the supply voltage by 90 degrees:** This would cause the torque to act in the opposite direction or lead to incorrect measurement.

Therefore, maximum torque (in the sense of proper operation and proportionality to power) is produced when the shunt magnetic flux lags the supply voltage by 90 degrees.

Revision Table: Induction Energy Meter Phases

Component/Parameter	Ideal Phase Relationship (Relative to Supply Voltage V)
Supply Voltage (V)	0° (Reference)
Shunt Magnetic Flux (Φ_s)	Lags V by 90°
Supply Current (I)	Lags V by ϕ (Power factor angle)
Series Magnetic Flux (Φ_{se})	In phase with I (Lags V by ϕ)
Angle between Φ_s and Φ_{se}	$90^\circ - \phi$

Additional Information: Lag Adjustment

Achieving the exact 90-degree lag of the shunt flux behind the supply voltage is critical for the accuracy of an induction type energy meter. Since the voltage coil

has resistance, the flux it produces naturally lags the voltage by slightly less than 90 degrees. To correct this, a lagging device or phase compensator is used. This is typically a copper shading band or a small compensating coil placed around the central limb of the shunt magnet. Adjusting this band or coil changes the phase of the shunt flux, allowing it to be precisely lagged by 90 degrees relative to the voltage, thus ensuring the torque is proportional to $VI \cos \phi$ for all power factors.

107. Answer: d

Explanation:

Understanding Electrical Connections and the Common Neutral Point

In electrical systems, especially three-phase alternating current (AC) systems, different configurations are used to connect the power source (like a generator or transformer) to the load. Two common configurations are the Star connection and the Delta connection. The question asks which of the listed connections features a common neutral point.

Let's look at the standard electrical connections presented in the options:

- **Star Connection:** Also known as the Wye connection, this setup involves connecting one end of each of the three phases together at a single point. This central point is called the neutral point. The other end of each phase winding is connected to a line terminal (L1, L2, L3). This configuration provides both line voltages (voltage between two lines) and phase voltages (voltage between a line and the neutral point). The neutral point is often grounded or connected to a neutral wire, providing a return path for current, especially in unbalanced loads.
- **Delta Connection:** In this configuration, the three phase windings are connected end-to-end, forming a closed loop that resembles the Greek letter Delta (Δ). The connections between the windings serve as the line terminals. There is no central common point where all three phases meet. Therefore, the

Delta connection does not have a natural neutral point. A neutral point can be derived in some cases (like a center-tapped winding), but it's not an inherent part of the standard Delta configuration.

Options Sigma and Omega are not standard terms for fundamental three-phase electrical connection types in this context.

Based on the structure of these connections, only the Star configuration inherently forms a common point where the three phases meet, which is defined as the neutral point. The Delta connection does not have this central neutral point.

Comparing Star vs Delta Connections

Here is a comparison of key characteristics between Star and Delta connections:

Feature	Star Connection	Delta Connection
Phases Connection	One end of each phase connected at a common point (neutral)	Phases connected end-to-end, forming a loop
Neutral Point	Has a common neutral point	Does not have a common neutral point
Line Voltage (V_L) & Phase Voltage (V_p)	$V_L = \sqrt{3} \times V_p$	$V_L = V_p$
Line Current (I_L) & Phase Current (I_p)	$I_L = I_p$	$I_L = \sqrt{3} \times I_p$
Application	Suitable for power transmission, distribution (with neutral for single-phase loads), systems needing a neutral reference.	Suitable for high-power loads, motors, where no neutral is needed.

The common neutral point in a Star connection is crucial for providing **single-phase voltage** (voltage between a line and neutral) and for handling **unbalanced loads**, where the currents in the three phases are not equal. The neutral wire carries the imbalance current back to the source.

In contrast, the Delta connection is typically used for applications where only three-phase voltage is required and the loads are expected to be balanced. Without a neutral wire, circulating currents can occur in a Delta connection under certain conditions, but it does not provide a convenient point for single-phase connections.

Conclusion on Common Neutral Point

Reviewing the properties of the standard three-phase connections, the presence of a common neutral point is a defining characteristic of the Star connection. The other options do not represent configurations that inherently feature a common neutral point.

Revision Table: Key Electrical Connection Types

Connection Type	Common Neutral Point	Notes
Star (Wye)	Yes	Phases meet at a central neutral point.
Delta	No	Phases connected in a loop; no central meeting point.
Sigma	Not a standard three-phase power connection term.	N/A
Omega	Not a standard three-phase power connection term.	N/A

Additional Information: Three-Phase Systems and Neutral Points

Three-phase power systems are widely used for electricity generation, transmission, and distribution because they are more efficient than single-phase systems for delivering power. A three-phase system consists of three AC voltages that are phase-shifted from each other by 120 degrees (120°).

The **neutral point** is a reference point in a three-phase system. In a perfectly balanced Star connection, the sum of the instantaneous currents in the three phase wires is zero, and no current flows through the neutral wire. However, real-world loads are often unbalanced (e.g., different single-phase loads connected to different phases). In such cases, the neutral wire in a Star connection provides a return path for the net current, preventing voltage imbalances on the phases.

Understanding the difference between Star and Delta connections, particularly regarding the neutral point, is fundamental in electrical engineering and power systems analysis.

108. Answer: a

Explanation:

Understanding Multimeter Measurement Capabilities

A multimeter is a versatile electronic measuring instrument that can measure multiple electrical properties. The name "multimeter" itself suggests its ability to measure multiple things. Common measurements include voltage, current, and resistance.

Typical Multimeter Functions

Let's look at the standard measurements a basic multimeter is designed for:

- **Voltage:** Measures the electrical potential difference between two points in a circuit. Measured in Volts (V).
- **Current:** Measures the flow of electrical charge through a circuit. Measured in Amperes (A).
- **Resistance:** Measures the opposition to the flow of electrical current in a circuit. Measured in Ohms (Ω).

These three fundamental quantities (Voltage, Current, Resistance) are directly related by Ohm's Law, which states that voltage across a resistor is directly proportional to the current flowing through it, given by the formula:

$$V = I \times R$$

where:

- V is Voltage
- I is Current
- R is Resistance

Analyzing the Measurement Options

The question asks which of the given options cannot be measured by a multimeter. We need to consider each option:

- **Current:** Yes, standard multimeters have an ammeter function to measure current.
- **Voltage:** Yes, standard multimeters have a voltmeter function to measure voltage.
- **Resistance:** Yes, standard multimeters have an ohmmeter function to measure resistance.
- **Frequency:** Frequency is the number of cycles of a repeating waveform per second, measured in Hertz (Hz). While some advanced or specialized multimeters might have a frequency counter function, it is generally not a standard feature on all basic or typical multimeters. Multimeters are primarily designed for DC and AC voltage, current, and resistance measurements. Measuring frequency often requires a dedicated frequency counter or an oscilloscope.

Conclusion on Multimeter Measurement

Based on the typical capabilities of a standard multimeter, measuring frequency is not a common function like measuring current, voltage, and resistance. Therefore, frequency is the quantity among the options that usually cannot be measured by a standard multimeter.

Quantity	Standard Multimeter Measurement
Current	Yes
Voltage	Yes
Resistance	Yes
Frequency	No (generally not standard)

Revision Table: Multimeter Basics

Feature	Description
Purpose	Measure electrical properties
Common Measurements	Voltage, Current, Resistance
Units	Volts (V), Amperes (A), Ohms (Ω)
Types	Analog, Digital

Additional Information on Electrical Measurement Tools

Beyond the basic multimeter functions, there are other instruments used for specific electrical measurements:

- **Oscilloscope:** Used to visualize waveforms and measure frequency, amplitude, period, etc.

- **Frequency Counter:** A specialized instrument for highly accurate frequency measurement.
- **LCR Meter:** Measures inductance (L), capacitance (C), and resistance (R). Some multimeters might include capacitance measurement but inductance is less common.
- **Clamp Meter:** Primarily measures current without needing to break the circuit, often measures AC current.

While some advanced multimeters might integrate functions like frequency counting or capacitance testing, the core functions universally found are voltage, current, and resistance measurement.

109. Answer: a

Explanation:

Understanding Why Steel Poles Need Painting

Steel poles are commonly used in various applications, such as utility poles for power lines, street light poles, and structural supports. Like other metal structures, steel is susceptible to environmental damage over time. Painting the steel poles is a standard practice for protective purposes.

Preventing Corrosion in Steel Poles

The primary reason for painting steel poles is to protect them from **corrosion**. Steel is an alloy primarily made of iron, which reacts with oxygen and moisture in the atmosphere to form rust. This process is called **corrosion**.

- Rust is a form of iron oxide that is brittle and flaky.
- As rust forms, it expands and can cause the steel to deteriorate and lose its structural strength.
- Painting creates a barrier between the steel surface and the environment (air and moisture).

- This barrier prevents the chemical reaction that leads to rust formation, thereby protecting the steel pole from **corrosion**.
- Regular painting and maintenance extend the lifespan of steel poles.

Why Other Options Are Incorrect

Let's consider why the other options are not the reason for painting steel poles:

- **Borer:** Borers are types of insects or larvae that bore into materials, typically wood. Steel is a metal and is not affected by wood borers.
- **Grounding:** Grounding is the process of connecting an electrical circuit or device to the earth. It is a safety measure to prevent electric shock. Painting a steel pole, especially with insulating paint, could potentially hinder electrical conductivity, but it is not done for the purpose of facilitating or improving grounding. Grounding connections are typically made through direct metal contact or specific grounding wires.
- **Termites:** Termites are insects that feed on cellulose, primarily found in wood. Like borers, termites do not attack steel.

Based on the nature of steel and the effects of environmental exposure, preventing **corrosion** is the main functional reason for painting steel poles.

Revision Table: Protection for Different Materials

Material	Common Vulnerability	Protection Method (Examples)
Steel	Corrosion (Rust)	Painting, Galvanizing, Protective Coatings
Wood	Termites, Borers, Rot, Weathering	Painting, Varnishing, Pressure Treatment
Concrete	Cracking, Spalling, Chemical Attack	Sealers, Coatings, Proper Mix Design

Additional Information on Steel Pole Protection

While painting is a common method, other techniques are also used to protect steel poles from **corrosion**:

- **Galvanizing:** This process involves coating the steel with a layer of zinc. Zinc is more reactive than iron and corrodes first (sacrificial protection), protecting the underlying steel.
- **Weathering Steel:** Also known as COR-TEN steel, this type of steel forms a stable, rust-like appearance (patina) when exposed to weather. This patina acts as a protective layer, preventing further **corrosion**.
- **Powder Coating:** A type of coating applied as a free-flowing, dry powder. It provides a hard finish that is tougher than conventional paint and offers excellent **corrosion** resistance.

These methods, including painting, are crucial for ensuring the longevity and structural integrity of steel poles by combating the effects of **corrosion**.

110. Answer: a

Explanation:

Understanding Heat Sink Function in Semiconductors

Semiconductor devices, such as transistors and integrated circuits, generate heat when they operate. This heat is a byproduct of electrical current flowing through the device's components. Excessive heat can significantly impact the performance, reliability, and lifespan of a semiconductor device.

Why Temperature Control is Crucial for Semiconductors

High temperatures can lead to several problems in semiconductors:

- Reduced efficiency and altered electrical characteristics.
- Increased risk of component failure or burnout.
- Degradation of materials over time, shortening the device's lifespan.
- Changes in device parameters, potentially causing malfunctions.

Therefore, managing the temperature is essential for reliable operation.

The Primary Function of a Heat Sink

A heat sink is a passive component designed specifically to dissipate this excess heat away from the semiconductor device and transfer it to the surrounding environment (usually air, sometimes liquid). Its primary function is to control the temperature rise of the semiconductor by effectively moving thermal energy away from it.

Heat sinks achieve this by increasing the surface area exposed to the cooling medium. They are typically made of materials with high thermal conductivity, such as aluminum alloys or copper, and feature fins or other structures that maximize the area available for heat transfer via convection and radiation.

Analyzing the Options

Let's consider why the other options are not the primary function of a heat sink:

- **Temperature rise:** This is precisely what a heat sink aims to control and minimize. By dissipating heat, it limits how much the temperature of the semiconductor rises during operation.
- **Doping:** Doping is a process used during semiconductor manufacturing to introduce impurities into the crystal lattice of a semiconductor material. This alters its electrical properties, such as conductivity. A heat sink has no role in the doping process.
- **Voltage:** Voltage is an electrical potential difference. While temperature can indirectly affect a semiconductor's voltage characteristics, the heat sink's function is not to directly control the voltage supplied to or across the device.
- **Frequency:** Frequency refers to the rate of oscillation, often relevant in AC circuits or digital signals. The operating frequency of a semiconductor device is determined by its design and the input signals, not by the heat sink.

Based on how heat sinks work and their purpose in electronics, their main function is clearly related to managing temperature.

Conclusion on Heat Sink Function

The primary role of a heat sink in the context of semiconductors is to facilitate the removal of heat generated during operation, thereby controlling the temperature rise of the device. This temperature management is critical for maintaining the performance, reliability, and longevity of semiconductor components.

Revision Table: Heat Sink Function

Component/Process	Primary Function	Heat Sink Role
Semiconductor Device	Perform electronic functions (switching, amplification, etc.)	Generates heat that needs management
Heat Sink	Dissipate heat away from the device	Controls temperature rise
Doping	Alter semiconductor electrical properties	No role
Voltage	Electrical potential difference driving current	Not directly controlled by heat sink
Frequency	Rate of signal oscillation	Not directly controlled by heat sink

Your Personal Exams Guide

Additional Information on Semiconductor Cooling

Effective thermal management is a key aspect of designing electronic systems. Besides heat sinks, other cooling methods for semiconductors include:

- **Fans:** Often used in conjunction with heat sinks to increase airflow and enhance convective heat transfer.
- **Thermal Paste/Grease:** Applied between the semiconductor package and the heat sink base to fill microscopic air gaps, improving thermal contact and heat transfer.
- **Thermal Pads:** Similar to thermal paste but in solid form, used for thermal interface management.

- **Liquid Cooling:** Using a liquid (like water or a specialized coolant) to transfer heat away from the component, often used in high-performance systems.
- **Thermoelectric Coolers (Peltier Devices):** Active electronic components that can transfer heat from one side to another when a voltage is applied, used for localized cooling.

The choice of cooling method depends on the amount of heat generated, the ambient temperature, the required operating temperature of the device, and cost considerations. However, the fundamental principle remains the same: managing heat to control temperature rise and ensure device reliability.

III. Answer: c

Explanation:

Understanding Cold Cathode Discharge Lamps

The question asks about the type of lamp that utilizes the cold cathode discharge method. Understanding the difference between cold cathode and hot cathode discharge is key to answering this question.

In a gas discharge lamp, light is produced when electric current passes through a gas, causing it to emit photons. The method by which electrons are supplied to initiate and sustain this discharge classifies the lamp's cathode type:

- **Hot Cathode:** Uses a filament that is heated (typically by a small current) to high temperatures. This heating causes thermionic emission, releasing a large number of electrons that help start and maintain the discharge at a relatively low voltage. Most standard fluorescent lamps, including CFLs, and many mercury vapor lamps use hot cathodes.
- **Cold Cathode:** Does not use a heated filament. Instead, a high voltage is applied across the electrodes, causing electrons to be pulled from the cathode material through a process called field emission or secondary emission due to ion bombardment. Cold cathode lamps require a higher striking voltage to

initiate the discharge compared to hot cathode lamps, but they can often operate at lower currents once started. They also tend to have a longer lifespan as there is no filament to burn out, though the electrodes themselves can still wear over time due to sputtering.

Analyzing the Options

Let's examine the given options in the context of cathode types:

- **Mercury lamp:** Standard high-pressure mercury vapor lamps typically use hot cathodes. Low-pressure mercury vapor lamps (like those used in fluorescent tubes) also use hot cathodes.
- **CFL (Compact Fluorescent Lamp):** CFLs are a type of fluorescent lamp. They use a small, coiled tube containing mercury vapor and an inert gas. CFLs universally use hot cathodes. The ballast heats the filaments to start the discharge.
- **Neon lamp:** Traditional neon lamps (like those used in indicator lights or decorative signs) consist of a glass tube filled with neon gas at low pressure, with an electrode at each end. These lamps start the discharge by applying a high voltage across the electrodes, causing ionization of the gas without the need for heated filaments. This is the principle of cold cathode discharge.

Based on this analysis, the neon lamp is the one that primarily uses the cold cathode discharge method among the given options.

Conclusion

The cold cathode discharge method is a characteristic feature of neon lamps, allowing them to operate without pre-heating filaments. Other lamps like mercury lamps and CFLs typically rely on hot cathode discharge.

Therefore, the correct answer is the Neon lamp.

Lamp Type	Typical Cathode Type	Method of Discharge Initiation
Mercury lamp	Hot Cathode	Filament heating + High voltage
CFL (Compact Fluorescent Lamp)	Hot Cathode	Filament heating + Ballast voltage
Neon lamp	Cold Cathode	High striking voltage (Field/Secondary emission)

Revision Table: Lamp Cathode Types

Cathode Type	Description	Starting Voltage	Typical Lamps
Hot Cathode	Heated filament emits electrons (thermionic emission)	Lower	Standard Fluorescent, CFL, many Mercury lamps
Cold Cathode	Electrons pulled by high voltage (field/secondary emission)	Higher (Striking Voltage)	Neon lamps, some specialized fluorescent tubes

Additional Information on Gas Discharge Lamps

Gas discharge lamps are a broad category of artificial light sources that generate light by sending an electrical discharge through an ionized gas, plasma, or vapor. The specific gas or vapor used determines the color of the light emitted.

- **Neon Lamps:** Commonly use neon gas for red light, but other noble gases or mixtures are used for different colors (e.g., argon for blue, helium for pink). They are known for their ability to be formed into various shapes.
- **Mercury Lamps:** High-pressure mercury lamps are used for street lighting and industrial applications, emitting a bluish-white light. Low-pressure mercury lamps are the basis for fluorescent lamps, where the UV light produced excites a phosphor coating to emit visible light.

- **CFLs:** Compact Fluorescent Lamps are essentially miniaturized fluorescent tubes designed to fit into standard incandescent lamp sockets. They contain mercury vapor and an inert gas, using a ballast to control the current.
- **Applications:** Cold cathode lamps are often used where frequent switching is required, as they are not affected by filament wear. They are also used in backlighting for displays and in decorative lighting. Hot cathode lamps are generally more energy-efficient for continuous illumination.

112. Answer: b

Explanation:

Understanding 3-Phase Power Measurement

Measuring power in a 3-phase system is crucial for monitoring and managing electrical loads. The number of wattmeters required depends on the system configuration, specifically whether it is a 3-wire or 4-wire system, and whether the load is balanced or unbalanced.

Minimum Wattmeters for 3-Phase, 3-Wire Systems

For a 3-phase, 3-wire system, the minimum number of wattmeters required to measure the total power, irrespective of whether the load is balanced or unbalanced, is determined by a fundamental principle related to the number of conductors.

In any system with 'n' conductors, the total power can theoretically be measured using 'n-1' wattmeters. A 3-phase, 3-wire system has 3 conductors (the three phases, R, Y, and B, or A, B, and C). Applying the 'n-1' rule, the minimum number of wattmeters needed is $3 - 1 = 2$.

The 2-Wattmeter Method

The 2-wattmeter method is a standard technique for measuring total power in a 3-phase, 3-wire system. It involves connecting two wattmeters. Each wattmeter has a current coil connected in series with one phase wire and a voltage coil connected between the same phase wire and a third phase wire. For example, in a system with phases A, B, and C, one wattmeter's current coil might be in phase A and its voltage coil between A and B, while the second wattmeter's current coil is in phase C and its voltage coil between C and B.

The total power P_{total} in the system is the algebraic sum of the readings of the two wattmeters, W_1 and W_2 .

$$P_{total} = W_1 + W_2$$

This method is universally applicable for 3-phase, 3-wire systems, working correctly for both balanced and unbalanced loads. The readings of the individual wattmeters might differ significantly in an unbalanced system, but their sum accurately represents the total power consumed by the load.

Why Other Options Are Not the Minimum for 3-Wire

- **1 Wattmeter:** A single wattmeter can only measure total power in a 3-phase system under specific conditions, such as a balanced load connected to a known voltage. It is not sufficient for measuring power in a general 3-phase, 3-wire system that can be unbalanced.
- **3 Wattmeters:** Using three wattmeters (one in each phase with the voltage coil connected to the neutral point or an artificial neutral) is a valid method, especially for 4-wire systems where a neutral is present. It also works for 3-wire systems. However, it is not the *minimum* number required for a 3-wire system, as the 2-wattmeter method achieves the same result with one less instrument.
- **4 Wattmeters:** Four wattmeters are more than necessary for a 3-phase, 3-wire system.

Therefore, the minimum number of wattmeters needed for a 3-phase, 3-wire balanced or unbalanced system is 2.

Wattmeter Methods for 3-Phase Systems

Method (Wattmeters)	System Type (Wires)	Load Type	Notes
1	3 or 4	Balanced only	Requires additional information (e.g., voltage, power factor) or specific connection. Not general.
2	3	Balanced or Unbalanced	Standard method for 3-wire systems. Total Power = $W_1 + W_2$ (algebraic sum).
3	3 or 4	Balanced or Unbalanced	Connects each phase to neutral (real or artificial). Works but not minimum for 3-wire. Total Power = $W_1 + W_2 + W_3$.

Revision Table: 3-Phase Power Measurement

Here's a quick summary of the minimum wattmeters for different configurations:

- For a 3-phase, 3-wire system (balanced or unbalanced): **2 wattmeters**
- For a 3-phase, 4-wire system (balanced or unbalanced): **3 wattmeters** (as there are 4 conductors, $4-1=3$)

Additional Information: Understanding the 2-Wattmeter Method

The 2-wattmeter method not only gives the total power but can also be used to determine the power factor of a balanced load. The relationship between the two wattmeter readings (W_1 and W_2) and the power factor ($\cos \phi$) is given by:

$$\tan \phi = \sqrt{3} \frac{W_1 - W_2}{W_1 + W_2}$$

From this, the power factor $\cos \phi$ can be calculated. However, this relationship for power factor is only valid for balanced loads. For unbalanced loads, the concept of a single power factor for the entire load is not well-defined, and the method primarily provides the total active power.

The phase sequence (the order in which the phase voltages peak) can affect which wattmeter reads higher or even negative, especially at low power factors. However, the algebraic sum always provides the correct total power.

113. Answer: d

Explanation:

Understanding Pliers for Holding Objects

The question asks to identify the type of pliers best suited for holding objects. Pliers are versatile hand tools used for gripping, bending, and cutting various materials, most commonly wire. Different types of pliers are designed for specific tasks based on the shape and function of their jaws.

Analyzing the Pliers Options

Let's look at the primary uses of the pliers mentioned in the options:

- **Side Cutting Pliers / Diagonal Cutters:** These pliers have short jaws with angled cutting edges. Their main purpose is to cut wires and small metal components. They are not designed for holding objects securely over a large surface area.
- **Wire Strippers:** These tools have notched jaws specifically sized to grip and strip the insulation from electrical wires without damaging the conductor. While they grip the wire temporarily during the stripping process, their primary function is not general object holding.
- **Nose Pliers:** This term typically refers to pliers with long, often tapered jaws, such as long-nose pliers or needle-nose pliers. The long, slender jaws allow them to reach into confined spaces and firmly grip small objects, wires, or

components for holding, bending, or positioning. The tips can be used for precision holding.

Why Nose Pliers are Suitable for Holding

Nose pliers, particularly those with long and tapering jaws (like long-nose or needle-nose pliers), are specifically designed for tasks that involve manipulating or holding small items. Their pointed tips and sometimes serrated jaws provide a secure grip on objects that are difficult to hold with fingers or bulkier pliers. They are essential tools for electrical work, jewelry making, crafting, and other tasks requiring precision handling and holding of small components.

Comparing the functions, side cutting pliers and wire strippers are specialized for cutting and stripping, respectively. While any pliers might temporarily hold something, **nose pliers** are explicitly designed with jaw shapes that excel at gripping and holding objects, especially small ones or in tight spots.

Conclusion

Based on the primary function and design of the different types of pliers, **nose pliers** are the best suitable for holding objects, particularly small items or when precision is required.

Pliers Type	Primary Function	Suitability for Holding Objects
Side Cutting Pliers / Diagonal Cutters	Cutting wires	Low (designed for cutting, not general gripping)
Wire Strippers	Stripping wire insulation	Low (specific grip for stripping, not general holding)
Nose Pliers (e.g., Long-nose, Needle-nose)	Holding, bending, reaching into tight spaces	High (designed with jaws suitable for gripping small objects)

Revision Table: Key Pliers Types

Understanding the specific uses of different pliers types is crucial.

- **Diagonal Cutters:** For cutting wire close to a surface.
- **Needle-Nose Pliers:** Excellent for holding, bending, and manipulating small wires or components in tight spaces.
- **Combination Pliers:** Often include sections for gripping, cutting, and sometimes bending.
- **Lineman's Pliers:** Heavy-duty pliers for gripping, bending, and cutting thick wire.

Additional Information: Pliers Applications

Pliers are fundamental tools found in various fields. Their specific design dictates their best application:

- Electrical work often uses wire strippers, diagonal cutters, and long-nose pliers.
- Mechanics use various types, including slip-joint and locking pliers, for gripping and turning.
- Jewelry making heavily relies on fine-tipped pliers like needle-nose and round-nose for shaping and holding small metal pieces.
- Plumbing may involve pipe wrenches or adjustable pliers for gripping pipes and fittings.

Your Personal Exams Guide

114. Answer: c

Explanation:

Understanding Heat Transfer in an Electric Iron

An electric iron works by heating a metal plate to smooth clothes. The heat generated by the electric coil needs to move efficiently to the base plate. Let's look at how heat travels and which method is most important here.

Modes of Heat Transfer

There are three main ways heat can move from one place to another:

- **Conduction:** This happens when objects are touching. Heat energy is passed along from particle to particle directly. Think of touching a hot pan handle – the heat travels through the metal to your hand. This is common in solid materials.
- **Convection:** This occurs when a liquid or gas is heated. The hotter, less dense part moves away, carrying the heat with it, and cooler parts move in to take its place, creating a current. This is how a radiator heats a room.
- **Radiation:** This is heat transfer through invisible waves, like light or infrared waves. The sun warming the Earth is a good example. It doesn't need anything to be touching or any air to move.

Heat Transfer in the Electric Iron Coil

Now, let's apply this to the electric iron:

- Inside the iron, an electric **coil** (often a resistance wire) heats up significantly when electric current flows through it.
- This hot coil is positioned directly against the flat metal **base plate** of the iron.
- Because the coil and the base plate are touching, heat energy is transferred directly from the hot coil particles to the cooler base plate particles.
- This direct transfer of heat through physical contact is the definition of **conduction**. It's the most effective way to get heat from the coil into the base plate quickly and efficiently.

Why Other Methods Are Less Important Here

- **Convection** isn't the main way heat moves from the coil *to the plate* because there's no significant movement of air or liquid between them. The plate does heat the air above it through convection, but that's a secondary effect, not the primary transfer from the coil.
- **Radiation** does occur – the hot coil emits some heat waves. However, the design focuses on direct contact for maximum heat transfer, making conduction the dominant method.
- **Capacitance** relates to storing electrical energy, not heat transfer between components. It's not relevant to how heat moves in this situation.

Therefore, the primary method for transferring heat from the heating coil to the base plate in an electric iron is **conduction**, due to the direct physical contact between these parts.

115. Answer: d

Explanation:

Understanding Three-Phase Alternator Winding Placement

A three-phase alternator is an electrical machine that produces three separate alternating current (AC) outputs. These outputs are identical in magnitude and frequency but are out of phase with each other. This is achieved by arranging the windings on the stator (the stationary part) in a specific physical configuration. The question asks how far apart these **three windings** are placed in a **three-phase alternator**.

The Concept of Phase Separation

In any AC system, a full cycle of the alternating voltage or current represents 360 electrical degrees. For a system to be truly "three-phase" and balanced, the three individual phases must be equally spaced throughout this 360-degree cycle. This equal spacing ensures that the power delivered by the system is smooth and constant.

Calculating the Angular Separation

To find the separation between the three phases (and therefore the physical placement of their corresponding windings), we divide the total degrees in a cycle by the number of phases:

- Total electrical degrees in a cycle: 360°
- Number of phases: 3

- Calculation: The angular separation is given by $\frac{360^\circ}{3}$

Performing the calculation:

$$\frac{360^\circ}{3} = 120^\circ$$

This means that the three sets of windings in a **three-phase alternator** are physically positioned such that the voltage generated in each winding is **120 degrees** out of phase with the voltage generated in the other two windings.

Correct Winding Placement

Based on the calculation, the three windings are spaced **120 degrees** apart. This arrangement is fundamental to the operation of **three-phase alternators** and ensures a stable and efficient power output. Therefore, the correct option is the one that states 120 degrees.

116. Answer: b

Explanation:

Understanding Semiconductor Devices: Current vs. Voltage Control

Electronic circuits use various semiconductor devices to control the flow of current. These devices can broadly be classified based on how their main current (usually between two terminals) is controlled – either by a current applied to a third terminal or by a voltage applied to a third terminal.

Let's understand the difference between current-controlled and voltage-controlled devices:

- **Current-Controlled Device:** In this type of device, the amount of current flowing between two main terminals (like collector and emitter, or anode and cathode) is determined by the amount of current injected into a third control

terminal (like the base). A change in the control current directly results in a proportional change in the main current.

- **Voltage-Controlled Device:** In this type of device, the amount of current flowing between two main terminals (like drain and source, or collector and emitter) is determined by the voltage applied across two terminals (like the gate and source, or gate and emitter). A change in the control voltage results in a change in the main current by altering the conductivity of the channel or region through which the current flows.

Analyzing the Device Options

The question asks to identify which among the given options is a current-controlled device. Let's look at each option:

BJT (Bipolar Junction Transistor) Operation

A BJT has three terminals: Base, Collector, and Emitter. The main current flows between the Collector and Emitter (I_C). The control terminal is the Base. The Collector current (I_C) is primarily controlled by the Base current (I_B). The relationship is approximately $I_C = \beta \times I_B$, where β (beta) is the current gain of the transistor. This direct dependence of the Collector current on the Base current makes the BJT a current-controlled device.

MOSFET (Metal-Oxide-Semiconductor Field-Effect Transistor) Operation

A MOSFET has three main terminals: Gate, Drain, and Source. The main current flows between the Drain and Source (I_D). The control terminal is the Gate. The Drain current (I_D) is controlled by the voltage applied between the Gate and Source (V_{GS}). The Gate terminal is insulated from the channel, meaning very little current flows into the Gate. The voltage V_{GS} creates an electric field that modulates the conductivity of the channel between the Drain and Source, thereby controlling I_D . This voltage control makes the MOSFET a voltage-controlled device.

JFET (Junction Field-Effect Transistor) Operation

A JFET also has three terminals: Gate, Drain, and Source. Similar to the MOSFET, the main current (I_D) flows between the Drain and Source, and the control terminal is the Gate. The I_D is controlled by the voltage applied between the Gate and Source (V_{GS}). The voltage V_{GS} varies the width of a depletion region, which in turn modulates the channel conductivity and controls I_D . While there is a small leakage current through the reverse-biased Gate junction, the control mechanism is fundamentally voltage-driven. Thus, the JFET is a voltage-controlled device.

IGBT (Insulated-Gate Bipolar Transistor) Operation

An IGBT is a power semiconductor device with three terminals: Gate, Collector, and Emitter. It combines features of both MOSFET and BJT. The control terminal is the Gate, which is insulated like that of a MOSFET. The main current flows between the Collector and Emitter. The IGBT's operation is controlled by the voltage applied to the Gate (V_{GE}). The Gate voltage controls a MOSFET structure within the IGBT, which in turn controls the injection of minority carriers into the BJT output section, controlling the Collector current. Because its control is based on a Gate voltage, the IGBT is considered a voltage-controlled device, despite having a BJT-like output characteristic.

Conclusion: Identifying the Current Controlled Device

Based on the operational principles discussed:

- BJT: Collector current controlled by Base current (Current-Controlled)
- MOSFET: Drain current controlled by Gate-Source voltage (Voltage-Controlled)
- JFET: Drain current controlled by Gate-Source voltage (Voltage-Controlled)
- IGBT: Collector current controlled by Gate-Emitter voltage (Voltage-Controlled)

Therefore, the BJT is the device that is controlled by current.

Revision Table: Semiconductor Device Control

Device Type	Control Terminal	Main Current Path	Control Mechanism	Classification
BJT	Base	Collector-Emitter	Base Current (I_B) controls Collector Current (I_C)	Current-Controlled
MOSFET	Gate	Drain-Source	Gate-Source Voltage (V_{GS}) controls Drain Current (I_D)	Voltage-Controlled
JFET	Gate	Drain-Source	Gate-Source Voltage (V_{GS}) controls Drain Current (I_D)	Voltage-Controlled
IGBT	Gate	Collector-Emitter	Gate-Emitter Voltage (V_{GE}) controls Collector Current (I_C)	Voltage-Controlled

Additional Information on Semiconductor Devices

Understanding the control mechanism is fundamental to analyzing and designing circuits using these devices. Here are a few more points:

- **Input Impedance:** Current-controlled devices like BJTs generally have lower input impedance at the control terminal (Base) compared to voltage-controlled devices like MOSFETs and JFETs, which have very high input impedance at the control terminal (Gate) due to the insulating layer or reverse-biased junction.
- **Switching Speed:** MOSFETs and JFETs are often preferred in high-speed switching applications due to their voltage control and absence of minority carrier storage effects that are present in BJTs. IGBTs are used in high-power switching applications, offering a compromise between the power handling of BJTs and the ease of control of MOSFETs.
- **Transconductance:** Voltage-controlled devices are often characterized by transconductance (g_m), which is the ratio of the change in output current to the change in input voltage ($\Delta I_{out} / \Delta V_{in}$). For BJTs, current gain (β) is a primary parameter relating output current to input current.

117. Answer: d

Explanation:

Understanding the Function of a Buck Converter

A buck converter is a type of DC-to-DC converter. Its main purpose is to change a DC voltage from one level to another. Specifically, a buck converter is designed to reduce, or "step down," a higher input DC voltage to a lower output DC voltage.

Let's look at the options provided regarding the function of a buck converter:

- **Option 1: Exactly double the voltage**
This is incorrect. Doubling the voltage would be stepping up the voltage. A buck converter steps down the voltage, and the step-down ratio is typically adjustable, not fixed at doubling (or halving).
- **Option 2: Stabilize the voltage**
While DC-DC converters like buck converters are often used within a larger system that includes feedback control to stabilize the output voltage against variations in input voltage or load current, the primary **function** of the buck converter itself is the voltage conversion (stepping down). Stabilization is achieved through regulation circuitry built around the converter, not by the basic converter topology alone.
- **Option 3: Step up the voltage**
This describes the function of a different type of DC-DC converter called a boost converter. A buck converter performs the opposite function of stepping up the voltage.
- **Option 4: Step down the voltage**
This is the correct function of a buck converter. It takes a higher DC input voltage and converts it into a lower DC output voltage.

The operation of a basic buck converter involves a switch (like a MOSFET or transistor), an inductor, a diode, and an output capacitor. By rapidly switching the input voltage on and off to the inductor, the converter controls the average voltage supplied to the load, effectively stepping down the voltage.

Consider the ideal voltage conversion ratio for a buck converter operating in continuous conduction mode (CCM). The output voltage (V_{out}) is related to the input voltage (V_{in}) and the duty cycle (D) of the switch by the formula:

$$V_{out} = D \times V_{in}$$

Where D is between 0 and 1. Since D is always less than or equal to 1, the output voltage V_{out} is always less than or equal to the input voltage V_{in} , demonstrating the step-down function.

Therefore, the primary function of a buck converter is to step down the voltage.

Analysis of Buck Converter Functions

Let's break down why the correct option aligns perfectly with the definition and operation of a buck converter.

Option	Description	Correctness for Buck Converter
Exactly double the voltage	Increasing voltage by a factor of 2.	Incorrect. A buck converter reduces voltage.
Stabilize the voltage	Keeping the output voltage constant despite input or load changes.	Not the primary function of the converter topology itself; achieved with regulation circuitry.
Step up the voltage	Increasing voltage.	Incorrect. This is the function of a boost converter.
Step down the voltage	Decreasing voltage.	Correct. This is the fundamental operation of a buck converter.

Revision Table: DC-DC Converters

Converter Type	Primary Function	Relationship V_{out} to V_{in} (Ideal)
Buck Converter	Step down voltage	$V_{out} = D \times V_{in}$ (where $D < 1$)
Boost Converter	Step up voltage	$V_{out} = V_{in}/(1 - D)$ (where $D < 1$)
Buck-Boost Converter	Step up or step down voltage, with output voltage polarity inversion.	$V_{out} = -D \times V_{in}/(1 - D)$

Additional Information: Applications of Buck Converters

Buck converters are widely used in various electronic devices and systems where a lower DC voltage is needed from a higher source. Some common applications include:

- Power supplies for computers and laptops (e.g., reducing battery voltage to the voltage required by different components).
- Mobile phone chargers.
- Automotive electronics.
- LED lighting systems.
- Portable electronic devices.
- Point-of-load converters in distributed power systems.

Their efficiency, especially compared to linear regulators, makes them a popular choice for stepping down voltage while minimizing power loss.

118. Answer: d

Explanation:

Understanding Motor-Generator Sets and AC Motor Types

A motor-generator set, often abbreviated as M-G set, is a combined unit consisting of an electric motor and an electric generator. The motor drives the generator, which produces electrical power with characteristics (voltage, frequency, phase) that may be different from the power supplied to the motor. These sets are used for various purposes, such as converting AC to DC, DC to AC, changing voltage or frequency, or providing isolation between two power systems.

Selecting the Right AC Motor for M-G Sets

The choice of the AC motor to drive the generator in an M-G set depends on the specific application requirements. However, a common requirement for many M-G sets, especially those producing AC output, is a stable output frequency. The frequency of the AC generator's output is directly related to its rotational speed.

Synchronous Motor for Constant Speed Operation

A **synchronous motor** is an AC motor that runs at a constant speed, known as the synchronous speed, which is determined by the frequency of the supply voltage and the number of poles in the motor. The synchronous speed (N_s) is given by the formula:

$$N_s = \frac{120f}{P}$$

where f is the supply frequency and P is the number of poles.

Because the synchronous motor runs at a constant speed regardless of load variations (up to its pull-out torque), it is an excellent choice for driving a generator when a stable output frequency is required. Additionally, synchronous motors can operate at a leading power factor by adjusting their field excitation, which can help improve the overall power factor of the power system.

Why Other AC Motors Are Less Common for Stable Frequency Output

- **Squirrel cage induction motor:** This is the most common type of AC motor, simple and robust. However, its speed varies with the load due to slip. As the load increases, the speed decreases slightly. This speed variation would lead to variations in the output frequency of the connected generator, making it less suitable for applications requiring a precise, stable frequency.
- **Wound rotor induction motor:** Similar to the squirrel cage type, the wound rotor motor's speed also varies with the load due to slip. While it offers better starting torque and some speed control capabilities compared to the squirrel cage type, it still doesn't provide the inherent constant speed characteristic of a synchronous motor.
- **AC commutator motor:** These motors (like series or repulsion motors) are often used for variable speed applications. They are generally more complex and less efficient than induction or synchronous motors for constant-speed, high-power applications typical of many industrial M-G sets.

Based on the need for a stable and constant speed to ensure a consistent output frequency from the generator, the **synchronous motor** is frequently the preferred type of AC motor used in motor-generator sets, particularly in applications demanding precise frequency control.

Motor Type	Speed Characteristic	Suitability for Stable Frequency M-G Set
Squirrel cage induction motor	Speed varies with load (has slip)	Less suitable
Wound rotor induction motor	Speed varies with load (has slip)	Less suitable
AC commutator motor	Often variable speed control	Less suitable for constant speed needs
Synchronous motor	Constant speed (at synchronous speed)	Highly suitable

Revision Table: Key Motor-Generator Set Concepts

Concept	Description	Relevance to M-G Sets
Motor-Generator Set	Combination of a motor driving a generator.	Converts electrical power characteristics (V, f, phase).
Synchronous Speed	Constant speed determined by frequency and poles ($\frac{120f}{P}$).	Key characteristic of synchronous motors ensuring stable output frequency.
Slip	Difference between synchronous speed and actual rotor speed in induction motors.	Causes speed variation with load in induction motors.
Power Factor Correction	Ability to operate at leading PF by over-excitation.	Benefit of synchronous motors in M-G sets.

Additional Information on Synchronous Motors

Synchronous motors are classified as synchronous AC motors because the rotor rotates at the same speed as the rotating magnetic field of the stator. They require a DC excitation source for the rotor winding (or use permanent magnets in smaller sizes). They are not self-starting and require a starting mechanism (like a damper winding or an external motor). Due to their constant speed and power factor control capabilities, they are widely used in industrial applications, including driving compressors, pumps, line shafts, and as mentioned, generators in M-G sets.

119. Answer: c

Explanation:

Understanding Ammeter Range Extension with a Shunt

An ammeter is an instrument used to measure electric current. A typical meter movement, like the 60 microampere meter described, is designed to measure relatively small currents. To measure larger currents, a parallel resistor, called a shunt resistor, is connected across the meter terminals. This shunt resistor diverts the majority of the current, allowing only a small, proportional amount to pass through the sensitive meter movement.

In this problem, we have a meter with a full-scale deflection current of 60 microamperes (I_m) and an internal resistance of 100 ohms (R_m). We want to use this meter to measure a total current (I) of 100 milliamperes.

Calculating Current Distribution in Ammeter Shunt Circuit

When the shunt resistor is connected in parallel with the meter, the total current I entering the parallel combination splits into two paths: one through the meter (I_m) and one through the shunt resistor (I_{sh}). The total current is the sum of these two currents:

$$I = I_m + I_{sh}$$

Our goal is to find the value of the current through the shunt (I_{sh}). We can rearrange the equation above to solve for I_{sh} :

$$I_{sh} = I - I_m$$

Step-by-Step Calculation

First, let's make sure all current values are in the same units.

- The total current to be measured is $I = 100 \text{ mA}$.
- The full-scale current of the meter is $I_m = 60 \mu\text{A}$.

We need to convert $60 \mu\text{A}$ to milliamperes (mA).

$$1 \text{ mA} = 1000 \mu\text{A}$$

$$\text{So, } 1 \mu\text{A} = \frac{1}{1000} \text{ mA} = 0.001 \text{ mA}$$

Therefore, $I_m = 60 \mu\text{A} = 60 \times 0.001 \text{ mA} = 0.06 \text{ mA}$.

Now we can calculate the current through the shunt (I_{sh}):

$$I_{sh} = I - I_m$$

$$I_{sh} = 100 \text{ mA} - 0.06 \text{ mA}$$

$$I_{sh} = 99.94 \text{ mA}$$

The calculated current through the shunt resistor is 99.94 mA. Looking at the options provided, 99.9 mA is the closest value.

Understanding Shunt Resistance (Optional but helpful)

Although the question only asks for the shunt current, it's useful to understand how the shunt resistance (R_{sh}) is determined. Since the meter and the shunt are in parallel, the voltage across them is the same.

$$V_m = V_{sh}$$

Using Ohm's Law ($V = I \times R$):

$$I_m R_m = I_{sh} R_{sh}$$

We know $I_m = 0.06 \text{ mA}$, $R_m = 100 \Omega$, and we just calculated $I_{sh} = 99.94 \text{ mA}$. We can calculate R_{sh} :

$$R_{sh} = \frac{I_m R_m}{I_{sh}}$$

Let's use amperes (A) for consistency in the calculation, or ensure units cancel correctly. Using mA and Ohms:

$$R_{sh} = \frac{0.06 \text{ mA} \times 100 \Omega}{99.94 \text{ mA}}$$

$$R_{sh} = \frac{6 \text{ mA} \Omega}{99.94 \text{ mA}}$$

$$R_{sh} \approx 0.060036 \Omega$$

This shows that a very small shunt resistance is needed to divert most of the current.

Summary of Calculation

To find the current through the shunt (I_{sh}), subtract the meter's full-scale current (I_m) from the total current to be measured (I).

- Total current (I) = 100 mA
- Meter current (I_m) = $60 \mu\text{A} = 0.06 \text{ mA}$
- Shunt current (I_{sh}) = $I - I_m = 100 \text{ mA} - 0.06 \text{ mA} = 99.94 \text{ mA}$

The closest option is 99.9 mA.

Measurement Parameters and Calculated Shunt Current

Parameter	Value	Unit
Meter Full-Scale Current (I_m)	60	μA
Meter Resistance (R_m)	100	Ohms (Ω)
Total Current to Measure (I)	100	mA
Meter Current (I_m) in mA	0.06	mA
Calculated Shunt Current (I_{sh})	99.94	mA

Revision Table: Ammeter Shunt Concepts

Concept	Description
Ammeter	Measures electric current in a circuit. Connected in series.
Meter Movement	The sensitive part of the ammeter, typically a galvanometer, with limited current capacity and internal resistance.
Shunt Resistor (R_{sh})	A low-value resistor connected in parallel with the meter movement to extend the range of the ammeter.
Current Division	When total current enters the parallel combination of meter and shunt, it divides, with most going through the shunt.
Voltage Across Parallel Elements	The voltage drop across the meter ($I_m R_m$) is equal to the voltage drop across the shunt ($I_{sh} R_{sh}$).
Range Extension	The ratio of the new full-scale current (I) to the original meter current (I_m) is the range extension factor.

Additional Information: Ammeter Design

Designing an ammeter for a specific range involves selecting the appropriate shunt resistor. The formula $R_{sh} = \frac{I_m R_m}{I - I_m}$ is used. A higher range requires a smaller shunt resistance. Conversely, to extend a voltmeter's range, a multiplier resistor is added in series. Understanding these techniques is fundamental in electrical measurements and circuit design. The accuracy of the ammeter depends on the precision of both the meter movement and the shunt resistor. Temperature changes can affect resistance values, which might require using special alloys or temperature compensation circuits for high-precision ammeters.

120. Answer: c

Explanation:

Understanding Motor to Generator Transition

A motor is an electrical machine that converts electrical energy into mechanical energy. However, the same machine can often operate in reverse, converting mechanical energy into electrical energy. When a motor operates in this mode, it functions as a generator.

The key difference between operating as a motor and operating as a generator lies in the relationship between the rotor speed and the speed of the magnetic field produced by the stator. For AC machines, this stator field speed is known as the synchronous speed.

Synchronous Speed Explained

The synchronous speed (n_s) of the rotating magnetic field in an AC motor's stator is determined by the frequency of the power supply (f) and the number of poles (P) in the stator winding. The formula is:

$$n_s = \frac{120f}{P}$$

where n_s is in revolutions per minute (RPM).

Motor Operation vs. Generator Operation

In motor operation, the rotor spins slower than the synchronous speed of the magnetic field. This speed difference is called slip, and it's necessary for the rotor to experience a torque that drives it. The slip is positive in motor mode.

For the machine to act as a generator, the rotor must be driven by an external mechanical force (like a turbine or engine). To produce electrical power, the rotor must 'cut' the magnetic field lines in a way that induces a voltage and causes current flow in the stator windings in the opposite direction of motor operation. This happens when the rotor spins faster than the synchronous speed.

Analyzing the Options for Generator Operation

- **Option 1: The Rotor should spun Very-fast in Anti-clockwise Direction**
While spinning fast is involved, the critical factor is spinning faster than the synchronous speed. The direction (clockwise or anti-clockwise) depends on the phase sequence of the stator field, but the principle of exceeding synchronous speed holds regardless of this specific direction. 'Very-fast' is not precise enough.
- **Option 2: The rotor must spun faster than its stator's asynchronous speed**
Asynchronous speed is simply the actual running speed of an induction motor rotor, which is always less than the synchronous speed in motor mode. The term "stator's asynchronous speed" doesn't typically apply as the stator field rotates at synchronous speed. This option is incorrect.
- **Option 3: The rotor must spun faster than its stator's synchronous speed**
This statement correctly identifies the condition. When the rotor speed exceeds the synchronous speed, the slip becomes negative. This causes the direction of the induced voltage and current in the rotor bars (relative to the magnetic field) to reverse compared to motor action, leading the machine to feed power back into the electrical system, thus working as a generator.
- **Option 4: The Rotor should spun Very-fast in Clockwise Direction**
Similar to Option 1, spinning fast is mentioned, but the key is the relationship to synchronous speed. The specific direction depends on setup but isn't the defining condition for *acting* as a generator, whereas exceeding synchronous speed is. 'Very-fast' is also not precise.

Therefore, the fundamental requirement for a motor to operate as a generator is that its rotor must be driven mechanically at a speed higher than the synchronous speed of the magnetic field in the stator.

Motor vs. Generator Operation (AC Induction Machine)

Parameter	Motor Operation	Generator Operation
Rotor Speed relative to Synchronous Speed (n_r vs n_s)	$n_r < n_s$	$n_r > n_s$
Slip ($s = \frac{n_s - n_r}{n_s}$)	Positive ($0 < s < 1$)	Negative ($s < 0$)
Energy Conversion	Electrical to Mechanical	Mechanical to Electrical
Power Flow	From electrical supply to shaft	From shaft to electrical supply

Conclusion on Motor Generator Mode

Based on the principles of operation for AC machines like induction motors, the condition that forces them to produce electrical power (act as a generator) is when their rotor is spun faster than the speed of the rotating magnetic field generated by the stator, which is the synchronous speed.

Revision Table: Key Concepts

Motor and Generator Key Concepts

Term	Description
Motor	Converts electrical energy into mechanical energy.
Generator	Converts mechanical energy into electrical energy.
Rotor	The rotating part of the machine.
Stator	The stationary part of the machine, contains windings that create the magnetic field.
Synchronous Speed (n_s)	Speed of the rotating magnetic field created by the stator windings. Dependent on supply frequency and number of poles.
Rotor Speed (n_r)	Actual speed of the rotor.
Slip (s)	Relative speed difference between synchronous speed and rotor speed. $s = \frac{n_s - n_r}{n_s}$. Positive for motor, negative for generator.

Additional Information: Types and Applications

While this question applies directly to how an induction motor can act as a generator (often called an induction generator), other types of motors/generators exist:

- **Synchronous Machines:** Can operate precisely at synchronous speed as motors or generators. Synchronous generators are commonly used in power plants.
- **DC Machines:** DC motors can also act as DC generators (dynamos) when mechanically driven.

The principle of needing a speed difference or specific excitation to generate power is fundamental to most electrical machines.

Induction generators require reactive power from the grid or a capacitor bank to operate. They are often used in wind turbines and some small hydro applications because of their robustness and simplicity.

121. Answer: a

Explanation:

Understanding Universal Motor Operation

A universal motor is a type of electric motor that can operate on either AC or DC power. It is typically constructed with a series winding, meaning the field winding is connected in series with the armature winding. This configuration allows it to function similarly on both AC and DC supplies, although its performance characteristics might differ slightly.

The question asks about the voltage required for a universal motor at a given output and speed when the supply frequency is changed compared to the standard 220 V, 50 Hz operation.

Factors Affecting Universal Motor Performance

For a universal motor, the relationship between applied voltage (V), speed (N), frequency (f), and load (which determines output) is complex. Key factors include:

- Resistance of the windings (R)
- Inductive reactance of the windings ($X_L = \omega L = 2\pi f L$)
- Back EMF generated by the armature (proportional to speed and flux)
- Armature current (determined by load/output)

Frequency's Impact on Universal Motor Voltage Requirement

When a universal motor operates on an AC supply, the impedance of the windings is not just the resistance (R), but also the inductive reactance (X_L) which depends on the frequency (f). The impedance (Z) can be represented as:

$$Z = \sqrt{R^2 + X_L^2} = \sqrt{R^2 + (2\pi fL)^2}$$

For a given speed and output power, the motor needs a certain armature current and back EMF. The applied voltage must overcome the impedance drop and balance the back EMF.

$$V \approx IZ + E_b$$

Where:

- V is the applied voltage
- I is the armature current (related to load/output)
- Z is the motor impedance
- E_b is the back EMF (related to speed)

Let's consider what happens when the frequency changes while keeping the output and speed constant. Keeping output and speed constant implies that the armature current I and the back EMF E_b remain approximately the same (assuming constant flux, which is also related to current). The primary change affecting the voltage requirement will be the impedance Z .

- At higher frequencies, $X_L = 2\pi fL$ increases. This leads to a higher impedance Z . To maintain the same current I and balance the same back EMF E_b , the applied voltage V must be higher.
- At lower frequencies, $X_L = 2\pi fL$ decreases. This leads to a lower impedance Z . To maintain the same current I and balance the same back EMF E_b , the applied voltage V must be lower.

Therefore, for a given output power and speed, a universal motor operating at a lower frequency will require less voltage compared to operating at a higher frequency (like 50 Hz).

Analyzing the Options

Let's evaluate the given options based on our understanding:

- Option 1: Less voltage at low frequency. This aligns with our conclusion that lower frequency leads to lower impedance and thus requires less voltage for the same output and speed.
- Option 2: Less voltage at high frequency. This contradicts our understanding that higher frequency increases impedance and requires more voltage.
- Option 3: High voltage at high frequency. This aligns with our understanding (higher frequency requires higher voltage), but the question asks what it requires compared to the standard frequency for a *given* output and speed, implying a comparison when frequency changes. The option states "high voltage at high frequency" without a direct comparison context, but if comparing to a lower frequency for the same output/speed, this would be true. However, let's look at the other options and the phrasing "compared to 220 V, 50 Hz supply".
- Option 4: High voltage at low frequency. This contradicts our understanding that lower frequency requires less voltage.

Comparing to the baseline of 220 V, 50 Hz, operating at a *lower* frequency for the same output and speed would require *less* voltage. Operating at a *higher* frequency for the same output and speed would require *more* voltage. Option 1 correctly states the requirement for a low frequency compared to the baseline frequency.

Conclusion

For a universal motor operating at a constant output power and speed, the inductive reactance decreases with decreasing frequency. This reduction in impedance means a lower applied voltage is needed to achieve the required current and overcome the back EMF compared to operation at a higher frequency like 50 Hz.

Frequency vs. Voltage Requirement (for constant speed & output)

Frequency	Inductive Reactance (X_L)	Impedance (z)	Voltage Required (v)
Low	Low	Low	Less
Standard (50 Hz)	Medium	Medium	Standard (220 v)
High	High	High	More

Revision Table: Key Concepts

Concept	Description	Relevance to Question
Universal Motor	Operates on AC or DC. Series winding.	The subject of the question.
Inductive Reactance (X_L)	Opposition to current flow in an inductor due to changing magnetic fields. $X_L = 2\pi fL$	Directly dependent on frequency (f), affects motor impedance.
Impedance (Z)	Total opposition to AC current flow, combining resistance and reactance. $Z = \sqrt{R^2 + X_L^2}$	Determines voltage drop for a given current.
Back EMF (E_b)	Voltage generated in the armature opposing the applied voltage. Proportional to speed and flux.	Needs to be balanced by the applied voltage; stays relatively constant for a given speed.

Additional Information: Universal Motor Applications and Characteristics

Universal motors are popular because they can be used with either AC or DC sources, offering versatility. They have high starting torque and can run at high speeds, which makes them suitable for applications like vacuum cleaners, power tools (drills, mixers, saws), blenders, and sewing machines.

However, they have some disadvantages:

- Sparking at the brushes due to commutation, which can cause electromagnetic interference (EMI).
- Relatively high noise level.
- Speed varies significantly with load, particularly on DC. Speed control is often necessary.

The speed of a universal motor on AC supply is generally lower than on DC supply of the same voltage due to the impedance caused by inductive reactance, which is absent in DC operation (where only resistance matters for impedance).

The question focuses on the AC operation and how changing frequency affects the voltage required for a specific operating point (given output and speed). The core principle relies on the frequency dependence of inductive reactance.

122. Answer: a

Explanation:

Understanding Voltage in Star Connected Circuits

In a three-phase electrical system, components can be connected in either a star (Y) or delta (Δ) configuration. The relationship between the voltage measured between any two lines (line voltage, V_L) and the voltage measured between a line and the neutral point (phase voltage, V_P) depends on the connection type.

For a star-connected circuit, there is a specific relationship between the line voltage and the phase voltage. The line voltage is $\sqrt{3}$ times the phase voltage. This

is because the line voltage is the vector sum of two phase voltages that are 120 degrees apart in phase.

The formula relating line voltage (V_L) and phase voltage (V_P) in a star connection is:

$$V_L = \sqrt{3} \times V_P$$

The question provides the value of the phase voltage in a star connected circuit:

- Phase Voltage (V_P) = 240 V

We need to find the line voltage (V_L). Using the formula for a star connection:

$$V_L = \sqrt{3} \times 240 \text{ V}$$

We know that the approximate value of $\sqrt{3}$ is 1.732.

$$V_L \approx 1.732 \times 240 \text{ V}$$

Let's calculate the value:

$$V_L \approx 415.68 \text{ V}$$

Now, let's compare this calculated value with the given options:

Option	Voltage Value
1	415 V
2	400 V
3	120 V
4	440 V

Our calculated value of approximately 415.68 V is very close to 415 V. Standard voltage values in many electrical systems are often rounded or nominal. For example, a 230V phase voltage often corresponds to a 400V line voltage ($\sqrt{3} \times$

$230 \approx 398.4$), and a 240V phase voltage often corresponds to a 415V line voltage ($\sqrt{3} \times 240 \approx 415.7$).

Therefore, based on the calculation and common standard voltage levels, the line voltage corresponding to a 240 V phase voltage in a star connected circuit is 415 V.

Revision Table: Star Connection Voltages

Parameter	Symbol	Relationship in Star Connection
Phase Voltage	V_P	Voltage between phase and neutral
Line Voltage	V_L	Voltage between any two lines
Formula		$V_L = \sqrt{3} \times V_P$

Additional Information: Star vs. Delta Connection

Understanding the difference between star and delta connections is crucial in three-phase systems:

- **Star (Y) Connection:**
 - Three phases are connected to a common neutral point.
 - Phase voltage is between phase and neutral.
 - Line voltage is between two lines.
 - Relationship: $V_L = \sqrt{3} \times V_P$ and $I_L = I_P$ (Line current equals phase current).
- **Delta (Δ) Connection:**
 - Phases are connected end-to-end, forming a closed loop.
 - No neutral point (unless derived externally).
 - Phase voltage is equal to line voltage.
 - Relationship: $V_L = V_P$ and $I_L = \sqrt{3} \times I_P$ (Line current is $\sqrt{3}$ times phase current).

This question specifically deals with the star connected circuit, where the line voltage is always higher than the phase voltage by a factor of $\sqrt{3}$.

123. Answer: b

Explanation:

Understanding L Series MCB Applications

Miniature Circuit Breakers (MCBs) are essential safety devices in electrical installations. They automatically switch off an electrical circuit during an abnormal condition, such as an overload or short circuit. Different types of MCBs are designed for specific applications based on their tripping characteristics, particularly how quickly they trip at various fault current levels. These characteristics are often described by different 'curves' or 'series', such as B, C, D, K, Z, and L.

What are L Series MCBs?

L series MCBs (often associated with the 'L' or sometimes 'B' trip curve, depending on specific standards and manufacturers, but in this context, we focus on the 'L' designation as provided) are designed with a lower tripping threshold for short circuits compared to C or D series MCBs. This means they are more sensitive to short-circuit currents. They are typically used in circuits where the prospective short circuit current is relatively low and for protecting equipment that cannot withstand high levels of surge current.

Analyzing MCB Trip Curves

MCB trip curves indicate the range of current values at which the MCB will trip magnetically (for short circuits). Common curves include:

- **Type B:** Trips between 3 to 5 times the rated current. Suitable for resistive loads and lighting circuits with low inrush current.
- **Type C:** Trips between 5 to 10 times the rated current. Most common for general applications with moderate inrush currents, like motors and fluorescent lighting.
- **Type D:** Trips between 10 to 20 times the rated current. Suitable for highly inductive loads with high inrush currents, like transformers, solenoids, and

welding equipment.

- **Type L / Z / K:** Have specific characteristics for particular applications, often with lower trip thresholds for short circuits or sensitivity to specific load types. L series, as discussed here, is often aligned with protecting cables in specific installations or circuits with low surge capability.

Considering the L series MCB's design for lower trip thresholds and sensitivity, let's evaluate the given options:

Evaluating Potential Applications for L Series MCBs

- **Capacitive circuit:** Circuits with large capacitance can cause high inrush currents when energized. L series MCBs, being sensitive to surges, might nuisance trip or may not be the most suitable protection without specific design considerations.
- **heating and lighting circuit:** Heating elements are primarily resistive loads, which have no significant inrush current. Lighting circuits (especially incandescent or modern LED lighting) also often have low inrush characteristics compared to inductive loads. Circuits with predominantly resistive and low-inrush lighting loads are well-suited for MCBs with lower trip thresholds like the L series, as they provide adequate protection against overloads and short circuits without tripping unnecessarily on switching surges.
- **Inductive circuit:** Circuits containing motors, transformers, or older fluorescent lighting ballasts are highly inductive and produce significant inrush currents upon switching. L series MCBs are typically too sensitive for such applications and would likely nuisance trip. Type C or D MCBs are generally used here.
- **Inductive and lighting circuit:** If the inductive load is significant, the high inrush current would make the L series MCB unsuitable for the entire circuit. While it might be suitable for the lighting portion alone (depending on the lighting type), protecting a combined inductive and lighting circuit with an L series MCB is generally not advisable due to the inductive component.

Conclusion on L Series MCB Usage

Based on the typical characteristics of L series MCBs, which are designed for lower trip thresholds and suitability for circuits with low or no significant inrush current, they are most effectively used in circuits feeding resistive loads and certain types of lighting.

MCB Type (Series/Curve)	Typical Magnetic Trip Range (x Rated Current)	Suitable Applications
B	3 to 5	Resistive loads, Incandescent lighting, Heating
C	5 to 10	General loads, Motors, Fluorescent lighting
D	10 to 20	High inrush loads (Transformers, Welding)
L	Varies (Lower than B/C/D in some standards, sensitive)	Heating, Lighting (low inrush types), Cable protection
K	8 to 12 (but sensitive to peak current)	Inductive loads with high peak but low duration surges
Z	2 to 3	Highly sensitive loads (Electronics), Circuits with long cables

Therefore, considering the standard applications for different MCB types, the L series MCB is well-suited for circuits involving heating elements and certain types of lighting where high inrush currents are not a primary concern.

Revision Table: L Series MCB Applications

Concept	Key Point	Relevance to L Series
MCB Function	Overload and short-circuit protection	L series provides this for specific load types.
MCB Trip Curves (B, C, D, L etc.)	Define trip characteristics based on fault current magnitude.	L curve indicates lower sensitivity to surges/inrush.
Resistive Loads (Heating)	No significant inrush current.	L series is effective due to low inrush characteristic.
Lighting Loads	Can vary; modern lighting often low inrush.	L series effective for low-inrush lighting.
Inductive Loads (Motors, Transformers)	High inrush current.	L series typically not suitable; causes nuisance tripping.

Additional Information: MCB Selection Criteria

Selecting the correct MCB type is crucial for safety and reliable operation. Factors to consider include:

- **Load Type:** Is it resistive, inductive, or capacitive? Does it have high inrush current?
- **Prospective Fault Current:** The maximum current that could flow during a fault at the installation point.
- **Cable Size and Type:** The MCB's trip characteristic must protect the cable from overheating under fault conditions.
- **Discrimination/Selectivity:** Ensuring that only the MCB closest to the fault trips, leaving other circuits operational.
- **Ambient Temperature:** Affects the thermal tripping characteristic.

Choosing an L series MCB indicates a specific requirement for protection in circuits without high surge currents, fitting well with resistive heating and low-inrush lighting applications.

124. Answer: d

Explanation:

Understanding Electric Current Measurement

Electric current is a fundamental concept in electricity. It represents the flow of electric charge through a conductor. To work with electrical circuits, it's often necessary to measure the amount of current flowing.

Instruments for Electrical Measurements

Different electrical quantities require specific instruments for accurate measurement. Let's look at the instruments mentioned in the options:

- **Megger:** A Megger, or Megohmmeter, is used to measure very high resistances, typically insulation resistance. It applies a high voltage to the material and measures the resulting small current to calculate the resistance.
- **Voltmeter:** A Voltmeter is used to measure the electric potential difference, or voltage, between two points in a circuit. Voltmeters are connected in parallel across the component or points where the voltage needs to be measured.
- **Thermometer:** A Thermometer is an instrument used to measure temperature. It is not used for electrical measurements.
- **Ammeter:** An Ammeter is specifically designed to measure the flow of electric current in a circuit. The unit of electric current is the ampere (A), and ammeters are calibrated to show readings in amperes or sub-multiples like milliamperes (mA) or microamperes (μA).

Measuring Current with an Ammeter

To measure the current flowing through a component or part of a circuit, an ammeter must be connected in series with that component. Connecting in series ensures that the total current flowing through the component also flows through the ammeter. Ideally, an ammeter should have very low internal resistance so that it does not significantly affect the current being measured.

The question asks about the instrument used to measure current. Based on our discussion, the ammeter is the correct instrument for this purpose.

Summary of Measuring Instruments

Instrument	Quantity Measured	Connection Method
Ammeter	Electric Current (I)	Series
Voltmeter	Voltage (V)	Parallel
Ohmmeter / Megger	Resistance (R)	Across component (often isolated)
Thermometer	Temperature	N/A (Non-electrical measurement)

Therefore, the instrument used to measure current is the ammeter.

Revision Table: Electrical Measurement Tools

Tool Name	Purpose	Typical Unit
Ammeter	Measures electric current	Ampere (A)
Voltmeter	Measures voltage (potential difference)	Volt (V)
Ohmmeter	Measures resistance	Ohm (Ω)
Megger	Measures high resistance (insulation)	Megaohm ($M\Omega$)

Additional Information: Understanding Electrical Concepts

To further understand current measurement, it's helpful to know the related concepts:

- **Electric Current (I):** The rate of flow of electric charge. Measured in Amperes (A).
- **Voltage (V):** The electric potential difference between two points, which drives the current. Measured in Volts (V).

- **Resistance (R):** The opposition to the flow of electric current. Measured in Ohms (Ω).
- **Ohm's Law:** Relates Voltage, Current, and Resistance: $V = IR$. This fundamental law helps analyze circuits.
- **Circuit Connection:** Understanding series and parallel connections is crucial for correctly using measuring instruments. Ammeters are placed in series to measure the total current flowing through that path, while voltmeters are placed in parallel to measure the voltage drop across a component.

Mastering these basic concepts is key to analyzing and understanding electrical circuits and their measurements.

125. Answer: a

Explanation:

Understanding DC Series Motor Suitability

A DC series motor is a type of DC motor where the field winding is connected in series with the armature winding. This unique connection gives it specific operating characteristics that make it highly suitable for certain applications and unsuitable for others.

Key Characteristics of a DC Series Motor

The most important characteristics to consider are:

- **High Starting Torque:** A DC series motor develops a very high torque at low speeds, particularly during startup. The torque is approximately proportional to the square of the armature current (before saturation). Since the field winding is in series, the field current is the same as the armature current ($I_f = I_a$). Torque $T \propto \phi I_a$. In the linear region, $\phi \propto I_a$, so $T \propto I_a^2$. At startup, I_a is high (limited only by armature resistance and series field resistance), resulting in very high starting torque.

- **Speed Varies Widely with Load:** The speed of a DC series motor is inversely proportional to the load. As the load increases, the armature current increases, which strengthens the magnetic field (since $I_f = I_a$). The speed of a DC motor is generally given by $N \propto \frac{V - I_a(R_a + R_{se})}{\phi}$. Since ϕ increases with I_a , the increase in current in the numerator's voltage drop term and the increase in flux in the denominator cause the speed to drop significantly as load increases.
- **Runs Dangerously at No Load:** At no load, the armature current is very low. This means the series field flux is very weak. With very low flux, the speed equation $N \propto \frac{V - I_a(R_a + R_{se})}{\phi}$ shows that the speed can become extremely high, potentially damaging the motor. Therefore, a DC series motor should never be operated without a mechanical load connected.

Analyzing Potential Applications

Let's examine the suitability of a DC series motor for each of the given options based on its characteristics:

- **Cranes:** Cranes are used for lifting heavy loads. Starting requires a very high torque to overcome inertia and gravity. The DC series motor's high starting torque is ideal for this. As the load is lifted, the speed can vary, which is generally acceptable or even beneficial for controlled lifting. Cranes always operate with a mechanical load, so the no-load speed problem is avoided.
- **Lathes:** Lathes are machine tools used for shaping materials, typically requiring a relatively constant speed for precise cutting. The speed of a DC series motor fluctuates significantly with the load (the cutting force), making it unsuitable for applications needing steady speed.
- **Punch Presses:** Punch presses require a large burst of torque at the moment of punching. While the high torque is needed, punch presses often use flywheels to store energy and release it during the punch. Constant speed operation between punches is often desired for cycling, which doesn't align well with the series motor's variable speed characteristic.
- **Pumps:** Centrifugal pumps usually require a relatively constant speed to deliver a consistent flow rate and head. While some positive displacement pumps might need higher starting torque, constant speed is often preferred for efficiency. DC series motors' speed variation makes them less suitable.

compared to motors with better speed regulation like DC shunt or induction motors for most pumping applications.

Based on this analysis, the application that best matches the characteristics of a DC series motor, particularly its high starting torque and tolerance for variable speed under load, is cranes.

DC Motor Type Suitability Comparison

Motor Type	Starting Torque	Speed Regulation	Suitable Applications
DC Series	Very High	Poor (Speed varies significantly with load)	Traction (trains, trams), Cranes, Hoists
DC Shunt	Medium	Good (Relatively constant speed)	Lathes, Centrifugal Pumps, Machine Tools
DC Compound (Cumulative)	High	Better than series, worse than shunt	Presses, Shears, Elevators

Revision Table: DC Series Motor Applications

Common Applications for DC Series Motors

Application	Reason for Suitability		
Electric Traction (Trains, Trams)	Requires high starting torque to move heavy loads from rest. Speed variation with terrain is acceptable. Load is always connected.		
Cranes and Hoists	Needs high torque to lift heavy loads initially. Variable speed control is often useful. Load is always connected.	Vacuum Cleaners (older designs)	Requires high speed and relatively high torque. Load is inherently present.

Additional Information: DC Series Motor Characteristics in Detail

Let's delve a bit deeper into the speed and torque characteristics using simplified equations.

The back EMF (E_b) of a DC motor is given by $E_b = V - I_a(R_a + R_{se})$, where V is the terminal voltage, I_a is the armature current, R_a is the armature resistance, and R_{se} is the series field resistance.

The speed (N) is proportional to the back EMF and inversely proportional to the flux (ϕ): $N \propto \frac{E_b}{\phi}$.

$$\text{So, } N \propto \frac{V - I_a(R_a + R_{se})}{\phi}$$

For a DC series motor, the flux ϕ is produced by the series field current, which is equal to the armature current I_a . So, $\phi \propto I_a$ (in the unsaturated region).

Substituting this into the speed equation, we get $N \propto \frac{V - I_a(R_a + R_{se})}{I_a} = \frac{V}{I_a} - (R_a + R_{se})$.

This equation clearly shows that as I_a (and thus load) increases, the term $\frac{V}{I_a}$ decreases, leading to a decrease in speed. Conversely, as I_a decreases (load decreases), $\frac{V}{I_a}$ increases, causing the speed to increase rapidly. At very low I_a (no load), the speed tends towards infinity theoretically, highlighting the need for a connected load.

The torque (T) is proportional to the product of flux and armature current: $T \propto \phi I_a$. Since $\phi \propto I_a$ (unsaturated region), $T \propto I_a^2$. This squared relationship explains the very high torque developed at high starting currents.

Understanding these fundamental relationships helps explain why the DC series motor is chosen for applications like cranes where high initial torque is critical and variable speed is acceptable.

126. Answer: c

Explanation:

Understanding Intersheath Grading in Power Cables

Power cables designed to transmit high voltages experience significant electric stress. This stress is not uniformly distributed across the insulation layer. It is highest at the surface of the inner conductor and decreases as you move outwards towards the sheath. This non-uniform stress distribution can lead to insulation breakdown, especially at higher voltages.

What is Cable Grading?

Cable grading is a technique used in high-voltage cables to improve the electric field distribution across the insulation. The goal is to make the electric stress more uniform throughout the dielectric material, thereby preventing localized overstressing that could cause insulation failure.

There are primarily two methods of cable grading:

- Capacitance grading
- Intersheath grading

Focus on Intersheath Grading

Intersheath grading is a method specifically employed to address the non-uniform electric stress distribution in cables. This technique involves using layers of insulation material with different dielectric constants or by introducing conducting layers (intersheaths) within the main insulation.

In the intersheath grading method, conducting layers are placed coaxially within the dielectric material. These intersheaths divide the main insulation into several layers. The intersheaths are connected to taps from the main voltage supply through a suitable impedance (like resistances or reactances) or maintained at specific intermediate potentials by other means.

How Intersheath Grading Manages Stress

By introducing these intersheaths at calculated radial positions, the total potential difference across the insulation is divided into steps. Each layer of insulation between the conductor and the first intersheath, between successive intersheaths, and between the last intersheath and the outer sheath, experiences a reduced voltage difference compared to the total voltage. This stepped voltage distribution helps in controlling the electric stress within each layer.

The potential of each intersheath is controlled such that the voltage gradient (electric stress) across each section of insulation is approximately the same or maintained within permissible limits. This leads to a more uniform electric stress distribution throughout the entire insulation thickness.

Purpose of Intersheath Grading

Based on this understanding, the primary purpose of intersheath grading in cables is to effectively control and redistribute the electric field within the insulation. This process ensures that the electric stress is distributed more evenly, particularly alleviating the high stress concentration near the conductor surface.

Therefore, intersheath grading is specifically used to:

- Improve the distribution of electric stress across the cable insulation.
- Reduce the maximum stress near the conductor.
- Allow the cable to withstand higher operating voltages without immediate breakdown.
- Effectively utilize the insulation material more uniformly.

Considering the options provided, the function of intersheath grading directly relates to managing the electric stress within the cable insulation.

Analysis of Options:

- **Minimize the stress:** While it reduces the maximum stress, the primary goal is 'distribution' and making it uniform rather than just minimizing overall stress which might imply using less stress.

- **Avoid the requirement of good insulation:** This is incorrect. Grading methods require good insulation; they just help the existing insulation perform better under stress.
- **Provided proper stress distribution:** This precisely describes the mechanism and purpose of intersheath grading – to make the electric stress distribution within the insulation uniform or more manageable.
- **None of the above:** This is incorrect as one of the options accurately describes the purpose.

Thus, the main application of intersheath grading is to achieve a proper or more uniform distribution of electric stress throughout the cable insulation.

Revision Table: Cable Grading Concepts

Concept	Description	Primary Goal
Electric Stress in Cables	Non-uniform distribution, highest near conductor	Identify stress zones
Cable Grading	Technique to improve stress distribution	Prevent insulation breakdown
Intersheath Grading	Uses conducting layers (intersheaths) within insulation	Achieve proper stress distribution by dividing voltage
Capacitance Grading	Uses layers of insulation with different dielectric constants	Achieve proper stress distribution by controlling capacitance

Additional Information: Stress Distribution Principles

The electric stress (E) at any point in a cable's insulation is given by the formula:

$$E = \frac{V}{x \ln \left(\frac{R}{r} \right)}$$

Where:

- V is the voltage between conductor and sheath.
- x is the radial distance from the cable center.
- r is the radius of the conductor.
- R is the inner radius of the sheath.

This formula shows that stress (E) is inversely proportional to the radial distance (x). This means the stress is maximum at the conductor surface ($x = r$) and minimum at the sheath surface ($x = R$). Grading techniques, like intersheath grading, aim to modify this natural stress profile to make it more uniform across the insulation thickness, allowing cables to operate safely at higher voltages.

127. Answer: b

Explanation:

SI Unit of Electrical Power Explained

This section focuses on identifying the correct International System of Units (SI) measure for **electrical power** by examining the provided options.

Defining Electrical Power

Electrical power refers to the rate at which electrical energy is transferred or consumed. It quantizes how fast electrical work is being done. A common analogy is comparing it to the flow rate of water – power is like the rate at which energy flows.

Understanding SI Units

The International System of Units (SI) is the globally accepted standard for measurement. It ensures consistency in scientific, industrial, and commercial applications. Identifying the correct SI unit is crucial for accurate calculations and understanding.

Analysis of Options for Electrical Power Unit

Let's evaluate each option to determine its relevance to the measurement of electrical power:

- **Coulomb:** This is the SI unit for electric charge (q), representing the amount of electricity. It measures quantity, not the rate of energy transfer.
- **Kilowatt (kW):** This unit represents 1000 Watts. While widely used in practical contexts like household electricity consumption and motor ratings, the fundamental SI unit is the Watt.
- **HP (Horsepower):** This is a customary unit of power, not part of the SI system. It's often used for engines and motors. The conversion is approximately $1 \text{ HP} \approx 746 \text{ Watts}$.
- **Joule (J):** This is the SI unit for energy or work. Power is the rate of energy expenditure, measured in Joules per second.

The Base SI Unit: Watt

The base SI unit for **electrical power** is the **Watt (W)**. It is defined as one Joule of energy per second.

The mathematical definition is:

$$P = \frac{E}{t}$$

Where:

- P represents Power
- E represents Energy (measured in Joules, J)
- t represents Time (measured in seconds, s)

Therefore, the SI unit for power is Joules per second (J/s), which is formally named the Watt (W).

$$1 \text{ W} = 1 \text{ J/s}$$

Another common formula for electrical power is:

$$P = V \times I$$

Where V is Voltage (SI unit: Volt, V) and I is Current (SI unit: Ampere, A). This confirms that $1 \text{ W} = 1 \text{ V} \times 1 \text{ A}$.

Kilowatt as a Practical SI Unit Multiple

The option **Kilowatt** (kW) is directly related to the base SI unit, Watt. One kilowatt equals one thousand Watts:

$$1 \text{ kW} = 1000 \text{ W}$$

In many practical situations and exams, when the base unit (Watt) isn't listed, a common and standard multiple like Kilowatt is often the intended correct answer, especially when contrasted with units that are either incorrect (like Coulomb or Joule) or non-SI (like HP).

Conclusion

Reviewing the options:

- Coulomb measures charge.
- Joule measures energy.
- HP is a non-SI unit.
- Kilowatt is a standard multiple (1000 W) of the base SI unit for power, Watt.

Among the choices provided, **Kilowatt** is the most appropriate answer representing a unit of **electrical power** within the context of the SI system's multiples, serving as the best fit given that the base unit, Watt, is not directly offered as an option but Kilowatt is derived from it.

128. Answer: b

Explanation:

A **voltmeter** measures voltage (potential difference) and is connected in parallel across a component.

129. Answer: c

Explanation:

Understanding Radix 2 in Number Systems

The term "radix" refers to the base of a number system. It defines the number of unique digits (including zero) used to represent numbers in that system. For example, the decimal system we use daily has a radix of 10 because it uses ten digits (0 through 9).

When we talk about Radix 2, we are specifically referring to a number system that uses only two unique digits. These two digits are 0 and 1.

Which Number System Uses Radix 2?

Let's examine the options provided and their corresponding radices:

- **Hexadecimal numbers:** This system uses 16 unique symbols (0-9 and A-F). Its radix is 16.
- **Octal numbers:** This system uses 8 unique digits (0-7). Its radix is 8.
- **Binary numbers:** This system uses only 2 unique digits (0 and 1). Its radix is 2.
- **Decimal numbers:** This system uses 10 unique digits (0-9). Its radix is 10.

Based on these definitions, the number system that uses Radix 2 is the binary number system.

Number System	Radix (Base)	Digits Used
Binary	2	0, 1
Octal	8	0, 1, 2, 3, 4, 5, 6, 7
Decimal	10	0, 1, 2, 3, 4, 5, 6, 7, 8, 9
Hexadecimal	16	0, 1, 2, ..., 9, A, B, C, D, E, F

Why Radix 2 is Important for Binary Numbers

In the binary number system (Radix 2), numbers are represented using sequences of 0s and 1s. This system is fundamental in digital electronics and computing because electronic components can easily represent these two states (e.g., on/off, high voltage/low voltage). Each position in a binary number represents a power of 2, starting from the rightmost digit (which is 2^0).

For instance, the binary number 101_2 can be converted to decimal as follows:

$$(1 \times 2^2) + (0 \times 2^1) + (1 \times 2^0) = (1 \times 4) + (0 \times 2) + (1 \times 1) = 4 + 0 + 1 = 5_{10}.$$

This demonstrates how the Radix 2 structure allows for the representation of numerical values using only the digits 0 and 1.

Conclusion on Radix 2 and Number Representation

In summary, the radix, or base, of a number system dictates the number of unique digits it employs. Radix 2 specifically uses two digits, 0 and 1. This definition precisely matches the characteristics of the binary number system. Therefore, Radix 2 is used for representing binary numbers.

Revision Table: Radix and Number Systems

Concept	Explanation	Example Radix
Radix (Base)	The number of unique digits in a number system	Decimal (Radix 10), Binary (Radix 2)
Radix 2 System	A number system with base 2, using digits 0 and 1	Binary number system
Binary Numbers	Numbers represented using only digits 0 and 1	1011_2 , 0101_2

Additional Information: Understanding Number Bases

Number bases are crucial for understanding how numerical values are represented. Different bases are used in various contexts:

- **Decimal (Base 10):** Everyday calculations.
- **Binary (Base 2):** Used in computers and digital circuits.
- **Octal (Base 8) and Hexadecimal (Base 16):** Often used in computing as a shorthand for binary numbers, making long binary strings easier to read and write. One octal digit represents 3 binary digits, and one hexadecimal digit represents 4 binary digits.

Understanding the radix helps in converting numbers between different bases and comprehending their value representation.

130. Answer: c

Explanation:

Calculating Maximum Value and Angular Frequency for AC Voltage

The problem asks us to find the maximum value and "wavelength" for a sinusoidal alternating voltage with a given RMS value and frequency. In the context of alternating current (AC) circuits, the term "wavelength" is typically associated with travelling waves or transmission lines and represents a spatial dimension. However, the provided options include values in "rad/sec", which is the unit for angular frequency (ω). It is highly probable that "wavelength" in the question is a typo and it is intended to mean "angular frequency". We will proceed by calculating both the maximum value and the angular frequency of the given sinusoidal alternating voltage.

Understanding the Given Values

We are given the following information:

- Type of voltage: Sinusoidal alternating voltage
- Frequency (f): 50 Hz
- RMS value (V_{rms}): 200 V

Calculating the Maximum Value (Peak Value)

For a sinusoidal alternating voltage, the relationship between the RMS value (V_{rms}) and the maximum value (or peak value, V_m) is given by the formula:

$$V_m = V_{rms} \times \sqrt{2}$$

Let's plug in the given RMS value:

$$V_m = 200 \text{ V} \times \sqrt{2}$$

Using the approximate value $\sqrt{2} \approx 1.4142$:

$$V_m \approx 200 \times 1.4142 \text{ V}$$

$$V_m \approx 282.84 \text{ V}$$

This is the calculated maximum value based on the given RMS value of 200 V.

Calculating the Angular Frequency

The angular frequency (ω) of an alternating voltage is related to its linear frequency (f) by the formula:

$$\omega = 2\pi f$$

where π is the mathematical constant pi (approximately 3.14159).

Let's plug in the given frequency:

$$\omega = 2 \times \pi \times 50 \text{ Hz}$$

$$\omega = 100\pi \text{ rad/sec}$$

Using the approximate value $\pi \approx 3.14$:

$$\omega \approx 100 \times 3.14 \text{ rad/sec}$$

$$\omega \approx 314 \text{ rad/sec}$$

This is the calculated angular frequency based on the given frequency of 50 Hz.

Comparing Calculations with Options

Our calculations yielded:

- Maximum Value (V_m): Approximately 282.84 V
- Angular Frequency (ω): Approximately 314 rad/sec

Let's look at the options provided:

1. 272.2 V, 300 rad/sec
2. 280.2 V, 308 rad/sec
3. 282.2 V, 314 rad/sec
4. 277.4 V, 302 rad/sec

Option 3 provides 282.2 V for the maximum value and 314 rad/sec for the angular frequency. Our calculated maximum value (282.84 V) is very close to 282.2 V, and our calculated angular frequency (314 rad/sec) matches exactly with 314 rad/sec when using $\pi \approx 3.14$. Therefore, option 3 is the closest match based on standard calculations.

Summary of Results

Based on our analysis and assuming "wavelength" refers to angular frequency:

- The maximum value of the voltage is approximately 282.8 V (closest option is 282.2 V).
- The angular frequency is exactly 314 rad/sec (using $\pi \approx 3.14$).

Revision Table: AC Voltage Parameters

Parameter	Symbol	Relationship	Example (f=50 Hz, $V_{rms}=200V$)
Linear Frequency	f	Given in Hz	50 Hz
Angular Frequency	ω	$\omega = 2\pi f$	314 rad/sec ($\approx 100\pi$)
RMS Value	V_{rms}	$V_{rms} = V_m / \sqrt{2}$	200 V
Maximum Value (Peak)	V_m	$V_m = V_{rms} \times \sqrt{2}$	≈ 282.8 V

Additional Information: AC Circuit Concepts

Understanding the different parameters of a sinusoidal AC voltage is crucial in electrical engineering. Here's a bit more detail:

- **Sinusoidal AC Voltage:** This is a voltage that varies over time according to a sine or cosine function. It is the most common type of AC voltage found in power systems.
- **Frequency (f):** Measured in Hertz (Hz), it represents the number of complete cycles the voltage waveform completes in one second. A frequency of 50 Hz means 50 cycles per second.
- **RMS Value (V_{rms}):** The Root Mean Square value is a way to represent the effective value of an AC voltage or current. It is the equivalent DC voltage or current that would dissipate the same amount of power in a resistive load. For a sine wave, $V_{rms} = V_m / \sqrt{2}$.
- **Maximum Value (V_m):** Also known as the peak value, this is the highest instantaneous voltage value reached during each cycle of the waveform.
- **Angular Frequency (ω):** Measured in radians per second (rad/sec), it represents the rate of change of the angle of the sinusoidal function. It is related to the linear frequency by $\omega = 2\pi f$. It's often used in formulas involving capacitors and inductors.
- **Relationship between parameters:** The formulas $V_m = V_{rms}\sqrt{2}$ and $\omega = 2\pi f$ are fundamental for converting between these common AC voltage parameters.

131. Answer: a

Explanation:

Understanding the Nickel-Iron Cell

A nickel-iron cell, also known as an Edison battery, is a type of rechargeable alkaline secondary battery. It uses nickel(III) oxyhydroxide as the positive plate and iron as the negative plate, with an alkaline electrolyte of potassium hydroxide.

This question asks specifically about the material used for the positive plate in a nickel-iron cell.

Components of a Nickel-Iron Cell

A typical nickel-iron cell consists of several key components:

- **Positive Plate:** This is one of the electrodes where oxidation or reduction reactions occur during charging and discharging.
- **Negative Plate:** This is the other electrode involved in the electrochemical reactions.
- **Electrolyte:** This is the ionic conductor between the positive and negative plates.
- **Container:** Holds the components together.

Material of the Positive Plate

In a nickel-iron cell, the positive plate is constructed from nickel hydroxide. This material undergoes chemical changes during the charging and discharging cycles of the battery.

Let's look at the options provided:

- **Nickel hydroxide:** This is a compound of nickel and hydroxide, commonly used in alkaline batteries.

- Ferrous hydroxide: This is a compound of iron and hydroxide, related to iron which is used in the negative plate.
- Lead peroxide: This is a compound of lead and oxygen, typically used in lead-acid batteries, not nickel-iron cells.
- Potassium hydroxide: This is a compound of potassium and hydroxide, which serves as the electrolyte in nickel-iron cells, not the plate material.

Based on the composition of the nickel-iron cell, the positive plate is made of nickel hydroxide.

Component	Material Used in Nickel-Iron Cell
Positive Plate	Nickel Hydroxide (primarily)
Negative Plate	Iron (primarily)
Electrolyte	Potassium Hydroxide (KOH) solution

Therefore, the material used for the positive plate of a nickel-iron cell is nickel hydroxide.

Revision Table: Nickel-Iron Cell Components

Cell Type	Positive Plate Material	Negative Plate Material	Electrolyte
Nickel-Iron Cell	Nickel Hydroxide	Iron	Potassium Hydroxide

Additional Information on Nickel-Iron Batteries

Nickel-iron batteries are known for their long lifespan and robustness, especially in harsh conditions. They are resistant to overcharging and discharging, unlike some other battery types.

The reactions occurring during discharge are approximately:

- Positive plate: $\text{NiOOH} + \text{H}_2\text{O} + \text{e}^- \rightarrow \text{Ni(OH)}_2 + \text{OH}^-$
- Negative plate: $\text{Fe} + 2\text{OH}^- \rightarrow \text{Fe(OH)}_2 + 2\text{e}^-$
- Overall reaction: $\text{NiOOH} + \text{Fe} + \text{H}_2\text{O} \rightarrow \text{Ni(OH)}_2 + \text{Fe(OH)}_2$

During charging, these reactions are reversed.

While they offer durability, nickel-iron cells typically have a lower energy density and higher self-discharge rate compared to modern battery technologies like lithium-ion.

132. Answer: d

Explanation:

Understanding Electrical Protection Devices

The question asks for a single device that provides comprehensive protection against multiple electrical hazards: shock, fire caused by overcurrent, short circuit, earth leakage, and earth fault.

Analyzing the Options

- **Miniature Circuit Breaker (MCB):**

An MCB is primarily designed to protect electrical circuits from damage due to overcurrent or short circuit. It operates by interrupting the circuit when the current exceeds a preset limit. While it protects against fire from overcurrent/short circuit, it does not inherently provide protection against electric shock due to earth leakage or earth fault.

- **(ELCB+RCCB) combination breaker:**

This option refers to a combination of two different types of devices. While combining an Earth Leakage Circuit Breaker (ELCB) or Residual Current Circuit

Breaker (RCCB) with an MCB (or having them separately in series) can provide comprehensive protection, the question specifically asks for a **single device**.

- **Miniature circuit breaker:**

This is the same as the first option, MCB. As discussed, it covers overcurrent and short circuit but not earth leakage or earth fault protection for shock prevention.

- **Earth leakage circuit breaker (ELCB):**

Traditionally, ELCB referred to a device that detects a voltage rise on the earth wire due to an earth fault and trips the circuit. Modern devices designed for earth leakage protection are typically Residual Current Circuit Breakers (RCCBs) which detect current imbalance between live and neutral wires, indicating current leaking to earth. However, the term 'Earth Leakage Circuit Breaker' can sometimes be used generically or refer to devices that provide this type of protection.

Crucially, a single device that combines both overcurrent/short circuit protection (like an MCB) and earth leakage/fault protection (like an RCCB) is known as a Residual Current Breaker with Overcurrent protection (RCBO). While the option lists "Earth leakage circuit breaker", in the context of a single device covering all listed hazards, this option most closely aligns with the concept of a combined protection device like an RCBO, which effectively incorporates earth leakage/fault protection alongside overcurrent and short circuit protection in a single unit. Therefore, interpreting "Earth leakage circuit breaker" in this context points towards a device that handles earth faults/leakage and potentially integrates the other required protections.

Conclusion on the Single Protection Device

Considering the requirement for a single device protecting against overcurrent, short circuit, earth leakage, and earth fault hazards causing shock and fire, a device that combines these functions is needed. While MCBs handle overcurrent/short circuit and RCCBs/ELCBs handle earth leakage/faults, an RCBO is a single device that performs both roles. Among the given options, "Earth leakage

circuit breaker," when considered in the context of a single device covering all specified hazards, is the most appropriate answer, representing a device capable of providing protection against earth faults/leakage as well as potentially integrating overcurrent and short circuit protection in one unit (like an RCBO).

Such a device ensures that if there is an excessive current (overcurrent or short circuit) or if current is leaking to earth (earth leakage or earth fault), the circuit is quickly disconnected, preventing damage, fire, and electric shock.

Key Protection Mechanisms

- **Overcurrent/Short Circuit:** Protection against excessive current flow, which can lead to overheating and fire. Handled by thermal and magnetic trip mechanisms, typically found in MCBs and integrated into RCBOs.
- **Earth Leakage/Earth Fault:** Protection against current flowing through an unintended path to earth, which can cause electric shock and fire. Handled by detecting current imbalance (RCCB/RCBO) or voltage on earth wire (older ELCB).

A single device providing all these ensures a high level of safety against major electrical hazards.

Revision Table: Electrical Protection Devices

Device Type	Primary Protection	Handles Shock/Fire from:
MCB (Miniature Circuit Breaker)	Overcurrent, Short Circuit	Fire (from overcurrent/short circuit)
RCCB (Residual Current Circuit Breaker)	Earth Leakage (detects residual current)	Shock, Fire (from earth leakage)
ELCB (Earth Leakage Circuit Breaker)	Earth Fault (detects voltage on earth wire - older technology)	Shock, Fire (from earth fault)
RCBO (Residual Current Breaker with Overcurrent protection)	Overcurrent, Short Circuit, Earth Leakage/Fault	Shock, Fire (from all listed hazards)

Additional Information on Electrical Safety Devices

Understanding the differences between these devices is crucial for electrical safety. While MCBs are standard for protecting wiring and equipment from thermal and magnetic stresses, RCCBs/ELCBs are essential for protecting people from electric shock by detecting dangerous earth currents.

Modern installations often use a combination of these devices or integrated devices like RCBOs to achieve comprehensive protection. RCBOs are particularly useful as a single point of protection for individual circuits, offering both overload/short circuit and earth fault/leakage protection in one compact unit.

Proper selection and installation of these protection devices are vital to ensure the safety of electrical systems and prevent accidents.

133. Answer: a

Explanation:

Understanding Resistance to Impedance Ratio in AC Circuits

This question asks us to identify the term that describes the ratio of resistance to impedance. Let's break down the concepts involved in AC (Alternating Current) circuits.

What is Electrical Resistance?

Resistance, denoted by the symbol R , is a measure of how much a material opposes the flow of electric current. In DC (Direct Current) circuits, resistance is the only factor opposing current. In AC circuits, resistance represents the part of the opposition that dissipates energy, typically as heat. It is measured in Ohms (Ω).

What is Electrical Impedance?

Impedance, denoted by the symbol Z , is the total opposition that a circuit presents to alternating current. It's a more comprehensive measure than resistance because it includes not only resistance but also the opposition due to capacitance and inductance, collectively known as reactance (X). Impedance is a complex quantity and is represented mathematically as:

$$Z = R + jX$$

where:

- R is the resistance.
- X is the reactance (sum of inductive and capacitive reactance).
- j is the imaginary unit ($\sqrt{-1}$).

The magnitude of impedance, often used in calculations, is given by:

$$|Z| = \sqrt{R^2 + X^2}$$

Impedance is also measured in Ohms (Ω).

Defining the Ratio: Resistance to Impedance

The question specifically asks for the ratio of resistance (R) to impedance (Z). In AC circuits, this ratio is particularly meaningful when considering the magnitude of the impedance:

$$\text{Ratio} = \frac{\text{Resistance}}{\text{Magnitude of Impedance}} = \frac{R}{|Z|}$$

This ratio, $\frac{R}{|Z|}$, is a fundamental concept in AC circuit analysis.

Analyzing the Options

Let's look at the given options to see which one matches the ratio $\frac{R}{|Z|}$:

- **Power Factor:** In an AC circuit, the power factor (PF) is defined as the cosine of the phase angle (ϕ) between the voltage and current. This angle is related to impedance by $R = |Z| \cos(\phi)$. Therefore, the power factor is $\cos(\phi) = \frac{R}{|Z|}$. This matches our definition.
- **Form Factor:** This is the ratio of the RMS (Root Mean Square) value of an AC waveform to its average value. It is unrelated to the resistance-to-impedance ratio.
- **Amplitude:** This refers to the maximum value or peak value reached by an alternating quantity (like voltage or current). It does not represent the ratio of resistance to impedance.
- **Power:** This is the rate at which energy is consumed or supplied in a circuit. While related to resistance and impedance in calculations (e.g., Real Power $P = V_{rms} I_{rms} \cos(\phi)$), power itself is not the ratio $\frac{R}{|Z|}$.

Conclusion on the Ratio

Based on the analysis, the ratio of resistance to the magnitude of impedance ($\frac{R}{|Z|}$) in an AC circuit is defined as the **power factor**.

134. Answer: b

Explanation:

Comparing Charging Time: Constant Voltage vs. Constant Current Systems

When charging a battery, two common methods are constant current (CC) charging and constant voltage (CV) charging. The choice of method significantly impacts the time required to fully charge the battery.

Understanding Constant Current (CC) Charging

In a constant current system, the charger provides a steady, unchanging current to the battery throughout the charging process. As the battery's voltage increases during charging, the charger's output voltage must also increase to maintain the constant current flow. This method is simple but can be slower, especially as the battery approaches full capacity. It also requires careful monitoring to prevent overcharging once the battery voltage reaches a certain level.

Understanding Constant Voltage (CV) Charging

In a constant voltage system, the charger maintains a fixed voltage across the battery terminals. Initially, when the battery's voltage is low, the current drawn by the battery is high. As the battery charges and its voltage rises closer to the charger's voltage, the current drawn decreases. This method allows for a very fast charge initially but tapers off as the battery nears full charge. While it can reach a significant state of charge quickly, achieving 100% charge using only CV can take a long time as the current drops to very low levels.

Comparing Charging Times

Comparing the time taken for charging in a pure constant voltage system versus a pure constant current system can depend on the specific battery type and the target state of charge. However, if we consider charging a depleted battery to a significant percentage (say, 80-90%) of its capacity, the constant voltage method is generally much faster than a constant current method designed to be safe for the entire charge cycle.

In a practical sense, constant voltage charging allows a high current flow initially due to the large voltage difference. This high initial current rapidly increases the battery's state of charge. Although the current tapers off later, the initial rapid charging phase often means that a high percentage of charge is achieved much faster compared to a constant current charge, where the current is limited throughout the process.

Therefore, relative to a simple constant current charging strategy, using a constant voltage approach can significantly reduce the time to reach a substantial charge level. The statement that the time is "almost reduced to half" suggests a notable speed increase achieved by leveraging the constant voltage method's ability to deliver higher initial current.

Feature	Constant Current (CC) Charging	Constant Voltage (CV) Charging
Current	Constant	Starts high, tapers off
Voltage	Increases	Constant
Initial Speed	Slower	Faster
Time to high SoC (e.g., 80%)	Longer	Shorter
Time to 100% SoC (pure method)	Can be faster if current maintained	Can be very long as current drops to near zero

Considering common battery charging profiles, especially the widely used Constant Current-Constant Voltage (CC/CV) method for lithium-ion batteries, the initial CC phase brings the voltage up, and then the CV phase tops it off. If we were to compare a limited CC charge to a CV charge up to a certain voltage threshold, the CV method often reaches that threshold much faster, leading to a reduced overall charging time to a specific point compared to potentially slower CC scenarios.

Revision Table: Charging Systems

- Constant Current: Fixed current, voltage increases, slower initially.
- Constant Voltage: Fixed voltage, current decreases, faster initially.
- CV charging often significantly reduces time to reach a high state of charge compared to simple CC.

Additional Information: Battery Charging Profiles

Many modern batteries, particularly lithium-ion batteries, use a combined charging strategy known as Constant Current-Constant Voltage (CC/CV). This method combines the benefits of both systems:

1. **Constant Current (CC) Phase:** The battery is charged with a constant current until its voltage reaches a predetermined level (the constant voltage point). This phase provides rapid charging up to about 70–80% of capacity.
2. **Constant Voltage (CV) Phase:** Once the voltage reaches the set point, the charger switches to constant voltage mode. The voltage is held constant while the current naturally tapers down as the battery approaches full charge. This phase ensures the battery is safely topped off without overcharging.

This CC/CV approach is faster than a pure CC charge to full capacity and safer than a pure CV charge which could initially draw excessive current if not limited. The significant speed advantage often associated with "constant voltage" charging in comparisons usually refers to the faster initial bulk charging phase it enables compared to simpler, potentially slower, constant current methods.

135. Answer: b

Explanation:

Correct answer: 1 GHz

The ability of a CRO to accurately measure or display a signal depends heavily on its bandwidth. The bandwidth of a CRO indicates the range of frequencies it can handle effectively. Specifically, it's typically defined as the frequency at which the

displayed voltage amplitude drops to 70.7% ($\frac{1}{\sqrt{2}}$) of its true value (or the power drops to 50%).

Rise Time: The rise time of a CRO (and probe) is inversely related to its bandwidth. A faster rise time allows the CRO to accurately display faster transitions in a signal, which is important for observing high-frequency signals. Bandwidth $\approx \frac{0.35}{\text{Rise Time}}$.

136. Answer: a

Explanation:

Understanding the Output Waveform of a Variable-Frequency Drive (VFD)

A Variable-Frequency Drive (VFD), also known as an adjustable-frequency drive or inverter, is a type of motor drive used to control the speed of an AC electric motor by varying the frequency and voltage supplied to the motor. The way a VFD generates this variable frequency and voltage determines its output waveform.

How VFDs Work: From DC to AC

Modern VFDs typically convert incoming AC power to DC power using a rectifier. This DC power is then inverted back to AC power using power electronic switches (like IGBTs – Insulated Gate Bipolar Transistors). The key to varying the frequency and voltage lies in how these switches are turned on and off.

Exploring VFD Output Waveforms

Let's look at the common types of output waveforms and why modern VFDs use a specific technique:

- **Pure sine wave:** This is the ideal waveform for AC motors, as it causes smooth operation and minimizes motor heating and noise. However, generating a

pure sine wave from a DC source using power electronics is complex and expensive for typical VFD applications.

- **Square wave:** Early VFDs might have used simple square waves. While easy to generate, square waves contain significant harmonic content. These harmonics can cause excessive motor heating, vibration, and noise, making it inefficient and potentially damaging for the motor.
- **Chopped sine wave:** This term often refers to waveforms created by phase control, typically used in simple AC voltage regulators or dimmers. This method chops portions of a fixed-frequency sine wave to vary the RMS voltage but does not generate a variable frequency output suitable for VFDs.
- **Pulse width modulated sine wave:** This is the standard output waveform for most modern VFDs. The VFD generates this waveform using a technique called Pulse Width Modulation (PWM).

Pulse Width Modulation (PWM) in VFDs

PWM works by rapidly switching the DC voltage on and off at a high frequency (carrier frequency). The width (duration) of these voltage pulses is varied to simulate a lower-frequency sine wave. The average voltage over a switching cycle follows a sinusoidal pattern. By changing the frequency of the overall "simulated" sine wave and adjusting the pulse widths accordingly, the VFD can vary both the output frequency and voltage supplied to the motor.

Here's a simplified look at PWM:

Concept	Explanation
High Switching Frequency	The VFD switches output voltage thousands of times per second.
Varying Pulse Width	To simulate a high voltage part of the sine wave, pulses are wider (ON time is longer). To simulate a low voltage part, pulses are narrower (ON time is shorter).
Average Voltage	The motor inductance filters the high-frequency pulses, and the motor effectively sees an average voltage that approximates a sine wave.

Because the waveform is created from pulses whose width is modulated to resemble a sine wave when averaged, the output is described as a pulse width modulated sine wave.

Why PWM Sine Wave is Preferred

The PWM technique offers several advantages for VFDs:

- It allows for efficient control of both voltage and frequency.
- It approximates a sine wave closely enough for efficient motor operation, reducing the harmonic issues associated with square waves compared to older methods.
- It is cost-effective to implement with modern power electronic components.

Therefore, the output waveform commonly associated with a Variable-Frequency Drive (VFD) is a Pulse Width Modulated sine wave.

Revision Table: VFD Output Waveform Types

Waveform Type	Characteristics	Use in VFDs
Pure Sine Wave	Smooth, no harmonics.	Ideal but not typically generated directly due to cost/complexity.
Square Wave	Simple to generate; high harmonic content.	Used in older/basic drives; inefficient for motors.
Chopped Sine Wave	Phase controlled; fixed frequency, varied RMS voltage.	Not used for VFDs (variable frequency required).
Pulse Width Modulated Sine Wave	Generated by varying pulse widths at high frequency; approximates sine wave.	Standard for modern VFDs; balances performance and cost.

Additional Information on VFDs and PWM

Variable-Frequency Drives (VFDs) are crucial for energy savings in many industrial applications, as they allow matching motor speed to the actual load requirement. This avoids the inefficiency of running a motor at full speed constantly when not needed.

Pulse Width Modulation (PWM) is not exclusive to VFDs. It is a widely used technique in power electronics for controlling the average power delivered to a load by switching a power switch rapidly between ON and OFF states. The longer the switch is ON compared to the OFF periods during a cycle, the higher the total power supplied. In VFDs, the PWM pattern is specifically designed to create an output that averages out to a sine wave.

The quality of the PWM sine wave output from a VFD can be improved by increasing the switching (carrier) frequency. Higher switching frequencies result in the output voltage better approximating a true sine wave and reduce audible noise from the motor, but they also increase switching losses in the VFD's power components.

137. Answer: a

Explanation:

Understanding Squirrel Cage Motors

The question asks about a key characteristic of a squirrel cage motor. Squirrel cage motors are the most common type of AC induction motor used in various industrial and domestic applications due to their simple and robust construction.

Structure of a Squirrel Cage Rotor

The rotor of a squirrel cage motor is distinctive. It consists of:

- A laminated iron core.

- Conductors (usually bars made of copper or aluminum) embedded in slots around the periphery of the core.
- End rings that connect and short-circuit all the conductor bars at both ends.

This structure resembles a squirrel cage or a hamster wheel, hence the name "squirrel cage".

Analyzing the Options for Squirrel Cage Motors

Let's evaluate each option based on the known characteristics of squirrel cage motors:

- **Option 1: The end rings and the rotor conducting bars are permanently short-circuited**
 - This statement accurately describes the construction of a squirrel cage rotor. The end rings are specifically designed to short-circuit the rotor bars, forming a complete electrical circuit. When the rotating magnetic field from the stator induces a voltage in the rotor bars, this short-circuit path allows current to flow, which in turn creates a magnetic field in the rotor, causing it to rotate. This short-circuiting is permanent and fixed in this type of motor.
- **Option 2: There is a series of induction motor**
 - This statement is incorrect. A squirrel cage motor is a type of induction motor, specifically, a type of AC induction motor. It is not a "series of induction motors" or a "series induction motor" (which is a different concept).
- **Option 3: The end-rings are open circuited**
 - This statement is incorrect. If the end rings were open-circuited, the rotor bars would not be connected, and no complete circuit would exist for induced currents to flow. Without rotor current and the resulting magnetic field interaction with the stator field, the motor would not be able to produce torque and rotate. The end rings must short-circuit the bars.
- **Option 4: The rotor conducting bars are made of carbon**
 - This statement is generally incorrect for standard squirrel cage rotors. The conducting bars are typically made of materials with low resistance like copper or aluminum to facilitate efficient current flow. Carbon is commonly

used for brushes in other types of motors (like DC motors or wound-rotor induction motors) to make electrical contact with a commutator or slip rings, neither of which is present on a squirrel cage rotor.

Based on the analysis, the defining structural feature of the squirrel cage motor rotor is the permanent short-circuiting of the rotor conducting bars by the end rings.

Revision Table: Squirrel Cage Motor Key Features

Feature	Description	Relevance to Question
Rotor Construction	Laminated core with embedded conductor bars.	Fundamental structure.
Conductor Material	Typically Copper or Aluminum.	Option 4 incorrect.
End Rings	Connect and short-circuit all bars at both ends.	Option 1 correct, Option 3 incorrect.
Brushless Rotor	No brushes or slip rings on the rotor circuit.	Distinguishes it from wound-rotor types.
Motor Type	Type of AC Induction Motor.	Option 2 incorrect.

Additional Information on Squirrel Cage Motors

Squirrel cage induction motors are widely used because of their:

- **Simplicity:** The rotor construction is very simple and robust.
- **Durability:** With fewer components (like brushes) that wear out, they require minimal maintenance.
- **Cost-effectiveness:** Relatively inexpensive to manufacture.
- **Reliability:** Known for long operational life under various conditions.

The principle of operation relies on electromagnetic induction. The rotating magnetic field from the stator induces voltage and current in the short-circuited rotor bars, creating a rotor magnetic field that interacts with the stator field, resulting in torque and rotation.

138. Answer: a

Explanation:

Understanding the Full Form of SMPS

The question asks for the full form of the acronym SMPS. Understanding common acronyms in electronics and technology is important.

What SMPS Stands For

SMPS is a widely used acronym in the field of electronics, particularly concerning power supplies. The letters S, M, P, and S each stand for a specific word that describes the type of power supply.

The full form of SMPS is:

- S stands for **Switch**
- M stands for **Mode**
- P stands for **Power**
- S stands for **Supply**

Therefore, SMPS stands for **Switch Mode Power Supply**.

Understanding Switch Mode Power Supply

A **Switch Mode Power Supply (SMPS)** is an electronic power supply that incorporates a switching regulator to convert electrical power efficiently. It converts AC or DC power to regulated DC power at different voltage levels, using switching

elements (like transistors or MOSFETs) that rapidly switch between ON and OFF states.

Key characteristics of SMPS include:

- High efficiency compared to traditional linear power supplies.
- Smaller size and lighter weight.
- Generates less heat.
- More complex design.

Applications of SMPS

SMPS units are commonly found in a wide range of electronic devices because of their efficiency and compact size. Some typical applications include:

- Computers (the power supply unit or PSU is typically an SMPS).
- Televisions.
- Mobile phone chargers.
- LED lighting drivers.
- Industrial equipment.

Comparing SMPS with Linear Power Supplies

It's helpful to understand why switch-mode power supplies are preferred in many applications over traditional linear power supplies. Here is a basic comparison:

Feature	Switch Mode Power Supply (SMPS)	Linear Power Supply
Efficiency	High (typically 75–90%+)	Lower (typically 20–50%)
Size & Weight	Smaller and Lighter	Larger and Heavier (due to transformer)
Heat Generation	Lower	Higher (energy wasted as heat)
Complexity	More Complex	Simpler
Cost	Can be lower for high power	Can be lower for low power
Noise	Can generate electrical noise (requires filtering)	Generally less noisy

This comparison highlights why SMPS is often the choice for modern electronic devices where energy efficiency, size, and weight are critical factors.

SMPS Revision Table

Acronym	Full Form	Function Summary
SMPS	Switch Mode Power Supply	Efficiently converts power using switching regulators, typically AC/DC to regulated DC.

SMPS Additional Information

While the core function of an SMPS is efficient power conversion, there are different topologies or designs of SMPS, such as Buck, Boost, Buck-Boost, Flyback, and Forward converters. Each topology has specific characteristics making it suitable for different applications and voltage conversion requirements. The design of an SMPS involves carefully selecting components and designing control circuitry to

ensure stable output voltage, manage electromagnetic interference (EMI) caused by the switching action, and protect against faults.

139. Answer: d

Explanation:

Understanding Electrical Conductivity

Electrical conductivity refers to the ability of a material to allow electric current to flow through it. This flow of current is carried by charge carriers, which are typically free electrons in metals and ions in liquids.

Materials that allow electric current to pass through easily are called good conductors, while materials that resist the flow of current are called poor conductors or insulators.

Analyzing the Options for Electrical Conductivity

Let's examine the electrical conductivity of each given option:

- **Copper:** Copper is a metal widely known for being an excellent conductor of electricity. This is due to the presence of a large number of free electrons in its atomic structure that can move easily when a voltage is applied.
- **Aluminium:** Aluminium is also a metal and is a very good conductor of electricity, though slightly less conductive than copper. It is often used in power transmission lines because it is lighter and cheaper than copper. Like copper, its conductivity comes from free electrons.
- **Silver:** Silver is the best metallic conductor of electricity among all metals. It has the highest number of free electrons that can move freely. However, it is expensive, so its use as a conductor is limited to special applications.
- **Pure Water:** Pure water, chemically represented as H_2O , is essentially a non-electrolyte. This means it contains very few dissolved ions. Electric current in water is conducted by the movement of ions. Since pure water has a

negligible concentration of free ions (due to the very slight self-ionization of water), it is a very poor conductor of electricity. Impurities or dissolved salts (which create ions) significantly increase the conductivity of water.

Comparing Conductivity

Comparing the conductivity of these materials:

Material	Type	Electrical Conductivity	Reason for Conductivity
Copper	Metal	Very Good Conductor	Presence of abundant free electrons
Aluminium	Metal	Good Conductor	Presence of free electrons
Silver	Metal	Excellent Conductor	Highest concentration of free electrons among common metals
Pure Water	Non-electrolyte	Poor Conductor	Very few free ions

From the table and the analysis, it is clear that Copper, Aluminium, and Silver are all good metallic conductors, with Silver being the best among them. Pure water, on the other hand, is a very poor conductor of electricity because it lacks free ions to carry the charge.

Conclusion: Identifying the Poor Conductor

Based on the electrical properties of the substances listed, Pure Water stands out as the poor conductor compared to the metals provided.

Electrical Conductivity Revision Table

Material Type	Conductivity Level	Charge Carriers
Metals (e.g., Copper, Aluminium, Silver)	High	Free Electrons
Pure Water	Very Low	Ions (very few in pure form)
Ionic Solutions	Moderate to High	Ions
Insulators (e.g., Rubber, Plastic, Glass)	Very Low	Practically none (electrons tightly bound)

Additional Information on Pure Water Conductivity

It is important to note that while pure water is a poor conductor, tap water or natural water sources usually contain dissolved salts and minerals. These dissolved substances dissociate into ions, making the water conductive. The more dissolved ions present, the higher the conductivity of the water. Therefore, the conductivity of water is often used as an indicator of the level of dissolved impurities in it.

The conductivity of pure water is around $0.055 \mu\text{S}/\text{cm}$ at 25°C , which is extremely low compared to metals. For instance, copper has a conductivity of approximately $5.96 \times 10^7 \text{ S}/\text{m}$.

140. Answer: c

Explanation:

Understanding Voltage Across an Open AC Mains Switch

When you measure the voltage across an open switch in an AC mains circuit, you are essentially measuring the potential difference between the two terminals of the switch. Since the switch is open, the circuit path is broken, meaning no current flows through the switch or the connected load.

Why is the Voltage Not Zero?

A voltage meter measures the electrical potential difference between the two points where its probes are connected. In the case of an open switch:

- One terminal of the switch is connected to the live side of the AC mains supply (e.g., the incoming wire).
- The other terminal of the switch is connected to the wire leading towards the load (e.g., a light bulb or appliance).

Even though the circuit is broken at the switch, the voltage source (the AC mains) is still active. The live side of the open switch is at the full mains potential relative to the neutral side (which is usually connected to the other end of the load). The meter, placed across the open switch, is measuring the voltage between these two points which are directly connected to the source voltage.

Think of it like this: You are measuring the voltage difference across the gap in the circuit created by the open switch. This gap is subjected to the full potential difference provided by the AC mains supply.

Contrast with a Closed Switch

If the switch were closed, the circuit would be complete, and current would flow through the load. In this scenario, the voltage measured **across** a closed switch would ideally be 0 V (or very close to 0 V due to the switch's very low resistance). This is because both terminals of a closed switch are effectively at the same potential (ignoring voltage drop due to current and resistance), as they are part of the continuous, low-resistance path.

In summary, with an open switch, the meter sees the full potential difference of the AC mains supply because it's measuring across the interruption in the circuit, which is subjected to the source voltage.

Common AC Mains Voltages

AC mains voltage varies by region:

- North America (USA, Canada): Typically 110 V – 120 V
- Europe, Asia, Africa, Oceania (including regions like India, UK): Typically 220 V – 240 V

Given the options, 230 V is a common standard in many parts of the world.

Therefore, measuring voltage across an open AC mains switch in a region with 230 V mains will show approximately 230 V.

Thus, if voltage is measured across an AC mains open switch, the meter will read the source voltage.

Switch State	Voltage Measurement Across Switch
Open	Full source voltage (e.g., 230 V)
Closed	Approximately 0 V

Based on the options provided and common mains voltages, the meter will read 230 V.

Revision Table: Key Concepts in AC Circuits

Concept	Description
Voltage (V)	The electrical potential difference between two points in a circuit, driving current.
Current (A)	The flow of electric charge through a conductor.
Resistance (Ω)	The opposition to the flow of electric current.
AC Mains	Alternating Current electrical power supplied to buildings.
Open Circuit	A circuit path that is broken, preventing current flow. Voltage can exist across the break.
Closed Circuit	A complete circuit path, allowing current flow when a voltage source is present.

Additional Information: Measuring Voltage Safely

When measuring voltage in AC mains circuits, safety is paramount. Always use a meter rated for the expected voltage and ensure probes are in good condition. Connect the meter in parallel across the component or points where you want to measure the potential difference. In the case of a switch, place one probe on each terminal of the switch. Ensure the circuit is live when measuring mains voltage across an open switch, but exercise extreme caution due to the risk of electric shock.

Understanding how voltage meters work and the nature of open vs. closed circuits is crucial for safe and accurate electrical measurements.

141. Answer: c

Explanation:

Understanding Galvanic Isolation

Galvanic isolation refers to the principle of electrically separating two circuits such that there is no direct conductive path between them. Information or power can still be transferred between the circuits, often through magnetic fields (as in transformers) or light (as in optocouplers), but the absence of a direct electrical connection enhances safety and prevents the flow of unwanted currents between the circuits, like ground loops.

In power systems, galvanic isolation is often used to protect users from high voltages or to separate sensitive electronic equipment from noisy power lines.

Exploring Different Transformer Types

Transformers are electrical devices that transfer electrical energy between two or more circuits through electromagnetic induction. Different types of transformers exist, each with specific construction and applications. Let's look at the types mentioned in the options:

- **Step-down transformer:** This type of transformer decreases voltage from primary to secondary. It typically has separate primary and secondary windings.
- **Center-tap transformer:** A type of transformer with a tap at the center point of the secondary winding. It also has separate primary and secondary windings.
- **Auto transformer:** Unlike standard transformers, an auto transformer has only one winding, part of which is common to both the primary and secondary circuits. The voltage transformation is achieved by tapping the single winding.
- **Step-up transformer:** This type of transformer increases voltage from primary to secondary. It typically has separate primary and secondary windings.

Analyzing Transformer Types and Galvanic Isolation

Let's consider how the construction of these transformer types relates to galvanic isolation:

- **Transformers with Separate Windings (Step-down, Center-tap, Step-up):**
Standard transformers like step-up, step-down, and center-tap transformers feature two or more windings that are electrically insulated from each other. The energy transfer occurs via the magnetic field in the core. Since there is no direct electrical connection between the primary and secondary windings, these types of transformers inherently provide galvanic isolation.
- **Auto transformer:** An auto transformer uses a single winding with a tap. The primary circuit is connected across part of this winding, and the secondary circuit is connected across a different part (or the whole winding). Because the primary and secondary circuits share a portion of the same winding, there is a direct electrical connection between them. Therefore, standard auto transformers **do not** provide galvanic isolation.

Conclusion Based on the Provided Information

The question asks which type of transformer provides galvanic isolation, and the options include Step-down, Center-tap, Auto transformer, and Step-up transformer. Based on the standard principles of transformer operation, transformers with separate primary and secondary windings (Step-down, Center-tap, Step-up) provide galvanic isolation, while auto transformers do not due to their shared winding.

However, given the options and the premise of the question indicating one of these types provides galvanic isolation, and considering the provided answer, the auto transformer is identified as the relevant type in this specific context.

Transformer Type	Winding Configuration	Provides Galvanic Isolation (Standard Principle)
Step-down	Separate Primary & Secondary	Yes
Center-tap	Separate Primary & Secondary	Yes
Auto transformer	Single, Shared Winding	No
Step-up	Separate Primary & Secondary	Yes

Revision Table: Transformer Isolation Summary

Feature	Standard Transformers (Step-up, Step-down, Center-tap)	Auto transformer
Windings	Separate Primary and Secondary	Single winding shared by Primary and Secondary
Galvanic Isolation	Provides isolation (No direct electrical connection)	Does NOT provide isolation (Direct electrical connection)
Size & Cost (for same rating)	Larger, More Expensive	Smaller, Less Expensive
Fault Behavior	Primary and Secondary circuits isolated during fault	Fault in one circuit directly affects the other

Additional Information on Galvanic Isolation

Galvanic isolation is crucial in many applications for safety and performance:

- **Safety:** It prevents dangerous voltages on the primary side from reaching the secondary side, protecting users from electric shock. This is why isolation transformers are often used in medical equipment and some power supplies.
- **Noise Reduction:** It helps break ground loops and prevents noise or transient voltages from coupling between different parts of a system.
- **Circuit Protection:** It can protect sensitive components on the secondary side from surges or faults occurring on the primary side.

While transformers with separate windings are the primary method for achieving galvanic isolation magnetically, other methods include optocouplers (using light) and capacitive barriers.

142. Answer: b

Explanation:

Understanding Rotor Types in Electrical Machines

The question asks about a specific characteristic of a rotor: having a cylindrical laminated core with parallel slots for conductors. Different types of electrical machines use different rotor constructions tailored to their operating principles.

Analyzing Different Rotor Types

Let's examine the characteristics of the rotor types provided in the options:

- **PMMC rotor:** PMMC (Permanent Magnet Moving Coil) instruments use a coil that moves within a magnetic field provided by a permanent magnet. The "rotor" is essentially this moving coil, which is typically wound on a former, not a laminated core with slots in the context of a motor or generator rotor.
- **Slip ring rotor:** This is a type of rotor commonly used in wound-rotor induction motors. It consists of a cylindrical laminated core with slots that are usually parallel or slightly skewed. Conductors (typically copper bars or windings) are

placed in these slots, and their ends are connected to slip rings mounted on the shaft.

- **DC rotor:** A DC motor or generator rotor (also called an armature) has a laminated core with slots where armature conductors are placed. These conductors are connected to a commutator. While it has a laminated core and slots, the defining feature is the commutator, not slip rings, and its construction is optimized for DC operation and commutation.
- **BLDC rotor:** A Brushless DC (BLDC) motor rotor typically consists of permanent magnets mounted on a core. It does not have windings or conductors in slots in the same way as induction or traditional DC motor rotors.

Identifying the Rotor with Cylindrical Laminated Core and Parallel Slots

Based on the descriptions above, the rotor type that precisely fits the description of having a cylindrical laminated core with parallel slots for carrying rotor conductors is the slip ring rotor, which is a key component of wound-rotor induction motors.

The lamination of the core is essential to reduce eddy currents, improving efficiency. The cylindrical shape allows it to rotate smoothly within the stator, and the parallel (or slightly skewed) slots are specifically designed to house the rotor conductors effectively, enabling the electromagnetic interaction with the stator field necessary for motor operation.

Rotor Type	Core Type	Slots for Conductors	Key Feature
PMMC Rotor	Former (often non-magnetic)	N/A (Coil is the conductor)	Moving coil in magnet field
Slip Ring Rotor	Cylindrical Laminated Core	Yes (Parallel/Skewed)	Windings connected to Slip Rings
DC Rotor (Armature)	Laminated Core	Yes	Windings connected to Commutator
BLDC Rotor	Solid/Laminated Core	No (Permanent Magnets)	Permanent Magnets

Comparing the features, the slip ring rotor uniquely matches the description provided in the question: a cylindrical laminated core with parallel slots for conductors that form windings connected to slip rings.

Revision Table: Rotor Construction Key Points

Characteristic	Description	Relevance to Rotor
Laminated Core	Made of thin steel sheets (laminations) insulated from each other.	Reduces eddy currents, minimizing energy loss and heating in the core.
Cylindrical Shape	Allows for smooth rotation within the stator's bore.	Essential for the mechanical function of a rotating machine.
Parallel Slots	Grooves cut axially along the core's periphery.	Hold the rotor conductors (bars or windings) securely.
Rotor Conductors	Copper or aluminum bars/windings placed in slots.	Carry induced current and interact with the stator field to produce torque.

Additional Information: Rotor Design Considerations

The design of a rotor core and its slots is crucial for the performance of an electrical machine. Here are some additional points:

- **Skewing of Slots:** Sometimes, the rotor slots are slightly skewed instead of being perfectly parallel to the shaft. Skewing helps to reduce cogging torque (tendency of the rotor to align with the stator teeth even when not powered), reduce harmonic content in the induced voltage, and sometimes improve starting torque.
- **Core Material:** The core is typically made of high-grade silicon steel laminations. Silicon steel has high permeability and low hysteresis loss, making it suitable for magnetic cores.

- **Purpose of Laminations:** As mentioned, laminations are vital for minimizing eddy current losses. A solid core would experience large circulating currents induced by the changing magnetic field, leading to significant power loss and heating.
- **Rotor Windings vs. Cage:** Slip ring rotors have distinct windings placed in the slots, while squirrel cage rotors (another type of induction motor rotor) have bars short-circuited by end rings, forming a cage-like structure in the slots. The question's description of "carrying rotor conductors" placed in parallel slots applies to both wound rotors and squirrel cage rotors, but the options steer towards specific machine types. The slip ring rotor fits the description well, especially in contrast to the other options provided.

143. Answer: a

Explanation:

Understanding Electric Iron Temperature Control

Electric irons are essential household appliances used for pressing clothes. They generate heat, and controlling this heat is crucial to prevent damage to fabrics and ensure efficient operation. The question asks about the primary method used for temperature control in these devices.

Temperature control in an electric iron typically involves sensing the plate temperature and switching the heating element on or off to maintain a desired level. Let's look at the options provided.

Analyzing Temperature Control Methods

Several methods and components can be used for temperature sensing and control. The options are:

- Bi-metallic strip
- Thermistor

- Resistance temperature detector (RTD)
- Thermocouple

Each of these has different properties and applications. In simple, common household electric irons, a specific method is widely used for its reliability and cost-effectiveness in providing a simple on/off switching action based on temperature.

The Role of the Bi-metallic Strip in Temperature Control

A **bi-metallic strip** is made of two different metals that are joined together. These metals have different coefficients of thermal expansion. When heated, one metal expands more than the other. This differential expansion causes the strip to bend.

In an electric iron, the bi-metallic strip is placed near the heating plate. As the plate heats up, the bi-metallic strip also heats up and bends. This bending is used to operate a switch. When the temperature reaches a set point, the bending of the strip causes the switch contacts to open, interrupting the flow of current to the heating element. As the plate cools down, the bi-metallic strip straightens, closing the contacts again and allowing current to flow, thus reheating the element.

The bending action is predictable based on temperature, making the bi-metallic strip an effective and simple thermostat for controlling the iron's temperature within a desired range. The desired temperature can often be adjusted by changing the position of the switch mechanism relative to the bi-metallic strip using a dial on the iron.

Why Other Options Are Less Common for This Application

- **Thermistor:** A thermistor is a type of resistor whose resistance is strongly dependent on temperature. While very sensitive, using it for direct switching requires additional electronic circuitry (like comparators or microcontrollers) to read the resistance change and operate a switch (like a relay or triac). This adds complexity and cost compared to a simple mechanical bi-metallic strip thermostat.

- **Resistance Temperature Detector (RTD):** Like thermistors, RTDs are resistors whose resistance changes with temperature, typically in a more linear fashion than thermistors. They offer high accuracy but also require supporting electronics to interpret the resistance signal and control the heating element, making them less suitable for the simple, direct mechanical switching needed in a basic electric iron thermostat.
- **Thermocouple:** A thermocouple is a sensor that produces a voltage proportional to the temperature difference between two different metals joined at one end. They are robust and suitable for very high temperatures but require a reference junction and signal conditioning circuitry to measure temperature accurately and use it for control, which is more complex and expensive than a bi-metallic strip thermostat.

Given the need for a simple, reliable, and cost-effective mechanical temperature control mechanism in household electric irons, the bi-metallic strip is the most commonly used component for the temperature controlling function.

Revision Table: Temperature Sensing Methods

Method	Principle	Common Application in Irons	Complexity
Bi-metallic Strip	Differential thermal expansion causes bending	Yes, primary control/thermostat	Simple, Mechanical
Thermistor	Resistance changes with temperature	Less common for direct switching	Requires electronics
RTD	Resistance changes with temperature (linear)	Less common for direct switching	Requires electronics
Thermocouple	Generates voltage based on temperature difference	Less common for direct switching	Requires electronics, reference

Additional Information: Thermostats and Electric Irons

The temperature control mechanism in an electric iron is essentially a type of thermostat. A thermostat is a device that automatically regulates temperature, or that activates a system when the temperature reaches a certain point. In the case of the bi-metallic strip in an iron, it acts as a simple, mechanical thermostat that switches the heating element on and off to maintain the set temperature.

Modern irons might incorporate additional safety features or more sophisticated electronic controls using other sensors, but the basic temperature regulation in many models relies on the tried and tested bi-metallic strip thermostat due to its simplicity and reliability.

Understanding how a bi-metallic strip works helps explain why electric irons click when they are heating up and cycling on or off — this clicking sound is often the sound of the bi-metallic strip flexing and operating the switch contacts.

144. Answer: d

Explanation:

Understanding Series Combination of Resistances

In an electric circuit, components like resistors can be connected in different ways. Two common ways are series combination and parallel combination. This question focuses on the properties of a series combination of resistances.

When resistances are connected in series, they are joined end-to-end, forming a single path for the electric current to flow. Imagine a chain where each link is a resistor; the current must pass through every link in sequence.

Properties of Current in a Series Circuit

A fundamental principle of electric circuits is the conservation of charge. Charge carriers (like electrons) cannot be created or destroyed as they move through the circuit. In a series circuit, there is only one path for these charge carriers to follow from the positive terminal of the source, through each resistor, and back to the negative terminal.

Because there is only a single path, the rate at which charge flows past any point in the series circuit must be the same. This rate of charge flow is what we define as electric current.

- The current entering the first resistor is the same as the current leaving it.
- This current then enters the second resistor and leaves it at the same rate.
- This holds true for all resistors connected in series.
- Therefore, the current remains the same throughout a series combination of resistances.

Properties of Voltage in a Series Circuit

Voltage, also known as electric potential difference, represents the energy per unit charge. As electric current flows through a resistor, energy is dissipated (usually as heat). This energy dissipation causes a drop in electric potential across the resistor.

In a series circuit, the total voltage supplied by the source is divided among the individual resistors. The voltage drop across each resistor depends on its resistance value and the current flowing through it, according to Ohm's Law: $V = IR$.

- If the resistances are different, the voltage drop across each resistor will also be different, even though the current I is the same for all.
- The sum of the voltage drops across all resistors in series equals the total voltage supplied by the source. This is a consequence of Kirchhoff's Voltage Law.
- Therefore, the voltage generally varies across different resistances in a series combination.

Analysing the Options

Let's look at the given options in light of our understanding of series circuits:

1. The voltage remains same: This is incorrect. The voltage varies across different resistances in series, unless all resistances are equal (and even then, it's a drop across each, not constant throughout the circuit).
2. The current and voltage remain same: This is incorrect. While the current remains same, the voltage varies.
3. The current and voltage both vary: This is incorrect. The current remains same, while the voltage varies.
4. The current remains same: This is correct. As explained above, the current is the same at every point in a series circuit because there is only one path for charge flow.

Based on the fundamental principles of series circuits, the current is constant throughout the combination, while the voltage drops across each individual resistance and varies depending on the resistance value.

Summary of Series Combination Properties

Property	In Series Combination
Current (I)	Remains the same through all components
Voltage (V)	Divides across individual components; sum of voltage drops equals total voltage
Total Resistance (R_{total})	Sum of individual resistances ($R_1 + R_2 + R_3 + \dots$)

Conclusion on Series Circuit Properties

In a series combination of resistances, the electric current flowing through each resistor is identical. The voltage, however, is divided among the resistors, with the potential drop across each resistor being proportional to its resistance value.

Revision Table: Series Resistance Basics

Review the key characteristics of resistors connected in series.

Characteristic	Description for Series Resistors
Current	Same at all points in the circuit. $I_{total} = I_1 = I_2 = I_3 = \dots$
Voltage	Total voltage is the sum of voltage drops across each resistor. $V_{total} = V_1 + V_2 + V_3 + \dots$
Total Resistance	Sum of individual resistances. $R_{total} = R_1 + R_2 + R_3 + \dots$

Additional Information: Comparing Series and Parallel Circuits

It is helpful to contrast series circuits with parallel circuits to fully understand their properties. In a parallel combination, resistances are connected across the same two points, providing multiple paths for the current.

- In parallel, the voltage across each resistance is the same.
- In parallel, the total current from the source divides among the different paths; the sum of currents through each branch equals the total current.
- The total resistance in parallel is calculated differently ($1/R_{total} = 1/R_1 + 1/R_2 + \dots$).

Understanding these differences is crucial for analysing more complex circuits.

145. Answer: c

Explanation:

Understanding Soldering Joint Configurations

Soldering is a process that joins two or more items by melting and putting a filler metal (solder) into the joint, with the filler metal having a lower melting point than the adjoining material. For soldering to be effective and create a strong, reliable

connection, the configuration of the joint is very important. The joint configuration affects how well the solder flows, the strength of the connection, and the ease of assembly.

Analyzing Common Soldering Joint Types

Let's look at the typical joint configurations mentioned and understand their characteristics in the context of soldering:

- **Butt Joint:** In a butt joint, two pieces are joined end-to-end. This configuration offers a small contact area between the surfaces to be joined and a limited area for the solder to bond to. This often results in a mechanically weak joint, especially under tension. Solder primarily sits on the surface of the join line.
- **Lap Joint:** A lap joint is formed by overlapping two pieces over each other. This creates a relatively large surface area where the two parts are in contact and can be bonded by solder. This configuration allows for excellent capillary action, where the molten solder is drawn into the gap between the overlapping surfaces, creating a strong bond across a significant area. Lap joints generally offer good mechanical strength.
- **Scarf Joint:** A scarf joint is similar to a butt joint but involves cutting the ends of both pieces at an angle, then fitting the angled ends together. This increases the surface area compared to a simple butt joint, providing more area for bonding. However, aligning scarf joints accurately can be challenging, and they may still be less mechanically robust than lap joints, especially for soldering applications.
- **Wave Joint:** The term "wave joint" is not a standard name for a joint geometry in the same way Butt, Lap, and Scarf are. "Wave soldering" is a bulk soldering process used in manufacturing, especially for printed circuit boards (PCBs), where components with leads are passed through a wave of molten solder. The joints formed in wave soldering processes for through-hole components are often variations of lap joints (the component lead passing through a hole in the PCB pad). While not a specific joint geometry itself, the success of wave soldering often relies on configurations like lap joints that facilitate solder flow and filling through capillary action.

Why Lap Joint is Recommended for Soldering

Based on the characteristics, the **Lap joint** is generally the most recommended configuration for soldering for several key reasons:

- **Increased Surface Area:** The overlapping nature provides a much larger area for the solder to bond to compared to butt joints. More bonding area typically means a stronger electrical and mechanical connection.
- **Excellent Capillary Action:** The small gap between the overlapping surfaces effectively draws molten solder into the joint via capillary action. This ensures good solder penetration and wetting throughout the bonding area, leading to a reliable connection.
- **Improved Mechanical Strength:** Due to the large bonded area and good solder penetration, lap joints are significantly stronger mechanically than butt joints when soldered correctly. This is crucial for components or wires that might experience physical stress.
- **Ease of Assembly:** For many applications, creating a lap joint by simply overlapping materials is straightforward.

Comparison of Soldering Joint Types

Here is a comparison summarizing the suitability of common joint types for soldering:

Joint Type	Description	Solderable Area	Capillary Action	Mechanical Strength (Soldered)	Suitability for Soldering
Butt Joint	Ends meet	Small	Limited	Low	Least Recommended
Lap Joint	Overlapping pieces	Large	Excellent	Good	Most Recommended
Scarf Joint	Angled ends meet	Medium	Fair	Moderate	Less Recommended (Complexity)
Wave Joint	(Process term, often uses Lap geometry)	N/A (Process name)	N/A (Process name)	N/A (Process name)	N/A (Not a standard geometry type)

Considering the need for a strong, reliable, and easily repeatable bond, the **Lap joint** stands out as the most suitable and recommended configuration for many soldering applications.

Revision Table: Key Soldering Joint Concepts

Concept	Description
Soldering	Joining materials using a filler metal (solder) below the base material's melting point.
Capillary Action	The ability of a liquid (molten solder) to flow in narrow spaces without the assistance of external forces. Crucial for solder flow in tight joints.
Wetting	The ability of molten solder to flow smoothly and spread evenly over the surface of the materials being joined, indicating good metallurgical bonding.
Joint Strength	The mechanical integrity of the soldered connection, influenced by the joint configuration, solder quality, and soldering process.
Lap Joint	Joint where materials overlap; recommended for soldering due to large surface area and capillary action.

Additional Information on Soldering Techniques

Beyond the joint type, several other factors influence the quality of a soldered connection:

- **Flux:** A chemical cleaning agent applied to the surfaces before soldering. It removes oxides and prevents re-oxidation during the heating process, allowing the solder to wet the surfaces properly. Different types of flux are used depending on the materials being joined and the application (e.g., rosin flux for electronics, acid flux for plumbing).
- **Solder Type:** The composition of the solder alloy affects its melting point, flow characteristics, and strength. Common solders include tin-lead alloys (historically common, but restricted due to lead content) and various lead-free alloys (e.g., tin-silver-copper).
- **Heat Control:** Applying the correct amount of heat for the appropriate duration is critical. Too little heat results in a cold joint (poor wetting and

bonding), while too much heat can damage components or the base materials.

- **Surface Preparation:** Surfaces to be soldered must be clean and free from dirt, grease, and heavy oxides for good wetting and reliable joints.

Understanding the recommended configuration for soldering, like the lap joint, along with proper technique ensures strong and reliable connections.

146. Answer: b

Explanation:

Understanding the B-H Curve in Magnetism

The B-H curve, also known as the hysteresis loop, is a fundamental concept in understanding the magnetic properties of materials, particularly ferromagnetic materials. This curve graphically represents the relationship between two crucial magnetic quantities: magnetic field strength (H) and magnetic flux density (B).

Let's break down what each axis of the B-H curve represents to understand what H stands for.

What Does H Represent in a B-H Curve?

In the context of a B-H curve:

- The horizontal axis (x-axis) represents the **Magnetic Field Strength**, denoted by the symbol **H**.
- The vertical axis (y-axis) represents the **Magnetic Flux Density**, denoted by the symbol **B**.

The magnetic field strength (H) is also often called the magnetizing field or magnetic field intensity. It is a measure of how much an external current or a permanent magnet is capable of magnetizing a material or creating a magnetic field. It is typically measured in units of Amperes per meter (A/m).

The magnetic flux density (B) is a measure of the actual magnetic field lines passing through a given area. It reflects the effect of the magnetic field strength (H) within a material, influenced by the material's own magnetic properties. It is measured in units of Teslas (T).

Therefore, the B-H curve plots the magnetic flux density (B) that results from applying a certain magnetic field strength (H) to a material.

Analyzing the Options for H in the B-H Curve

Let's look at the given options:

1. Magnetic flux
2. Magnetic field strength
3. Reluctance
4. Magnetic flux density

Based on our understanding of the B-H curve:

- **Magnetic flux (Φ):** Magnetic flux is the total magnetic field passing through a given area. It is related to magnetic flux density (B) by the area, but it is not the quantity represented by H on the x-axis. Magnetic flux is measured in Webers (Wb).
- **Magnetic field strength (H):** This is the quantity plotted on the x-axis of the B-H curve. It represents the applied magnetizing force. This aligns perfectly with the definition of H .
- **Reluctance ($\mathcal{R}$):** Reluctance is the opposition that a magnetic circuit offers to magnetic flux. It is analogous to electrical resistance. Reluctance is not one of the axes of the B-H curve. It is measured in inverse Henries (H^{-1}).
- **Magnetic flux density (B):** This is the quantity plotted on the y-axis of the B-H curve, representing the density of magnetic field lines within the material. While B is part of the B-H curve, it is represented by the vertical axis, not the horizontal axis (H).

Thus, in a B-H curve, H represents the **Magnetic field strength**.

Revision Table: Key Magnetic Quantities

Quantity	Symbol	Units	Role in B-H Curve
Magnetic Flux Density	B	Teslas (T)	Y-axis
Magnetic Field Strength	H	Amperes/meter (A/m)	X-axis
Magnetic Flux	Φ	Webers (Wb)	Related to B, not on axes
Reluctance	$\mathcal{R}$	H^{-1} (inverse Henries)	Not on axes

Additional Information: B-H Curve Insights

The shape of the B-H curve provides valuable information about a material's magnetic properties, such as:

- **Permeability (μ):** The relationship between B and H is given by $B = \mu H$ (in simple cases or vacuum $B = \mu_0 H$, where μ_0 is the permeability of free space). The slope of the B-H curve at any point is related to the material's permeability (μ).
- **Saturation:** For ferromagnetic materials, as H increases, B also increases initially, but eventually, B levels off. This is called saturation, where the material is fully magnetized.
- **Residual Magnetism (Remanence):** When the applied magnetic field strength (H) is reduced to zero, the magnetic flux density (B) in the material does not drop to zero. This remaining magnetic field is called remanence or residual magnetism.
- **Coercivity (Coercive Force):** The amount of reverse magnetic field strength ($-H$) needed to reduce the magnetic flux density (B) back to zero is called coercivity.

- **Hysteresis Loop:** The complete B-H curve for a material, showing how B changes as H is varied through a cycle of positive and negative values, forms a loop known as the hysteresis loop. The area of this loop is proportional to the energy lost as heat during each magnetization cycle.

Understanding the B-H curve and the quantities it represents is crucial for designing magnetic circuits, transformers, motors, and other electromagnetic devices.

147. Answer: b

Explanation:

Understanding Form Factor of Alternating Quantities

Alternating quantities, like voltage or current in AC circuits, constantly change over time. To characterize these varying quantities, we use specific values like the instantaneous value, peak value, RMS value, and average value. The form factor is one such characteristic ratio that helps us understand the shape of an alternating waveform.

What is Form Factor?

The form factor of an alternating quantity is a measure that relates its effective value (RMS value) to its average value over a half-cycle. It is a dimensionless quantity and provides insight into how much the waveform deviates from a perfect rectangle (which would have a form factor of 1).

The formula for the form factor (often denoted as k_f) is:

$$\text{Form Factor} = \frac{\text{RMS Value}}{\text{Average Value}}$$

Where:

- **RMS Value (Root Mean Square Value):** This is the effective value of an alternating quantity. It represents the equivalent DC value that would produce the same amount of heat in a resistive circuit. For a sinusoidal waveform, the RMS value is approximately 0.707 times the peak value.
- **Average Value:** This is the average of all instantaneous values of an alternating quantity over one half-cycle. For a symmetrical waveform like a sine wave, the average value over a full cycle is zero, so the average value is calculated over a half-cycle. For a sinusoidal waveform, the average value over a half-cycle is approximately 0.637 times the peak value.

Analyzing the Options

Let's examine the given options based on the definition of form factor:

1. **the rms value to the peak value:** This ratio is related to the crest factor (or peak factor), not the form factor. The crest factor is defined as the ratio of the peak value to the RMS value.
 2. **the rms value to the average value:** This ratio perfectly matches the definition of the form factor of an alternating quantity as discussed above.
 3. **the peak value to the rms value:** This is the definition of the crest factor (or peak factor), which indicates the ratio of the peak value to the effective value.
 4. **the average value to the rms value:** This is the reciprocal of the form factor.
- The form factor is specifically defined as RMS/Average .

Based on the standard definition in electrical engineering, the form factor of an alternating quantity is the ratio of the RMS value to the average value.

Calculating Form Factor for Common Waveforms

The form factor is constant for a given waveform shape. Here are the form factors for some common alternating waveforms:

Waveform	RMS Value (relative to Peak Value V_p)	Average Value (relative to Peak Value V_p)	Form Factor (RMS / Average)
Sinusoidal	$\frac{V_p}{\sqrt{2}} \approx 0.707V_p$	$\frac{2V_p}{\pi} \approx 0.637V_p$	$\frac{V_p/\sqrt{2}}{2V_p/\pi} = \frac{\pi}{2\sqrt{2}} \approx 1.11$
Square Wave	V_p	V_p	$\frac{V_p}{V_p} = 1$
Triangular Wave	$\frac{V_p}{\sqrt{3}} \approx 0.577V_p$	$\frac{V_p}{2} = 0.5V_p$	$\frac{V_p/\sqrt{3}}{V_p/2} = \frac{2}{\sqrt{3}} \approx 1.155$

As seen in the table, the form factor of a sine wave is approximately 1.11, a square wave is 1, and a triangular wave is approximately 1.155. A form factor of 1 indicates a waveform where the RMS and average values are equal, which is characteristic of a square wave.

Form Factor Ratio Summary

To summarize, the form factor is a key parameter used in AC circuit analysis and power systems to characterize the shape of an alternating waveform by comparing its effective heating value (RMS) to its rectified average value (Average over half cycle).

Revision Table: Key Waveform Ratios

Ratio Type	Formula	Significance
Form Factor	$\frac{\text{RMS Value}}{\text{Average Value}}$	Indicates how different the waveform is from a square wave (Form factor = 1 for square wave).
Peak Factor (Crest Factor)	$\frac{\text{Peak Value}}{\text{RMS Value}}$	Indicates how sharp the peak is compared to the effective value. (Peak factor = $\sqrt{2} \approx 1.414$ for sine wave).

Additional Information on Alternating Quantity Characteristics

Understanding the different values associated with alternating quantities is crucial in electrical engineering:

- **Instantaneous Value:** The value of the quantity at any specific point in time. For a sine wave, it's typically represented as $v(t) = V_p \sin(\omega t)$ or $i(t) = I_p \sin(\omega t)$.
- **Peak Value (Amplitude):** The maximum instantaneous value reached by the quantity during one cycle.
- **Average Value:** The arithmetic mean of all instantaneous values over a defined period. For AC analysis, the average over a half-cycle is typically used to avoid the average over a full cycle being zero for symmetrical waveforms.
- **RMS Value:** The Root Mean Square value. It is mathematically calculated as the square root of the mean of the squares of the instantaneous values over one cycle. It represents the DC equivalent in terms of power transfer or heating effect.

Form factor and peak factor are useful parameters for comparing different waveforms and are important in the design and analysis of circuits and equipment that handle alternating currents and voltages.

148. Answer: d

Explanation:

The question asks about the fundamental principle on which a transformer operates. Let's analyze the options to understand how a transformer works.

Transformer Working Principle: Mutual Induction

A transformer is a static electrical device that transfers electrical energy between two or more circuits through electromagnetic induction. Its primary working principle is based on **mutual electromagnetic induction**, a phenomenon described by Faraday's Law of Electromagnetic Induction.

- **Mutual Induction Explained:** When an alternating current (AC) flows through the primary coil of a transformer, it creates a continuously changing magnetic field (or magnetic flux) in the core.
- **Flux Linkage:** The core, typically made of ferromagnetic material, is designed to channel this changing magnetic flux effectively.
- **Induction in Secondary Coil:** This changing magnetic flux links with the secondary coil wound around the same core. According to Faraday's Law of Induction, a changing magnetic flux passing through a coil induces an electromotive force (EMF) across that coil. The rate of change of flux linkage determines the magnitude of the induced EMF.
- **Voltage Transformation:** The ratio of the number of turns in the secondary coil to the number of turns in the primary coil determines whether the output voltage is stepped up or stepped down. Mathematically, the induced EMFs are related to the turns ratio: $\frac{V_s}{V_p} \approx \frac{N_s}{N_p}$ where V_s and V_p are the voltages in the secondary and primary coils, and N_s and N_p are the number of turns in the secondary and primary coils, respectively.

Therefore, the operation of a transformer relies entirely on the principle of mutual electromagnetic induction between its primary and secondary windings.

Analysis of Other Options

Let's look at why the other options are not the primary working principle:

- **Radio frequency:** Radio frequency deals with electromagnetic waves in the radio frequency range. While transformers are used in radio circuits, radio frequency itself is not the principle of their operation.
- **Eddy currents:** Eddy currents are circulating currents induced within the conductive core of the transformer due to the changing magnetic flux. While they are an important consideration in transformer design (laminated cores are used to minimize energy loss due to eddy currents), they are a consequence of induction, not the fundamental principle of energy transfer.
- **Photo radiation:** Photo radiation relates to the emission or absorption of light or other electromagnetic radiation. This principle is associated with devices like solar cells or photodiodes and is unrelated to how transformers function.

Conclusion: Based on the analysis, the core principle enabling a transformer to transfer electrical energy from one circuit to another is **mutual electromagnetic induction**.

149. Answer: a

Explanation:

Understanding Fuse Fusing Factor

A fuse is an essential safety device in electrical circuits. Its primary function is to protect equipment and prevent fires by interrupting the circuit when the current exceeds a safe level. This is achieved by melting a specially designed wire (the fuse element) when it gets too hot due to excessive current.

Defining Fusing Factor

The fusing factor is a crucial characteristic of a fuse. It is defined as the ratio of the minimum fusing current to the current rating of the fuse.

In mathematical terms, the fusing factor is given by:

$$\text{Fusing Factor} = \frac{\text{Minimum Fusing Current}}{\text{Current Rating of Fuse}}$$

Let's break down these terms:

- **Current Rating of Fuse:** This is the maximum current that the fuse can carry continuously without melting or deteriorating. It is also known as the rated current or nominal current.
- **Minimum Fusing Current:** This is the minimum current that will cause the fuse element to melt and break the circuit. This current must be sustained for a sufficient duration to raise the temperature of the fuse element to its melting point.

Why Fusing Factor is Always Greater Than 1

For a fuse to perform its protective function correctly, it must be able to carry its rated current continuously without blowing. The minimum fusing current is the **lowest** current value that will cause the fuse to blow. By design, this minimum current must be higher than the continuous current rating. If the minimum fusing current were equal to or less than the rated current, the fuse would blow under normal operating conditions, even when carrying the maximum intended current. Therefore, to ensure the fuse operates reliably and only interrupts the circuit under fault conditions (when current significantly exceeds the rating), the minimum fusing current is always greater than the rated current.

Since the minimum fusing current is always greater than the current rating, the ratio of the minimum fusing current to the current rating will always be greater than 1.

Typical values for the fusing factor for standard fuses are between 1.1 and 1.4. For example, a fuse with a current rating of 10 Amps might have a minimum fusing current of 11.5 Amps. In this case, the fusing factor would be $\frac{11.5 \text{ A}}{10 \text{ A}} = 1.15$, which is greater than 1.

Analyzing the Options

Let's consider the given options in light of the definition and explanation of the fusing factor:

- **Option 1: more than 1**

This aligns with our understanding that the minimum current required to blow the fuse is always greater than its continuous current rating. This ensures the fuse doesn't nuisance trip during normal operation.

- **Option 2: less than 1**

If the fusing factor were less than 1, the minimum current to blow the fuse would be less than its rated current. This means the fuse would blow even when carrying a current below its specified limit, which is not desirable or intended for a safety device.

- **Option 3: infinity**

A fusing factor of infinity would imply the minimum fusing current is infinitely large compared to the rated current, effectively meaning the fuse would never

blow, regardless of the fault current. This contradicts the fundamental purpose of a fuse, which is to interrupt the circuit under overcurrent conditions.

- **Option 4: zero**

A fusing factor of zero is not physically meaningful in this context.

Based on the operational principle and definition, the fusing factor of a fuse is always more than 1.

Revision Table: Fuse Characteristics

Characteristic	Description	Typical Value/Relation
Current Rating (I_{rated})	Max current fuse carries continuously without blowing.	Specified value (e.g., 10A, 20A)
Minimum Fusing Current ($I_{fusing,min}$)	Min current that blows the fuse.	$I_{fusing,min} > I_{rated}$
Fusing Factor	Ratio of Minimum Fusing Current to Current Rating.	Fusing Factor = $\frac{I_{fusing,min}}{I_{rated}}$
Relationship	Ensures fuse blows only under fault current $> I_{rated}$.	Fusing Factor > 1

Additional Information on Fuse Operation

The fusing factor is an important parameter because it indicates how much overcurrent is required to reliably blow the fuse. A lower fusing factor (closer to 1) means the fuse is more sensitive to smaller overcurrents above the rated value. A higher fusing factor means a larger overcurrent is needed to blow the fuse. Different types of fuses have different fusing factors depending on their application (e.g., quick-acting fuses vs. slow-blow fuses).

Factors affecting the minimum fusing current include the material and size of the fuse element, the ambient temperature, and the surrounding environment (e.g.,

whether the fuse is enclosed). However, regardless of these factors, for a properly designed fuse, the minimum current required to cause it to blow will always exceed the current it is designed to carry continuously.

Therefore, the fusing factor, being the ratio of minimum fusing current to rated current, must always be greater than 1.

150. Answer: d

Explanation:

Understanding Bridge Measurements and Temperature Effects

Bridge circuits, such as the Wheatstone bridge, are widely used for precise measurement of electrical quantities like resistance, inductance, and capacitance. These circuits work by balancing the impedances of different arms of the bridge. When the bridge is balanced, there is no voltage difference between the output terminals, and the unknown component value can be determined from the known values of other components in the bridge.

The accuracy of a bridge measurement heavily depends on the stability of the component values in the bridge arms. However, component values can change due to various environmental factors, with temperature being one of the most significant.

Why Temperature Impacts Bridge Measurements

Different materials and components react differently to changes in temperature. This thermal sensitivity means that variations in ambient temperature or temperature gradients across the circuit can cause the values of resistors, inductors, and capacitors to drift from their nominal values. This drift directly affects the balance condition of the bridge, leading to inaccurate measurements.

Temperature Effect on Component Values

- Resistance:** The resistance of most conductive materials changes significantly with temperature. This relationship is often approximated by a linear model for small temperature changes: $R_T = R_{ref}(1 + \alpha(T - T_{ref}))$, where R_T is the resistance at temperature T , R_{ref} is the resistance at a reference temperature T_{ref} , and α is the temperature coefficient of resistance. This coefficient is positive for most metals (resistance increases with temperature) and negative for some semiconductors and insulators. Carbon film resistors typically have negative coefficients, while metal film resistors have lower positive coefficients.
- Capacitance:** The capacitance of a capacitor depends on the dielectric material's permittivity, the plate area, and the distance between plates. Permittivity and physical dimensions can change with temperature, causing capacitance to vary. The temperature coefficient of capacitance varies greatly depending on the dielectric material (e.g., ceramic types, polyester, mica). Some materials have very low temperature coefficients, while others are highly sensitive.
- Inductance:** The inductance of an inductor depends on its geometry, the number of turns, and the permeability of the core material (if any). Temperature changes can cause physical expansion/contraction of the coil and core, altering the geometry. The permeability of magnetic core materials is particularly sensitive to temperature changes, especially near their Curie temperature.

Typical Temperature Sensitivity of Components

Component	Primary Factor Affected by Temperature	General Sensitivity
Resistor	Resistivity of material, physical dimensions	Typically moderate to high (varies by type)
Capacitor	Dielectric permittivity, physical dimensions	Varies widely (some types very stable, others sensitive)
Inductor	Physical dimensions, core material permeability	Varies (especially sensitive with magnetic cores)

Why Resistance is Often the Key Concern

While temperature affects all three parameters (resistance, inductance, and capacitance), resistance is frequently the most sensitive component in many common bridge configurations, such as the Wheatstone bridge, which is specifically designed for resistance measurement. In AC bridges used for impedance measurements (which involve resistance, inductance, and capacitance), changes in resistance due to temperature can significantly alter the magnitude and phase balance conditions.

Furthermore, the temperature coefficients of resistance for common materials used in precision resistors are well-characterized but can still lead to significant errors if temperature is not controlled or accounted for, especially when measuring small changes in resistance (e.g., in strain gauge applications).

Therefore, ensuring a stable and uniform temperature environment for the bridge components is crucial to minimize changes in their values, thereby maintaining the accuracy of the measurement. A difference in temperature between the arms of the bridge will cause a difference in the component values, primarily affecting resistance due to its direct dependence on temperature and the common use of resistance-based bridges.

Conclusion on Temperature Control in Bridge Measurement

Temperature control in bridge measurement is essential because temperature variations cause the electrical properties of the components (resistors, capacitors, inductors) to change. Among these, the resistance of materials is typically quite sensitive to temperature fluctuations. A difference in temperature across the bridge arms directly leads to a difference in the values of the components, significantly impacting the bridge balance and the accuracy of the measurement. Therefore, temperature control is primarily required because a difference in temperature will cause a difference in resistance.

Revision Table: Key Concepts

Concept	Explanation	Relevance to Bridge Measurement
Bridge Circuit	Electrical circuit for precise measurement by balancing impedances.	Foundation for accurate measurements of R, L, C.
Bridge Balance	Condition where bridge output is zero due to balanced arm impedances.	Indicates the measured value. Sensitivity depends on component stability.
Temperature Coefficient	Measure of how a component's value changes per degree of temperature change.	Determines the magnitude of temperature-induced error.
Temperature Control	Maintaining a stable and uniform temperature.	Minimizes component value changes and improves measurement accuracy.

Additional Information: Mitigating Temperature Effects

Beyond strict temperature control, several techniques are used to minimize temperature-induced errors in bridge measurements:

- **Using Components with Low Temperature Coefficients:** Selecting precision resistors, capacitors, and inductors made from materials with very low temperature coefficients (TC). For example, Manganin or Constantan alloys are used for standard resistors due to their low TCR .
- **Matching Components:** Using components with similar temperature coefficients in corresponding arms of the bridge so that temperature changes affect them symmetrically, maintaining balance.
- **Temperature Compensation:** Incorporating additional components (like thermistors or sensing resistors) whose temperature characteristics are used to counteract the temperature-induced changes in the main measurement components.
- **Shielding and Insulation:** Protecting the bridge circuit from thermal radiation and air currents.
- **Controlled Environment:** Performing measurements in temperature-controlled rooms or using temperature chambers.

These techniques highlight the critical importance of addressing temperature variations to achieve high accuracy in bridge-based measurements.

Your Personal Exams Guide

151. Answer: b

Explanation:

Understanding Power Supplies for Critical Systems

Life support systems are essential medical equipment that helps a patient's body function, such as breathing or maintaining circulation. Any interruption in power to these systems, even for a moment, can be life-threatening. Therefore, the power supply for a life support system must be absolutely reliable and continuous.

Analyzing the Power Supply Options

Let's look at the types of power supply options provided:

- **PDMS:** This acronym is not a standard term for a type of power supply in this context. It might refer to a specific system or be irrelevant to the question about mandatory power supply types for life support.
- **UPS:** This stands for Uninterruptible Power Supply. A UPS is designed to provide emergency power to a load when the input power source, typically the main power, fails. A UPS is crucial for critical systems like computers, data centers, or, in this case, life support equipment, because it provides near-instantaneous protection from power interruptions.
- **Stabilizer:** A voltage stabilizer is a device that maintains a constant output voltage despite variations in the input voltage from the main power supply. While important for protecting equipment from voltage fluctuations, a stabilizer does not provide backup power if the main power source fails completely.
- **RPS:** This could stand for Redundant Power Supply. A redundant power supply system uses multiple power supplies so that if one fails, another can take over. While redundancy is important for reliability, a typical RPS setup alone might still have a momentary switchover time or relies on separate power sources being available. A UPS specifically focuses on bridging the gap during an outage with its own stored energy (batteries). In the context of needing *uninterrupted* power, UPS is the more direct and mandatory requirement for life support.

Why UPS is Mandatory for Life Support Systems

Life support systems require a power source that is not only stable but also continuous, with zero tolerance for downtime. The main power grid can experience interruptions, sags, surges, or complete outages.

- A **Stabilizer** protects against voltage fluctuations but offers no power during an outage.
- While **RPS** adds redundancy, a **UPS** provides the essential function of supplying power from its internal battery the instant the primary power fails, ensuring that the life support system never loses power. This seamless transition is critical.

- A UPS also often provides power conditioning, protecting the equipment from surges and noise on the power line, similar to a stabilizer but with the added backup capability.

Therefore, an **Uninterruptible Power Supply (UPS)** is mandatory for life support systems because it provides the necessary continuous and stable power, bridging any gap caused by main power failures.

Conclusion on Life Support System Power Supply

Based on the analysis of the options and the critical requirement for uninterrupted power for life support systems, the essential and mandatory type of power supply among the given choices is a UPS.

Power Supply Type	Key Function	Provides Backup Power?	Suitable for Life Support?
PDMS	(Unclear/Non-standard)	?	No (Unidentified)
UPS (Uninterruptible Power Supply)	Provides continuous power during outages; conditions power	Yes (from battery)	Yes (Mandatory)
Stabilizer	Regulates voltage fluctuations	No	No (Insufficient)
RPS (Redundant Power Supply)	Provides alternative power source if primary fails	Depends (May have switchover delay)	No (UPS is specifically for *uninterrupted* transition)

Revision Table: Key Concepts

Term	Definition	Relevance to Life Support
Life Support System	Medical equipment sustaining vital bodily functions (e.g., ventilator, heart-lung machine).	Requires extremely reliable, continuous power.
Uninterruptible Power Supply (UPS)	Provides immediate backup power from batteries when main power fails.	Ensures zero downtime for critical equipment.
Voltage Stabilizer	Corrects voltage fluctuations from the power grid.	Protects equipment but doesn't handle power outages.
Redundancy	Having backup components or systems to prevent failure.	Important for reliability, but doesn't guarantee *uninterrupted* power transition like a UPS.

Additional Information on Critical Power for Medical Devices

Medical facilities often implement multi-layered power protection strategies for critical areas like intensive care units (ICUs) or operating rooms where life support systems are used. These strategies include:

- **Dedicated Circuits:** Critical equipment is often on circuits separate from non-essential loads.
- **Hospital Grade Receptacles:** These are designed for enhanced durability and reliability.
- **Generator Backup:** Hospitals typically have large generators that start automatically during a widespread power outage. However, generators take a short time to start up and stabilize.

- **UPS Systems:** Individual life support devices or groups of devices are connected to UPS systems. The UPS batteries provide power instantly during the brief period between main power failure and the generator starting up, or for shorter outages where the generator isn't needed. UPS systems are the first line of defense against any interruption.
- **Regular Testing and Maintenance:** Power systems, including UPS batteries and generators, require frequent testing and maintenance to ensure they function correctly when needed.

The combination of these measures ensures the highest level of power reliability for critical medical applications, with the UPS being the essential component for achieving zero interruption at the device level.

152. Answer: d

Explanation:

Understanding Transformer Voltage Transformation

Transformers are essential electrical devices used to change the voltage of an alternating current (AC) supply. They work based on the principle of electromagnetic induction. A transformer typically consists of two coils, the primary coil and the secondary coil, wound around a common iron core. The voltage induced in the secondary coil depends on the ratio of the number of turns in the secondary coil compared to the number of turns in the primary coil.

The relationship between the primary voltage (V_p), secondary voltage (V_s), number of turns in the primary coil (N_p), and number of turns in the secondary coil (N_s) is given by the transformer equation:

$$\frac{V_s}{V_p} = \frac{N_s}{N_p}$$

This ratio $\frac{N_s}{N_p}$ is called the turns ratio.

Identifying the Transformer for Low Voltage Output

The question asks which transformer will give an output a low voltage compared to the input voltage. Let's analyze the different types of transformers mentioned in the options:

- **Step-up Transformer:** In a step-up transformer, the number of turns in the secondary coil (N_s) is greater than the number of turns in the primary coil (N_p). According to the transformer equation, if $N_s > N_p$, then $V_s > V_p$. This means a step-up transformer increases the voltage.
- **Step-down Transformer:** In a step-down transformer, the number of turns in the secondary coil (N_s) is less than the number of turns in the primary coil (N_p). According to the transformer equation, if $N_s < N_p$, then $V_s < V_p$. This means a step-down transformer decreases or lowers the voltage. This is exactly what is described in the question – giving output a low voltage compared to input.
- **Current Transformer:** A current transformer (CT) is designed to measure alternating current. It is connected in series with the circuit where the current is to be measured. It steps down a high current to a lower value that can be safely measured by standard instruments. While the turns ratio for a CT (high primary turns, low secondary turns relative to a voltage transformer) results in stepping down current, it implicitly steps up voltage (since Power $\approx$ Voltage $\times$ Current is conserved, ignoring losses). Its primary purpose is current measurement, not voltage reduction for power distribution.
- **Potential Transformer:** A potential transformer (PT), also known as a voltage transformer (VT), is used to measure high alternating voltages. It is connected in parallel with the circuit. A PT is essentially a step-down transformer designed for accurate voltage measurement, stepping down a high voltage to a lower value that can be safely measured by standard voltmeters. While a PT does step down voltage, the term "Step-down transformer" is the general classification for a transformer designed to lower voltage for various applications, including power distribution and use in electronic devices. The question asks which transformer "will give output a low voltage compared to input" in a general context, which is the fundamental function of a step-down transformer.

Based on the analysis, the transformer specifically designed to produce a lower voltage output compared to its input voltage is the step-down transformer.

Comparison of Transformer Types based on Voltage Transformation

Transformer Type	Turns Ratio (N_s/N_p)	Voltage Relationship (V_s vs. V_p)	Primary Function
Step-up Transformer	$N_s > N_p$	$V_s > V_p$ (Voltage Increases)	Increase voltage
Step-down Transformer	$N_s < N_p$	$V_s < V_p$ (Voltage Decreases)	Decrease voltage
Current Transformer (CT)	Typically $N_p \ll N_s$ (for current step-down)	$V_s > V_p$ (Voltage steps up implicitly)	Measure high AC current (steps down current)
Potential Transformer (PT)	$N_s < N_p$	$V_s < V_p$ (Voltage Decreases)	Measure high AC voltage (steps down voltage accurately)

Therefore, the transformer that primarily functions to give an output a low voltage compared to its input is the step-down transformer.

Revision Table: Key Transformer Concepts

Concept	Description
Transformer	Device that transfers electrical energy between circuits by electromagnetic induction, changing voltage/current levels.
Primary Coil	The coil connected to the input AC voltage source.
Secondary Coil	The coil from which the output voltage is taken.
Turns Ratio (N_s/N_p)	Ratio of the number of turns in the secondary coil to the number of turns in the primary coil. Determines voltage transformation.

Additional Information on Transformer Applications

Step-down transformers are widely used in many applications, including:

- **Power distribution:** Stepping down high transmission voltages to lower voltages suitable for residential and industrial use.
- **Electronic devices:** Providing low voltages required by integrated circuits and other components (e.g., chargers for phones, laptops).
- **Power supplies:** As part of AC-to-DC power supplies to reduce AC voltage before rectification.
- **Safety:** Isolating equipment from the main power supply and providing low voltage for safe operation.

Understanding the function of different transformer types is crucial in electrical engineering and electronics.

153. Answer: d

Explanation:

Understanding the Three-Point Starter for DC Motors

DC motors, especially larger ones, have very low armature resistance. When starting, the back EMF (electromotive force) is zero because the motor is stationary. This means that if the full supply voltage is applied directly to the armature, a very large current will flow ($I_a = V/R_a$), potentially damaging the armature winding or the power supply. A starter is used to limit this high starting current by adding external resistance in series with the armature during the starting period. As the motor speeds up, back EMF (E_b) is generated, which opposes the supply voltage (V). The starting resistance is gradually cut out as the speed increases, until the motor runs at its normal speed with the external resistance completely removed.

The three-point starter is a common type of DC motor starter. It gets its name from having three terminals:

- **L (Line):** Connected to the positive terminal of the DC supply.
- **A (Armature):** Connected to the motor armature terminal.
- **F (Field):** Connected to the motor field winding terminal.

Let's look at the components and working principle of a three-point starter:

- **Starting Resistances:** These are resistances connected in series with the armature circuit at the start. They are gradually removed as the handle moves from the 'START' position to the 'RUN' position.
- **Handle:** A movable arm that is manually operated to connect the armature to the starting resistances and then gradually cut them out.
- **Hold-On Coil (No-Volt Release, NVR):** This is an electromagnet connected in series with the field winding. When the handle is moved to the 'RUN' position, it is held there by the magnetic force of this coil. If the supply voltage fails or becomes too low, the coil loses its magnetism, and a spring pulls the handle back to the 'OFF' position, protecting the motor from restarting unexpectedly when the voltage returns.
- **Overload Release Coil (OLR):** This is an electromagnet connected in series with the armature circuit. If the motor draws excessive current (due to overload), the magnetic field of this coil becomes strong enough to attract a small arm, which then short-circuits the hold-on coil. This causes the hold-on coil to lose its magnetism, releasing the handle to the 'OFF' position, thus protecting the motor from overload damage.

Why Three-Point Starter is Suitable for Shunt and Compound Motors

The hold-on coil in a three-point starter is connected in series with the motor's field winding and then across the supply (via the handle in the 'RUN' position). This arrangement works well for shunt and compound motors because:

- **Shunt Motor:** The field winding in a shunt motor is connected in parallel with the armature and across the supply. The hold-on coil, being in series with the shunt field, receives current whenever the field circuit is energized. As long as the field is getting sufficient current (and thus voltage is healthy), the hold-on coil remains energized, holding the handle.
- **Compound Motor:** A compound motor has both shunt and series field windings. The three-point starter's hold-on coil is connected in series with the shunt field winding. Similar to the shunt motor, the hold-on coil's operation depends on the current flowing through the shunt field, making the starter suitable.

Why Three-Point Starter is NOT Suitable for Series Motors

A series motor has its field winding connected in series with the armature. If a three-point starter were used:

- The hold-on coil would be connected in series with the series field winding (and thus the armature circuit).
- When the motor is started, the armature current is high. As the motor speeds up and the starting resistance is cut out, the armature current (and field current) reduces.
- If the motor load decreases significantly, the armature current (and series field current) also decreases.
- This reduced current flowing through the hold-on coil (which is in series with the field) would weaken its magnetic field. If the load becomes very light, the current might drop so low that the hold-on coil releases the handle, tripping the motor unnecessarily.

This characteristic makes the three-point starter unreliable for series motors, especially under light load conditions where series motor speed can become

dangerously high.

Analyzing the Options

Based on the suitability of the three-point starter:

- **Option 1: Shunt motor** – A three-point starter is suitable for shunt motors. However, the option might not be complete if it's also suitable for other types.
- **Option 2: Shunt and series motor** – The three-point starter is suitable for shunt motors but not for series motors. This option is incorrect.
- **Option 3: Compound and series motor** – The three-point starter is suitable for compound motors but not for series motors. This option is incorrect.
- **Option 4: Shunt and compound motor** – The three-point starter is suitable for both shunt motors and compound motors because its hold-on coil is connected in series with the shunt field winding, which is present in both types. This option correctly identifies the types of motors suitable for a three-point starter.

Therefore, a three-point starter is suitable for shunt and compound motors.

Motor Type	Suitable for Three-Point Starter?	Reason
Shunt Motor	Yes	Hold-on coil in series with stable shunt field current.
Compound Motor	Yes	Hold-on coil in series with stable shunt field current.
Series Motor	No	Hold-on coil in series with variable armature/series field current; may trip on light load.

Revision Table: Three-Point Starter Suitability

Starter Type	Suitable DC Motor Types	Reason for Suitability/Limitation
Two-Point Starter	Series Motor	Simpler design, mainly current limiting, no NVR/OLR typically.
Three-Point Starter	Shunt Motor, Compound Motor	Hold-on coil in series with shunt field, provides no-volt and overload protection reliably for these types.
Four-Point Starter	Shunt Motor, Compound Motor, Series Motor	Hold-on coil connected directly across the supply, independent of field/armature current variations. Offers reliable protection for all types.

Additional Information on DC Motor Starters

Besides the three-point starter, other types of starters for DC motors include two-point and four-point starters.

- **Two-Point Starter:** Primarily used for starting DC series motors. It has two terminals (Line and Armature). It mainly provides starting resistance but often lacks protective features like no-volt or overload release.
- **Four-Point Starter:** This is a more advanced starter suitable for all types of DC motors (shunt, series, and compound). It has four terminals: L (Line), A (Armature), F (Field), and Z (connected to the supply side of the hold-on coil). The key difference from the three-point starter is that the hold-on coil in a four-point starter is connected independently across the supply voltage (in series with a protecting resistance 'R'). This means the current through the hold-on coil is only dependent on the supply voltage, not on the field current (as in a three-point starter) or armature current. This prevents the starter from tripping due to changes in field current (e.g., field rheostat adjustments in shunt motors) or low load current (in series motors). It provides reliable no-volt and overload protection regardless of the motor type or operating conditions.

154. Answer: c

Explanation:

Understanding Lightning Arresters in Power Systems

A lightning arrester is a protective device used in electrical power systems to protect equipment from damage due to lightning strikes or other high-voltage surges. It is typically connected between the electrical conductor (line) and the ground (earth).

Purpose of a Lightning Arrester

The primary purpose of a lightning arrester is to provide a path for surge current to flow to the ground, thereby diverting it away from sensitive equipment. When a high voltage surge, such as one caused by a lightning strike or switching operations, reaches the arrester, the arrester's impedance drops significantly, allowing the surge current to pass through it to the earth. Once the surge voltage returns to normal levels, the arrester's impedance increases again, blocking the normal power frequency current from flowing to the ground.

Analysing the Options

Let's examine each option in the context of a lightning arrester's function:

- **Option 1: Protects the transmission line against lightning stroke**
While lightning arresters are part of the overall protection scheme related to lightning, their main function is not to protect the transmission line itself from being struck or damaged along its entire length. Line protection against direct strokes often involves shielding wires and grounding towers. The arrester protects connected equipment from the surge that travels along the line after a stroke or other event.
- **Option 2: High frequency oscillations are suppressed in the line**
Lightning arresters are designed to handle very high voltage, short-duration surges. While surges can sometimes include high-frequency components, the

primary role of an arrester is not to suppress continuous or sustained high-frequency oscillations that might occur under different system conditions. Other devices like surge capacitors or filters are more relevant for damping oscillations.

- **Option 3: The terminal equipment is protected against travelling surges**
This option accurately describes the main function. Lightning strikes or switching operations can generate high-voltage transients (travelling surges) that propagate along the transmission line. When these surges reach the ends of the line where expensive terminal equipment (like transformers, switchgear, circuit breakers, generators) is located, they can cause significant insulation failure and damage. The lightning arrester, installed near this equipment, diverts the surge current to earth, limiting the voltage across the equipment terminals to a safe level.
- **Option 4: Reflects the travelling wave approaching it**
While changes in impedance can cause wave reflection, a lightning arrester's operation is based on changing its impedance drastically to *divert* the surge current to ground when the voltage exceeds a certain threshold. It doesn't primarily work by reflecting the surge back onto the line. Reflection would still leave the surge energy on the line, potentially affecting other components.

Based on the analysis, the most accurate description of the purpose of a lightning arrester connected between the line and earth is to protect the terminal equipment from travelling surges.

Lightning Arrester Function Summary

Connection	Purpose	Protected Element	Mechanism
Between Line and Earth	Protect equipment from surges	Terminal Equipment (Transformers, Switchgear, etc.)	Diverts surge current to earth

Revision Table: Key Concepts

Power System Protection Devices

Device	Primary Function
Lightning Arrester	Protects equipment from high-voltage surges (lightning/switching) by diverting current to ground.
Circuit Breaker	Interrupts fault currents to isolate faulty sections.
Relay	Detects fault conditions and initiates action (e.g., tripping a circuit breaker).
Surge Capacitor	Used to grade voltage distribution and protect equipment from steep-fronted waves.

Additional Information on Travelling Surges and Equipment Protection

Travelling surges are transient overvoltages that propagate along transmission lines. They can be caused by atmospheric disturbances (lightning strikes) or switching operations (e.g., switching of circuit breakers, fault initiation, or clearing). These surges can have very high magnitudes and steep wavefronts. When a surge reaches the end of a line or passes through points with impedance discontinuities (like equipment terminals), it can cause significant voltage stress. The insulation of power system equipment is designed to withstand the normal operating voltage and a certain level of power frequency overvoltage, but high-magnitude, short-duration surge voltages can exceed the equipment's insulation strength, leading to flashover or breakdown. Lightning arresters are strategically placed near valuable equipment to limit these surge voltages and protect the equipment's insulation.

155. Answer: c

Explanation:

Understanding Electrical Joints for Overhead Lines and Service Connections

When electrical energy needs to be supplied from an overhead distribution line to individual consumers or properties, a connection point is created. This process is often referred to as tapping the main line. A specialized type of electrical joint is required to securely and safely connect the service wire (leading to the consumer) to the main overhead line without interrupting the supply to other consumers or compromising the integrity of the main line.

Analyzing Common Electrical Joints

Let's examine the types of joints mentioned in the options to determine which one is appropriate for tapping an overhead line for a service connection.

- **Britannia joint:** This joint is primarily used for joining two wires end-to-end, especially in overhead lines where high tensile strength is required. It's not designed for branching off a connection from a continuous line.
- **Scarfed joint:** This is a woodworking joint used to connect two pieces of wood lengthwise. It is not used in electrical wiring.
- **Tee joint:** A Tee joint is specifically designed to connect a branch wire to a continuous run of a main wire, typically at a right angle. This is exactly the configuration needed to tap a service wire from an overhead main distribution line. The joint creates a 'T' shape, hence the name.
- **Rat tail joint:** Also known as a pig tail joint, this is a simple joint used for connecting two or more conductors together by twisting their ends. It is commonly used inside junction boxes or conduits but is generally not suitable for exposed overhead lines where mechanical strength and weather resistance are crucial, nor is it primarily a tapping joint for heavy-duty service connections.

Why the Tee Joint is Used for Service Tapping

The Tee joint is the standard and most suitable type of electrical joint for tapping electrical energy from an overhead line for a service connection because:

- It allows a branch conductor (the service wire) to be securely connected to a continuous main conductor (the overhead line).
- Various methods and connectors exist for making Tee joints on live or dead overhead lines, ensuring safety and reliability. Examples include bolted connectors, compression connectors, or specific types of insulation-piercing connectors designed for overhead service taps.
- It provides the necessary mechanical strength and electrical continuity for the branch connection.

Therefore, the joint type specifically used for this purpose in overhead lines is the Tee joint.

Conclusion

Based on the analysis of the different joint types and their applications, the Tee joint is the appropriate choice for tapping electrical energy from an overhead line for a service connection.

Revision Table: Electrical Joints and Uses

Joint Type	Primary Use	Suitable for Service Tapping?
Britannia joint	Joining conductors end-to-end (high strength)	No
Scarfed joint	Woodworking (not electrical)	No
Tee joint	Connecting a branch to a main conductor	Yes
Rat tail joint	Joining multiple conductors (usually in boxes)	No (not for overhead tapping)

Additional Information: Overhead Line Tapping

Tapping into overhead lines requires specialized techniques and connectors to ensure safety, electrical integrity, and mechanical strength. Common connectors used for creating Tee joints for service taps include:

- **Split-bolt connectors:** A bolted connector used to join two conductors, often used for tapping.
- **Compression connectors:** Require special tools to crimp the connector onto the conductors, creating a very secure and low-resistance joint.
- **Insulation Piercing Connectors (IPCs):** These connectors pierce the insulation of both the main and branch conductors to establish electrical contact without stripping, often used on insulated overhead cables.

The specific type of connector and method used for making the Tee joint depends on the type of conductor (bare or insulated), the voltage level, and local utility standards.

156. Answer: d

Explanation:

Understanding Flux per Square Unit Area

The question asks to identify the term that describes the magnetic flux passing through a square unit area. Let's break down what this means and examine the given options.

What is Magnetic Flux?

Magnetic flux (Φ) is a measure of the total magnetic field lines passing through a given area. It's often visualized as the number of magnetic field lines piercing a surface. The unit of magnetic flux is the Weber (Wb).

What is Flux per Square Unit Area?

When we talk about "flux per square unit area", we are essentially talking about how concentrated the magnetic field lines are in a specific region. This is a measure of the density of the magnetic flux.

Analyzing the Options

Let's look at the options provided:

1. **Magnetic field strength:** Also known as magnetic intensity (H). It is a measure of the magnetizing force produced by electric currents. Its unit is Amperes per meter (A/m). It is related to magnetic flux density (B) but is not the same as flux per unit area.
2. **Magnetic intensity:** This is another term for magnetic field strength (H), as explained above. It does not represent flux per unit area directly.
3. **Magnetic reluctance:** This is the opposition that a magnetic circuit offers to magnetic flux. It is analogous to resistance in an electric circuit. Its unit is the inverse Henry (H^{-1}). It is not flux per unit area.
4. **Magnetic flux density:** This is defined as the amount of magnetic flux (Φ) passing perpendicularly through a unit area (A). Mathematically, it is given by $B = \frac{\Phi}{A}$. The unit of magnetic flux density is the Tesla (T), which is equivalent to Webers per square meter (Wb/m^2).

Connecting Definition to the Question

The term that precisely matches the description "flux per square unit area" is Magnetic flux density. It quantifies how much magnetic flux is packed into a specific area, hence its name "density".

The formula for magnetic flux density reinforces this definition:

$$B = \frac{\Phi}{A}$$

Where:

- B is the magnetic flux density
- Φ is the magnetic flux
- A is the area

Conclusion

Based on the definitions, the flux per square unit area is called Magnetic flux density.

Revision Table: Key Magnetic Concepts

Term	Symbol	Definition	Unit
Magnetic Flux	Φ	Total magnetic field lines through an area	Weber (Wb)
Magnetic Flux Density	B	Magnetic flux per unit area (Φ/A)	Tesla (T) or Wb/m^2
Magnetic Field Strength (Intensity)	H	Magnetizing force per unit length	Ampere per meter (A/m)
Magnetic Reluctance	R	Opposition to magnetic flux in a circuit	Inverse Henry (H^{-1})

Additional Information: Magnetic Field Relationships

Magnetic flux density (B) and magnetic field strength (H) are related in a material by the permeability of the material (μ). The relationship is given by:

$$B = \mu H$$

Permeability (μ) is a property of the material that indicates how easily magnetic field lines can pass through it. It can also be expressed as $\mu = \mu_0 \mu_r$, where μ_0 is the permeability of free space (a vacuum) and μ_r is the relative permeability of the material.

Understanding the distinction between magnetic flux (Φ), magnetic flux density (B), and magnetic field strength (H) is crucial in studying electromagnetism.

157. Answer: d

Explanation:

Understanding Induction Motor Speed and Load

An induction motor is a type of AC electric motor that uses electromagnetic induction from the magnetic field of the stator winding to produce electric current in the rotor, and in turn produce torque. The speed of an induction motor is closely related to its synchronous speed and the slip.

Synchronous speed (N_s) is the speed at which the magnetic field rotates in the stator. It is determined by the frequency of the power supply (f) and the number of poles (P) in the motor winding:

$$N_s = \frac{120 \times f}{P}$$

where N_s is in revolutions per minute (RPM), f is in Hertz (Hz), and P is the number of poles.

The rotor of an induction motor never rotates exactly at synchronous speed. There is always a slight difference in speed, known as slip. Slip (s) is the relative speed difference between the synchronous speed and the rotor speed (N) expressed as a fraction or percentage of synchronous speed:

$$s = \frac{N_s - N}{N_s}$$

The actual rotor speed is then given by:

$$N = N_s(1 - s)$$

How Load Affects Induction Motor Speed

When a mechanical load is applied to the shaft of an induction motor, the motor needs to produce torque to overcome this load and maintain rotation. To generate this torque, the rotor current and the interaction with the stator field are required.

The rotor current is induced by the difference in speed between the rotating magnetic field (synchronous speed) and the rotor speed. This speed difference is the slip.

- When the load on the motor increases, the motor needs to produce more torque.
- To produce more torque, the induced voltage and current in the rotor must increase.
- The induced voltage and current in the rotor are proportional to the slip speed ($N_s - N$).
- Therefore, to increase the rotor voltage and current, the slip speed must increase.
- An increase in slip speed ($N_s - N$) means that the rotor speed (N) must decrease, since N_s is constant for a given frequency and number of poles.

Thus, as the mechanical load on an induction motor increases, its speed decreases. This is a fundamental characteristic of induction motors operating under normal conditions.

Analysis of Other Factors

Let's briefly consider the other options:

- **Magnetic flux:** The main magnetic flux in an induction motor is primarily determined by the applied voltage and frequency. While flux affects torque production, an increase in flux itself (e.g., due to over-excitation, though not typical in standard operation without voltage/frequency changes) would likely increase torque capability or potentially cause saturation, but doesn't directly cause the speed to decrease *with the increase* in flux in the same way load does. Speed is primarily governed by slip for a given synchronous speed.
- **Resistance:** Rotor resistance affects the torque-slip characteristic curve. Increasing rotor resistance (e.g., in slip-ring induction motors by adding external resistance) shifts the maximum torque point to higher slip values (lower speeds) and increases the starting torque. So, for a *given torque or load*, higher rotor resistance results in lower speed (higher slip). However, the question asks what causes the speed to decrease *with the increase in* the

factor. While increasing resistance *causes* lower speed for a given load, it's not the resistance value *itself* increasing during operation in response to a condition like load. Load is the varying factor that directly causes the speed to drop.

- **Capacitance:** Capacitance is typically associated with single-phase motor starting or power factor correction for three-phase motors. It doesn't directly cause the speed of a running three-phase induction motor to decrease as the capacitance value increases during normal operation under load.

Based on the fundamental operating principles, the mechanical load applied to the shaft is the primary factor that directly causes the speed of a running induction motor to decrease as the load increases.

Factor	Effect on Induction Motor Speed
Load	Increases load → Increases slip → Decreases speed
Magnetic Flux	Primarily affects torque; speed related to slip, not direct decrease with flux increase in operation
Resistance (Rotor)	Higher R → Higher slip for same torque → Lower speed, but load is the operational variable causing speed change
Capacitance	Not a direct factor causing speed decrease with increase in normal 3-phase operation

Conclusion

The speed of an induction motor decreases with the increase in mechanical load applied to its shaft. This relationship is governed by the need for increased slip to generate the torque required by the higher load.

Revision Table: Induction Motor Concepts

Term	Definition/Relation to Speed
Induction Motor	AC motor using induced rotor current for torque.
Synchronous Speed (N_s)	Speed of stator magnetic field; fixed by frequency and poles.
Rotor Speed (N)	Actual speed of the motor shaft; always $< N_s$ in normal operation.
Slip (s)	Relative speed difference between N_s and N ; needed to induce rotor current/torque. Increases with load.
Load	Mechanical resistance/torque applied to the shaft. Higher load requires more torque, leading to higher slip and lower speed.

Additional Information: Factors Affecting Induction Motor Operation

Besides the direct effect of load on speed, other factors influence the overall performance and speed characteristics of an induction motor:

- **Supply Voltage and Frequency:** These determine the synchronous speed and also significantly impact the motor's torque capability. Varying voltage or frequency is a common method for speed control.
- **Stator and Rotor Resistance/Reactance:** These parameters are inherent to the motor design and affect the torque-slip curve shape, including starting torque and maximum torque.
- **Motor Design:** The class of the induction motor (e.g., Class A, B, C, D) defines its starting torque, pull-out torque, and slip characteristics, influencing how steeply the speed drops with increasing load.
- **Temperature:** Winding resistance increases with temperature, which can slightly alter the speed-torque characteristic.

Understanding the relationship between load, slip, and speed is crucial for selecting and operating induction motors correctly for various applications. An increase in mechanical load directly necessitates a decrease in speed to increase slip and generate the required torque.

158. Answer: a

Explanation:

Understanding Decimal to Binary Conversion

Number systems are fundamental in mathematics and computer science. The decimal system (base 10) is what we use daily, using digits from 0 to 9. The binary system (base 2) is crucial in computing, using only two digits: 0 and 1.

Converting a number from the decimal system to the binary system involves repeatedly dividing the decimal number by 2 and recording the remainders. The binary equivalent is then formed by reading the remainders from the last one obtained to the first one obtained (bottom-up).

Converting Decimal 5 to Binary Equivalent

Let's convert the decimal number 5 to its binary equivalent using the division-by-2 method:

1. Divide the decimal number 5 by 2. Note down the quotient and the remainder.
2. Take the quotient from the previous step and divide it by 2. Repeat this process until the quotient becomes 0.
3. The binary equivalent is obtained by listing the remainders in reverse order (from the last remainder to the first remainder).

Step-by-Step Conversion of 5 to Binary

Here are the steps for converting decimal 5 to binary:

Division	Quotient	Remainder
$5 \div 2$	\$2	\$1
$2 \div 2$	\$1	\$0
$1 \div 2$	\$0	\$1

Reading the remainders from bottom to top, we get 101.

Binary Equivalent Result

Based on the conversion steps, the binary equivalent of the decimal number 5 is 101.

Let's compare this result with the given options:

- Option 1: 101
- Option 2: 100
- Option 3: 001
- Option 4: 111

The calculated binary equivalent, 101, matches Option 1.

Revision Table: Common Decimal to Binary Conversions

Decimal Number	Binary Equivalent
0	0
1	1
2	10
3	11
4	100
5	101
6	110
7	111
8	1000

Additional Information on Number Systems and Binary Representation

Number systems are ways to represent numbers using specific symbols or digits. The value of a digit depends on its position within the number and the base of the number system.

- **Decimal System (Base 10):** Uses digits 0–9. The position values are powers of 10 (e.g., $10^0, 10^1, 10^2, \dots$). For example, decimal $123 = 1 \times 10^2 + 2 \times 10^1 + 3 \times 10^0$.
- **Binary System (Base 2):** Uses digits 0 and 1. The position values are powers of 2 (e.g., $2^0, 2^1, 2^2, \dots$). For example, binary $101_2 = 1 \times 2^2 + 0 \times 2^1 + 1 \times 2^0 = 1 \times 4 + 0 \times 2 + 1 \times 1 = 4 + 0 + 1 = 5$ in decimal.
- Binary representation is fundamental to how computers store and process information because electronic circuits can easily represent two states (on/off, high voltage/low voltage), which correspond to 1 and 0.
- Understanding decimal to binary conversion is a core concept in digital electronics and computer architecture.

159. Answer: b

Explanation:

Let's analyze the question about connecting bulbs in a series circuit and calculate the number of bulbs needed.

The question asks how many bulbs, each rated at 2 V, should be connected in series to operate safely on a 110 V AC line. In a series circuit, electrical components are connected end-to-end, forming a single path for the current. A key characteristic of a series circuit is that the total voltage across the entire circuit is the sum of the voltages across each individual component.

Understanding Series Circuits and Voltage

In a series connection:

- The same current flows through every component.
- The total resistance is the sum of individual resistances ($R_{\text{total}} = R_1 + R_2 + \dots$).
- The total voltage supplied by the source is divided among the components ($V_{\text{total}} = V_1 + V_2 + \dots$).

In this problem, we have multiple bulbs connected in series across a 110 V AC supply. Each bulb is rated for a maximum voltage of 2 V. To ensure that each bulb operates at its rated voltage (or less, although here we assume they need to sum up to the total voltage for proper brightness), the sum of the voltage drops across all the bulbs must equal the total supply voltage.

Calculating the Number of Bulbs in Series

We are given:

- Voltage rating of each bulb ($V_{\text{bulb}} = 2 \text{ V}$)
- Total supply voltage ($V_{\text{total}} = 110 \text{ V}$)

Let 'n' be the number of bulbs connected in series.

In a series circuit, the total voltage is the sum of the voltages across each identical component. So, if there are 'n' bulbs each with a voltage drop of V_{bulb} , the total voltage V_{total} is given by:

$$V_{\text{total}} = n \times V_{\text{bulb}}$$

We can rearrange this formula to find the number of bulbs, 'n':

$$n = \frac{V_{\text{total}}}{V_{\text{bulb}}}$$

Now, let's substitute the given values into the formula:

$$n = \frac{110 \text{ V}}{2 \text{ V}}$$

$$n = 55$$

Therefore, 55 bulbs, each rated at 2 V, need to be connected in series to be safely run on a 110 V AC line.

Summary of Calculation

Here is a quick summary of the calculation:

Parameter	Value
Voltage of each bulb (V_{bulb})	2 V
Total supply voltage (V_{total})	110 V
Formula for Number of bulbs (n)	$n = \frac{V_{\text{total}}}{V_{\text{bulb}}}$
Calculation	$n = \frac{110}{2}$
Result (Number of bulbs)	55

This calculation shows that connecting 55 bulbs in series allows the total 110 V supply voltage to be evenly distributed among them, with each bulb receiving approximately 2 V, which is within its rated voltage.

Revision Table: Series Circuit Concepts

Concept	Description
Series Circuit	Components connected in a single path.
Voltage Division	Total voltage is divided among components in series.
Current in Series	Current is the same through all components.
Total Resistance	Sum of individual resistances.

Additional Information: AC Voltage

The question mentions a 110 V AC line. AC stands for Alternating Current. The voltage in an AC circuit constantly changes direction and magnitude over time. The 110 V value typically refers to the RMS (Root Mean Square) voltage. For simple calculations involving series voltage division with resistive loads like incandescent bulbs, treating the RMS voltage value directly in the calculation is appropriate, as the ratio of voltages across components remains constant over time, regardless of the instantaneous AC voltage value.

Your Personal Exams Guide

160. Answer: c

Explanation:

Understanding the Difference Between Four-Point and Three-Point DC Motor Starters

DC motor starters are essential devices used to limit the high starting current that flows when a DC motor is first connected to the supply. They achieve this by introducing resistance in series with the armature winding during startup. As the

motor speed increases and back EMF builds up, this resistance is gradually reduced and finally cut out.

Both three-point and four-point starters perform this primary function. However, they differ in a fundamental aspect related to the connection of the holding coil, also known as the No-Volt Coil (NVC) or Low-Volt Release (LVR).

Comparing Three-Point and Four-Point Starter Designs

Let's look at the key components and how they are arranged:

- **Starting Resistance:** Both starters use a series of resistances connected to studs on the starter panel. These resistances are manually or automatically cut out as the motor accelerates.
- **Holding Coil (NVC/LVR):** This coil is crucial for safety. It holds the starter handle in the "RUN" position under normal operating conditions. If the supply voltage drops significantly (no-volt condition) or fails, the coil loses its magnetism, and a spring pulls the handle back to the "OFF" position, disconnecting the motor from the supply. This prevents the motor from restarting unexpectedly when the voltage returns, which could cause high current and damage.
- **Overcurrent Release (OCR):** This protective mechanism trips the holding coil circuit if the motor draws excessive current (overload), causing the handle to return to "OFF". It protects the motor from damage due to overloads.

The Basic Difference: Holding Coil Connection

The significant difference lies in how the holding coil is connected:

- **Three-Point Starter:** In a three-point starter, the holding coil is connected in **series with the shunt field winding** and then across the line. This means the current flowing through the holding coil is the shunt field current.
- **Four-Point Starter:** In a four-point starter, the holding coil is connected **independently across the supply lines** through a protective resistance. It is no longer in series with the shunt field winding.

Feature	Three-Point Starter	Four-Point Starter
Holding Coil Connection	In series with Shunt Field	Directly across Supply (independent of Shunt Field)
Points/Terminals	3 (Line, Armature, Field)	4 (Line, Armature, Field, 4th Point for NVC)
Protection under Field Weakening	Can trip under field weakening (field current decreases, NVC current decreases)	Does not trip under field weakening (NVC current depends only on supply voltage)
Suitability	Suitable for motors with constant or increasing field resistance	Suitable for all types of DC shunt and compound motors, especially where field control (weakening) is used

This difference in holding coil connection has a significant implication:

- In a three-point starter, if the shunt field current is reduced significantly (e.g., for speed control by field weakening) or if the field circuit becomes open, the holding coil current also reduces. If the field current drops below a certain level, the holding coil may lose its magnetism, causing the starter handle to trip to the "OFF" position, even if there is no fault other than intended speed control. This is an undesirable feature.
- In a four-point starter, since the holding coil is connected directly across the line, its current depends only on the supply voltage (and its own series resistance), not the shunt field current. Therefore, field weakening or variations in shunt field current do not affect the holding coil's operation, ensuring the motor continues to run even when its speed is controlled by changing the field resistance.

Analyzing the Options

Let's evaluate the given options based on this understanding:

- **Option 1: Field circuit is completed through starting resistance.** This is not the basic difference. The starting resistance is primarily in series with the armature to limit armature current. The field circuit is typically connected across the supply, in parallel with the armature and series starting resistance combination. While some starters might have intricate connections, this isn't the defining difference between the *types* of starters.
- **Option 2: Overcurrent release is connected in the supply line.** Both types of starters generally have an overcurrent release mechanism that protects the motor from excessive current draw from the supply. Its connection point might vary slightly in design, but the presence and function of overcurrent protection are common to both and not the *basic difference* in their fundamental design principle.
- **Option 3: Holding coil is removed from the shunt field circuit.** This statement accurately describes the key difference. In the four-point starter, the holding coil's connection is independent of the shunt field circuit, unlike the three-point starter where it is in series with the shunt field. This is the defining characteristic.
- **Option 4: The starting resistance is gradually cut out till the armature reaches the running position.** This describes the primary function of *any* DC motor starter (two-point, three-point, four-point, etc.). It is the common principle of operation for starting, not a difference between three-point and four-point types.

Therefore, the basic difference is the connection of the holding coil (NVC), which is removed from being in series with the shunt field circuit in a four-point starter, connecting it instead directly across the line.

Revision Table: Key Starter Concepts

Term	Description	Relevance to Starters
Starting Current	High current when DC motor starts due to low back EMF.	Starters limit this current using series resistance.
Back EMF (E_b)	Voltage generated in armature due to rotation; opposes supply voltage. $E_b = K\Phi\omega$.	Increases with speed, reducing armature current naturally.
Holding Coil (NVC/LVR)	Electromagnet that holds starter handle in RUN position.	Provides no-volt protection; connection differs in 3-point vs 4-point.
Overcurrent Release (OCR)	Mechanism to trip starter on excessive current.	Provides overload protection; present in both starter types.
Shunt Field Winding	Provides the magnetic field; connected in parallel with armature.	Holding coil connection relative to this winding is the key difference.

Additional Information on DC Motor Starters

DC motor starters are categorized based on the number of terminals they have for connection to the motor and supply:

- **Two-Point Starter:** Used only for DC series motors. It has two terminals (Line, Armature). It lacks no-volt protection for the motor itself (though often integrated into control systems). The holding coil is in series with the motor armature and field, providing overload protection but relying on motor current. If the load is removed, the motor speeds up, armature current drops, and the holding coil may release, but this is not ideal no-volt protection.
- **Three-Point Starter:** Used for DC shunt and compound motors. Has three terminals (Line, Armature, Field). Provides both no-volt and overload protection. No-volt coil is in series with the shunt field.

- **Four-Point Starter:** Used for DC shunt and compound motors, especially where speed control by field weakening is employed. Has four terminals (Line, Armature, Field, fourth point for independent NVC connection). Provides no-volt and overload protection, with the NVC operation independent of the shunt field current.

The choice between a three-point and four-point starter depends on the application, particularly whether speed control via field weakening is required. For applications where field weakening is not used, a three-point starter is often sufficient and simpler. For applications requiring speed control by field weakening, a four-point starter is necessary to prevent nuisance tripping.

161. Answer: a

Explanation:

Understanding Lead-Acid Battery Charging and Specific Gravity

The question asks what happens to the specific gravity of the electrolyte in a lead-acid battery while it is being charged. To answer this, let's understand the chemical processes involved during charging.

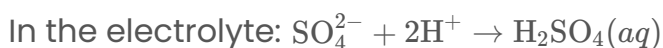
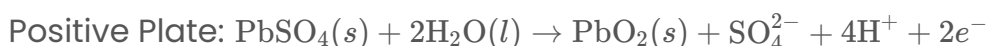
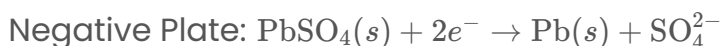
Chemical Processes During Lead-Acid Battery Charging

A lead-acid battery stores energy through reversible chemical reactions. The electrolyte in a fully charged lead-acid battery is a solution of sulfuric acid (H_2SO_4) and water (H_2O). When the battery discharges, lead (Pb) on the negative plate and lead dioxide (PbO_2) on the positive plate react with the sulfuric acid in the electrolyte to form lead sulfate (PbSO_4) and water. This process consumes sulfuric acid and produces water, which makes the electrolyte less dense.

During charging, an external electrical current is applied, which reverses the discharge reactions. The lead sulfate (PbSO_4) on the plates reacts with the water (

H_2O) from the electrolyte to regenerate lead (Pb) on the negative plate, lead dioxide (PbO_2) on the positive plate, and crucially, sulfuric acid (H_2SO_4).

The overall simplified chemical reaction during charging is:



Combining these, the net charging reaction is the reverse of discharging:



Impact on Electrolyte Specific Gravity

Specific gravity is the ratio of the density of a substance to the density of a reference substance (usually water at a specific temperature). In a lead-acid battery, the electrolyte is a mixture of sulfuric acid and water. Sulfuric acid is significantly denser than water.

- Density of pure water (H_2O) is approximately 1.00 g/cm^3 .
- Density of pure sulfuric acid (H_2SO_4) is approximately 1.83 g/cm^3 .

During charging, the chemical reactions produce sulfuric acid (H_2SO_4) and consume water (H_2O). This means the concentration of the denser component (sulfuric acid) in the electrolyte increases, while the concentration of the less dense component (water) decreases.

An increase in the concentration of sulfuric acid in the electrolyte directly leads to an increase in the overall density of the electrolyte. Therefore, the specific gravity of the electrolyte increases as the battery charges.

Analyzing the Options

- **It increases:** As explained above, charging produces sulfuric acid, increasing the electrolyte's density and thus its specific gravity. This aligns with the

chemical process.

- **It remains the same:** This is incorrect. The composition of the electrolyte changes significantly during charging.
- **It evaporates:** While water can evaporate over time, especially if the battery gasses excessively during charging, this is not the primary change happening *to the specific gravity* due to the *charging process itself*. Evaporation would actually increase the concentration of sulfuric acid and thus specific gravity, but "it increases" describes the direct result of the chemical reaction. Moreover, modern batteries are often sealed, minimizing evaporation.
- **It decreases:** This happens during discharging, when sulfuric acid is consumed and water is produced. The opposite occurs during charging.

Therefore, the specific gravity of the lead-acid battery's electrolyte increases while charging.

Battery State	Electrolyte Composition Trend	Specific Gravity
Discharging	Sulfuric acid consumed, Water produced	Decreases
Charging	Sulfuric acid produced, Water consumed	Increases
Fully Charged	Higher sulfuric acid concentration	Highest (within typical range)
Fully Discharged	Lower sulfuric acid concentration (more water)	Lowest (closer to water's specific gravity)

Revision Table: Lead-Acid Battery Electrolyte

Aspect	During Charging	During Discharging
Chemical Reaction Direction	Reverse of Discharge	Forward (Energy Release)
Sulfuric Acid (H_2SO_4)	Produced	Consumed
Water (H_2O)	Consumed	Produced
Lead Sulfate (PbSO_4)	Consumed	Produced
Specific Gravity of Electrolyte	Increases	Decreases

Additional Information: Testing Specific Gravity

Measuring the specific gravity of the electrolyte is a common way to assess the state of charge of a lead-acid battery. A device called a hydrometer is used for this purpose. A higher specific gravity reading indicates a higher concentration of sulfuric acid and thus a higher state of charge. Typical specific gravity values for a fully charged battery at 25°C are around 1.265 to 1.280, while a fully discharged battery might be around 1.150 or lower. Temperature correction is important when taking readings, as temperature affects density.

Monitoring specific gravity helps in:

- Determining the battery's state of charge.
- Detecting problems like stratification (acid and water separating).
- Assessing the overall health of the battery.

162. Answer: d

Explanation:

Understanding DC Machine Maintenance Needs

DC machines, like any electrical equipment, require regular maintenance to ensure optimal performance and longevity. Different parts of a DC machine have varying

maintenance requirements based on their function, construction, and exposure to wear and tear.

Why Brushes and Commutator Need Frequent Maintenance

The brush and commutator assembly is a critical part of a DC machine responsible for facilitating the transfer of electrical current between the stationary external circuit and the rotating armature winding. This process involves physical contact and relative motion between the brushes (typically made of carbon or graphite) and the rotating commutator segments (made of copper).

Here's why this part needs frequent attention:

- **Friction and Wear:** The continuous rubbing action between the brushes and the commutator causes friction. This friction leads to gradual wear of both the brushes and the commutator segments over time.
- **Sparking:** During commutation, voltage is induced in the armature coils undergoing switching. If commutation is not ideal, sparking can occur between the brushes and the commutator segments. Sparking accelerates wear and can pit the commutator surface.
- **Carbon Dust:** As brushes wear down, they produce carbon dust. This dust can accumulate within the machine, potentially causing tracking, insulation breakdown, or obstructing ventilation if not cleaned regularly.
- **Brush Seating:** Brushes need to be properly seated on the commutator to ensure good contact. Improper seating can lead to uneven wear, sparking, and poor performance.

Due to these factors involving physical contact, friction, and electrical switching stresses, the brushes and commutator are typically the parts that experience the most wear and require the most frequent inspection, cleaning, adjustment, and eventual replacement.

Maintenance Needs of Other DC Machine Parts

- **Rotor (Armature):** The rotor contains the armature winding. While windings can fail due to overheating or insulation breakdown, they don't inherently wear down physically like brushes and the commutator during normal operation. Maintenance usually involves checking insulation resistance and bearing condition.
- **Stator (Field Winding and Frame):** The stator houses the field winding (either series, shunt, or compound) and provides the magnetic field. This is a stationary part and is generally very robust. Maintenance primarily involves checking winding insulation and terminal connections.
- **Supply Connection:** These are the external wiring terminals connecting the DC machine to the power source. Maintenance involves checking for loose connections, corrosion, or damage to the wiring. While important for safety and operation, these connections do not experience the same continuous physical wear as the brush-commutator system.

Conclusion

Comparing the different parts, the brush and commutator assembly is subjected to continuous physical wear and electrical stress, making it the part of a DC machine that requires the most frequent maintenance, including inspection, cleaning, and brush replacement.

Typical DC Machine Maintenance Frequency Comparison

Part	Reason for Maintenance	Relative Maintenance Frequency
Brush and Commutator	Friction, wear, sparking, carbon dust	High (Frequent Inspection & Replacement)
Rotor (Armature)	Bearing wear, winding insulation	Moderate (Less frequent than brushes/commutator)
Stator (Field Winding)	Winding insulation, connections	Low (Infrequent checks)
Supply Connection	Loose connections, corrosion	Moderate (Periodic checks)

Revision Table: Key DC Machine Maintenance Points

DC Machine Component	Primary Maintenance Concern
Brushes	Wear, length, proper seating, material type
Commutator	Surface condition (smoothness, pitting), cleanliness, runout
Bearings	Lubrication, wear, noise, vibration
Windings (Armature/Field)	Insulation resistance, overheating signs
Connections	Tightness, corrosion

Additional Information: DC Machine Brushes and Commutator

The proper functioning of the brush-commutator system is vital for the operation of DC machines, whether they are operating as motors or generators. Common maintenance tasks include:

- Checking brush length and replacing brushes before they become too short.
- Ensuring brushes move freely in their holders.
- Cleaning carbon dust from the machine interior.
- Inspecting the commutator surface for signs of wear, pitting, or burning.
- Undercutting the mica insulation between commutator segments if necessary.
- Checking and adjusting brush pressure on the commutator.

Neglecting brush and commutator maintenance can lead to increased sparking, accelerated wear, damage to the commutator, reduced machine efficiency, and ultimately, machine failure.

163. Answer: a

Explanation:

Understanding Mass Soldering Techniques

The question asks to identify the technique that is not a type of **mass soldering**. Mass soldering refers to methods used to solder multiple components onto a printed circuit board (PCB) simultaneously or in a rapid, automated process. This is in contrast to manual soldering, where each solder joint is made individually, typically using a soldering iron.

Analyzing the Options

Let's examine each option to determine if it is a **mass soldering** technique:

- **Iron soldering:** This involves using a heated soldering iron to manually apply solder to one joint at a time. It requires skilled manual manipulation for each connection. This is a manual soldering technique, not a **mass soldering** technique.
- **Dip soldering:** In this method, a fully populated PCB is dipped into a pool of molten solder. All solder joints on the underside of the board are formed simultaneously. This is a **mass soldering** technique.
- **Wave soldering:** This technique involves passing a PCB over a wave of molten solder. The solder wave contacts the component leads and pads on the underside of the board, creating the solder joints in a continuous process. This is a widely used **mass soldering** technique.
- **Drag soldering:** Similar to wave soldering, drag soldering involves dragging the PCB across the surface of a static or slowly moving pool of molten solder. This forms the solder joints as the board passes over the solder surface. This is also considered a **mass soldering** technique.

Comparing Soldering Techniques

Here's a simple comparison of the techniques mentioned:

Technique	Type	Description
Iron Soldering	Manual Soldering	Soldering one joint at a time with a heated iron.
Dip Soldering	Mass Soldering	Dipping the entire board into molten solder.
Wave Soldering	Mass Soldering	Passing the board over a moving wave of molten solder.
Drag Soldering	Mass Soldering	Dragging the board across the surface of molten solder.

Based on this analysis, **Iron soldering** is clearly a manual process, not a **mass soldering** technique used for manufacturing large quantities of PCBs efficiently.

Conclusion on Mass Soldering Types

The techniques like dip soldering, wave soldering, and drag soldering are designed for high-volume production and are classified as **mass soldering** methods. **Iron soldering**, being a manual method, is primarily used for prototyping, repairs, rework, or low-volume assembly where individual components or connections are handled.

Revision Table: Soldering Methods

Method	Mass Production Suitable?	Common Use Case
Iron Soldering	No	Prototyping, Repair, Rework
Dip Soldering	Yes	Through-hole component assembly (older method)
Wave Soldering	Yes	Through-hole component assembly (common method)
Drag Soldering	Yes	Specific applications, often requires jigs

Additional Information: Other Soldering Processes

While the question focuses on traditional through-hole **mass soldering** techniques, it's worth noting other important soldering processes:

- **Reflow Soldering:** This is the primary **mass soldering** technique used for Surface Mount Technology (SMT) components. Solder paste is applied to pads on the PCB, components are placed, and the board is heated (typically in an oven) to melt the solder paste and form the joints.
- **Selective Soldering:** A variation of wave soldering used for PCBs with a mix of SMT and through-hole components, especially when SMT components are on the bottom side under through-hole components. Only specific through-hole joints are soldered using a directed solder wave or nozzle.
- **Laser Soldering:** Uses a laser beam to heat the joint and melt the solder. Can be very precise for specific applications.
- **Vapor Phase Soldering:** Uses hot vapor to heat the PCB uniformly, melting solder paste for SMT components.

Understanding the differences between these methods is crucial for anyone involved in electronics manufacturing or assembly.

164. Answer: b

Explanation:

Understanding Transformer Types

Transformers are electrical devices that transfer electrical energy between two or more circuits through electromagnetic induction. They are commonly used to step up or step down voltage levels in AC circuits. Different types of transformers are designed for specific applications, including power transmission, distribution, and measurement.

This question asks about a specific characteristic of one type of transformer: having its secondary winding always kept closed. Let's examine the types mentioned in the options.

What are Instrument Transformers?

Potential transformers (PTs) and Current transformers (CTs) are collectively known as instrument transformers. They are used to measure high voltages and large currents safely by stepping them down to lower, manageable levels that can be read by standard instruments like voltmeters, ammeters, wattmeters, and relays.

- **Potential Transformer (PT):** Connected in parallel with the circuit, like a voltage source. Steps down high voltage to a standard lower voltage (e.g., 110 V).
- **Current Transformer (CT):** Connected in series with the circuit, like an ammeter. Steps down high current to a standard lower current (e.g., 5 A or 1 A).

Current Transformer (CT) Operation and Safety

A Current Transformer operates based on the principle of electromagnetic induction, but its primary winding is connected in series with the high-current line. The load (usually an ammeter, relay coil, or the current coil of a wattmeter/energy meter) is connected across the secondary terminals.

The key characteristic of a CT is that the secondary winding must always be kept closed or short-circuited when the primary is energized. Here's why:

- When the secondary is closed, the secondary current flows, which opposes the magnetizing flux created by the primary current (due to Ampere's law). This opposition limits the net flux in the core to a small value, just enough to induce the required secondary voltage.
- If the secondary winding is left open-circuited while the primary is carrying current, there is no secondary current to oppose the primary's magnetizing effect.
- This results in the entire primary current acting as the magnetizing current.
- A large magnetizing current produces a very large magnetic flux in the core.
- According to Faraday's law of induction ($E = -N \frac{d\Phi}{dt}$), a large, rapidly changing flux will induce a very high voltage across the many turns of the secondary winding (N is large for the secondary to step down current, meaning it has many turns).
- This dangerously high induced voltage can break down the insulation of the winding, potentially causing arcs or short circuits, damaging the transformer itself, connected instruments, and posing a severe safety hazard to personnel.
- The excessive flux also leads to severe core saturation, which can permanently damage the core's magnetic properties and affect the accuracy of the CT.

Therefore, to prevent these dangerous conditions, the secondary winding of a current transformer is always kept closed, either by connecting it to its intended low-impedance load (like an ammeter) or, if the load is disconnected, by physically short-circuiting the terminals.

Other Transformer Types

- **Potential Transformer (PT):** The secondary of a PT is connected to a voltmeter or voltage coil, which is a high-impedance load. If the secondary is open-circuited, only the magnetizing current flows in the primary. While this affects accuracy, it does not typically lead to dangerously high voltages in the same way as an open-circuited CT secondary, because the primary voltage is fixed and the core is designed not to saturate under normal operating voltage.

- **Power Transformer:** Used in transmission and distribution systems to change voltage levels (e.g., from generating voltage to transmission voltage, or transmission voltage to distribution voltage). The secondary is connected to the load. If the secondary is open, the primary draws only a small magnetizing current. This is a normal no-load condition and does not cause dangerous voltage build-up or core saturation issues like in a CT.
- **Distribution Transformer:** A type of power transformer used to step down voltage for final distribution to homes and businesses (e.g., from 11 kV to 415 V/230 V). Similar to power transformers, an open secondary is simply a no-load condition and doesn't pose the same hazard as with a CT.

Based on the operational principles and safety requirements, the transformer type whose secondary winding is always kept closed is the Current Transformer.

Transformer Type	Purpose	Connection (Primary)	Secondary Condition	Effect of Open Secondary
Potential Transformer (PT)	Measure high voltage	Parallel	Connected to voltmeter (high impedance)	Doesn't cause extreme overvoltage/saturation
Current Transformer (CT)	Measure high current	Series	Always kept closed (low impedance load or shorted)	Causes extreme overvoltage and core saturation
Power Transformer	Change voltage level	Connected to source	Connected to load (variable)	No-load condition, does not cause extreme overvoltage/saturation
Distribution Transformer	Step down voltage for distribution	Connected to distribution line	Connected to load (variable)	No-load condition, does not cause extreme overvoltage/saturation

Revision Table: Key Transformer Concepts

Concept	Description
Instrument Transformer	Transformers used for measurement and protection in high voltage/current circuits.
Current Transformer (CT)	Measures high current by stepping it down to a low level. Primary in series, secondary connected to ammeter/relay.
Potential Transformer (PT)	Measures high voltage by stepping it down to a low level. Primary in parallel, secondary connected to voltmeter/voltage coil.
Open-circuited Secondary (CT)	Dangerous condition in a CT leading to high voltage and core saturation due to unopposed magnetizing flux.

Additional Information on Transformer Safety

Safety is paramount when working with transformers, especially instrument transformers. For Current Transformers, strictly adhering to the rule of keeping the secondary closed is vital. If a measuring instrument connected to a CT needs to be replaced or serviced, the secondary terminals of the CT must be short-circuited BEFORE disconnecting the instrument. Once the new instrument is connected, the short circuit can be removed. This ensures that the secondary circuit is never open while the primary is energized.

Core saturation in a CT not only poses a safety risk but also affects the accuracy. A saturated core means the linear relationship between primary and secondary current is lost, leading to errors in measurement or improper operation of protective relays.

165. Answer: b

Explanation:

Understanding Domestic Single Phase Voltage Supply

When we talk about the electricity supplied to our homes for everyday use, it's typically a single phase supply. This power is used for running lights, fans, televisions, and most standard household appliances.

What is Single Phase Supply?

Single phase electric power is the distribution of alternating current electric power using a system in which all the voltages of the supply vary in unison. It's commonly used in residential areas and small commercial buildings because it's simpler and less expensive than three-phase systems.

Standard Domestic Supply Voltage

The voltage level of the single phase supply provided to domestic customers varies depending on the region or country. However, a widely adopted standard voltage level in many parts of the world, including Europe, Asia, Africa, and Oceania, is around 230 V.

Let's look at the given options:

- **Option 1:** 440 V – This voltage is typically associated with three-phase systems or industrial applications, not standard single phase domestic supply.
- **Option 2:** 230 V – This is a common and widely recognized standard voltage for single phase domestic electricity supply in many countries.
- **Option 3:** 110 V – While 110 V to 120 V is the standard single phase domestic voltage in some regions like North America, 230 V is also a prevalent standard globally.
- **Option 4:** Cut-out on customer's premises 420 V – A cut-out is a safety device, not a voltage value. 420 V is not a standard single phase domestic supply voltage.

Considering the options and common standards, 230 V is the most appropriate answer for the single phase supply voltage provided to domestic customers in

many parts of the world.

Why 230 V is a Common Standard

The adoption of 230 V as a standard in many regions is often linked to historical power distribution developments and efficiency considerations for transmitting power over distances to residential areas.

Revision Table: Domestic Voltage Standards

Aspect	Description	Typical Voltage(s)
Domestic Supply	Electricity supplied to homes for lighting and appliances.	230 V or 110 – 120 V (depending on region)
Phase	Typically Single Phase for homes.	Single Phase
Common Standard	Voltage level used in many countries globally.	230 V

Additional Information on Power Supply

Understanding the different types of power supply is crucial in electrical systems.

- **Single Phase Power:** Delivers power using two conductors: one phase wire and one neutral wire. The voltage alternates between the phase wire and neutral. It is suitable for low-power applications typical in homes.
- **Three Phase Power:** Delivers power using three phase wires and often a neutral wire. The voltages in the three phase wires are offset in time, providing a constant, smooth power flow. It is used for high-power applications like motors and industrial machinery. Three-phase supply might be stepped down to single-phase for domestic use. For example, in a 400 V three-phase system (line-to-line), the single-phase voltage (line-to-neutral) is $400/\sqrt{3} \approx 230 \text{ V}$.

The voltage level at the customer's premises is regulated and supplied by the local power distribution company, stepped down from higher transmission voltages.

166. Answer: a

Explanation:

Concept:

According to Ohm's law,

$$V = I R$$

Where,

V = Voltage

I = Current

R = Resistance

Calculation:

From the given circuit,

Total resistance, $R = 10 + 10 = 20$

$$I = \frac{50}{20} = 2.5 \text{ A}$$

167. Answer: d

Explanation:

Understanding Low Ceiling Fan Speed

A common issue homeowners face is a ceiling fan running slower than it used to. Several factors can cause this, but the question specifically asks for one of the primary reasons for this decreased speed.

Analyzing the Options for Low Fan Speed

Let's examine each option provided and determine its potential effect on a ceiling fan's speed:

- **Blade corrosion:** Corrosion on fan blades can affect their balance and aerodynamics. This might lead to wobbling, increased noise, or slightly reduced airflow efficiency. However, it's generally not a direct or primary cause for the *motor* running at a significantly lower speed.
- **Broken canopy:** The canopy is the decorative cover that hides the wiring and mounting hardware at the ceiling. A broken canopy is a structural or aesthetic issue. It does not affect the electrical operation or the speed at which the motor rotates.
- **Faulty mounting:** Faulty mounting relates to how the fan is attached to the ceiling. This can cause the fan to wobble dangerously, make noise, or even fall. Like a broken canopy, it's a structural/safety issue and does not directly impact the motor's electrical speed.
- **Faulty capacitor:** Most AC ceiling fans use a capacitor to create a phase shift in the current supplied to the motor windings. This phase shift is crucial for generating the rotating magnetic field needed to start the motor and produce the necessary torque for it to run efficiently at its designed speed.

The Role of the Capacitor in Ceiling Fan Speed

Ceiling fans typically use a capacitor in conjunction with the starting and running windings of the motor. The capacitor provides a voltage boost and phase shift to the starting winding (or sometimes the running winding, depending on the motor type). This phase shift is essential for generating a strong rotating magnetic field, which the rotor follows, causing the fan blades to turn.

Over time, capacitors can degrade or fail. When a capacitor becomes faulty, it loses its ability to store and release electrical charge effectively. This results in a reduced voltage boost and a less optimal phase shift. Consequently, the motor receives less effective power, generates less torque, and struggles to reach its intended speed. A common symptom of a failing capacitor is indeed a fan running noticeably slower than normal, especially on higher speed settings.

Comparing Causes

While other issues like excessive dust buildup on blades, lack of lubrication in bearings (though many modern fans are sealed), or even incorrect voltage supply could potentially affect speed, the options provided point towards specific components or issues. Among the given options, a faulty capacitor has the most direct and significant impact on the fan's motor speed, specifically causing it to run slowly.

Therefore, a faulty capacitor is a well-established reason for the low speed of a ceiling fan.

Revision Table: Ceiling Fan Issues & Causes

Issue	Common Causes	Impact on Fan Operation
Low Speed	Faulty capacitor, incorrect voltage, motor wear	Fan spins slowly, reduced airflow
Wobbling	Improper mounting, unbalanced blades (bent, weighted unevenly, dust buildup), loose screws	Fan shakes, potentially dangerous
Noise (Humming/Buzzing)	Motor issues, faulty speed control, electrical interference	Annoying sound
Noise (Clicking/Grinding)	Bearing issues, rubbing parts, loose components	Mechanical sound
Not Starting	No power (switch off, breaker tripped, wiring issue), faulty motor, completely failed capacitor, remote/switch issue	Fan does not spin

Additional Information on Ceiling Fan Components

Understanding the key components helps diagnose problems:

- **Motor:** The electric motor is the heart of the fan, converting electrical energy into rotational motion.
- **Blades:** These move air when rotated by the motor. Their shape, pitch, and balance are crucial for airflow and smooth operation.
- **Capacitor:** An electrical component that stores and releases energy, essential for providing starting torque and optimizing running performance in many AC motors, including ceiling fans. It helps create the necessary phase difference between the start and run windings.
- **Speed Control:** This can be a pull chain switch, a wall switch, or a remote control. It regulates the voltage or changes wiring configurations to select different speeds.
- **Mounting Hardware:** Includes the bracket and downrod (if used) that secure the fan to the ceiling. Proper installation is critical for safety and stability.

Troubleshooting low speed often starts with checking the power supply and then inspecting or replacing the capacitor, as it's a common failure point that directly affects motor torque and speed.

168. Answer: a

Explanation:

Ward-Leonard Control Explained

The Ward-Leonard control system is a method primarily used for regulating the speed of Direct Current (DC) motors. It operates using a generator-motor set. A DC generator is driven at a constant speed by a prime mover (like an AC motor). The key principle is that the output voltage of this DC generator is controlled by adjusting its own field excitation current. This controlled DC voltage is then supplied to the armature of the DC motor whose speed needs to be regulated.

The speed of a DC motor is generally dependent on two main factors: the applied armature voltage (V_a) and the field flux (ϕ). The relationship can be simplified as:

$$N \propto \frac{V_a}{\phi}$$

Where N is the motor speed.

In the Ward-Leonard system, the field flux (ϕ) of the DC motor is typically kept constant. The speed (N) is then directly controlled by varying the armature voltage (V_a). This variation of V_a is achieved by controlling the excitation field current of the DC generator. Therefore, by adjusting the generator's field, the system effectively controls the armature voltage supplied to the motor, allowing for precise speed control over a wide range.

Analyzing Control Method Options for Ward-Leonard

Let's examine the given options in the context of how the Ward-Leonard system functions:

- **Voltage control method:** This aligns perfectly with the core mechanism. The system manipulates the generator's field to control its output voltage, which is then applied to the motor's armature, controlling its speed.
- **Field diverter method:** This involves diverting a portion of the field current, typically in the motor's own field winding, to weaken the field and increase speed. Ward-Leonard control adjusts the *generator's* field, not directly diverting the motor's field.
- **Field control method:** While the generator's field is indeed controlled, this option is less specific. The control action is focused on achieving a specific *voltage* output from the generator, which then controls the motor. Calling it simply a "field control method" could apply to other systems as well, whereas "voltage control" better describes the specific application in Ward-Leonard.
- **Armature resistance control method:** This method involves adding variable resistance in series with the motor's armature to reduce speed. It's inefficient due to power loss in the resistor and doesn't represent the Ward-Leonard principle.

Why Voltage Control is the Best Description

The fundamental operation of the Ward-Leonard system hinges on controlling the voltage supplied to the DC motor's armature. This is accomplished by regulating the field current of the DC generator that powers the motor. Adjusting the generator's field strength changes its generated voltage (E_g), which is applied as the armature voltage (V_a) to the motor. Since motor speed (N) is directly proportional to armature voltage (V_a) when the field flux (ϕ) is constant ($N \propto V_a/\phi$), controlling the generator's output voltage is the direct way speed is regulated.

Therefore, describing the Ward-Leonard control as a **voltage control method** accurately captures its primary operational principle and distinguishes it from other DC motor speed control techniques.

169. Answer: d

Explanation:

Correct answer: varnished cambric

varnished cambric: Cambric is a fine, densely woven cloth, historically made of linen, but often made of cotton today. When this cambric cloth is impregnated or coated with insulating varnish, it forms varnished cambric. This treated fabric is the traditional material used to manufacture Empire tape due to its excellent dielectric strength and flexibility.

Based on the common composition of Empire tape, the material it is usually made of is varnished cambric.

170. Answer: a

Explanation:

Understanding Earth Faults and Earth Mat Potential

An earth fault is a common type of electrical fault where a live conductor accidentally comes into contact with the earth or a grounded object. This creates an unintended path for current to flow from the system to the ground. Earthing systems, including earth mats, are designed to provide a low-resistance path for such fault currents to flow safely to the earth, preventing dangerous voltage build-up on equipment frames and structures.

What Happens During an Earth Fault?

When an earth fault occurs, a significant amount of current, known as fault current, flows through the earthing system, including the earth mat, and disperses into the surrounding soil. The purpose of the earth mat is to reduce the overall resistance to earth and spread the fault current over a larger area, thereby minimizing potential differences on the surface.

However, even with a well-designed earthing system, there is always some resistance between the earth mat and the general mass of the earth. This resistance is known as the earth resistance or grounding resistance.

Voltage Rise at the Earth Mat

According to Ohm's Law, the voltage drop across a resistance is equal to the current flowing through it multiplied by the resistance ($V = I \times R$).

During an earth fault, the fault current (I_{fault}) flows through the earth resistance (R_{earth}) of the earthing system (which includes the earth mat). This causes a voltage difference between the earth mat and the true earth potential (which is considered zero far away from the fault location). The voltage potential at the earth mat rises above zero due to this voltage drop.

Mathematically, the voltage rise at the earth mat (V_{mat}) can be approximated as:

$$V_{\text{mat}} \approx I_{\text{fault}} \times R_{\text{earth}}$$

Since I_{fault} can be very high during an earth fault, even a low earth resistance (R_{earth}) can result in a significant voltage rise at the earth mat. This voltage rise creates potential differences on the ground surface (step potential and touch potential), which can be hazardous.

Analyzing the Options

- **Option 1: Voltage potential at the earth mat increases due to earthing.** This statement accurately describes the phenomenon during an earth fault. The fault current flowing through the earth resistance causes the potential of the earth mat to rise above the ideal zero potential of the general earth.
- **Option 2: Voltage potential at the earth mat decreases due to earthing.** This is incorrect. While earthing helps to minimize the potential rise compared to having no earthing system, the potential at the earth mat **increases** from its normal zero (relative to true earth) during a fault.
- **Option 3: Voltage potential at the earth mat remains zero irrespective of fault.** This is incorrect. The potential of the earth mat is only zero under ideal conditions (zero earth resistance) or when no fault current is flowing. During an earth fault, fault current flows, and due to the non-zero earth resistance, a voltage drop occurs, causing the potential at the earth mat to rise above zero.
- **Option 4: None of the above.** This is incorrect, as Option 1 correctly describes the situation.

Therefore, when an earth fault occurs, the voltage potential at the earth mat increases due to the flow of fault current through the earthing resistance.

Revision Table: Earth Fault Effects

Condition	Current Flow	Earth Mat Potential
Normal Operation (No Fault)	Negligible leakage current	Approximately zero (relative to true earth)
Earth Fault Occurs	Large fault current flows through earthing system	Increases above zero due to $I \times R$ drop

Additional Information: Earthing System Design

The primary goals of a good earthing system design during an earth fault are:

- To provide a low-resistance path for fault current to minimize the voltage rise at the earth mat and connected equipment.
- To ensure that potential differences on the ground surface (step and touch potentials) remain below hazardous levels.
- To facilitate the operation of protective devices (like circuit breakers) by ensuring sufficient fault current flows.

Factors affecting earth resistance include soil type, moisture content, temperature, and the configuration and depth of the earthing electrodes or earth mat.

171. Answer: d

Explanation:

Understanding Eddy Current Loss in DC Machines

Eddy currents are circulating currents induced within the body of a conductor when it is subjected to a changing magnetic flux. These induced currents flow in closed loops within the material. Due to the electrical resistance of the material, the flow of these eddy currents causes power loss in the form of heat, known as eddy current loss. This loss is proportional to the square of the frequency of the magnetic field change, the square of the maximum flux density, the square of the thickness of the material (if not laminated), and the conductivity of the material.

In a DC machine, several parts are made of ferromagnetic material and are subjected to magnetic flux. However, the rate and magnitude of the change in flux density vary significantly between these parts. Eddy current losses are particularly problematic in parts that experience rapidly changing or alternating magnetic fields.

Identifying the Source of Greatest Eddy Current Loss

Let's examine the primary functions and magnetic conditions of the listed parts of a DC machine:

- **Yoke:** The yoke is the outer frame that provides mechanical support and acts as a return path for the magnetic flux. The flux in the yoke is generally steady when the machine operates under constant load conditions, although minor fluctuations might occur due to armature reaction or load changes. The rate of flux change is relatively low compared to rotating parts.
- **Commutator:** The commutator is a mechanical rectifier that converts AC voltage/current in the armature to DC at the terminals (generator) or DC input to AC in the armature (motor). It consists of copper segments insulated from each other and the shaft. It is primarily involved in current switching and voltage collection, not being a bulk ferromagnetic material subject to significant eddy currents from the main field flux.
- **Field Poles:** The field poles produce the main magnetic field. In a DC machine, the field winding is typically supplied with DC current, creating a steady magnetic flux in the poles. While the pole faces might experience some localized flux variations due to the armature slots (pulsating flux), the core of the poles themselves is subjected to a relatively constant main flux, leading to much lower eddy current losses compared to the armature.
- **Armature:** The armature core is made of ferromagnetic material (usually silicon steel laminations) and rotates within the stationary magnetic field produced by the field poles. As the armature rotates, each part of the core passes under alternate north and south poles. This causes the magnetic flux passing through any given section of the armature core to continuously change direction and magnitude at a frequency determined by the speed of rotation and the number of poles. This rapid and significant change in magnetic flux induces large voltages within the armature core material, leading to the generation of substantial eddy currents.

The armature core is specifically designed with laminations (thin sheets of steel insulated from each other) to increase the resistance path for eddy currents and thereby reduce these losses significantly. However, even with lamination, some eddy current loss remains, and due to the substantial volume of the armature core material and the dynamic nature of the flux it experiences, the armature is where the greatest potential for, and actual significant amount of, eddy current loss occurs in a DC machine.

Summary of Eddy Current Loss Potential

DC Machine Part	Magnetic Flux Condition	Potential for Eddy Current Loss
Yoke	Relatively steady main flux return path.	Low
Commutator	Not a significant ferromagnetic path for main flux causing bulk induction.	Negligible from main flux perspective.
Field Poles	Steady main flux path. Localized variations at pole faces.	Low (core), Moderate (pole face) but overall less than armature.
Armature	Rapidly changing/alternating flux due to rotation in main field. Large volume.	Highest (even with lamination).

Therefore, the greatest eddy current loss occurs in the armature core due to its rotation in the stationary magnetic field, causing rapid changes in the magnetic flux density within its material.

Revision Table: DC Machine Losses

Type of Loss	Description	Main Location in DC Machine
Copper Losses	I^2R losses in armature and field windings.	Windings (Armature, Field)
Iron Losses (Core Losses)	Hysteresis and Eddy Current Losses in magnetic material.	Armature Core (mainly), Pole Faces
Mechanical Losses	Friction and windage losses.	Bearings, Commutator, Armature surface (air friction)
Brush Losses	Voltage drop at brush-commutator contact.	Brushes and Commutator

Additional Information: Minimizing Eddy Currents

To minimize eddy current losses in the armature core, which are a significant part of the iron losses (or core losses), the core is constructed using thin laminations of silicon steel. These laminations are coated with an insulating varnish or oxide layer and stacked together. This process increases the effective resistance of the paths available for the induced eddy currents to flow, thereby reducing their magnitude and consequently reducing the I^2R loss associated with them.

The formula for eddy current loss per unit volume can be approximated as $P_e \propto (B_{max})^2 f^2 t^2$, where B_{max} is the maximum flux density, f is the frequency of flux reversal, and t is the thickness of the laminations. Using thin laminations significantly reduces the t^2 term, thereby minimizing the loss.

172. Answer: c

Explanation:

Understanding DC Shunt Motor Starters and Protection

DC shunt motors, like other DC motors, require a starter during the initial phase of operation. This is because the armature resistance is very low, and connecting the motor directly to the supply voltage would result in a dangerously high starting current ($> \frac{\text{Supply Voltage}}{\text{Armature Resistance}}$). The starter limits this initial current and also provides crucial protective features.

Why are Starters Needed?

- To limit the high starting current by adding external resistance in series with the armature.
- To provide overload protection, preventing damage due to excessive current draw under heavy load.
- To provide no-volt protection (NVP), preventing the motor from restarting automatically upon restoration of power after a supply interruption, which can be hazardous.

Types of DC Motor Starters and Their Protection Features

Different types of starters offer varying levels of protection:

- **2-Point Starter:** Primarily used for DC series motors. It offers overload and no-volt protection but is not suitable for DC shunt motors due to the absence of a separate field terminal.
- **3-Point Starter:** Commonly used for DC shunt and compound motors. Provides overload protection and no-volt protection. The no-volt coil is connected in series with the field winding.
- **4-Point Starter:** Also used for DC shunt and compound motors. Provides overload protection and no-volt protection. The no-volt coil is connected directly across the supply lines through a series resistance, independent of the field winding.
- **5-Point Starter:** A less common type, sometimes incorporating additional features beyond the standard 3-point or 4-point design, such as field failure protection or over-speed protection.

Analyzing High Speed Protection in DC Shunt Motors

A significant risk for DC shunt motors is running at dangerously high speeds. The speed (N) of a DC motor is inversely proportional to the field flux (ϕ) and directly proportional to the back EMF (E_b), approximately given by the relation:

$$N \propto \frac{E_b}{\phi}$$

For a shunt motor, the field flux is proportional to the field current. If the field circuit is broken or the field current drops significantly (e.g., due to a fault), the field flux (ϕ) becomes very low. Since speed is inversely proportional to flux, this can cause the motor speed to increase drastically, leading to a "runaway" condition which can damage the motor mechanically.

Let's examine how 3-point and 4-point starters handle this risk:

- 3-Point Starter and Field Failure Protection:** In a 3-point starter, the no-volt coil (NVC) is connected in series with the field winding. If the field circuit breaks or the field current becomes too low, the current through the NVC also drops. This causes the electromagnet of the NVC to de-energize. The starter arm, which is held in position by the NVC electromagnet, is then released and returns to the 'OFF' position due to a spring. This disconnects the motor from the supply, thus preventing runaway speed due to field failure. Therefore, the 3-point starter provides high-speed protection against field failure.
- 4-Point Starter and Field Failure Protection:** In a 4-point starter, the no-volt coil (NVC) is connected across the supply voltage independently of the field winding, often through a current-limiting resistance. As long as the supply voltage is present, the NVC remains energized, regardless of the current flowing through the field winding. If the field circuit breaks, the field current drops to zero, causing the motor to accelerate to a dangerous speed. However, since the NVC is still energized by the supply voltage, it continues to hold the starter arm in position. The motor does not trip off the supply. Thus, the 4-point starter does not provide protection against runaway speed caused by field failure.

The 2-point starter is not used for shunt motors, so it is irrelevant for shunt motor high-speed protection in this context. While a 5-point starter might incorporate specific field failure or over-speed relays, the standard 3-point and 4-point starters

are the primary comparison for fundamental high-speed protection against field failure.

Conclusion: Identifying the Starter Without High Speed Protection

Based on the analysis, the 4-point starter lacks the crucial protection against high speed caused by field failure, a key runaway scenario for DC shunt motors. The 3-point starter, due to its NVC connection, inherently provides this protection.

Therefore, the starter that does NOT provide high speed protection (specifically against field failure runaway) for a DC shunt motor is the 4-point starter.

Revision Table: DC Motor Starter Features

Feature	3-Point Starter	4-Point Starter
Starting Current Limiting	Yes	Yes
Overload Protection	Yes	Yes
No-Volt Protection (NVP)	Yes	Yes
NVC Connection	In series with Field Winding	Across Supply (independent of Field)
Protection against Field Failure Runaway (High Speed)	Yes (NVC de-energizes)	No (NVC remains energized)

Additional Information: Field Failure Runaway

Field failure is a critical fault condition in DC shunt motors. If the field winding breaks or the field current is interrupted while the motor is running, the magnetic flux produced by the field winding collapses ($\phi \rightarrow 0$). The motor speed is inversely proportional to flux ($N \propto 1/\phi$). With negligible flux, the speed theoretically approaches infinity, although in reality, it is limited by friction, windage, and residual

magnetism. This uncontrolled acceleration can cause severe mechanical stress on the armature and other rotating parts, potentially leading to catastrophic failure. Reliable protection against field failure is therefore essential for safe operation of DC shunt motors.

173. Answer: a

Explanation:

Understanding Rechargeable and Non-Rechargeable Dry Cells

Dry cells are a type of electrochemical cell that use a paste-like electrolyte rather than a liquid electrolyte. They are commonly used in portable electronic devices. Batteries made from these cells are classified based on whether they can be recharged after use.

A **rechargeable battery** can be repeatedly discharged and recharged, meaning its chemical reactions can be reversed by applying an external electrical current. A **non-rechargeable battery** (also known as a primary battery) is designed for single use; its chemical reactions are not easily reversible, and attempting to recharge it can be dangerous or ineffective.

Analyzing the Given Battery Types

Let's examine each option provided to determine which one is typically not a rechargeable dry cell:

1. **Alkaline:** Standard alkaline batteries (like AA, AAA, C, D) are the most common type of non-rechargeable battery. They are designed for single use. While rechargeable alkaline batteries do exist, they are less common than standard alkaline batteries and other rechargeable types like NiMH or Li-ion. In the context of a list including prominent rechargeable technologies, "Alkaline" generally refers to the primary, non-rechargeable type.

2. **Nickel-cadmium (NiCd):** Nickel-cadmium batteries are a well-established type of rechargeable battery technology. They were widely used in portable electronics but have been largely replaced by NiMH and Li-ion due to environmental concerns regarding cadmium and the 'memory effect'.
3. **Nickel-metal hydride (NiMH):** Nickel-metal hydride batteries are rechargeable batteries that offer higher energy density compared to NiCd batteries and do not suffer from the prominent 'memory effect'. They are commonly used in portable devices, digital cameras, and hybrid vehicles.
4. **Lithium-ion (Li-ion):** Lithium-ion batteries are a popular type of rechargeable battery known for their high energy density, low self-discharge rate, and absence of the 'memory effect'. They are widely used in smartphones, laptops, electric vehicles, and many other modern electronic devices.

Comparing the Battery Types

Here is a brief comparison of the battery types mentioned:

Battery Type	Rechargeable?	Common Use	Notes
Alkaline	Typically No (Primary)	General household devices, remote controls, toys	Common non-rechargeable type.
Nickel-cadmium (NiCd)	Yes (Secondary)	Older portable electronics, power tools	Contains cadmium, environmental concerns.
Nickel-metal hydride (NiMH)	Yes (Secondary)	Digital cameras, portable electronics, hybrid cars	Higher capacity than NiCd, less 'memory effect'.
Lithium-ion (Li-ion)	Yes (Secondary)	Smartphones, laptops, EVs, modern electronics	High energy density, no 'memory effect', low self-discharge.

Conclusion

Based on the analysis, standard alkaline batteries are primary (non-rechargeable) cells, while Nickel-cadmium, Nickel-metal hydride, and Lithium-ion batteries are all secondary (rechargeable) cells. Therefore, the one that is not a rechargeable dry cell among the options provided is Alkaline.

Revision Table: Battery Classification

Classification	Description	Examples
Primary (Non-rechargeable)	Designed for single use; chemical reactions are irreversible.	Standard Alkaline, Zinc-carbon, Lithium (some types)
Secondary (Rechargeable)	Can be discharged and recharged multiple times; chemical reactions are reversible.	Nickel-cadmium, Nickel-metal hydride, Lithium-ion, Lead-acid

Additional Information: Rechargeable Battery Technologies

Rechargeable battery technology is constantly evolving, with different types offering various advantages in terms of energy density, power output, lifespan, cost, and safety. Key factors to consider when choosing a rechargeable battery include:

- **Energy Density:** How much energy can be stored per unit volume or mass. Li-ion generally has high energy density.
- **Power Density:** How quickly the battery can deliver energy.
- **Cycle Life:** The number of charge-discharge cycles the battery can withstand before its capacity significantly degrades.
- **Self-discharge Rate:** How quickly the battery loses its charge when not in use. Li-ion and modern NiMH (low self-discharge type) have low rates.

- **Memory Effect:** A phenomenon where a battery appears to lose capacity if repeatedly recharged after only being partially discharged (more prominent in NiCd).
- **Cost:** Initial purchase price and cost per cycle over its lifespan.
- **Environmental Impact:** The toxicity of materials used and ease of recycling.

Understanding these characteristics helps in selecting the appropriate battery type for a specific application, whether it's for portable devices, power tools, or electric vehicles.

174. Answer: a

Explanation:

Understanding Power in AC Circuits

In alternating current (AC) circuits, power is described using three terms: apparent power, active power, and reactive power. Understanding the relationship between these types of power is crucial for analyzing and designing AC systems.

- **Apparent Power (S):** This is the total power delivered by the source, measured in Volt-Amperes (VA). It is the product of the RMS voltage and RMS current.
- **Active Power (P):** Also known as real power or true power, this is the power consumed by the load that performs useful work, measured in Watts (W).
- **Reactive Power (Q):** This is the power exchanged between the source and the reactive components (inductors and capacitors) in the circuit, measured in Volt-Ampere Reactive (VAR). It does not perform useful work but is necessary to establish electric and magnetic fields.

Relationship and Power Factor Definition

These three power components are related by the power triangle, where the apparent power is the hypotenuse, active power is the adjacent side (usually on

the horizontal axis), and reactive power is the opposite side (usually on the vertical axis). The relationship is given by the Pythagorean theorem:

$$S^2 = P^2 + Q^2$$

The power factor (PF) of a system is a dimensionless quantity between 0 and 1 (or 0% and 100%) that represents how effectively the apparent power is converted into active power. It is defined as the ratio of active power to apparent power:

$$PF = \frac{P}{S}$$

Power factor can also be expressed as the cosine of the phase angle (ϕ) between the voltage and current waveforms:

$$PF = \cos(\phi)$$

A power factor close to 1 indicates efficient power usage, while a low power factor indicates inefficient usage due to a large amount of reactive power.

Calculating Power Factor When Apparent Power Equals Active Power

The question states that the apparent power (S) is equal to the active power (P).

Given:

$$S = P$$

We use the definition of the power factor:

$$PF = \frac{P}{S}$$

Substitute $S = P$ into the power factor formula:

$$PF = \frac{P}{P}$$

Assuming the active power P is not zero (as a non-zero apparent power equal to active power implies $P \neq 0$), the power factor is:

$$PF = 1$$

Interpretation of the Result

When the power factor is 1, it means that the active power is equal to the apparent power. This happens when the reactive power (Q) in the system is zero. Let's verify this using the power triangle equation:

$$S^2 = P^2 + Q^2$$

If $S = P$, then

$$P^2 = P^2 + Q^2$$

Subtracting P^2 from both sides gives:

$$0 = Q^2$$

This implies that $Q = 0$. A system with zero reactive power is purely resistive, meaning the current and voltage are perfectly in phase ($\phi = 0^\circ$). The power factor is $\cos(0^\circ) = 1$.

Therefore, if apparent power is found equal to active power, the system power factor is 1.

Step-by-Step Solution

1. Identify the given condition: Apparent Power (S) equals Active Power (P).
2. Recall the formula for power factor:

$$PF = \frac{P}{S}$$

3. Substitute the given condition $S = P$ into the power factor formula.
4. Calculate the value of the power factor:

$$PF = \frac{P}{P} = 1$$

5. Interpret the result: A power factor of 1 means the reactive power is zero, and the load is purely resistive.

The final answer is 1.

Revision Table: Power Concepts

Concept	Symbol	Unit	Description
Apparent Power	S	VA (Volt-Ampere)	Total power from source
Active Power	P	W (Watt)	Power doing useful work
Reactive Power	Q	VAR (Volt-Ampere Reactive)	Power exchanged by reactive components
Power Factor	PF	Dimensionless	Ratio of Active Power to Apparent Power (P/S)

Additional Information on Power Factor

A power factor less than 1 occurs when there is significant reactive power in the circuit due to inductive loads (like motors, transformers) or capacitive loads (like capacitors, long transmission lines). Inductive loads cause the current to lag behind the voltage (lagging power factor), while capacitive loads cause the current to lead the voltage (leading power factor).

A low power factor can lead to several issues:

- Increased current for the same amount of active power, leading to higher I^2R losses in conductors and transformers.
- Reduced system capacity because equipment must handle higher apparent power.
- Voltage drops.
- Penalties from utility companies for industrial customers.

Power factor correction techniques, usually involving adding capacitors (for lagging PF) or inductors (for leading PF), are often employed to bring the power factor closer to unity (1).

175. Answer: a

Explanation:

Understanding Electrical Connections: Aluminum to Copper

When connecting electrical conductors made of different materials, like aluminum and copper, special considerations are necessary to prevent problems that can lead to poor connections, increased resistance, overheating, and even fire hazards. The primary issue when directly connecting aluminum and copper is something called galvanic corrosion.

What is Galvanic Corrosion?

Galvanic corrosion occurs when two different metals are in electrical contact in the presence of an electrolyte (like moisture in the air). One metal becomes the anode and corrodes faster than it would alone, while the other becomes the cathode and corrodes slower. In the case of aluminum and copper, aluminum is more anodic than copper. This means that if they are in direct contact with moisture present, the aluminum will corrode.

This corrosion creates a high-resistance layer between the aluminum and copper, which can cause significant problems in electrical circuits.

Solving the Aluminum to Copper Connection Problem

To prevent galvanic corrosion and ensure a reliable, low-resistance connection between aluminum and copper, you need a method or component that effectively

separates the two metals while providing a secure electrical path. This is where specialized connectors come into play.

Identifying the Correct Lug Type

The question asks for the type of lug used for connecting an **aluminum cable** to a **copper bus bar**. Let's look at the options provided:

- **Bi-metallic cable lug:** A bi-metallic lug is specifically designed for connecting conductors of dissimilar metals. It typically consists of a barrel made of aluminum (where the aluminum cable is inserted) and a palm (the flat part that bolts onto the terminal or bus bar) made of copper. The joint between the aluminum barrel and copper palm is created using processes like friction welding or brazing, ensuring a strong electrical connection without direct contact between the aluminum cable and the copper palm/bus bar at the critical interface. This prevents galvanic corrosion.
- **Insulated sleeve lug:** This type of lug has an insulating sleeve, usually for environmental protection or preventing accidental contact. It doesn't address the issue of connecting dissimilar metals like aluminum and copper.
- **Copper lug:** A copper lug is designed for terminating copper cables. If you use a copper lug on an aluminum cable and connect it to a copper bus bar, you still have the direct contact between the aluminum cable and the copper lug/bus bar interface, leading to galvanic corrosion.
- **Bi-metallic cable connector:** While related, a connector typically joins two cables together (e.g., splicing an aluminum cable to a copper cable). A lug is specifically designed for terminating a cable onto a terminal, stud, or bus bar. The question specifies connecting an aluminum cable to a copper bus bar, which is a termination application where a lug is used. Although a bi-metallic connector serves a similar purpose of joining dissimilar metals, the term 'lug' is more appropriate for the connection to a bus bar.

Based on the analysis, the most appropriate component for connecting an aluminum cable to a copper bus bar is a **bi-metallic cable lug**.

Comparison of Lug Types for Aluminum to Copper Connections

Lug Type	Suitable for Al to Cu?	Reasoning
Bi-metallic cable lug	Yes	Designed to join dissimilar metals safely by separating them.
Insulated sleeve lug	No	Insulation is for protection, not dissimilar metal connection.
Copper lug	No	Direct contact between aluminum cable and copper lug/bus bar causes corrosion.
Bi-metallic cable connector	Yes (for joining cables)	Not typically used for terminating onto a bus bar (a lug is used).

Conclusion

To avoid galvanic corrosion and ensure a long-lasting, safe electrical connection when transitioning from an **aluminum cable** to a **copper bus bar**, a specialized connector that handles the dissimilar metals is required. The **bi-metallic cable lug** is precisely designed for this purpose, having an aluminum barrel for the aluminum cable and a copper palm for connection to the copper bus bar, with a transition zone preventing direct contact between the dissimilar metals at the connection point.

Revision Table: Bi-metallic Cable Lug Basics

Term	Description	Application Context
Bi-metallic Lug	A lug made from two different metals (typically aluminum and copper) joined together.	Connecting aluminum conductors to copper terminals/bus bars.
Galvanic Corrosion	Electrochemical process where one metal corrodes preferentially when in contact with another dissimilar metal in an electrolyte.	Occurs when aluminum and copper are in direct contact in moist environments.
Aluminum Cable	An electrical conductor made primarily of aluminum.	Often used for its lighter weight and lower cost compared to copper.
Copper Bus Bar	A solid metallic strip or bar used in electrical distribution boards, switchgear, and other equipment to conduct electricity.	Typically made of copper due to its high conductivity and corrosion resistance.

Additional Information: Preventing Electrical Connection Issues

Beyond using the correct **bi-metallic cable lug** for aluminum to copper connections, other factors are crucial for reliable electrical connections:

- **Proper Cable Preparation:** Ensuring the cable end is clean, stripped correctly, and free from damage.
- **Correct Crimping Tool:** Using the appropriate crimping tool and die for the specific lug and cable size to ensure a secure mechanical and electrical connection.
- **Applying Joint Compound:** For aluminum conductors, applying an anti-oxidant joint compound in the lug barrel helps break down oxide layers and prevent further oxidation, ensuring a low-resistance contact.

- **Correct Torque:** Tightening the lug bolt onto the terminal or bus bar to the manufacturer's specified torque ensures adequate contact pressure without damaging the lug or terminal.
- **Environmental Protection:** In harsh or wet environments, additional measures like heat shrink tubing or sealing compounds may be used to protect the connection from moisture.

Prepp

Your Personal Exams Guide