

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 kJ = 1000 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:

$$KE = \frac{1}{2}Mv^2$$

Where:

- 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

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 $\frac{3}{6} = \frac{1}{2}$ and $\frac{6}{12} = \frac{1}{2}$.
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 $SP < CP$. $\text{Loss} = CP - SP$.
- **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 $CP = \frac{SP}{(1 - \frac{L}{100})}$:

1. Substitute the given values into the formula: $CP = \frac{32}{(1 - \frac{20}{100})}$
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 = $(\frac{\text{Loss}}{\text{CP}}) \times 100\% = (\frac{8}{40}) \times 100\%$
- Loss Percentage = $(\frac{1}{5}) \times 100\% = 20\%$

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 % = $(\frac{\text{Profit}}{\text{CP}}) \times 100$	N/A
Loss %	Loss % = $(\frac{\text{Loss}}{\text{CP}}) \times 100$	N/A
SP (with Profit P%)	N/A	$SP = CP \times (1 + \frac{P}{100})$
SP (with Loss L%)	N/A	$SP = CP \times (1 - \frac{L}{100})$
CP (with Profit P%)	N/A	$CP = \frac{SP}{(1 + \frac{P}{100})}$
CP (with Loss L%)	N/A	$CP = \frac{SP}{(1 - \frac{L}{100})}$

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

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

Length left = 50 m – 47.25 m = 2.75 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

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 (x y z) b b (x y z) c c c (x 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 'xbycy'.

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 $60\ \Omega$.

Verifying the Result

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

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

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:

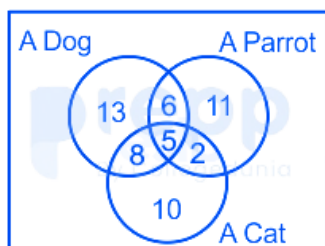
- **Series Connection:** Components are connected end-to-end, forming a single path for the current. The total resistance is the sum of individual resistances ($R_{eq} = R_1 + R_2 + R_3 + \dots$). The current is the same through all components.
- **Parallel Connection:** Components are connected across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

Your Personal Exams Guide

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

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

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

- 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:

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

- 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)
- 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:

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

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.

- Find the difference between the first and second numbers.
- 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).

Your Personal Exams Guide

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 ($Jkg^{-1}K^{-1}$) or Joules per kilogram per degree Celsius ($Jkg^{-1}C^{-1}$). For steel, it's given as $500Jkg^{-1}K^{-1}$.
- **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, $c = 500 \text{ Jkg}^{-1}K^{-1}$

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

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

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: $200Jkg^{-1}K^{-1}$ - This unit ($Jkg^{-1}K^{-1}$) is for specific heat capacity, not heat capacity (JK^{-1}). Also, the value is incorrect.
- Option 2: $1,250JK^{-1}$ - This matches our calculated value and has the correct unit for heat capacity.
- Option 3: $125JK^{-1}$ - This value is incorrect. It seems like a calculation error might have occurred (e.g., $2.5 \times 50 = 125$).
- Option 4: $20Jkg^{-1}K^{-1}$ - This unit ($Jkg^{-1}K^{-1}$) 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 $1250JK^{-1}$.

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 1250 JK^{-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

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:

Decimal value = $256 + 0 + 64 + 32 + 16 + 0 + 4 + 2 + 0$

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,

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

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

$$\Rightarrow 16 - 3$$

$\therefore$ 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

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 $\rho \approx 1$ 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 ($\mathbf{B}$) to magnetic field strength ($\mathbf{H}$).

The magnetic flux density $\mathbf{B}$ is measured in Tesla (T), and the magnetic field strength $\mathbf{H}$ 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:

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

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 $\frac{C/m^2}{V/m} = \frac{C}{m^2} \times \frac{m}{V} = \frac{C}{Vm}$.

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

Substituting this into the unit expression:

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

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.

Your Personal Exams Guide

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:

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

Now, solve for N_{driver} :

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

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

$$N_{driver} = 12$$

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).

$$\text{Combined Rate} = \text{Rate of A} - \text{Rate of B}$$

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):

$$\text{Net Rate} = \text{Rate of A} - \text{Rate of B}$$

$$2 = 5 - r$$

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

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:

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.

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 * 30°/hour = 90° 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 Jg⁻¹K⁻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 Jg⁻¹K⁻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 kJ = 1000 J
So, $Q = 3.6 \text{ kJ} = 3.6 \times 1000 \text{ J} = 3600 \text{ J}$.
- The temperature rise is given in degrees Celsius (°C). A change in temperature in Celsius is the same as a change in temperature in Kelvin.
So, $\Delta T = 45 \text{ }^{\circ}\text{C} = 45 \text{ K}$.
- 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)	$Q = \frac{kA\Delta T}{d}$	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	$-0.3 + 0.6 = 0.3$

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

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).
- **&empty::** 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 $\cos \theta/\sin \theta$)

Your Personal Exams Guide

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 ÷) becomes $4 +$
- 15 (replace -) becomes $15 \div$
- 3 (replace x) 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.

$$32 + 5 = 37$$

$$\text{Then, } 37 - 1 = 36$$

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)
+	×
÷	+
-	÷
×	-

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

Volume (V) = Length × Width × 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:

Mass = Density × 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: Volume = length × width × height.
3. Use the density formula (Density = Mass/Volume) to find the mass by rearranging it to: Mass = Density × 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	Length × Width × Height
Kerosene Density	800 kg/m^3	Given
Kerosene Mass	8000 kg	Density × 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.

$$\text{Profit} = \text{SP} - \text{CP}$$

$$\text{Profit} = ₹64 - ₹25 = ₹39$$

The vendor makes a profit of ₹39 on 40 bananas.

Calculating the Percentage Profit

The percentage profit is calculated using the formula:

$$\text{Percentage Profit} = \left(\frac{\text{Profit}}{\text{CP}} \right) \times 100\%$$

Using the values we found:

$$\text{Percentage Profit} = \left(\frac{39}{25} \right) \times 100\%$$

$$\text{Percentage Profit} = 39 \times \left(\frac{100}{25} \right) \%$$

$$\text{Percentage Profit} = 39 \times 4\%$$

$$\text{Percentage Profit} = 156\%$$

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: $\text{SP} - \text{CP} = ₹64 - ₹25 = ₹39$.
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 = $\left(\frac{\text{Profit per banana}}{\text{CP per banana}}\right) \times 100\%$
- 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 $1/16$ Area of A0 = $1/16\text{ m}^2 = 0.0625\text{ m}^2 = 625\text{ cm}^2$.

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

Additional Information: Properties of Regular Hexagons

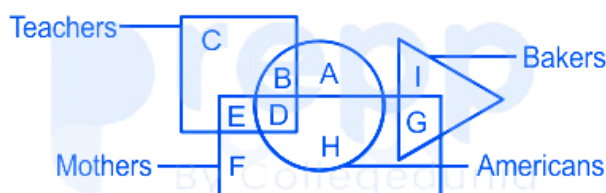
A regular hexagon is a fascinating shape with several interesting properties:

- 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):

$$V = l \times w \times h$$

$$V = 0.06 \text{ m} \times 0.08 \text{ m} \times 0.05 \text{ m}$$

$$V = (6 \times 10^{-2}) \text{ m} \times (8 \times 10^{-2}) \text{ m} \times (5 \times 10^{-2}) \text{ m}$$

$$V = (6 \times 8 \times 5) \times 10^{(-2-2-2)} \text{ m}^3$$

$$V = 240 \times 10^{-6} \text{ m}^3$$

$$V = 2.4 \times 10^{-4} \text{ m}^3$$

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/s^2 , but for simpler calculations in problems like this, it is often given as 10 m/s^2 .

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)	$P = I^2 R$, $P = VI$
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)	$P = I^2 R$	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)

Your Personal Exams Guide

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 $E = I^2 Rt$, it is a valid relation

derived from fundamental principles.

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

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 $\times$ 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.

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.

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 "_____ $^\circ\text{Celsius} = 167 \text{ Fahrenheit}$ " 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 $g_{\text{jupiter}} = 2.5 \times g_{\text{earth}}$. Substitute the given value for g_{earth} .

$$g_{\text{jupiter}} = 2.5 \times 10 \text{ m/s}^2$$

$$g_{\text{jupiter}} = 25 \text{ m/s}^2$$

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: $W_{\text{jupiter}} = m \times g_{\text{jupiter}}$. Substitute the mass of the satellite and the calculated gravity on Jupiter.

$$W_{\text{jupiter}} = 250 \text{ kg} \times 25 \text{ m/s}^2$$

$$W_{\text{jupiter}} = 6250 \text{ N}$$

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})	$2.5 \times g_{\text{earth}}$	Given relationship
Calculated g_{jupiter}	$2.5 \times 10 = 25 \text{ m/s}^2$	Step 1 result
Weight on Jupiter (W_{jupiter})	$m \times g_{\text{jupiter}}$	Formula for weight
Calculated W_{jupiter}	$250 \times 25 = 6250 \text{ N}$	Step 2 result

Your Personal Exams Guide

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:

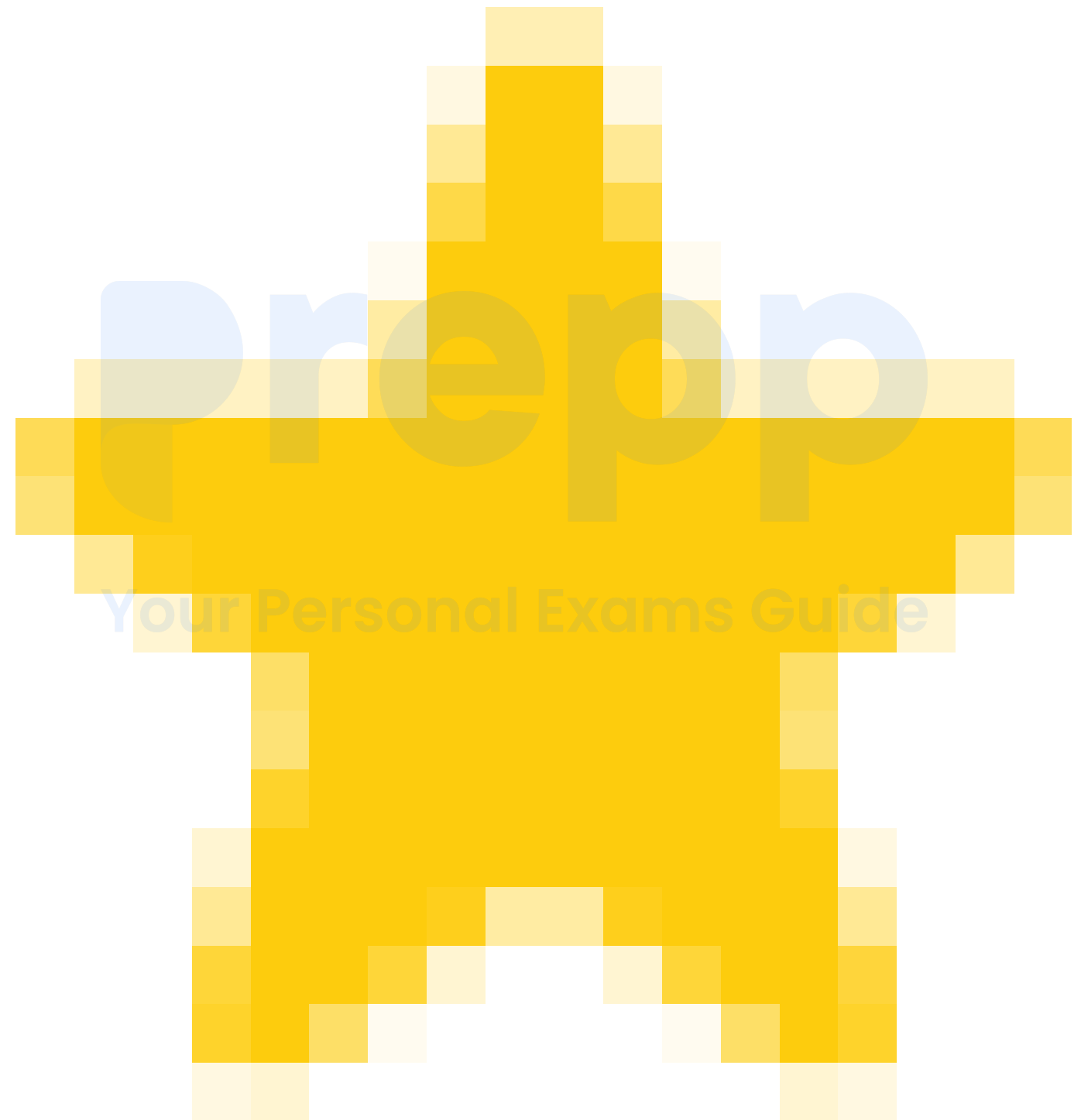
In option 1, 2 and 3, all the letters are same, while in option '4', H is missing and 'I' is present which is not there in other three options.

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 kΩ.
- 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:

$$R = 500 / 5 \times 10^{-3} \text{ (since the current is given in milliamperes and 'milli' refers to } 10^{-3}\text{)}$$

$$R = 100 \times 10^3 \Omega$$

$$R = 100 \text{ k}\Omega \text{ (since } 10^3 \text{ refers to 'kilo').}$$

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

Amount at the end of the first year = $P + I_1$

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.

Your Personal Exams Guide

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.
- **Axonometric 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.
- $2 \times$: 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 ' $2 \times$ ' 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 $2 \times \varnothing 6$ 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 $\text{Distance} = \text{Speed} \times \text{Time}$:

$$\text{Distance} = 64 \text{ km/h} \times \frac{5}{6} \text{ hours}$$

$$\text{Distance} = \frac{64 \times 5}{6} \text{ km} = \frac{320}{6} \text{ km} = \frac{160}{3} \text{ km}$$

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 $\text{Speed} = \text{Distance} / \text{Time}$:

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

$$\text{New Speed} = \frac{\frac{160}{3} \text{ km}}{\frac{2}{3} \text{ hours}}$$

$$\text{New Speed} = \frac{160}{3} \times \frac{3}{2} \text{ km/h}$$

$$\text{New Speed} = \frac{160}{2} \text{ km/h} = 80 \text{ km/h}$$

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:

⇒ Total degree of marks = 360°

⇒ Marks scored in Science = $\frac{82}{360} \times 1800 = 82 \times 5 = 410$

⇒ Marks scored in Mathematics = $\frac{84}{360} \times 1800 = 84 \times 5 = 420$

⇒ Difference = $420 - 410 = 10$

Therefore, the difference between the marks scored in Mathematics and Science is 10.

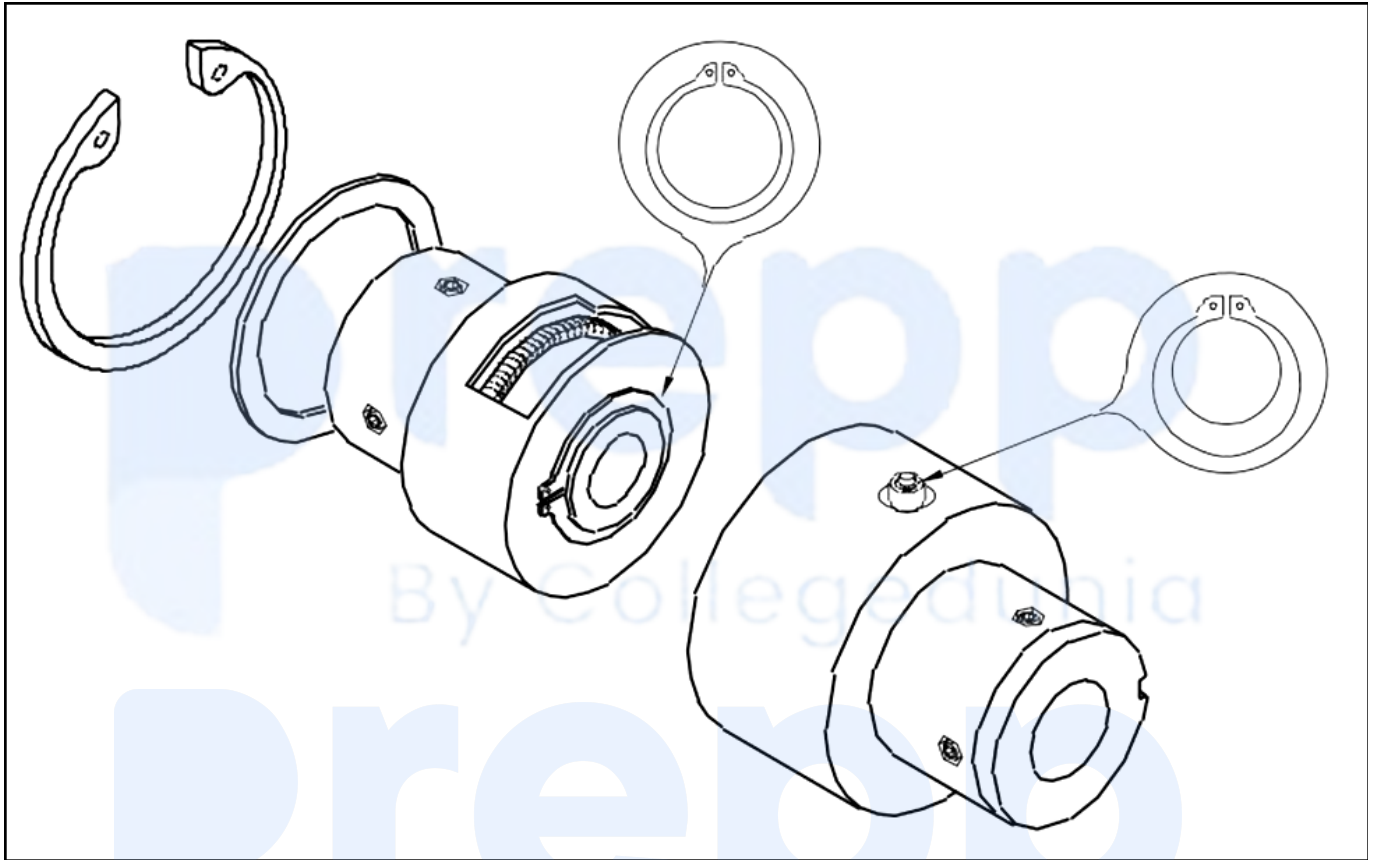
101. Answer: c

Explanation:

Explanation:

Snap ring/Circlips:

- Circlips are fastening devices used to provide shoulders for positioning or limiting the movement of parts in an assembly.
- Circlips are also called Retaining rings.
- The rings are generally made of materials having good spring properties so that the fastener may be deformed elastically to a considerable degree and still spring back to its original shape.
- This permits the circlips to spring back into a groove or other recess in a part or they may be seated on a part in a deformed condition so that they grip the part by functional means.
- Circlips are manufactured from spring steel with high tensile and yield strength.



Types:

Internal circlips:

- This type of ring is assembled in holes, bores, or housing.

External circlip:

- These types of rings are installed on shafts, pins, studs, and similar parts.

Material:

- Because retaining rings depend for their function largely on their ability to be deformed elastically during assembly and disassembly, the materials must have good spring properties.
- Circlips are manufactured from spring steel with high tensile and yield strength.

102. Answer: a

Explanation:

Understanding Variable Valve Timing (VVT)

Variable Valve Timing (VVT) is an automotive engine technology that allows the lift, duration, or timing (angle) of the intake or exhaust valves to be adjusted while the engine is in operation. This allows the engine to optimize performance, fuel economy, and emissions across different operating conditions (like idle, low speed, high speed, heavy load).

Purpose of Variable Valve Timing (VVT)

The main purpose of Variable Valve Timing (VVT) systems is to improve engine efficiency and performance by controlling how and when the engine valves open and close. Different VVT systems achieve this in various ways:

- Some systems primarily alter the valve timing, changing when the intake and exhaust valves open and close relative to the crankshaft position.
- Other more advanced systems can also alter the valve lift, which is the maximum distance the valve opens from its closed position.
- Some systems can alter the valve duration, which is the length of time the valve remains open.

Analyzing the Options

Let's look at the provided options in the context of Variable Valve Timing (VVT):

1. **Alter the valve lift:** Some Variable Valve Timing (VVT) systems, particularly the more sophisticated ones like those involving variable valve lift technology, are specifically designed to change how far the valve opens. This directly affects the amount of air-fuel mixture entering the cylinder or exhaust gases leaving it, optimizing combustion under different loads and speeds. Therefore, altering valve lift is indeed one of the purposes of Variable Valve Timing (VVT).
2. **Variation in valve size:** Variable Valve Timing (VVT) systems control the movement of the existing valves. They do not physically change the size or dimensions of the valves themselves while the engine is running. This option is incorrect.
3. **High vibration on valve head:** High vibration on the valve head is not a desired purpose of Variable Valve Timing (VVT). If anything, excessive vibration could indicate a problem with the valve train. This option is incorrect.
4. **No change in valve lift:** While some simpler Variable Valve Timing (VVT) systems might only adjust timing or duration, more advanced systems explicitly include the capability to change valve lift. Stating there is "no change in valve lift" is inaccurate for VVT systems designed to alter lift. Since altering valve lift is a known function of some VVT systems, this option is not universally true and contradicts a primary function of certain VVT implementations.

Considering the functions of various Variable Valve Timing (VVT) technologies, altering valve lift is a significant purpose achieved by many VVT systems to enhance engine performance and efficiency.

Conclusion

Based on the analysis of the options and the known capabilities of Variable Valve Timing (VVT) systems, altering the valve lift is a key purpose of some VVT implementations.

VVT Purpose	Description
Alter Valve Timing	Change when valves open/close relative to crankshaft.
Alter Valve Lift	Change how far valves open.
Alter Valve Duration	Change how long valves stay open.

Revision Table: Variable Valve Timing Aspects

Aspect	Explanation
Valve Timing	Refers to the specific crank angle at which a valve opens or closes. VVT can advance or retard timing.
Valve Lift	The maximum distance the valve moves away from its seat. Higher lift generally allows more flow.
Valve Duration	The amount of time (measured in crank degrees) a valve is open. Longer duration allows more flow over time.

Additional Information: Benefits of Variable Valve Timing

Implementing Variable Valve Timing (VVT) technology offers several benefits for internal combustion engines:

- **Improved Performance:** VVT can optimize valve operation for better power delivery across different engine speeds, especially improving torque at low RPM and horsepower at high RPM.
- **Increased Fuel Efficiency:** By precisely controlling airflow into the engine, VVT can help the engine burn fuel more efficiently, leading to better gas mileage.
- **Reduced Emissions:** Optimized combustion through VVT can help reduce the production of harmful exhaust emissions.
- **Smoother Idle:** VVT can adjust valve timing at idle for a more stable and smooth engine operation.

Different manufacturers have developed various Variable Valve Timing (VVT) systems, often with proprietary names like VVT-i (Toyota), VANOS (BMW), VTEC (Honda), VVEL (Nissan), etc. These systems may focus primarily on timing, or incorporate the ability to alter lift and duration as well.

103. Answer: d

Explanation:

Understanding the Automotive Cooling System

The cooling system in a car is crucial for maintaining the engine's operating temperature. It prevents the engine from getting too hot, which could cause significant damage. This system circulates coolant, a

mixture of antifreeze and water, through passages in the engine to absorb heat. The hot coolant is then sent to the radiator to cool down before being circulated back into the engine.

Role of Hoses in the Cooling System

Flexible hoses are used to connect the different components of the cooling system. These hoses must be able to withstand varying temperatures and pressures. The two main hoses connected to the radiator are typically the upper hose and the lower hose.

Identifying the Lower Hose Pipe

The location and function of the upper and lower radiator hoses are distinct:

- The **upper hose pipe** typically connects the engine's coolant outlet (often near the thermostat housing) to the inlet tank at the top of the radiator. This hose carries hot coolant from the engine to the radiator to be cooled.
- The **lower hose pipe** connects the outlet tank at the bottom of the radiator to the inlet of the water pump. This hose carries the cooled coolant from the radiator back to the water pump, which then pumps it back into the engine.

Analyzing the Lower Hose Pipe Connection Options

Let's examine the given options based on the known connections in a standard cooling system:

1. **thermostat and radiator:** The thermostat is located near the engine's coolant outlet. The upper hose connects the thermostat housing (engine outlet) to the radiator inlet. The lower hose does not connect the thermostat directly to the radiator.
2. **thermostat and bypass valve:** The bypass valve is usually internal to or associated closely with the thermostat housing and allows coolant to bypass the radiator when the engine is cold. Hoses typically connect major components, not internal parts or specific valves within a component assembly like the thermostat.
3. **thermostat and water pump:** The thermostat is on the hot side of the system (engine outlet), while the water pump inlet receives cooled coolant from the radiator. These two components are not directly connected by a hose in this manner.
4. **water pump inlet and radiator outlet tank:** This describes the correct path of the cooled coolant exiting the radiator. The coolant leaves the bottom tank (outlet) of the radiator and is drawn into the inlet side of the water pump, ready to be pumped back through the engine. The lower hose pipe makes this connection.

Correct Connection for the Lower Hose Pipe

Based on the analysis of the cooling system's flow path and the function of the lower hose, the lower hose pipe is connected between the **water pump inlet** and the **radiator outlet tank**.

Revision Table: Automotive Cooling System Hoses

Hose Type	Connects From	Connects To	Coolant Temperature
Upper Hose	Engine Outlet (often via Thermostat Housing)	Radiator Inlet Tank (Top)	Hot coolant from engine
Lower Hose	Radiator Outlet Tank (Bottom)	Water Pump Inlet	Cooled coolant from radiator

Additional Information on Cooling System Operation

The cooling system operates as a continuous loop. The water pump circulates coolant through the engine. The hot coolant then flows through the upper hose to the radiator. Inside the radiator, heat is transferred from the coolant to the air flowing through the radiator fins. The cooled coolant then flows out of the radiator's lower tank and is drawn back to the water pump via the lower hose pipe, completing the cycle. The thermostat controls when coolant flows through the radiator, ensuring the engine reaches its optimal operating temperature efficiently.

104. Answer: d

Explanation:

Understanding Crankshaft Movement in Engines

In an internal combustion engine, the crankshaft is a critical component that converts the linear motion of the pistons into rotational motion. While primarily rotating, the crankshaft also experiences forces that can cause it to move axially (forward and backward) within the engine block. This axial movement needs to be controlled to ensure proper operation and prevent damage to other engine parts.

The Role of the Main Thrust Bearing

The component specifically designed to limit this axial movement of the crankshaft is the **main thrust bearing**. Main bearings support the crankshaft's rotation, and one or more of these main bearings are designed as thrust bearings. These special bearings have surfaces that bear against thrust surfaces on the crankshaft and the engine block (or bearing cap), preventing excessive forward or backward sliding.

- The **main thrust bearing** is typically located at a specific main bearing journal, often near the center or at one end of the crankshaft.
- It contains flanges or separate thrust washers that engage with corresponding surfaces on the crankshaft cheeks or counterweights.
- These surfaces limit how far the crankshaft can travel axially, maintaining its correct position relative to the connecting rods, pistons, and other engine components.

Therefore, the **main thrust bearing** effectively limits the distance the crankshaft can travel axially and allows it to slide within a controlled range forward or backward in the block.

Analyzing Incorrect Options

Let's look at why the other options are not correct in this context:

- **Plain washer:** A plain washer is a simple flat ring used primarily to distribute load under a bolt or nut or act as a spacer. It is not designed to handle or limit the significant axial thrust forces of a crankshaft.
- **Dummy plates:** This term is not standard engine terminology related to crankshaft axial control. It might refer to temporary covers or plates used during assembly or manufacturing, but not a functional component that limits crankshaft movement.
- **Spring washer:** A spring washer (like a lock washer) is designed to provide tension to prevent fasteners from loosening under vibration. It has no role in limiting the axial travel of a rotating crankshaft within its bearings.

Conclusion on Crankshaft Limitation

Based on the function of engine components, the component that limits the axial distance of the crankshaft and controls its forward/backward sliding is the **main thrust bearing**.

Component Function Comparison		
Component	Primary Function	Limits Crankshaft Axial Movement?
Main Bearing (Radial)	Supports crankshaft rotation (radial load)	No
Main Thrust Bearing	Limits crankshaft axial movement (axial load)	Yes
Plain Washer	Load distribution, spacing	No
Spring Washer	Prevents fastener loosening	No

Revision Table: Key Engine Bearing Concepts

Engine Bearing Types and Roles

Bearing Type	Location	Primary Load Supported
Main Bearings	Between crankshaft journals and engine block/caps	Radial load from combustion forces
Connecting Rod Bearings	Between connecting rod journals and connecting rod big end	Radial load from piston/combustion forces
Thrust Bearings (part of main bearings)	Often one main bearing location on crankshaft	Axial load (thrust) from clutch/torque converter, helical gears, etc.

Additional Information: Why Axial Movement Occurs

Axial thrust forces on the crankshaft can arise from several sources:

- **Manual Transmission Clutch:** Depressing the clutch pedal applies a significant thrust load to the crankshaft.
- **Automatic Transmission Torque Converter:** While less severe than a clutch, the torque converter can also exert some axial force.
- **Helical Gears:** Gears with angled teeth (like timing gears or oil pump drive gears) create axial thrust when they mesh.
- **Engine Operation:** Imbalances or specific engine designs can also contribute to minor axial forces.

The **main thrust bearing** is crucial for managing these forces and keeping the crankshaft properly positioned, preventing damage to main bearings, connecting rod bearings, and the engine block itself.

105. Answer: d

Explanation:

Understanding the Full Form of CRDi in Engines

The question asks for the full form of the acronym **CRDi**, which is commonly associated with modern diesel engines.

Let's look at the options provided:

1. Common Rod Direct Inductive
2. Common Road Direct Injection
3. Common Road Drive Inspection
4. Common Rail Direct Injection

The acronym **CRDi** stands for **Common Rail Direct Injection**.

Common Rail Direct Injection is a modern fuel injection system used in diesel engines. It differs from older diesel injection systems where fuel was supplied to each cylinder individually by a pump. In a common rail system, a single, high-pressure pump pressurizes the fuel and stores it in a common rail or pipe, which acts as an accumulator. This common rail then distributes the high-pressure fuel to the injectors located in each cylinder. The injectors, which are typically electronically controlled (often using solenoids or piezoelectric actuators), precisely meter and spray the fuel directly into the combustion chamber at the optimal time and pressure.

This system offers several advantages over older systems, including:

- Higher injection pressures, leading to better fuel atomization and more efficient combustion.
- More precise control over injection timing and quantity, allowing for multiple injections per combustion cycle (e.g., pilot, main, and post-injections).
- Reduced emissions (particulate matter and NOx).
- Improved fuel efficiency.
- Lower engine noise and vibration.

Therefore, the correct full form of CRDi is **Common Rail Direct Injection**.

Revision Table: Key Diesel Engine Concepts

Term	Full Form	Brief Description
CRDi	Common Rail Direct Injection	A modern diesel fuel injection system using a common high-pressure fuel rail.
TDi	Turbocharged Direct Injection	Refers to diesel engines that are both turbocharged and use direct injection (can include CRDi or older pump-injector systems).
MPFI	Multi-Point Fuel Injection	A fuel injection system for petrol engines where fuel is injected into the intake port of each cylinder.

Additional Information on Common Rail Direct Injection Systems

The development of **Common Rail Direct Injection (CRDi)** technology has significantly advanced diesel engine performance and environmental compliance. Early diesel engines used mechanical injection systems, which were less precise. The advent of electronic control units (ECUs) made systems like CRDi possible.

Key components of a typical **CRDi** system include:

- **Fuel Tank:** Stores the diesel fuel.
- **Low-Pressure Fuel Pump:** Draws fuel from the tank and sends it to the high-pressure pump.
- **Fuel Filter:** Removes contaminants from the fuel.
- **High-Pressure Fuel Pump:** Pressurizes the fuel to very high levels (up to 2500+ bar in modern systems).

- **Common Rail:** A reinforced pipe that stores the high-pressure fuel and distributes it to the injectors. It dampens pressure fluctuations.
- **Fuel Injectors:** Electronically controlled valves that inject the precise amount of fuel into the cylinder at the right time and pressure.
- **Engine Control Unit (ECU):** The "brain" of the system, receiving data from sensors (engine speed, load, temperature, etc.) and controlling the high-pressure pump and injectors.

The precision offered by **CRDi** systems is crucial for meeting stringent emission standards and improving overall engine efficiency compared to older diesel technologies.

106. Answer: c

Explanation:

Understanding Valve Assembly Components

Valve assemblies are crucial parts of many mechanical systems, particularly in engines where they control the flow of gases into and out of the combustion chamber. A typical valve assembly involves several components working together. The question asks which of the given options is generally considered **not** part of the valve assembly.

Analyzing Valve Assembly Components

Let's examine the options provided and their typical roles in a valve system:

- **Stem:** The stem is the long, cylindrical part of the valve. It connects the valve head to the operating mechanism (like a camshaft or rocker arm). The stem slides within a valve guide, helping to keep the valve aligned. The stem is an integral part of the valve itself.
- **Bearing:** Bearings are typically used to support rotating or sliding shafts in mechanisms, reducing friction. While bearings are essential in the valve operating mechanism (e.g., camshaft bearings, rocker arm bearings), they are generally **not** considered a component of the valve assembly itself. The valve assembly primarily consists of the valve and its immediate associated parts like the guide, seat, spring, retainer, and locks.
- **Spring:** Valve springs are commonly used in many valve applications, especially in internal combustion engines. Their primary function is to return the valve to its closed position after it has been opened by the operating mechanism. Although critical for the proper function of the valve system, the spring is a component that acts *on* the valve rather than being an integral *part of* the valve itself (like the stem or head). In some contexts, it might be considered a separate, though necessary, component of the valve train rather than the valve assembly strictly defined as the valve and its immediate seating/guiding parts.
- **Head:** The head is the large, flat or slightly dished part at one end of the valve stem. This is the part that seats against the valve seat in the cylinder head to seal the port and prevent flow when the valve is closed. The head is an essential and integral part of the valve itself.

Comparing the Options

Considering the typical definition of a valve assembly, the stem and head are fundamental components of the valve itself. The spring is a separate component that assists the valve's operation by providing the force to close it. The bearing is typically part of the operating mechanism that moves the valve, not part of the valve assembly itself.

However, based on the options and the component commonly considered separate from the core valve parts (stem and head), the spring is presented as the item not part of the valve assembly in this context, distinguishing it from the valve's integral structure (stem, head).

Conclusion

Based on the analysis of the components and considering that stem and head are integral parts of the valve, and while springs are crucial for valve operation, they can sometimes be considered external to the valve assembly itself in certain classifications, the component listed that is **not** part of the valve assembly is the spring.

Component	Typical Role	Part of Valve Assembly?
Stem	Connects head to mechanism, slides in guide	Integral part of the valve
Bearing	Supports rotating/sliding parts in mechanism	Typically not part of the valve assembly itself
Spring	Returns valve to closed position	Often acts on the valve, sometimes considered separate from the valve itself (stem + head)
Head	Seals against the valve seat	Integral part of the valve

Thus, the component that is generally considered not part of the valve assembly among the given options is the spring.

Valve Assembly Revision Table

Reviewing the key components associated with a valve assembly helps solidify understanding.

- **Valve:** Consists of the stem and head.
- **Valve Guide:** Tube that the valve stem slides in, fitted into the cylinder head.
- **Valve Seat:** Machined surface in the cylinder head (or insert) that the valve head seals against.
- **Valve Spring:** Provides force to close the valve.
- **Spring Retainer (or Collar):** Sits on top of the spring.
- **Valve Locks (or Keepers):** Small split pieces that fit into grooves on the valve stem and hold the retainer, compressing the spring.

While all these parts work together in a valve train, the question specifically asks what is **not** part of the 'valve assembly', which, based on the options, points towards components that are either external mechanisms (like bearings) or components that act upon the valve but aren't physically part of the valve itself (like springs, in this specific context).

Additional Information on Valve Components and Function

Valves are critical for controlling fluid or gas flow. In internal combustion engines, intake valves allow air-fuel mixture (or just air) into the cylinder, and exhaust valves allow burnt gases out. Their precise timing is controlled by the camshaft or other actuation mechanisms.

The stem and head are forged or machined as a single piece in most engine valves. The material is chosen to withstand high temperatures and pressures.

Valve springs are designed to provide the correct closing force to ensure the valve seals properly against the seat, even at high engine speeds where inertia can cause the valve to 'float' or bounce.

Understanding the function of each component helps in diagnosing issues within the valve train.

107. Answer: d

Explanation:

Understanding Maximum Fuel Pressure in Unit Injector Systems

The question asks about the **maximum fuel pressure** achieved in a **unit injector system**. Unit injector systems, also known as Pumpe Düse, are a type of diesel fuel injection system that integrates the fuel pump and the injector into a single component located in the cylinder head.

High fuel pressure is critical in diesel engines for several reasons:

- To overcome the high pressure already present in the combustion chamber.
- To effectively atomize the fuel, breaking it down into very fine droplets for efficient mixing with air.
- To ensure proper penetration of the fuel spray into the combustion chamber.

The **unit injector system** is known for its ability to generate very high injection pressures directly at the point of injection for each cylinder. This allows for better fuel atomization and control over the injection process compared to older systems.

Typical **fuel pressure** ranges in modern diesel injection systems vary significantly. Older mechanical systems operated at much lower pressures, while modern common rail and unit injector systems operate at very high pressures.

In a **unit injector system**, the fuel pressure is generated individually for each cylinder by a cam-driven plunger within the injector unit. This design allows for extremely high pressures to be achieved.

The maximum pressure achieved in a **unit injector system** can be very high. Let's consider the typical ranges:

- Older diesel systems: often below 1000 bar.
- Common Rail systems: typically range from 1000 bar up to over 2500 bar in the latest generations.
- Unit Injector systems (Pumpe Düse): known for achieving pressures up to 2000 bar or even higher.

Considering the options provided, the possible maximum pressures are 1600 bar, 400 bar, 800 bar, and 2400 bar. Comparing these values with the typical capabilities of **unit injector systems**:

- 400 bar and 800 bar are too low for modern high-pressure diesel injection systems like unit injectors.
- 1600 bar is a plausible high pressure, but unit injector systems are capable of going higher.
- 2400 bar represents the upper limit of what can be achieved by advanced **unit injector systems**. While 2000 bar is a commonly cited figure, some later generation systems did push towards higher pressures.

Therefore, among the given options, **2400 bar** represents the **maximum fuel pressure** that can be achieved in a high-performance **unit injector system**.

Revision Table: Diesel Injection Pressure Comparison

Injection System Type	Typical Maximum Fuel Pressure
Older Mechanical Systems	~800 – 1000 bar
Early Common Rail / Unit Injector	~1350 – 1600 bar
Modern Common Rail / Unit Injector	Up to 2000 – 2500+ bar

Additional Information on Unit Injector Systems

The **unit injector system** (Pumpe Düse) was primarily used by the Volkswagen Group (Audi, VW, Skoda, Seat) for many years before they switched to common rail systems. Its main advantage was the ability to generate extremely high injection pressures independently for each cylinder, leading to efficient combustion and lower emissions for its time.

Key characteristics of a **unit injector system**:

- Pump and injector are combined into a single unit.
- Each cylinder has its own unit injector.
- The high pressure is generated right at the cylinder, minimizing pressure losses in fuel lines.
- Injection timing and quantity are controlled electronically.
- Achieves very high injection pressures, contributing to fine fuel atomization.

Despite their high performance in terms of pressure, **unit injector systems** had drawbacks, including noise, complexity in the cylinder head, and limitations in achieving multiple, flexible injection events compared to common rail systems, which became necessary for meeting stricter emissions standards.

108. Answer: c

Explanation:

Understanding Piston Rings in Engines

Piston rings are crucial components found in the pistons of internal combustion engines. They have several important functions, primarily sealing the combustion chamber, controlling oil consumption, and transferring heat from the piston to the cylinder wall.

A typical piston assembly uses several rings placed in grooves around the piston's circumference. These rings work together to ensure the engine operates efficiently and reliably.

Common Types of Piston Rings

There are typically two main categories of piston rings:

- **Compression Rings:** Located near the top of the piston, these rings form a seal between the piston and the cylinder wall. This seal prevents combustion gases from escaping into the crankcase (blow-by), maintaining cylinder pressure for efficient power generation.
- **Oil Rings:** Located below the compression rings, the oil ring's primary function is to scrape excess lubricating oil from the cylinder walls as the piston moves. This prevents oil from entering the combustion chamber, where it would burn and cause smoke, deposits, and reduce engine performance.

Analyzing the Options and their Relation to Piston Rings

Let's examine each option provided:

1. **Oil ring:** As discussed above, an oil ring is a standard type of piston ring. Its function is directly related to managing lubrication on the cylinder wall in conjunction with the piston's movement. Therefore, the oil ring is related to piston rings.
2. **Wiper ring:** A wiper ring is essentially a type of oil ring, sometimes specifically referring to the lower portion of a multi-piece oil control ring or a secondary compression ring designed with a scraping edge. Its function is also to control oil film on the cylinder wall. Thus, a wiper ring is related to piston rings.
3. **Clip ring:** A clip ring, also known as a circlip or snap ring, is a type of fastener. In an engine piston assembly, clip rings are often used to secure the piston pin (also called a gudgeon pin) within the piston bosses. The piston pin connects the piston to the connecting rod. Clip rings fit into grooves inside the piston boss bore to prevent the piston pin from sliding out sideways and damaging the cylinder walls. While used in the piston assembly, a clip ring does not go around the circumference of the piston and does not perform the sealing or oil control functions of a piston ring. Therefore, a clip ring is a fastener, not a piston ring.

4. Compression ring: As explained earlier, a compression ring is one of the fundamental types of piston rings. Its primary role is to seal the combustion chamber. Therefore, the compression ring is directly related to piston rings.

Conclusion

Based on the analysis of each option, the Oil ring, Wiper ring (as a type of oil ring/control ring), and Compression ring are all types or related to piston rings that fit into grooves around the piston's outer diameter. The Clip ring, however, is a fastener used to retain the piston pin and is not a ring that seals against the cylinder wall or controls oil in the same manner as piston rings.

Therefore, the item not related to piston ring function in the context of sealing and oil control on the cylinder wall is the Clip ring.

Ring Type	Relation to Piston Ring?	Primary Function
Oil ring	Yes	Scrape excess oil from cylinder wall.
Wiper ring	Yes (Type of Oil Ring)	Scrape excess oil from cylinder wall.
Clip ring	No	Retain piston pin in piston boss.
Compression ring	Yes	Seal combustion chamber gases.

Revision Table: Engine Rings Summary

Let's quickly summarize the key types of rings discussed:

- **Piston Rings:** Circular rings fitted into grooves on the piston circumference for sealing and oil control. Includes Compression and Oil rings.
- **Compression Ring:** Seals combustion gases.
- **Oil Ring:** Controls oil on cylinder walls.
- **Wiper Ring:** Often a term for an oil control ring or a secondary ring with scraping capability.
- **Clip Ring (Circlip/Snap Ring):** Fastener used to secure components like the piston pin within a bore or groove. Not a piston ring.

Additional Information: Piston Assembly Components

Understanding the context of piston rings involves knowing other parts of the piston assembly:

- **Piston:** The moving component that transfers force from expanding gases in the combustion chamber to the crankshaft via the connecting rod.
- **Piston Skirt:** The lower part of the piston that helps guide it in the cylinder bore.
- **Piston Crown:** The top surface of the piston exposed to combustion pressure.
- **Piston Grooves:** Channels machined into the piston circumference where the piston rings are fitted.
- **Piston Bosses:** The reinforced sections inside the piston where the piston pin is inserted.

- **Piston Pin (Gudgeon Pin):** A hollow steel pin connecting the piston to the small end of the connecting rod.
- **Connecting Rod:** Links the piston to the crankshaft, converting reciprocating motion to rotary motion.
- **Crankshaft:** Converts the linear motion of the piston into rotational motion.

Clip rings are specifically used to hold the piston pin in place within the piston bosses, preventing its axial movement. This is distinct from the function of piston rings which seal against the cylinder wall.

109. Answer: d

Explanation:

Understanding Vehicle Ignition Systems and Components

A vehicle's ignition system is crucial for starting and running the engine. Its primary function is to generate a high-voltage electrical spark that ignites the fuel-air mixture inside the engine's cylinders at precisely the right moment.

Components of a Typical Ignition System

Let's look at the options provided and see which one does not belong to the ignition system:

- **Spark Plug:** This is a vital component of the ignition system. The spark plug receives high voltage and creates the spark that ignites the fuel. It's the final step in the ignition process within the cylinder.
- **Ignition Coil:** The ignition coil is an electrical transformer that takes the relatively low voltage from the vehicle's battery (usually 12V) and converts it into a much higher voltage (tens of thousands of volts) required to jump the gap in the spark plug and create a spark. It is definitely a core part of the ignition system.
- **Distributor:** In older vehicles, the distributor is a mechanical device that directs the high voltage from the ignition coil to the correct spark plug in the correct firing order. Modern vehicles often use distributor-less ignition systems or coil-on-plug systems, but the distributor is historically and functionally part of the ignition system in many vehicles.
- **Rocker Arm:** The rocker arm is a component found in the valve train of an internal combustion engine. It pivots to transfer the lifting motion from the camshaft lobe (or a pushrod actuated by the camshaft) to open the intake or exhaust valves in the cylinder head. Rocker arms are essential for engine operation and combustion timing by controlling airflow, but they are not part of the *ignition* system, which is responsible for sparking the fuel, not controlling the valves.

Identifying the Non-Ignition System Component

Based on the functions described, the spark plug, ignition coil, and distributor are all directly involved in generating and delivering the spark needed for combustion. The rocker arm, however, is part of the valve train mechanism that controls gas flow in and out of the cylinders.

Therefore, the term that is not related to the ignition system of a vehicle is the rocker arm.

Component	Related to Ignition System?	Function
Spark plug	Yes	Creates the spark to ignite fuel.
Ignition coil	Yes	Increases voltage for the spark.
Distributor	Yes (in older systems)	Directs high voltage to correct spark plug.
Rocker arm	No	Opens/closes engine valves.

Revision Table: Key Engine Systems

System	Primary Function	Example Components
Ignition System	Ignite fuel-air mixture	Spark plug, Ignition coil, Distributor
Fuel System	Deliver fuel to engine	Fuel pump, Fuel injectors, Fuel filter
Exhaust System	Remove combustion gases	Exhaust manifold, Muffler, Catalytic converter
Cooling System	Regulate engine temperature	Radiator, Water pump, Thermostat
Lubrication System	Reduce friction and wear	Oil pump, Oil filter, Oil pan
Valve Train	Control air/gas flow via valves	Camshaft, Pushrods, Rocker arms, Valves

Additional Information: Modern Ignition Systems

While the distributor is part of traditional ignition systems, modern vehicles often use more advanced setups:

- **Distributor-less Ignition Systems (DIS):** These systems use multiple ignition coils (often one for every two cylinders) controlled electronically by the engine control unit (ECU).
- **Coil-on-Plug (COP) Systems:** This is the most common modern type, where each spark plug has its own dedicated ignition coil mounted directly on top of it. This eliminates the need for spark plug wires and allows for very precise ignition timing controlled by the ECU.

Understanding these different systems helps in diagnosing ignition-related issues in vehicles.

110. Answer: a

Explanation:

Explanation:

Oil Filter:

- Impurities can be introduced into a system as a result of mechanical wear, and external environmental influences.
- For this reason filters are installed in the hydraulic circuit to remove dirt particles from the hydraulic oil.
- The reliability of the system also depends on the cleanliness of the oil.
- An oil filter is used to remove the dust particles/contaminants from the oil before it is feeding to all the parts of the engine.

Types of Oil Filters:

There are 3 types of Oil Filters:

1. Cartridge Type Oil Filter
2. Edge Type Oil Filter
3. Centrifugal Type Oil Filter

Cartridge Type Oil Filter:

- Cartridge-type oil filter consists of filtering elements that are placed in the metallic casing for removing the impurities present in the lubricating oil and is mostly used in automobile engines.
- The filter elements with fine pores have been employed which has made it practicable to stop or arrest the particles of size down to within the region of 5 microns.
- In the filter, the oil enters the filter at the top of the casing and passes through the filter elements.
- The pure oil has to be passed through the porous metallic tube from where the oil goes to the outlet for circulation.
- A drain plug is also provided.

Edge Type Oil Filter:

- Edge-type oil filter is also called a stack-type Oil filter.
- In this oil filter, the oil is made to pass through several closely spaced discs which are mounted on the center spindle as well as the square rod alternately.
- The clearance or gap between two successive discs is a few microns only.
- The oil is allowed to pass through these spaces between the discs and due to the small spaces involved in between the discs, the impurities are left on the disk periphery itself from where they are removed periodically using operating the central knob.

Centrifugal Type Oil filter:

- In this Oil Filter, the impure oil from the engine enters the hollow Central spindle and has holes around its periphery.
- The impure/dirty oil comes out of these holes and fills the rotor casing after which it passes through the tubes at the ends to which Jets are attached.
- The oil under pressure passes through these jets and due to the reaction of which, it gives motion to the rotor casing in the opposite direction so that it starts rotating.

- The oil impinges on the outer stationary casing under heavy pressures where the impurities are retained there itself and cleaned oil falls below from where it is to be passed to other parts of the engine.

III. Answer: d

Explanation:

Understanding Valve Timing Diagrams

A valve timing diagram is a fundamental concept when studying the operation of internal combustion engines. It serves as a visual guide to understand precisely when the valves in an engine open and close during the engine cycle.

What a Valve Timing Diagram Represents

Specifically, a valve timing diagram is a graphical representation that illustrates the timing of the opening and closing events for the engine's intake valve and exhaust valve. This timing is shown in relation to the position or angle of the crankshaft, which dictates the position of the piston within the cylinder and the overall phase of the engine cycle.

The diagram plots the valve events (opening and closing) against the crankshaft rotation, typically measured in degrees. Key points on the diagram include:

- Intake Valve Opening (IO)
- Intake Valve Closing (IC)
- Exhaust Valve Opening (EO)
- Exhaust Valve Closing (EC)
- Top Dead Center (TDC)
- Bottom Dead Center (BDC)

These diagrams help engineers and mechanics visualize the duration each valve is open and how these durations overlap or separate during the 720 degrees of crankshaft rotation in a four-stroke engine cycle.

Analyzing the Options

Let's look at the provided options based on our understanding:

- Graphical representation of crankshaft position: While the diagram uses crankshaft position as a reference, it's not just about the crankshaft position itself, but rather valve events plotted against it.
- Graphical representation of camshaft location: The camshaft's rotation and lobe shape determine the valve timing, but the diagram shows the *result* of the camshaft's action on the valves' timing, not the camshaft's physical location.

- Graphical representation of piston position: Piston position is directly linked to crankshaft angle. However, the diagram's primary focus is on when the valves open and close relative to this position, not just plotting piston position.
- Graphical representation of the opening and closing of intake and exhaust valve of engine: This option accurately describes the core purpose and content of a valve timing diagram. It shows precisely when the intake valve and exhaust valve operate during the engine cycle as measured by crankshaft rotation.

Therefore, the most accurate description of a valve timing diagram is its depiction of the intake valve and exhaust valve timing.

Importance of Valve Timing

Accurate valve timing is critical for the performance, efficiency, and emissions of an internal combustion engine. The exact timing and duration that the intake valve and exhaust valve are open affect how well the cylinder is filled with the air-fuel mixture and how effectively the burnt gases are expelled.

Revision Table: Engine Timing Components

Component	Role in Valve Timing Diagram
Crankshaft	Provides the angular reference (X-axis) for the diagram; measures the engine cycle progress.
Intake Valve	Its opening and closing points are key events shown on the diagram. Controls mixture entry.
Exhaust Valve	Its opening and closing points are also key events shown on the diagram. Controls exhaust exit.
Camshaft	Mechanism that <i>*causes*</i> the valve opening and closing events shown; rotates synchronously with the crankshaft (at half speed in a 4-stroke).
Piston	Its position (TDC, BDC) corresponds to specific crankshaft angles, which are reference points on the diagram for valve events.

Additional Information on Engine Valve Timing

Valve timing diagrams often show periods where both the intake valve and exhaust valve are open simultaneously. This is known as valve overlap. Valve overlap typically occurs around TDC at the end of the exhaust stroke and beginning of the intake stroke. It helps with scavenging (pushing out remaining exhaust gases with the incoming mixture) at higher engine speeds, although too much overlap can lead to inefficiencies at lower speeds.

Modern engines often use Variable Valve Timing (VVT) systems. These systems can alter the valve timing diagram dynamically based on engine operating conditions like speed and load. This allows engineers to

optimize engine performance, fuel economy, and emissions across a wider range of operations than fixed valve timing.

112. Answer: d

Explanation:

Understanding Three-Way Catalytic Converters

A three-way catalytic converter is a crucial component in modern vehicle exhaust systems. Its primary function is to reduce the emission of harmful pollutants produced during the combustion process in the engine. These pollutants include Hydrocarbons (HC), Carbon Monoxide (CO), and Nitrogen Oxides (NOx).

How Three-Way Catalytic Converters Work

The "three-way" in the name refers to the fact that it simultaneously performs three main reactions:

1. Reduction of Nitrogen Oxides (NOx) to nitrogen (N₂).
2. Oxidation of Carbon Monoxide (CO) to carbon dioxide (CO₂).
3. Oxidation of unburnt Hydrocarbons (HC) to carbon dioxide (CO₂) and water (H₂O).

These reactions occur as the hot exhaust gases pass over a ceramic structure coated with precious metals like platinum, palladium, and rhodium, which act as catalysts.

Conversion Products of HC, CO, and NOx

Let's look at the specific conversions that take place for each pollutant:

- **Hydrocarbons (HC):** These are unburnt fuel molecules. The catalytic converter oxidizes them.
Reaction: $\text{HC} + \text{O}_2 \rightarrow \text{H}_2\text{O} + \text{CO}_2$
- **Carbon Monoxide (CO):** This is a poisonous gas resulting from incomplete combustion. The converter oxidizes it.
Reaction: $\text{CO} + \text{O}_2 \rightarrow \text{CO}_2$
- **Nitrogen Oxides (NOx):** These are formed at high temperatures in the engine. The converter reduces them. NOx represents various oxides like NO and NO₂.
Reaction: $\text{NO}_x \rightarrow \text{N}_2 + \text{O}_2$ (simplified)

Considering the combined effect of these reactions, the harmful gases (HC, CO, and NOx) are converted into much less harmful substances: water (H₂O), nitrogen gas (N₂), and carbon dioxide (CO₂).

Three-Way Catalytic Converter Conversions

Input Pollutant	Conversion Process	Output Product(s)
Hydrocarbons (HC)	Oxidation	Water (H ₂ O), Carbon Dioxide (CO ₂)
Carbon Monoxide (CO)	Oxidation	Carbon Dioxide (CO ₂)
Nitrogen Oxides (NOx)	Reduction	Nitrogen (N ₂)

Therefore, in a three-way catalytic converter, HC, CO, and NOx gases are primarily converted into water, nitrogen, and carbon dioxide.

Revision Table: Catalytic Converter

Key Aspect	Description
Function	Reduces harmful exhaust emissions (HC, CO, NOx)
Type	Three-way catalytic converter
Input Gases	Hydrocarbons (HC), Carbon Monoxide (CO), Nitrogen Oxides (NOx)
Output Gases	Water (H ₂ O), Nitrogen (N ₂), Carbon Dioxide (CO ₂)
Catalysts	Platinum, Palladium, Rhodium

Additional Information on Emission Control

Catalytic converters are a critical part of the overall emission control system in vehicles. Their effectiveness depends on the engine running at the correct air-fuel ratio (stoichiometric ratio) because the reduction and oxidation reactions require a balance of oxygen.

- Modern vehicles use oxygen sensors in the exhaust system to monitor the oxygen levels and provide feedback to the engine control unit (ECU). This allows the ECU to adjust the fuel injection and maintain the optimal air-fuel ratio for the catalytic converter to function efficiently.
- While carbon dioxide (CO₂) is produced, it is a greenhouse gas, but it is significantly less toxic than carbon monoxide or unburnt hydrocarbons. The primary goal of the catalytic converter is to eliminate the immediate, poisonous air pollutants.
- NOx formation is favoured at high temperatures and pressures inside the engine cylinder. Other technologies like Exhaust Gas Recirculation (EGR) are also used alongside catalytic converters to reduce NOx by lowering combustion temperatures.

113. Answer: b

Explanation:

Understanding MOSFET: Full Form Explained

The question asks for the full form of the acronym MOSFET. MOSFETs are crucial components in modern electronics, widely used in digital circuits and power electronics. Understanding their full name helps clarify their basic structure and operation.

What is the Full Form of MOSFET?

The acronym MOSFET stands for **Metal Oxide Semiconductor Field Effect Transistor**. This name describes the key parts and operating principle of the device.

- **Metal:** Refers to the gate electrode material, which was historically metal but is now often polysilicon.
- **Oxide:** Refers to the insulating layer (usually silicon dioxide, SiO_2) beneath the gate electrode.
- **Semiconductor:** Refers to the substrate material (typically silicon) where the current flows.
- **Field Effect Transistor (FET):** Indicates that the current flow between the source and drain terminals is controlled by an electric field produced by the voltage applied to the gate terminal.

Analyzing the Given Options

Let's look at the options provided and determine which one correctly expands the acronym MOSFET.

- **Option 1: Metal Oxide Semiconductor Final End Transistor** - This option is incorrect. "Final End" is not part of the standard terminology for this type of transistor.
- **Option 2: Metal Oxide Semiconductor Field Effect Transistor** - This option correctly represents the full form of MOSFET, detailing its construction (Metal, Oxide, Semiconductor) and its operating principle (Field Effect Transistor).
- **Option 3: Metal Oxide Semiconductor Fusion Effect Transistor** - This option is incorrect. "Fusion Effect" is not related to the operation of a MOSFET.
- **Option 4: Metal Offset Flow Effect Thermistor** - This option is incorrect. A thermistor is a temperature-sensitive resistor, which is a completely different type of electronic component from a MOSFET. Also, "Offset Flow" is not standard terminology.

Based on the standard definition and the analysis of the options, the correct full form of MOSFET is Metal Oxide Semiconductor Field Effect Transistor.

Revision Table: Key Terms

Term	Definition/Explanation
MOSFET	Metal Oxide Semiconductor Field Effect Transistor
Transistor	A semiconductor device used to amplify or switch electronic signals and electrical power.
Semiconductor	A material with electrical conductivity between that of a conductor and an insulator. Silicon is a common example.
Field Effect	The control of current flow in a semiconductor by an electric field.

Additional Information about MOSFETs

MOSFETs are voltage-controlled devices, meaning a voltage applied to the gate controls the current between the drain and source. This is different from Bipolar Junction Transistors (BJTs), which are current-controlled devices.

Key characteristics and uses of MOSFETs include:

- High input impedance, making them suitable for cascading multiple stages without significant signal loss.
- Used extensively in digital logic gates (like CMOS technology) due to their low power consumption in static states.
- Used in power electronics for switching applications due to their high switching speed.
- Available in two main types: Enhancement mode and Depletion mode.
- Available in two channel types: n-channel and p-channel.

Your Personal Exams Guide

114. Answer: c

Explanation:

Understanding Brazing Filler Metals

Brazing is a metal-joining process where two or more metal items are joined together by melting and flowing a filler metal into the joint. The filler metal has a lower melting point than the adjacent base metals. Unlike welding, the base metals themselves do not melt during brazing. The molten filler metal flows into the gap between the base metals by capillary action and, upon cooling, forms a strong metallurgical bond.

The choice of filler metal is crucial in brazing as it determines the strength, corrosion resistance, and other properties of the final joint. Brazing filler metals are typically alloys, designed to have specific melting ranges and flow characteristics suitable for different base metals and applications.

Composition of Common Brazing Alloys

Brazing filler metals come in various compositions, but some common types are widely used. Among these, alloys based on copper are very prevalent. When copper is alloyed with zinc, it forms brass, and specific brass compositions are used as brazing filler metals, often referred to as brass brazing rods or spelter.

Copper and zinc alloys are popular choices for brazing steel, copper, brass, and bronze. They offer good joint strength and are relatively cost-effective. The proportion of copper and zinc can vary, influencing the melting point and flow properties of the alloy.

Other common brazing filler metals include:

- Silver-based alloys (often contain copper, zinc, cadmium, tin) – lower melting point, high strength, good flow.
- Copper-phosphorus alloys (for copper to copper or copper to brass) – self-fluxing on copper.
- Aluminum-silicon alloys (for aluminum).
- Nickel alloys (for stainless steels and nickel alloys).

However, the question asks about filling metals used *in an alloy* during brazing, and copper and zinc are indeed primary constituents in a significant class of brazing filler alloys (brass brazing alloys).

Why Copper and Zinc are Used in Brazing Alloys

Copper-zinc alloys (brass) are effective as brazing filler metals because their melting points are significantly lower than the melting points of common base metals like steel or copper, but higher than those used in soldering (like tin-lead alloys). This allows the base metals to remain solid while the filler metal melts and flows.

Key reasons include:

- Appropriate melting range for many common brazing applications.
- Good flow characteristics, allowing capillary action.
- Ability to form strong metallurgical bonds with various base metals.
- Relatively lower cost compared to silver-based alloys.

Comparison of Joining Filler Metals

It's helpful to compare brazing filler metals with those used in soldering.

Process	Common Filler Metal Components	Typical Melting Point Range
Soldering	Tin, Lead, Silver, Copper, Bismuth	Below 450°C (840°F)
Brazing	Copper, Zinc, Silver, Phosphorus, Nickel, Silicon	Above 450°C (840°F) and below base metal melting points

Based on the common compositions of brazing alloys, specifically brass brazing alloys, copper and zinc are widely used as the primary filling metals.

Brazing Alloys Revision Table

Alloy Type	Primary Components	Typical Use Cases
Brass Brazing Alloys	Copper, Zinc	Brazing steel, copper, brass, bronze, cast iron
Silver Brazing Alloys	Silver, Copper, Zinc, Cadmium, Tin, Manganese	Brazing steel, copper alloys, nickel alloys, carbides; often used for critical joints due to strength and low brazing temperature
Copper-Phosphorus Alloys	Copper, Phosphorus, Silver (optional)	Brazing copper to copper or copper to brass; self-fluxing on copper

Additional Information on Brazing

Brazing requires a filler metal that melts above 450°C but below the melting point of the base materials being joined. This differentiates it from soldering, which uses filler metals melting below 450°C.

A flux is often used in brazing to clean the surfaces, prevent oxidation during heating, and promote the flow of the molten filler metal. The flux needs to be active at the brazing temperature range of the chosen filler metal.

Common brazing methods include torch brazing, furnace brazing, induction brazing, resistance brazing, and dip brazing. The choice depends on the application, part size, and required production volume.

115. Answer: d

Explanation:

Understanding Engine Strokes and Valve Operation

Internal combustion engines, like those found in cars and motorcycles, work by completing a cycle of events to convert fuel into power. A common type is the four-stroke engine, which performs four distinct movements or strokes for every power delivery.

Each stroke involves specific movements of the piston inside the cylinder and the opening or closing of the intake and exhaust valves. The question asks about which stroke(s) have both valves closed.

The Four Strokes of an Engine Cycle

Let's break down the four strokes:

1. **Intake Stroke:** The piston moves downwards. The intake valve is open to allow the fuel-air mixture (or just air in diesel engines) to enter the cylinder. The exhaust valve is closed.
2. **Compression Stroke:** The piston moves upwards. Both the intake valve and the exhaust valve are closed. This movement compresses the fuel-air mixture into a smaller volume, increasing its pressure and temperature.
3. **Power Stroke:** Ignition occurs (spark plug fires in gasoline engines or fuel is injected and auto-ignites in diesel engines). The expanding hot gases push the piston downwards. Both the intake valve and the exhaust valve are closed during the main power generation phase.
4. **Exhaust Stroke:** The piston moves upwards. The intake valve is closed. The exhaust valve is open to allow the burnt gases to exit the cylinder.

Valve Positions During Engine Strokes

To clearly see when both valves are closed, let's summarize the valve positions in a table:

Engine Stroke	Intake Valve	Exhaust Valve
Intake	Open	Closed
Compression	Closed	Closed
Power	Closed	Closed
Exhaust	Closed	Open

As the table shows, both the intake and exhaust valves are closed during the Compression stroke and the Power stroke.

Analyzing the Question and Options

The question asks which stroke is produced when both valves are closed. Based on our understanding of the four-stroke cycle:

- **Compression:** Both valves are closed. This matches.
- **Power:** Both valves are closed. This also matches.
- **Exhaust:** The exhaust valve is open. This does not match.
- **Both A & B:** This refers to both Compression and Power strokes. Since both of these strokes occur when both valves are closed, this option is correct.

Therefore, both the Compression stroke and the Power stroke are produced when both valves are closed in a typical four-stroke engine.

Revision Table: Engine Stroke Summary

Stroke	Piston Movement	Intake Valve	Exhaust Valve	Cylinder Action
Intake	Downward	Open	Closed	Draws in fuel-air mixture/air
Compression	Upward	Closed	Closed	Compresses mixture
Power	Downward	Closed	Closed	Ignition & expansion push piston
Exhaust	Upward	Closed	Open	Pushes out burnt gases

Additional Information on Engine Operation

Understanding the timing of the valves is crucial for efficient engine operation. The camshaft controls when the intake and exhaust valves open and close, synchronized with the piston's movement via the crankshaft.

While the basic description states valves are fully closed during Compression and Power, in reality, valves might open slightly before the piston reaches the very top or bottom of its stroke (valve overlap) to improve performance, especially at higher speeds. However, for the fundamental definition of the strokes, Compression and Power are characterized by having both valves closed during the critical phases.

The closed valves during Compression are necessary to trap the mixture and allow it to be compressed. During Power, they remain closed to contain the explosion and direct all force downwards onto the piston.

116. Answer: d

Explanation:

Understanding the Engine Starting Motor

The question asks to identify the component responsible for providing high torque to turn the engine flywheel. This specific task, which involves overcoming the initial resistance and compression inside the engine cylinders to get the engine rotating, requires a powerful electric motor designed for high torque output.

Analyzing the Options for Engine Starting

Let's look at the different options provided:

- **Dynamo:** A dynamo is an electrical generator that produces direct current (DC). Its primary function is to convert mechanical energy into electrical energy, typically used to charge the battery and power the vehicle's electrical system while the engine is running. It is not designed to turn the engine flywheel.
- **Alternator:** Similar to a dynamo, an alternator is also an electrical generator, but it produces alternating current (AC), which is then converted to DC for the vehicle's system. It performs the same

function as a dynamo (charging the battery, powering accessories) but more efficiently, especially at higher engine speeds. It does not provide the initial torque to start the engine.

- **Reversible motor:** A reversible motor is a motor that can operate in both clockwise and counter-clockwise directions. While some components in a vehicle might use reversible motors (like window motors), the term "reversible motor" itself does not specifically describe the high-torque motor used for starting an engine. The starting motor typically only needs to turn the engine in one direction.
- **Starting motor:** Also known as a starter motor or starter, this is a powerful electric motor designed specifically to rotate the internal combustion engine to initiate its operation. It draws a large current from the battery to produce the high torque necessary to crank the engine's crankshaft via the flywheel or ring gear. Its sole purpose is the brief period of engine cranking.

Identifying the High Torque Component

Based on the functions described, the component specifically designed to provide the high torque needed to turn the engine flywheel and start the engine is the starting motor.

Revision Table: Automotive Electrical Components

Component	Primary Function	Torque Requirement
Dynamo	Generate DC power (charging/accessories)	Low (driven by engine)
Alternator	Generate AC power (charging/accessories)	Low (driven by engine)
Reversible Motor	Operate in two directions (general term)	Varies depending on application
Starting Motor	Crank the engine (turn flywheel)	High (to overcome engine compression)

Additional Information on Engine Starting Systems

The starting system in most vehicles consists of the battery, the ignition switch, the starter solenoid, and the starting motor. When the ignition key is turned to the "start" position, it energizes the starter solenoid. The solenoid then performs two main actions: it pushes a gear (called the pinion gear) on the starting motor shaft to mesh with the ring gear on the engine flywheel, and it closes a heavy-duty electrical contact to send a large current directly from the battery to the starting motor. This allows the starting motor to generate immense torque momentarily to spin the engine until it starts running on its own power. Once the engine starts, the ignition key is released, which de-energizes the solenoid, causing the pinion gear to disengage from the flywheel and the power supply to the starting motor to be cut off.

117. Answer: d

Explanation:

Understanding Room Temperature Vulcanizing (RTV) Sealer

The question asks for another common name for Room Temperature Vulcanizing (RTV) sealer. RTV sealers are a type of adhesive or sealant that cures at room temperature.

Identifying RTV Sealer's Common Name

RTV sealers are widely known by a specific material name. Let's look at the options provided to determine the correct alternative name:

- **Tin sealer:** Tin is a metal and not typically used as a primary component in common sealers of this type.
- **Rubber sealer:** While RTV sealers result in a rubber-like material after curing (vulcanizing), "rubber sealer" is a very general term and not the specific common name for RTV type sealers.
- **Steel sealer:** Steel is a metal alloy and not used as a sealer material in this context.
- **Silicone sealer:** This is the most common and accurate alternative name for Room Temperature Vulcanizing (RTV) sealer. RTV sealers are a type of silicone sealant that cures via a chemical reaction upon exposure to air or moisture at room temperature.

Therefore, Room Temperature Vulcanizing (RTV) sealer is also commonly called Silicone sealer.

Why is it Called RTV Silicone Sealer?

The name "Room Temperature Vulcanizing" (RTV) describes the curing process. "Vulcanizing" generally refers to the process of hardening rubber or similar polymers, often by cross-linking the polymer chains. In the case of RTV silicone sealers, this hardening or curing happens at room temperature without the need for heat or other external catalysts (beyond atmospheric moisture or internal activators, depending on the specific type).

Since these sealers are based on silicone polymers and cure at room temperature through vulcanization, the term "RTV Silicone" or simply "Silicone sealer" is the most accurate and widely used descriptor.

Common Uses of RTV Silicone Sealer

RTV silicone sealers are versatile and used in many applications, including:

- Automotive gaskets and seals
- Industrial sealing and bonding
- Building and construction joints
- Electronic potting and sealing
- Household repairs and sealing gaps

Revision Table: Types of Sealers

Sealer Type	Key Characteristic	Common Examples
RTV Sealer	Cures at Room Temperature	Silicone Sealers
Silicone Sealer	Flexible, durable, good temperature resistance	Often RTV type
Acrylic Sealer	Water-based, paintable, less flexible than silicone	Caulking for gaps, trim
Polyurethane Sealer	Strong adhesion, durable, paintable	Construction joints, flooring

Additional Information on RTV Silicone Curing

RTV silicone sealers typically cure through one of two main mechanisms:

- **Acetoxo Cure:** Releases acetic acid (vinegar smell) during curing. Fast curing but can be corrosive to some materials.
- **Neutral Cure:** Releases non-acidic compounds (like alcohol or oxime) during curing. Slower curing but less corrosive and better adhesion to a wider range of substrates.

Both types result in a durable, flexible silicone rubber when fully cured, ideal for sealing and bonding applications requiring flexibility and environmental resistance.

118. Answer: b

Explanation:

Understanding File Types for Shaping Internal Curves

Files are essential hand tools used for shaping and smoothing materials, commonly metal or wood. Different types of files have distinct shapes and cuts, making them suitable for specific tasks. The question asks about the type of file used for shaping **internal curved surfaces**.

Analysis of File Types and Their Uses

Let's examine the options provided and determine which file is best suited for **internal curved surfaces**.

- **Triangular file:** This file has a triangular cross-section. It is typically used for filing internal angles, especially those greater than 60 degrees, and for sharpening saw teeth. It is not suitable for curved surfaces.
- **Half round file:** As the name suggests, this file has one flat side and one curved side. The curved side is specifically designed for working on concave surfaces, such as the inside of rings, curved openings, or other **internal curved surfaces**. The flat side can be used for flat surfaces or convex curves. This file is highly versatile for both internal and external curves, but particularly effective for internal ones due to its curved face.

- **Round file:** This file has a circular cross-section. It is primarily used for enlarging circular holes or filing concave curves with a specific, constant radius, such as notches. While it works on curves, the **half round file** offers more flexibility for various internal curved shapes that are not perfectly circular holes.
- **Square file:** This file has a square cross-section. It is used for filing square or rectangular holes and internal corners at 90 degrees. It is not suitable for curved surfaces.

Comparing File Shapes and Uses

File Type	Shape	Primary Use	Suitable for Internal Curved Surfaces?
Triangular file	Triangular	Internal angles, saw teeth	No
Half round file	Flat on one side, curved on the other	Internal and external curves, concave/convex surfaces	Yes
Round file	Round	Enlarging circular holes, specific concave curves	Limited suitability (mainly for holes)
Square file	Square	Square/rectangular holes, internal corners	No

Based on the shapes and common uses of these files, the **half round file** is explicitly designed with a curved surface that makes it the most appropriate tool for filing and shaping various **internal curved surfaces**.

Conclusion on Filing Internal Curves

To effectively shape or smooth an **internal curved surface**, you need a file with a corresponding curved profile. The **half round file** provides this profile with its curved face, making it the ideal choice among the given options for this task.

Revision Table: File Types for Shaping

File Name	Key Feature	Best Use Case
Triangular file	3 flat sides	Angles > 60°, saw teeth
Half round file	One flat, one curved side	Internal and external curved surfaces
Round file	Round profile	Circular holes, curved notches
Square file	4 flat sides	Square/rectangular holes, 90° corners

Additional Information on Files and Filing

Files come in various lengths, cuts (the spacing and depth of the teeth, e.g., rough, bastard, second cut, smooth), and types of teeth (single cut, double cut, rasp cut, curved cut). The choice of cut depends on the amount of material to be removed and the desired finish. A coarser cut removes material quickly but leaves a rougher finish, while a finer cut removes less material but provides a smoother finish.

When filing **internal curved surfaces** with a **half round file**, it's important to use smooth, even strokes, following the curve of the surface. Proper technique ensures an even finish and avoids creating flat spots or uneven profiles.

119. Answer: b

Explanation:

Understanding Piston Pin Connections

The question asks about the connection point for the small end of the piston pin within an engine assembly. To answer this correctly, let's understand the components involved in connecting the piston to the crankshaft.

What is a Piston Pin?

The piston pin, also known as a gudgeon pin, is a hollow or solid metal pin that connects the piston to the connecting rod. Its main function is to allow the connecting rod to pivot as the piston moves up and down in the cylinder, transmitting the force generated by combustion to the crankshaft.

Components of the Assembly

The key components in this part of the engine are:

- **Piston:** Moves up and down within the cylinder.
- **Piston Pin (Gudgeon Pin):** Connects the piston to the connecting rod. It passes through bosses (supports) in the piston skirt.
- **Connecting Rod:** A link between the piston and the crankshaft. It has two ends:
 - **Small End:** The end that attaches to the piston via the piston pin.
 - **Big End:** The end that attaches to the crankshaft journal.
- **Crankshaft:** Rotates and converts the linear motion of the piston into rotary motion. The big end of the connecting rod is attached to a crank pin on the crankshaft.

Analyzing the Connection Point

The piston pin fits snugly through the piston bosses and the small end of the connecting rod. This forms a hinged joint that allows the connecting rod to swing back and forth relative to the piston as the crankshaft

rotates.

Therefore, the small end of the piston pin is directly connected to the small end of the connecting rod.

Evaluation of Options

Let's look at the given options:

- **Big end of the journal:** The "journal" typically refers to the part of the crankshaft or camshaft that rotates within a bearing. The big end of the connecting rod connects to the crankshaft journal (specifically, the crank pin journal), not the small end of the piston pin.
- **Small end of the connecting rod:** This is the correct connection point for the piston pin. The piston pin passes through the piston and the small end bore of the connecting rod.
- **Cam shaft:** The camshaft controls the opening and closing of the engine's valves. It is not directly connected to the piston pin or the piston-connecting rod assembly.
- **Crank shaft:** The crankshaft receives the power from the piston via the connecting rod's big end. The piston pin connects to the small end of the connecting rod, which is then connected to the crankshaft by its big end, but the piston pin itself does not connect directly to the crankshaft.

Based on the function and arrangement of engine components, the small end of the piston pin connects to the small end of the connecting rod.

Piston-Connecting Rod Assembly Connections

Component	Connects to	Via
Piston	Connecting Rod (Small End)	Piston Pin
Connecting Rod (Small End)	Piston	Piston Pin
Connecting Rod (Big End)	Crankshaft	Crank Pin Journal

Conclusion

The piston pin serves as a critical link between the piston and the connecting rod. Its small end fits within the piston and connects to the small end of the connecting rod.

Revision Table: Piston Pin Knowledge

Term	Definition/Role	Key Connection
Piston Pin (Gudgeon Pin)	Connects piston to connecting rod, allowing pivot.	Piston & Small End of Connecting Rod
Connecting Rod Small End	Upper end of connecting rod.	Connects to Piston via Piston Pin
Connecting Rod Big End	Lower end of connecting rod.	Connects to Crankshaft Journal

Additional Information: Types of Piston Pin Connections

There are a few ways the piston pin can be fitted:

- **Full Floating:** The pin is free to rotate in both the piston bosses and the small end of the connecting rod. Retaining rings (circlips) prevent it from moving sideways and contacting the cylinder wall.
- **Semi-Floating (Semi-Rigid):** The pin is secured (usually clamped) in the connecting rod's small end and is free to rotate in the piston bosses.
- **Fixed:** The pin is fixed rigidly in the piston bosses and is free to rotate only in the small end bush of the connecting rod. This type is less common in modern engines.

Regardless of the fitting type, the physical connection point for the piston pin is between the piston and the small end of the connecting rod.

120. Answer: d

Explanation:

Understanding the Concept of Oiliness in Lubricants

The question asks about a property that is a combination of wetting, surface tension, and slipperiness. Let's analyze the terms and options provided to identify the correct concept.

The properties mentioned are:

- **Wetting:** The ability of a liquid to maintain contact with a solid surface, resulting from intermolecular interactions when the two are brought together.
- **Surface tension:** The tendency of liquid surfaces to shrink into the minimum surface area possible. High surface tension liquids tend to bead up, while low surface tension liquids spread out (wetting).
- **Slipperiness:** The quality of being slippery or causing things to slide easily. This relates to the friction-reducing ability of a substance.

We are looking for a term that encompasses how well a substance spreads on a surface (related to wetting and surface tension) and how effectively it reduces friction (slipperiness).

Analyzing the Options

Let's look at each option:

1. **Fire point:** This is the temperature at which a substance will continue to burn after being ignited by an open flame. This property is related to flammability, not the combination of wetting, surface tension, and slipperiness.
2. **Volatility:** This refers to how easily a substance evaporates at a given temperature. While related to the substance's composition, it is not directly related to its ability to wet a surface, its surface tension, or its slipperiness in the context described.
3. **Flash point:** This is the lowest temperature at which a volatile substance can vaporize to form an ignitable mixture with air. Like fire point, this relates to flammability and ignition properties, not the combination of wetting, surface tension, and slipperiness.
4. **Oiliness:** This term, especially in the context of lubricants, describes the property that causes differences in the friction-reducing ability of liquids, even when they have similar viscosities. It is often related to the presence of polar molecules that can adhere strongly to surfaces (wetting) and form a thin, slippery film. This adhesion and spreading (influenced by surface tension) contributes to the lubricant's ability to provide slipperiness under boundary lubrication conditions.

Connecting Properties to Oiliness

Oiliness is precisely the property that describes the combined effect of a lubricant's ability to adhere to a surface (wetting), its surface tension (which affects spreading), and its resulting slipperiness or friction-reducing capability, especially under conditions where a full fluid film is not maintained (boundary lubrication).

A good lubricant with high oiliness will wet the metal surface effectively, forming a strong, thin film that provides slipperiness and prevents metal-to-metal contact, even under pressure.

Summary Table of Concepts

Term	Primary Property	Relation to Wetting, Surface Tension, Slipperiness
Fire point	Flammability	None directly related
Volatility	Evaporation rate	None directly related
Flash point	Flammability/Ignition	None directly related
Oiliness	Lubrication/Boundary Friction Reduction	Directly involves the combination of wetting, surface tension (affecting spreading/adhesion), and slipperiness (friction reduction).

Conclusion

Based on the analysis, the term that encompasses the combination of wetting, surface tension, and slipperiness, particularly in the context of lubrication and surface interaction, is oiliness.

Revision Table: Key Lubricant Properties

Property	Description	Importance in Lubrication
Viscosity	Resistance to flow	Determines film thickness under hydrodynamic conditions.
Viscosity Index (VI)	How viscosity changes with temperature	Indicates suitability for varying temperatures.
Pour Point	Lowest temperature at which it flows	Ensures lubricant flows in cold conditions.
Flash Point	Temperature vapours ignite	Safety measure against fire hazards.
Oiliness	Ability to reduce friction (especially boundary) via wetting/adhesion	Crucial for lubrication where full film isn't maintained.

Additional Information on Oiliness

Oiliness is a complex property and is not easily measured by a single standard test like viscosity or flash point. It is often evaluated indirectly through friction and wear tests under specific conditions, such as boundary lubrication.

The chemical composition of a lubricant significantly affects its oiliness. Lubricants containing polar molecules, such as fatty acids or esters, tend to exhibit higher oiliness because their polar ends are attracted to and can form stronger bonds with metal surfaces. This strong adhesion helps the lubricant film resist being squeezed out, even under high pressure, contributing effectively to the slipperiness and friction reduction.

Understanding oiliness is important in selecting lubricants for applications involving heavy loads, slow speeds, or reciprocating motion, where maintaining a full hydrodynamic film is difficult, and boundary lubrication conditions are likely to occur.

121. Answer: a

Explanation:

Understanding Spark Plug Gap Adjustment

Setting the correct **spark plug gap** is essential for ensuring your engine runs efficiently. The gap is the critical space between the electrodes on the spark plug, where the spark ignites the fuel-air mixture. This gap needs to be precisely measured and adjusted to match the specific requirements of the vehicle's engine.

The Role of a Feeler Gauge in Setting Spark Plug Gaps

The standard and most effective tool for measuring and setting the **spark plug gap** is a **feeler gauge**. Here's a breakdown of why:

- **What is a Feeler Gauge?** A **feeler gauge** is a precision tool that comes with a set of thin, long metal strips of varying, calibrated thicknesses. Each strip is accurately marked to indicate its specific thickness, usually in millimeters (mm) or inches (in).
- **How it's Used:** To **adjust** the gap, you select the feeler gauge strip corresponding to the manufacturer's specified gap measurement for the spark plug. This strip is then carefully inserted between the spark plug's center and ground electrodes.
- **Achieving the Correct Gap:** The ideal result is that the correct feeler gauge strip slides into the gap with slight resistance or friction. If the strip slides in too easily, the gap is too wide, and if it doesn't fit, the gap is too narrow. The electrodes can then be gently bent or tapped to achieve the desired measurement.

Why Other Tools Are Not Suitable

Let's consider the other options provided:

- **Go Gauge:** A go gauge is typically used in manufacturing to verify if a part dimension meets a minimum size tolerance. It confirms whether something 'goes' through or fits, but it's not designed for the fine adjustment or measurement of the small electrode gap on a spark plug.
- **Blade Gap Adjustment:** This phrase describes the gap itself or the components involved (electrodes can be seen as blades), but it is not the name of a tool used for measurement or adjustment. The process involves adjusting the gap *between* the blades (electrodes).
- **Scale:** While a scale (like a ruler) measures length, it generally lacks the high precision required for setting spark plug gaps, which are often very small (e.g., 0.04 inches or 1.0 mm). Feeler gauges are specifically manufactured for these precise measurements.

Summary of Spark Plug Gap Tool Choice

In summary, the precision and specific design of a **feeler gauge** make it the correct tool for accurately measuring and ensuring the **spark plug gap** is set according to manufacturer specifications, which is vital for engine performance and fuel economy.

122. Answer: b

Explanation:

This question asks us to identify the range representing the **finest abrasive grain** size used in grinding wheels.

Abrasive Grain Size Classification for Grinding Wheels

Grinding wheels are made of abrasive grains bonded together. The size of these abrasive grains, often referred to as the 'grit', is crucial for determining the wheel's cutting action and the finish it produces on the workpiece. Abrasive grains are typically graded numerically, where a higher number indicates a smaller, finer grain, and a lower number indicates a larger, coarser grain.

Understanding the different grit size classifications helps in selecting the appropriate grinding wheel for a specific task:

- **Coarse Grit (e.g., 6 – 24):** These larger grains are used for rapid material removal, heavy grinding, and on soft materials. They leave a rougher surface finish.
- **Medium Grit (e.g., 30 – 60):** These are suitable for general-purpose grinding operations where a balance between material removal rate and surface finish is needed.
- **Fine Grit (e.g., 70 – 200):** Finer grains are used for lighter grinding tasks, achieving better surface finishes, and grinding harder materials.
- **Very Fine / Super Fine Grit (e.g., 220 – 1000 and above):** These extremely small grains are used for precision grinding, polishing, honing, and achieving very smooth, mirror-like finishes.

Grit Size Ranges and Their Applications

Grit Size Range	Classification	Typical Use
6 – 24	Coarse	Rapid material removal, heavy grinding
30 – 60	Medium	General purpose grinding
70 – 200	Fine	Precision grinding, improved surface finish
220 – 1000+	Very Fine / Super Fine	Polishing, fine finishing, honing

Identifying the Finest Abrasive Grain Size

Based on the standard classification system for abrasive grain sizes:

- Option 1 (6 – 24) represents **coarse** grit sizes.
- Option 3 (70 – 200) represents **fine** grit sizes.
- Option 4 (30 – 60) represents **medium** grit sizes.
- Option 2 (220 – 1000) represents **very fine to super fine** grit sizes.

Therefore, the range **220 – 1000** includes the finest abrasive grains among the choices provided, suitable for applications requiring exceptionally smooth finishes.

123. Answer: c

Explanation:

Understanding Crane Hook Materials

Crane hooks are critical components in lifting operations, designed to bear heavy loads and ensure safety. The material used for a crane hook must possess specific mechanical properties to withstand the significant forces and stresses involved, including tensile stress, shear stress, and potential shock loads. Selecting the right material is paramount for preventing failure, which could lead to serious accidents.

Key Properties for Crane Hook Material

A suitable material for crane hooks should ideally have the following characteristics:

- **High Tensile Strength:** Ability to withstand pulling forces without breaking.
- **Ductility:** Ability to deform significantly under tensile stress before fracturing. This allows the hook to show visible signs of deformation before catastrophic failure, providing a safety warning.
- **Toughness:** Ability to absorb energy and plastically deform without fracturing, especially under impact or sudden loads.
- **Wear Resistance:** To handle contact and friction with lifting slings and other equipment.
- **Fatigue Strength:** Ability to withstand repeated loading cycles.

Analyzing the Options

Let's consider the suitability of the materials provided in the options for manufacturing crane hooks:

Your Personal Exams Guide

Material	Typical Properties	Suitability for Crane Hooks
Cast Iron	High compressive strength, low tensile strength, brittle.	Not suitable. Its brittleness makes it prone to sudden fracture under tensile or shock loads.
Cast Iron Steel	This term is not a standard material classification. If referring to Cast Steel, properties vary, but many grades lack the ductility and toughness of wrought steel like mild steel for this application.	Generally less preferred than wrought steel for standard hooks due to potential variations in casting and lower ductility compared to mild steel.
Mild Steel	Good tensile strength, high ductility, good toughness, weldable, relatively inexpensive.	Highly suitable. Its combination of strength, high ductility, and toughness allows it to withstand expected loads and deform visibly before failure under overload conditions.
High Speed Steel	Very high hardness, wear resistance, used for cutting tools. Extremely brittle.	Completely unsuitable. Its brittleness would lead to catastrophic failure under typical crane hook loads.

Why Mild Steel is Preferred for Crane Hooks

Based on the required properties and the characteristics of the materials, mild steel emerges as the most suitable choice for general-purpose crane hooks. Its defining feature is its high ductility. While stronger steels exist, their reduced ductility means they might fracture suddenly under shock loads or excessive stress without prior warning deformation. Mild steel, on the other hand, will yield and deform plastically before breaking, providing a crucial safety indicator. Furthermore, mild steel offers a good balance of strength, toughness, cost-effectiveness, and ease of manufacturing (like forging or casting followed by appropriate heat treatment).

Therefore, crane hooks are generally made of mild steel.

Revision Table: Crane Hook Material Properties

Property	Mild Steel	Cast Iron	High Speed Steel
Tensile Strength	Good	Low	High (but brittle)
Ductility	High	Very Low	Very Low
Toughness	Good	Low	Very Low
Suitability for Hooks	High	Very Low	Very Low

Additional Information on Crane Hook Design and Materials

While mild steel is common, some high-capacity or specialized crane hooks might use alloy steels that are heat-treated to achieve a specific balance of strength, toughness, and ductility. However, for typical forged hooks, mild steel is the standard. The design of the hook itself is also critical, featuring a large radius at the bend to minimize stress concentration and a trapezoidal or triangular cross-section for optimal stress distribution. Crane hooks are often manufactured through forging, which refines the grain structure and improves the material's strength and toughness. Regular inspection of crane hooks for deformation or cracks is essential for safety, and the ductility of mild steel aids in identifying potential issues through visual inspection.

124. Answer: d

Explanation:

Understanding Drill Point Angle for Zinc Alloys

When drilling any material, selecting the correct drill point angle is crucial for efficient cutting, good surface finish, and tool longevity. The drill point angle is the angle formed at the tip of the drill bit between the cutting edges. Different materials require different point angles due to their varying hardness, ductility, and abrasive properties.

Importance of Drill Point Angle

The drill point angle affects several aspects of the drilling process:

- **Chip Formation and Evacuation:** The angle influences how chips are formed and move away from the cutting area.
- **Cutting Forces:** It impacts the thrust force (downward force) and torque required for drilling.
- **Centering:** A sharper point angle centers more easily, while a flatter angle requires more force or a pilot hole to prevent walking.
- **Tool Life:** The correct angle minimizes wear and prevents chipping or breaking of the cutting edge.

Drilling Zinc Alloys

Zinc alloys are generally considered soft and ductile materials. Drilling soft, ductile materials can be challenging because the material tends to deform and stick to the drill bit (built-up edge), and chips can be long and stringy, leading to poor chip evacuation and potential clogging. Selecting the appropriate drill point angle, along with other drill geometry features like rake angle and helix angle, helps mitigate these issues.

Analyzing Common Drill Point Angles

Let's look at typical drill point angles and their general applications:

- **Sharper angles ($< 118^\circ$):** These angles are generally used for softer materials like plastics or very soft metals. They cut more freely but can be weaker and prone to grabbing in ductile materials. Examples include 90° or even smaller angles for specific applications.
- **Standard general-purpose angle (118°):** This is a very common angle used for a wide range of materials, including mild steel, cast iron, and some non-ferrous metals. It offers a good balance between centering ability and strength.
- **Flatter angles ($> 118^\circ$):** Angles like 135° are used for harder or tougher materials like stainless steel or titanium. They require less thrust force and are more resistant to wear, but they have poorer centering capability.

Recommended Drill Point Angle for Zinc Alloys

For soft and ductile materials like zinc alloys, while a sharper angle might sometimes be used with modifications, the 118° point angle is often considered suitable or is the most common standard angle used effectively, especially when combined with appropriate feeds, speeds, and lubrication. Although specific recommendations can vary based on the exact alloy and drilling conditions, 118° provides a reasonable compromise for many non-ferrous materials, including zinc alloys among the given options.

Here is a general guide for drill point angles based on material type:

Material Type	Typical Drill Point Angle
Soft Steels, Cast Iron	118°
Harder Steels, Stainless Steel	135°
Brass, Bronze	118° or 90° – 118°
Aluminum Alloys	118° or 90° – 118°
Zinc Alloys	118° is a common choice; sometimes sharper angles with modified geometry.
Plastics	90° – 118° (sometimes modified geometry)

Based on standard practices for drilling materials with properties similar to zinc alloys and the options provided, 118° is the most appropriate drill point angle listed.

Revision Table: Drill Point Angle for Zinc Alloys

Let's summarise the key points regarding the drill point angle for zinc alloys.

Concept	Explanation
Drill Point Angle	Angle at the tip of the drill bit between cutting edges.
Purpose	Influences cutting action, chip formation, forces, and tool life.
Zinc Alloy Properties	Soft and ductile. Prone to grabbing and chip clogging.
Recommended Angle (from options)	118° is a common and suitable choice for many non-ferrous metals including zinc alloys.
Other Considerations	Rake angle, helix angle, feed/speed, lubrication also critical for drilling zinc alloys.

Additional Information on Drilling Zinc Alloys

Beyond the point angle, other factors significantly affect drilling performance in zinc alloys:

- **Rake Angle:** A smaller or negative rake angle can help prevent the drill from "grabbing" the soft material.
- **Helix Angle:** A larger helix angle helps improve chip evacuation, which is important for ductile materials like zinc alloys.
- **Feeds and Speeds:** Using appropriate cutting speeds and feeds is essential to control chip size and prevent heat buildup.
- **Lubrication/Coolant:** Using a suitable coolant or lubricant reduces friction, prevents built-up edge, and aids chip evacuation.
- **Drill Material:** High-speed steel (HSS) drills are commonly used, but coated drills can offer improved performance and tool life.

Choosing the right combination of drill geometry and cutting parameters is key to successful drilling of zinc alloys.

125. Answer: b

Explanation:

Understanding the Reamer and its Parts

A reamer is a type of rotary cutting tool used in metalworking. Its primary purpose is to enlarge and finish a pre-existing hole to a precise size, often with a smooth surface finish. Reamers are known for their ability to achieve tight tolerances and high-quality surface finishes compared to drills.

Like most cutting tools used in machines, a reamer has several distinct parts, each serving a specific function in the cutting process and in how the tool is held and operated.

Identifying the Held and Driven Portion of a Reamer

The question asks to identify the portion of the reamer that is held and driven by the machine or tool holder. This part must be strong and accurately shaped to ensure the reamer rotates correctly and applies force effectively during the reaming operation.

Let's look at the common parts of a reamer:

- **Shank:** This is the end of the reamer opposite to the cutting end. The shank is designed to fit into the spindle or tool holder of the machine tool (like a drill press, milling machine, or lathe). It transmits the rotary motion and torque from the machine to the reamer's cutting edges. Shanks can be cylindrical (parallel) or tapered (like Morse taper shanks) to suit different types of tool holders and provide secure gripping.
- **Body:** This is the main part of the reamer, extending from the start of the cutting flutes up to the shank. It contains the cutting edges and flutes along its length. The body performs the actual cutting and sizing of the hole.
- **Cutting Edges/Flutes:** These are the sharp edges that remove material from the hole wall. The flutes are grooves that provide channels for chips to escape and cutting fluid to reach the cutting edges.
- **Axis:** This is the imaginary central line around which the reamer rotates. It is not a physical part of the tool that is held or driven.
- **Recess:** A recess might be a groove or reduced diameter section on some tools, but it is not the primary part used for holding and driving the reamer.

Based on the functions of these parts, the portion of the reamer that is specifically designed to be held and driven by the machine is the shank. It is the interface between the reamer and the power source/holding mechanism of the machine tool.

Analyzing the Options

Let's evaluate the given options in the context of the reamer's structure and function:

- **Axis:** The axis is a theoretical line of rotation, not a physical part that is held. Incorrect.
- **Shank:** The shank is the part inserted into the machine's spindle or holder to be held and driven. This perfectly matches the description in the question. Correct.
- **Recess:** A recess is typically an indentation or groove, not the primary holding/driving part. Incorrect.
- **Body:** The body contains the cutting flutes and performs the cutting, but it is not the part that is held and driven by the machine; the shank is. Incorrect.

The shank is the critical component for mounting the reamer securely in the machine and transferring the necessary power for the reaming operation. The mention that it can be parallel or tapered further points towards the shank, as these are common types of shanks found on cutting tools.

Reamer Parts and Functions

Part	Function	Held & Driven?
Shank	Inserted into tool holder; receives power/torque; held securely.	Yes
Body	Contains cutting edges and flutes; performs cutting.	No
Axis	Imaginary line of rotation.	No
Recess	(If present) Groove or indentation for specific purposes (e.g., clearance).	No

Conclusion

The portion of the reamer that is held and driven, and can be parallel or tapered, is the shank.

Revision Table: Key Reamer Terms

Term	Definition
Reamer	A cutting tool used for finishing and sizing holes accurately.
Shank	The part of a cutting tool that is gripped by the tool holder or machine spindle.
Parallel Shank	A cylindrical shank for holding in collets or chucks.
Tapered Shank	A conical shank (e.g., Morse taper) for self-holding in matching tapered spindles.
Body	The main portion of the reamer containing the cutting flutes.
Flutes	Grooves on the reamer body that form the cutting edges and allow chip evacuation.

Your Personal Exams Guide

Additional Information on Reamer Shanks

The type of shank on a reamer is crucial for its application and the type of machine tool it can be used with. Here are some common types:

- **Straight/Parallel Shanks:** These are cylindrical and require a chuck (like a drill chuck or collet chuck) to hold them securely in the machine spindle. They are common on smaller reamers.
- **Taper Shanks:** The most common is the Morse taper shank, which fits directly into a matching tapered bore in the machine spindle or a tapered adapter sleeve. The taper provides a self-holding friction grip, capable of transmitting significant torque. Taper shanks are often found on larger reamers and tools requiring high rigidity.
- **Square End Shanks:** Some reamers, particularly hand reamers, might have a square end on the shank to allow them to be turned with a tap wrench. Machine reamers typically have round or tapered shanks for power rotation.

Choosing the correct shank type is essential to ensure the reamer is properly centered, securely held, and capable of withstanding the cutting forces encountered during the reaming process.

126. Answer: c

Explanation:

Correct answer: lock the reading after setting the job

Option 3: lock the reading after setting the job - This aligns perfectly with the function described. Once the micrometer is set to the job and the measurement is established, the lock nut is tightened to prevent the spindle from moving, thus locking the reading. This option is correct.

Therefore, the main purpose of the lock nut in a micrometer is to lock the reading after setting the job on the micrometer.

127. Answer: c

Explanation:

Understanding How to Clean Oily Workshop Floors

An oily workshop floor presents a safety hazard, making it crucial to clean it effectively. Different cleaning methods are available, but their suitability depends on the substance being cleaned and the surface.

Analyzing Cleaning Methods for Oil Spills

Let's evaluate the options provided for cleaning an oily workshop floor:

- Option 1: by washing with water

Oil and water generally do not mix. Washing an oily floor with water can spread the oil, making the floor even more slippery and dangerous. It doesn't effectively remove the oil itself.

- Option 2: by cotton waste

Cotton waste can absorb some oil, but it is not the most efficient method for large spills or general floor cleaning. It can become saturated quickly and may leave a residue. Also, oily cotton waste is a fire hazard and needs proper disposal.

- Option 3: with the help of wooden dust and sand

Wooden dust (sawdust) and sand are porous materials that are excellent absorbents. When spread over an oil spill, they soak up the oil, converting the liquid mess into a solid or semi-solid waste that

can be easily swept up. This method is effective, relatively inexpensive, and helps make the floor less slippery quickly.

- **Option 4: by spraying Carbon dioxide**

Carbon dioxide (CO_2) is primarily used as a fire extinguishing agent. It works by displacing oxygen or cooling the fuel source. Spraying CO_2 will not absorb or remove oil from a floor. It is completely ineffective for cleaning oil spills.

Selecting the Best Method for Oily Floors

Comparing the methods, using absorbent materials like wooden dust and sand is the most practical and effective way to clean an oily workshop floor. These materials quickly absorb the oil, making cleanup easier and improving safety by reducing slipperiness.

Therefore, cleaning an oily workshop floor is best achieved with the help of wooden dust and sand.

Revision Table: Oily Floor Cleaning Methods

Method	Effectiveness for Oil	Safety Considerations
Washing with water	Poor (spreads oil)	Increases slip hazard
Cotton waste	Moderate (absorbs some)	Fire hazard, leaves residue
Wooden dust and sand	Good (absorbs effectively)	Reduces slip hazard, requires cleanup
Spraying Carbon dioxide	None (fire extinguishing)	Ineffective for cleaning

Additional Information: Workshop Safety and Cleanup

Maintaining a clean and safe workshop environment is crucial to prevent accidents. Oily floors are a significant hazard, leading to slips and falls. Prompt cleanup using appropriate absorbent materials like sand, sawdust, or commercial oil absorbents is essential. Proper storage and disposal of oily waste are also important to prevent fires and environmental contamination. Regular maintenance and checking for leaks can help prevent oil spills from happening in the first place.

128. Answer: c

Explanation:

Explanation:

Spanners:

- Spanners are used for operating threaded fasteners, bolts, and nuts.

- They are made with jaws or opening that fit square on hexagonal nuts and bolts and screw heads.
- They are made of high tensile or alloy steel.
- They are drop-forged and heat-treated for strength.
- Finally they are given a smooth surface finish for ease of gripping.

The basic types of spanners are:

- Open end spanners
- tube or tubular box spanners
- Socket spanners
- Ring spanners

Ratchet spanners:

- It is a type of socket spanner.
- Ratchet spanners feature a closed head with a multi-toothed ring featuring a ratchet function, allowing the user to tighten or loosen nuts and bolts without having to remove and re-position the tool.

Socket spanners:

- The socket is one of the fastest and most convenient of all the spanners.
- Sockets come in two sizes; standard and deep.
- Standard sockets will handle most of the work, while the extra reach of the deep socket is occasionally needed.

129. Answer: a

Explanation:

Understanding Internal Combustion Engines and Fuels

An internal combustion engine (often shortened to ICE) is a heat engine where the combustion (burning) of a fuel occurs with an oxidizer (usually air) in a confined space called a combustion chamber. This combustion creates high-temperature and high-pressure gases, which expand and exert force on components like pistons or turbine blades, making them move and perform useful work.

Different types of internal combustion engines are designed to operate on specific types of fuel. Let's look at the options provided and see which one is commonly used in an internal combustion engine.

Analyzing Fuel Options for Internal Combustion Engines

Here's an analysis of each fuel option:

- **Diesel:** Diesel fuel is a liquid petroleum distillate fuel. Diesel engines are a major type of internal combustion engine. In a diesel engine, air is compressed to a very high temperature, and then diesel

fuel is injected into the hot compressed air, where it ignites spontaneously. This makes diesel a standard fuel for many internal combustion engines, especially in trucks, buses, ships, and some cars.

- **Hydrogen:** Hydrogen can be used as a fuel in internal combustion engines, often with modifications. Hydrogen ICEs are being researched and developed, and some prototype vehicles exist. However, hydrogen is not as widely or commonly used in standard internal combustion engines compared to traditional liquid fuels like petrol or diesel.
- **Wood:** Wood is a solid fuel. While wood can be used to generate energy (e.g., by burning it in a furnace or boiler, or through gasification), it is not suitable for direct use as a fuel in the combustion chamber of typical internal combustion engines like those found in cars or trucks. These engines require fuels that can be easily atomized or mixed with air as a gas or vapor.
- **Coal:** Coal is also a solid fuel. Like wood, coal is primarily used in external combustion systems, such as large power plants, where it is burned to heat water and produce steam to drive turbines. Coal is not used directly in standard internal combustion engines because of its solid state and the difficulty in controlling its combustion within the small, high-speed environment of an ICE cylinder.

Conclusion on Fuel Type

Based on the common and widespread use of the listed fuels in typical internal combustion engines, diesel is the most appropriate answer among the given options. While hydrogen is a potential future fuel for ICEs, diesel engines are a well-established and prevalent type of internal combustion engine.

Fuel Type	Common Engine Type	Used in Internal Combustion Engine (Typical)
Diesel	Diesel Engine (ICE)	Yes
Hydrogen	Hydrogen ICE (Developing)	Limited/Developing
Wood	Furnaces, Boilers (External Combustion/Gasification)	No (in typical ICE)
Coal	Power Plants (External Combustion)	No (in typical ICE)

Revision Table: Internal Combustion Engine Fuels

Fuel	State	Common Application	Suitability for Typical ICE
Diesel	Liquid	Cars, Trucks, Buses, Ships	High
Hydrogen	Gas	Future/Alternative Vehicles	Emerging/Requires modification
Wood	Solid	Heating, Power Generation (External)	Low (not direct)
Coal	Solid	Power Generation (External)	Low (not direct)

Additional Information: More About ICE Fuels

Besides diesel, other common fuels used in internal combustion engines include:

- **Petrol (Gasoline):** Widely used in spark-ignition internal combustion engines found in most cars and motorcycles.
- **LPG (Liquefied Petroleum Gas):** Used in modified petrol engines, primarily in vehicles.
- **CNG (Compressed Natural Gas):** Used in modified petrol or dedicated natural gas internal combustion engines, especially in buses and commercial vehicles.
- **Ethanol/Methanol:** Alcohols used alone or blended with petrol in some internal combustion engines.

The design of the internal combustion engine dictates which fuel or range of fuels it can efficiently and safely burn.

130. Answer: d

Explanation:

Understanding the Role of the Clutch in a Vehicle

The clutch is a vital component in vehicles with manual transmissions. Its primary function is to connect and disconnect the engine's rotation from the transmission and, subsequently, the drive wheels. This connection allows the driver to change gears smoothly and also to stop the vehicle without stopping the engine.

Think of the clutch as a bridge between the engine (power source) and the wheels (what makes the vehicle move). When the clutch is engaged, the bridge is 'up', allowing power to flow. When it's disengaged (clutch pedal pressed), the bridge is 'down', breaking the connection.

Impact of a Malfunctioning Clutch

If the clutch is not functioning correctly, it often means this connection or disconnection isn't happening as it should. A common failure is the clutch not being able to transmit the engine's power to the gearbox effectively. Let's examine what happens in this scenario by looking at the given options.

Analyzing the Given Options Regarding a Non-Functioning Clutch

Option 1: Engine gets overheated

Engine overheating is usually related to issues with the cooling system (like coolant levels, radiator, thermostat, or fan) or the engine running under extreme load for prolonged periods. While a slipping clutch can cause the engine RPM to increase without a corresponding increase in vehicle speed, potentially leading to some increased engine temperature due to wasted energy, it is not the primary or guaranteed

outcome of a clutch that is "not functioning" in the sense of failing to engage or transfer power. The most direct consequence relates to the vehicle's movement.

Option 2: Engine will not start

The engine starting process is largely independent of the clutch's ability to engage the transmission. When starting a vehicle with a manual transmission, it is standard practice (and often a safety feature requiring the pedal to be pressed) to fully disengage the clutch. This disconnects the engine from the transmission, reducing the load on the starter motor. A faulty clutch might make it difficult to shift gears once started, but it typically won't prevent the engine from starting itself.

Option 3: Fly wheel is at rest

The flywheel is directly attached to the engine's crankshaft. It rotates whenever the engine is running. The clutch is designed to connect to or disconnect from the flywheel. Therefore, if the engine is running, the flywheel will be rotating, regardless of the clutch's condition or whether it is engaged or disengaged. A non-functioning clutch does not cause the flywheel to stop if the engine is running.

Option 4: Vehicle will not move

This is the most direct and probable outcome if the clutch is not functioning in a way that prevents the transfer of power from the engine to the transmission and wheels. If the clutch cannot engage the transmission effectively (e.g., due to severe wear, hydraulic problems, or a broken component), the engine might run, the flywheel might spin, but the power will not reach the wheels. The vehicle will remain stationary even if the engine is revved up. This is a classic symptom of a failed clutch.

Component	Function in Vehicle Movement	Effect if Clutch Not Functioning (Cannot Engage)
Engine	Generates power (rotates)	Engine can still run
Flywheel	Attached to engine, spins with engine	Flywheel spins if engine runs
Clutch	Connects engine to transmission	Fails to transfer engine power
Transmission/Gearbox	Receives power, allows gear selection	Receives no power from engine
Drive Wheels	Receive power from transmission to move vehicle	Receive no power, vehicle does not move

Based on the analysis, if the clutch is not functioning, particularly if it cannot engage to transfer power, the vehicle will not move. The engine might run, but its power cannot reach the wheels.

Revision Table: Common Clutch Problems & Symptoms

Problem	Symptom
Slipping Clutch	Engine RPM increases but vehicle speed does not; loss of acceleration; burning smell.
Clutch Grab/Chatter	Vehicle shakes when engaging the clutch.
Hard Pedal	Requires excessive force to press the clutch pedal.
Clutch does not disengage (Dragging Clutch)	Difficult to shift gears, gears grind.
Clutch does not engage	Engine runs but vehicle does not move.

Additional Information on Vehicle Clutch Systems

Clutch systems are essential for manual transmission vehicles. They allow drivers to interrupt the flow of power from the engine to the transmission temporarily. This is crucial for:

- Starting the vehicle from a standstill.
- Changing gears smoothly while the vehicle is in motion.
- Idling the engine without the vehicle moving.

Modern clutches often use hydraulic systems, though older vehicles might use mechanical linkages. Over time, the friction material on the clutch disc wears out, leading to slipping and eventually failure to engage. Proper clutch operation, avoiding riding the clutch pedal unnecessarily, helps extend its lifespan.

131. Answer: d

Explanation:

Understanding the Engine Timing Belt

A timing belt is a critical component in many internal combustion engines. It is essentially a reinforced rubber belt with teeth that mesh with sprockets or pulleys on different engine shafts. Its primary function is to maintain synchronization between the rotation of specific parts of the engine.

Think of it like the conductor of an orchestra, ensuring that different sections play their parts in perfect time relative to each other for the music (or engine) to work correctly.

What Components Does the Timing Belt Connect?

The main components that the timing belt connects are the **camshaft** and the **crankshaft**. This connection is vital for the engine's operation.

- The **crankshaft** is connected to the pistons and is responsible for converting their up-and-down motion into rotational motion.
- The **camshaft** is responsible for operating the engine's valves (intake and exhaust valves), controlling when they open and close.

Why is this Connection Important? Engine Synchronization

For an engine to run efficiently and reliably, the opening and closing of the valves must be perfectly timed with the movement of the pistons. The timing belt ensures this synchronization by linking the rotation of the crankshaft and the camshaft.

If the timing belt were to fail or slip, the synchronization would be lost. This can lead to the pistons colliding with the valves, causing significant engine damage, especially in what are known as "interference engines".

Analyzing the Timing Belt Connection Options

Let's look at the options provided and determine which one correctly identifies the components connected by the timing belt:

- **Option 1:** Piston and connecting rod. The piston and connecting rod are directly connected to each other within the cylinder. The connecting rod connects the piston to the crankshaft. The timing belt does not connect these two components directly.
- **Option 2:** Alternator and engine. The alternator is typically driven by an accessory belt, often called a serpentine belt or V-belt, which is separate from the timing belt system. This belt powers accessories like the alternator, power steering pump, and air conditioning compressor.
- **Option 3:** Axle and wheel. The axle and wheel are part of the vehicle's drivetrain and suspension system, responsible for transferring power from the engine/transmission to the wheels. They are not directly connected by the engine's timing belt.
- **Option 4:** Camshaft and crankshaft. As discussed, the timing belt connects the camshaft and crankshaft. This is the essential link for synchronizing valve operation with piston movement, which is fundamental to the internal combustion engine cycle.

Based on the function and location of the timing belt in an engine, the correct connection is between the camshaft and the crankshaft.

Revision Table: Key Engine Components and Timing

Component	Primary Function	Connected by Timing Belt?
Crankshaft	Converts piston linear motion to rotation	Yes (to camshaft)
Camshaft	Operates engine valves	Yes (to crankshaft)
Piston	Moves up/down in cylinder during combustion	No (connected to crankshaft via connecting rod)
Connecting Rod	Links piston to crankshaft	No
Valves	Control flow of air/fuel and exhaust gases	No (operated by camshaft)
Alternator	Generates electrical power	No (typically driven by accessory belt)

Additional Information on Engine Timing

Some engines use a timing chain instead of a timing belt. Timing chains are typically more durable than belts but can be noisier. Both timing belts and chains serve the same fundamental purpose: to maintain the precise timing between the crankshaft and the camshaft to ensure smooth and efficient engine operation. Regular inspection and replacement of the timing belt (as per manufacturer recommendations) are crucial maintenance steps to prevent potentially catastrophic engine failure.

132. Answer: c

Explanation:

Understanding Vehicle Dashboard Indicators

Vehicle dashboards contain various indicators that provide drivers with essential information about the vehicle's operation. These indicators help monitor performance, safety, and maintenance needs. The question asks to identify which of the given options is a speed indicator found on the dashboard.

Analyzing the Options

Let's examine each option to determine its function:

- **Pressure gauge:** A pressure gauge measures the pressure of a fluid or gas within a system, such as engine oil pressure or fuel pressure. It is not a speed indicator.
- **AMP meter:** An AMP meter, or ammeter, measures the electrical current (amperes) flowing through a circuit. It typically indicates the charging system's performance or battery status. It is not a speed indicator.
- **RPM (Revolutions per minute):** RPM stands for Revolutions Per Minute and measures the rotational speed of the engine's crankshaft. This is commonly displayed on a tachometer on the dashboard.

While it measures engine speed, it is directly related to the vehicle's speed, as higher engine RPMs in a given gear generally result in higher vehicle speed. It indicates how hard the engine is working to achieve speed. Therefore, it functions as a crucial speed-related indicator.

- **Voltmeter:** A voltmeter measures the electrical voltage of the vehicle's battery and charging system. It indicates the health of the electrical system. It is not a speed indicator.

Why RPM is Considered a Speed Indicator Among the Options

While the primary indicator for vehicle speed over ground is the speedometer (which measures speed in MPH or Km/h), the RPM gauge (tachometer) indicates engine speed. Engine speed is directly linked to vehicle speed through the transmission ratio. Monitoring RPM helps drivers understand the engine load and select appropriate gears for efficient and powerful acceleration or deceleration. In the context of the provided options, RPM is the only one that directly reflects a rate of motion (engine rotation) that is fundamental to generating vehicle speed.

Here is a summary of the indicators:

Indicator	Measures	Related to Speed?
Pressure gauge	Fluid/gas pressure	No
AMP meter	Electrical current	No
RPM (Revolutions per minute)	Engine rotational speed	Yes (indirectly, related to engine output for speed)
Voltmeter	Electrical voltage	No

Conclusion

Based on the functions of the given dashboard indicators, RPM (Revolutions per minute) is the indicator that relates to the speed of the engine's operation, which in turn affects the vehicle's speed. While not the ground speed indicator, it is a fundamental speed-related metric displayed on many dashboards.

Revision Table: Key Dashboard Indicators

This table helps summarize the main functions of common dashboard indicators often encountered in vehicles.

Indicator	Purpose
Speedometer	Measures vehicle speed (MPH or Km/h).
Tachometer (RPM)	Measures engine rotational speed (Revolutions per Minute).
Fuel Gauge	Indicates the amount of fuel in the tank.
Temperature Gauge	Monitors engine coolant temperature.
Oil Pressure Gauge	Indicates engine oil pressure.
Voltmeter	Shows battery voltage and charging system status.

Additional Information: RPM and Driving

Understanding RPM is important for efficient and safe driving. Keeping the engine within the appropriate RPM range for the current gear helps optimize fuel economy and engine performance. Operating the engine at very high RPMs for extended periods can increase wear and tear, while operating at very low RPMs in a high gear can lug the engine. The RPM gauge is a vital tool, especially in manual transmission vehicles, to know when to shift gears.

133. Answer: b

Explanation:

Locating Components on the Crankshaft

This question requires identifying which specific engine part is attached to the **rear end** of a **crankshaft**. The crankshaft is a vital rotating shaft that converts the reciprocating motion of pistons into rotational motion. Different components are attached to its ends to perform specific functions.

The Flywheel's Position and Purpose

A **flywheel** is a large, heavy rotating disc connected to the **crankshaft**. Its main jobs include:

- Storing energy to ensure smoother rotation between the engine's power pulses.
- Providing a surface for the clutch or torque converter, enabling the engine's power to be transferred to the transmission.

The standard location for the **flywheel** is on the **rear end** of the **crankshaft**, which is the end typically facing away from the timing components and towards the transmission.

Comparison with Other Crankshaft Parts

Let's consider why the other options are not typically mounted on the rear end:

- **Vibration Damper:** Often called a harmonic balancer, this component is usually found at the *front* end of the crankshaft. It's designed to absorb torsional vibrations originating from the engine, preventing damage to the crankshaft.
- **Timing Sprocket:** This gear or sprocket is also mounted on the *front* end of the crankshaft. It meshes with the timing belt or chain to synchronize the crankshaft's rotation with the camshaft(s), ensuring correct valve timing.
- **Counterweight:** Counterweights are not separate components mounted on the end; they are integral parts of the crankshaft itself, built into the crank webs. Their function is to balance the mass of the pistons and connecting rods, reducing vibrations caused by uneven rotating forces.

Identifying the Correct Crankshaft Component

Considering the specific function and common placement within an engine's structure, the **flywheel** is the component that is mounted on the **rear end** of the **crankshaft**.

134. Answer: d

Explanation:

Understanding the Splash System in Vehicles

In automotive engineering, various systems work together to ensure a vehicle operates smoothly and efficiently. One such mechanism relates to how different parts are kept lubricated. Let's explore the function of a splash system and its connection to the core systems within a vehicle.

Splash System's Role in Lubrication

A **splash system**, also known as splash lubrication, is a simple method used primarily in engines, especially older designs or simpler engines like those in lawnmowers or small generators. Its main purpose is to lubricate moving parts that might not have direct oil lines.

- **Mechanism:** In this system, components like the crankshaft or connecting rods dip into a reservoir of oil (the oil sump). As they rotate, they splash the oil upwards, throwing it onto other engine parts like cylinder walls, pistons, and camshafts.
- **Function:** The primary goal is to reduce friction and wear between moving metal parts. By coating these surfaces with a thin layer of oil, the system prevents direct metal-to-metal contact, which could cause significant damage and overheating.
- **Cooling Effect:** While its main job is lubrication, the oil splashed around also helps to carry away some heat from the parts it contacts, contributing indirectly to temperature management.

Comparing Splash System with Other Vehicle Systems

It's important to differentiate the splash system from other vital vehicle systems:

- **Exhaust System:** This system is responsible for collecting and expelling burnt gases (exhaust) from the engine cylinders. It has no direct connection to oil splashing mechanisms.
- **Electrical System:** This system manages the vehicle's battery, alternator, starter motor, ignition, and lights. It deals with electrical energy and has no relation to the splash system's mechanical lubrication function.
- **Cooling System:** This system's primary role is to prevent the engine from overheating by circulating coolant (like antifreeze and water) through the engine block and radiator. While lubrication also helps manage heat, the cooling system uses a dedicated fluid and circulation method (water pump, radiator, fan) entirely separate from splash lubrication.

The Critical Connection: Splash System and Lubrication

The **splash system** is fundamentally a method of delivering lubrication to engine components. It's a specific technique within the broader **lubrication system** of a vehicle.

- The **lubrication system** as a whole ensures all crucial moving parts receive adequate oil. This oil reduces friction, prevents wear, cleans the engine, and helps dissipate heat.
- Splash lubrication is one way the lubrication system achieves this, particularly for parts that are difficult to reach with pressurized oil jets or galleries.
- Therefore, the splash system is an integral part of, or a method employed by, the vehicle's overall lubrication system.

135. Answer: a

Explanation:

Understanding the Solenoid Switch in Automotive Circuits

A solenoid switch is an electrical component that uses an electromagnet to operate a mechanical switch. In vehicles, solenoid switches play a crucial role, especially in systems that require handling high electrical currents, like starting the engine.

Role of Solenoid Switch in the Starting Circuit

The primary function of a solenoid switch in a vehicle is found within the starting circuit. The starter motor, which is responsible for cranking the engine to start it, draws a very large amount of current from the battery. A standard ignition switch inside the passenger cabin cannot directly handle this high current without risking damage or overheating.

This is where the solenoid switch comes in. It acts as a heavy-duty relay. When you turn the ignition key to the 'start' position, a small amount of current flows from the battery through the ignition switch to activate the solenoid's electromagnet. This relatively low current signal energizes the solenoid coil.

Once energized, the solenoid performs two main actions:

- It moves a plunger or core which pushes the starter motor's pinion gear forward to mesh with the engine's flywheel ring gear.
- It closes a set of large electrical contacts within the solenoid itself. These contacts are designed to handle the high current needed by the starter motor.

By closing these heavy contacts, the solenoid directly connects the battery's positive terminal to the starter motor's main power terminal, allowing the large current flow required to turn the engine over.

Therefore, the solenoid switch acts as a bridge between the low-current ignition switch and the high-current starter motor, safely controlling the power delivery needed to start the engine. This specific function is essential for the operation of the vehicle's starting system.

Why Other Circuits Are Less Likely for This Primary Use

Let's consider the other options:

- **Charging circuit:** The charging circuit, involving the alternator and battery, primarily deals with generating and regulating electrical power to recharge the battery and power the vehicle's electrical systems while the engine is running. While relays are used, a heavy-duty solenoid switch like the one for the starter is not typically the main component in this circuit's core function.
- **Ignition circuit:** The ignition circuit provides the spark to ignite the fuel mixture in gasoline engines. This involves components like the ignition coil, spark plugs, and ignition control module. This circuit generally operates at much lower currents than the starter motor circuit, and while it is controlled by the ignition switch, it does not require a heavy-duty solenoid switch for its primary power switching function.
- **Electrical circuit:** This is a very broad term that could refer to any circuit in the vehicle. While solenoid switches (and other types of solenoids) might be used in various electrical applications (like door locks or fuel injectors, though often integrated differently), their most prominent and essential role as a high-current switch for starting is in the starting circuit specifically.

Based on its function of handling high current and engaging the starter motor gear, the solenoid switch is most famously and critically used in the vehicle's starting circuit.

Summary: Solenoid Switch Application

The table below summarizes the typical use of a solenoid switch in the given options.

Circuit Type	Typical Use of Solenoid Switch
Starting circuit	Used as a high-current relay to power the starter motor and engage the pinion gear. This is its primary application.
Charging circuit	Not typically used as a high-current switch for the main charging function.
Ignition circuit	Not typically used as a high-current switch for the main ignition function.
Electrical circuit	A broad term; solenoids may have other uses, but the high-current switch application is key in starting.

Therefore, the solenoid switch is most commonly and critically used in the starting circuit.

Revision Table: Key Concepts

Term	Definition/Function
Solenoid Switch	An electromagnetic switch used to handle high electrical currents and often actuate a mechanical component.
Starting Circuit	The electrical system in a vehicle responsible for cranking the engine using the starter motor.
Starter Motor	An electric motor that rotates the engine's crankshaft to start it. Draws high current.
Ignition Switch	The main switch operated by the driver to control the vehicle's electrical systems, including the starting circuit (via the solenoid).
Pinion Gear	A small gear on the starter motor that engages with the flywheel ring gear to crank the engine.

Additional Information: Solenoid Switch Operation

The solenoid switch works on the principle of electromagnetism. When current flows through a coil of wire wrapped around a metallic core, it creates a magnetic field. In the case of a starter solenoid, this magnetic field pulls a metal plunger. This plunger movement serves a dual purpose:

1. Mechanical Action: The plunger is linked to a lever that pushes the starter pinion gear forward to engage the flywheel ring gear.
2. Electrical Switching: The same plunger often has a contact plate that bridges two large terminals, completing the high-current path between the battery and the starter motor windings.

When the ignition key is released from the 'start' position, the current to the solenoid coil is cut off. The magnetic field collapses, a return spring pulls the plunger back, disengaging the pinion gear from the

flywheel and opening the high-current contacts, stopping the starter motor.

This efficient and robust design allows a low-current control signal from the driver to activate a powerful motor that requires a large amount of electrical energy.

136. Answer: c

Explanation:

Understanding Automotive Hoist Equipment

Automotive workshops and garages use various types of equipment for vehicle maintenance and repair. These tools serve different purposes, from lifting vehicles to diagnosing engine problems or changing tyres. The question asks us to identify which piece of equipment falls under the category of a 'hoist'.

A hoist, in a general sense, is a mechanical device used for lifting or lowering a load by means of a drum or lift-wheel around which rope or chain wraps. In the automotive context, a hoist is primarily used to lift vehicles to allow mechanics access to the undercarriage for inspection, maintenance, or repair.

Let's examine the given options:

- **Wheel aligner:** This equipment is used to measure and adjust the angles of the wheels (camber, caster, toe) relative to the vehicle and the road. It does not lift the vehicle itself, although the vehicle might be placed on a lift or pit for the alignment process.
- **Tyre changer:** This machine is used to remove old tyres from wheels and mount new ones. It works directly with the wheel and tyre assembly and does not lift the entire vehicle.
- **Two post lift:** This is a common type of vehicle lift found in garages. It consists of two vertical posts with movable arms that support and lift the vehicle off the ground. This equipment directly performs the function of lifting a vehicle, fitting the definition of a hoist or lift in the automotive context.
- **Engine scanner:** This is a diagnostic tool that connects to a vehicle's onboard computer system to read fault codes and access data related to the engine and other systems. It is used for diagnosis and troubleshooting, not for lifting the vehicle.

Based on the functions described, the equipment that fits the description of a hoist in the automotive category is the Two post lift, as its primary purpose is to lift the vehicle.

Analysis of Equipment Functions

Equipment	Primary Function	Fits Hoist Category?
Wheel aligner	Measures and adjusts wheel angles	No
Tyre changer	Removes and mounts tyres on wheels	No
Two post lift	Lifts the entire vehicle off the ground	Yes
Engine scanner	Diagnoses vehicle electronic systems	No

The analysis clearly shows that the Two post lift is the only equipment among the options whose function is to lift a vehicle, aligning it with the purpose of a hoist in a garage setting.

Conclusion on Automotive Hoists

In the context of automotive repair equipment, a hoist is essentially a lift used to elevate a vehicle. Out of the given options, the Two post lift is designed specifically for this purpose, allowing mechanics access to the vehicle's undercarriage.

Revision Table: Automotive Equipment

Equipment Type	Key Use
Hoists/Lifts	Lifting vehicles for undercarriage access (e.g., Two post lift, Four post lift, Scissor lift)
Wheel & Tyre Service	Working on wheels and tyres (e.g., Tyre changer, Wheel balancer, Wheel aligner)
Diagnostic Equipment	Identifying and troubleshooting issues (e.g., Engine scanner, Multimeter, Oscilloscope)
Hand Tools	Manual tasks (e.g., Wrenches, Sockets, Screwdrivers)

Additional Information on Vehicle Lifts

Vehicle lifts, often referred to as hoists, come in various types depending on the intended use, space availability, and lifting capacity. Some common types include:

- **Two Post Lifts:** Ideal for general repair and maintenance, allowing clear access to wheels and undercarriage.
- **Four Post Lifts:** Often used for vehicle storage, parking, or alignment services. Vehicles drive onto ramps.
- **Scissor Lifts:** Lift vehicles straight up via a scissor mechanism. Can be surface-mounted or in-ground, good for limited space.
- **In-Ground Lifts:** Installed directly into the floor, offering a very clean workspace.

- **Mobile Column Lifts:** Portable lifts often used for heavy-duty vehicles like trucks and buses.

Each type of lift offers distinct advantages for different types of automotive work.

137. Answer: b

Explanation:

Understanding Engine Oil Consumption

Engine oil is crucial for lubricating the internal moving parts of an engine, reducing friction, and transferring heat. While a small amount of oil consumption is normal in most engines, excessively high engine oil consumption can indicate a problem within the engine.

Factors Affecting High Engine Oil Consumption

Several issues can lead to increased engine oil consumption. We will analyze the given options to identify the primary cause among them.

- **Compression pressure is weak:** Weak compression can be caused by various issues like worn piston rings, worn cylinder bore, damaged valves, or head gasket problems. While weak compression might correlate with some causes of high oil consumption (like worn rings), weak compression itself isn't the direct cause of the oil being consumed.
- **Piston rings are worn out:** Piston rings form a seal between the piston and the cylinder wall. There are typically compression rings (to seal the combustion chamber) and oil control rings (to scrape excess oil off the cylinder walls). If these rings, especially the oil control rings, are worn, they cannot effectively scrape the oil away. This allows oil to remain on the cylinder walls during the power stroke, where it gets exposed to the heat of combustion and burns. Burning oil is a major reason for high engine oil consumption.
- **Cylinder bore is small:** The size of the cylinder bore is a design specification. A cylinder bore being 'small' relative to its design doesn't directly cause high oil consumption. Problems arise when the bore is worn, scored, or out-of-round, which would then prevent the piston rings from sealing properly, leading to oil consumption.
- **Oil level is low:** A low oil level is a *symptom* of high oil consumption or a leak, not a *cause* of high oil consumption. High consumption means the rate at which oil is being used up is high, leading to the level dropping quickly if not topped up.

Why Worn Piston Rings Cause High Oil Consumption

Let's focus on why worn piston rings are a significant cause of high engine oil consumption.

- The piston rings are designed to fit snugly against the cylinder wall and in their grooves on the piston.
- Oil control rings specifically have a design (like slots or holes) that allows scraped oil to drain back into the oil pan.

- When these rings wear down, their ability to maintain a tight seal against the cylinder wall diminishes.
- This allows engine oil that is splashing or being thrown onto the cylinder walls to pass by the rings and enter the combustion chamber above the piston.
- In the combustion chamber, this excess oil burns during the power stroke, producing exhaust gases that may appear bluish.
- The constant burning of oil results in the oil level in the sump dropping rapidly, indicating high engine oil consumption.

Therefore, among the given options, worn-out piston rings directly and significantly lead to high engine oil consumption because they fail to control and scrape oil effectively from the cylinder walls, allowing it to burn in the combustion chamber.

Revision Table: Engine Oil Issues

Issue	Effect on Engine Oil Consumption	Explanation
Worn Piston Rings	High Consumption	Fail to scrape oil from cylinder walls; oil burns in combustion chamber.
Weak Compression	Can be related (e.g., worn rings/valves) but not a direct cause of oil burning itself.	Indicates poor sealing, potentially allowing combustion gases past rings or intake/exhaust leaks, but doesn't directly put oil into combustion chamber like worn oil rings.
Small Cylinder Bore	Indirect (if worn)	Bore dimensions are design; wear/damage to bore surface causes issues with rings sealing, leading to consumption.
Low Oil Level	Symptom, not Cause	Level drops because oil is being consumed or leaked, doesn't *cause* the high rate of consumption.

Additional Information on Engine Oil Consumption and Piston Rings

Understanding the role of piston rings is key to diagnosing many engine problems, including high engine oil consumption.

- **Types of Piston Rings:** Most pistons have two or three rings. The top rings are typically compression rings, designed to seal the combustion pressure. The bottom ring is usually the oil control ring.
- **Causes of Ring Wear:** Piston rings can wear out due to high mileage, improper lubrication, contaminated oil, or overheating. Cylinder wall wear also contributes to poor ring seal.
- **Symptoms of High Oil Consumption:** Besides a frequently low oil level, other symptoms include blue smoke from the exhaust (especially on startup or acceleration), decreased engine performance, reduced fuel economy, and fouled spark plugs.
- **Prevention:** Regular oil changes with the correct type and viscosity of oil, maintaining proper coolant levels to prevent overheating, and addressing any engine issues promptly can help prevent excessive

wear on piston rings and cylinder walls.

138. Answer: a

Explanation:

Understanding the Four-Stroke Diesel Engine Cycle

A four-stroke diesel engine, like other internal combustion engines, operates through a cycle involving four distinct strokes of the piston within the cylinder. These strokes are: Suction (or Intake), Compression, Power (or Expansion), and Exhaust. Each stroke corresponds to half a revolution of the crankshaft, completing the cycle in two full revolutions.

The Suction Stroke in a Diesel Engine

The question asks specifically about the **suction stroke** in a four-stroke **diesel engine**. During the suction stroke:

- The piston moves downwards from the Top Dead Centre (TDC) towards the Bottom Dead Centre (BDC).
- The inlet valve opens, while the exhaust valve remains closed.
- As the piston moves down, it increases the volume inside the cylinder, creating a vacuum or lower pressure compared to the intake manifold pressure.
- This pressure difference causes air (and **only air** in a diesel engine, unlike a petrol engine which draws in an air-fuel mixture) to be drawn into the combustion chamber through the open inlet valve.

Therefore, the primary event during the **suction stroke** of a four-stroke **diesel engine** is the intake of fresh air into the cylinder's combustion chamber.

Analyzing the Given Options

Let's examine each option in the context of a four-stroke diesel engine cycle:

- **Option 1: Air enters into the combustion chamber**

This accurately describes what happens during the **suction stroke** of a **diesel engine**, as explained above. Only air is drawn in because the fuel is injected later during the compression stroke.

- **Option 2: Air fuel mixture compresses**

This option describes the compression stroke, but it is incorrect for a diesel engine because an air-fuel mixture is not present during compression. In a diesel engine, only air is compressed during the compression stroke, and fuel is injected near the end of this stroke.

- **Option 3: Power is produced**

Power is produced during the power stroke (also known as the expansion stroke). This stroke occurs after combustion, when the hot, expanding gases push the piston down, generating work.

- **Option 4: Exhaust gases are exhausted**

This option describes the exhaust stroke. During this final stroke, the exhaust valve opens, the piston moves upwards from BDC to TDC, and the burnt gases are pushed out of the cylinder.

Based on the analysis of the four-stroke diesel engine cycle, the event that occurs during the **suction stroke** is the entry of air into the combustion chamber.

Four-Stroke Diesel Engine Cycle Summary

Stroke	Piston Movement	Valves	Event	Working Fluid
Suction (Intake)	TDC > BDC	Inlet Open, Exhaust Closed	Air enters cylinder	Air
Compression	BDC > TDC	Inlet Closed, Exhaust Closed	Air compressed, Fuel injected	Air
Power (Expansion)	TDC > BDC	Inlet Closed, Exhaust Closed	Combustion, High-pressure gases expand	Burnt Gases
Exhaust	BDC > TDC	Inlet Closed, Exhaust Open	Burnt gases expelled	Exhaust Gases

Conclusion on the Suction Stroke

In summary, the **suction stroke** is the first step in the four-stroke cycle of a **diesel engine**, where the piston's downward movement draws in fresh air. This air is then compressed in the subsequent stroke, setting the stage for combustion when fuel is injected.

Revision Table: Four-Stroke Diesel Engine Concepts

Key Actions in Diesel Engine Strokes

Stroke	What Happens?
Suction	Air enters the cylinder.
Compression	Air is compressed; Fuel is injected near the end.
Power	Combustion occurs; Expanding gases push piston.
Exhaust	Burnt gases are expelled.

Additional Information: Comparing Diesel and Petrol Engines

While both are four-stroke engines, a key difference lies in the **suction stroke** and ignition method:

- **Diesel Engine Suction:** Draws in only air. Ignition by compression ignition (high temperature of compressed air ignites fuel).
- **Petrol Engine Suction:** Draws in an air-fuel mixture (either from a carburetor or port/direct injection during intake). Ignition by spark plug.

Understanding this difference is crucial when studying internal combustion engines and the events of the **suction stroke**.

139. Answer: b

Explanation:

Understanding the Manifold Absolute Pressure Sensor Placement

The manifold absolute pressure (MAP) sensor is a crucial component in modern internal combustion engines, particularly those with electronic fuel injection. Its primary function is to measure the absolute pressure inside the intake manifold.

This pressure reading is vital because it indicates the load on the engine. At idle, the engine is drawing little air, creating a high vacuum (low absolute pressure) in the intake manifold. Under full throttle, the throttle plate is wide open, and the manifold pressure is close to atmospheric pressure (higher absolute pressure).

The engine control unit (ECU) uses the MAP sensor's signal, along with information from other sensors like the throttle position sensor and engine speed sensor, to calculate the correct amount of fuel to inject and sometimes adjust ignition timing.

Location of the MAP Sensor in an Engine

Based on its function – measuring the pressure within the intake manifold – the manifold absolute pressure sensor must be located where it can directly sense this pressure. The intake manifold is the network of tubes or passages that distributes the air (or air-fuel mixture in some older systems) from the throttle body (or carburetor) to the intake ports of the engine's cylinders.

Therefore, the logical and correct placement for the manifold absolute pressure sensor is on the intake manifold itself.

Analyzing the Placement Options

Let's examine the given options for the placement of the manifold absolute pressure sensor:

1. crankshaft: The crankshaft is part of the engine's rotating assembly and is located inside the engine block. It is related to engine speed and piston movement, not intake manifold pressure. A sensor located here would typically measure crankshaft position or speed, not manifold pressure.
2. inlet manifold: Also known as the intake manifold, this is the system that brings air into the engine cylinders. Placing the MAP sensor here allows it to measure the absolute pressure within this system, which is directly related to engine load and airflow. This is the correct location.
3. camshaft: The camshaft is part of the valve train, responsible for opening and closing the intake and exhaust valves. It is located either in the engine block or cylinder head. A sensor here measures camshaft position or speed, not manifold pressure.
4. exhaust manifold: The exhaust manifold collects exhaust gases from the cylinders and directs them out of the engine. The pressure and temperature in the exhaust manifold are vastly different from those in the intake manifold, and measuring exhaust pressure is not the function of a standard MAP sensor. Sensors in the exhaust system typically measure oxygen content or exhaust gas temperature.

Why the Intake Manifold is the Correct Location

The MAP sensor needs to measure the pressure that determines how much air is entering the combustion chambers. This pressure exists within the intake manifold between the throttle body and the intake valves. Changes in this pressure directly reflect changes in engine load. A low pressure (high vacuum) indicates low load (like idling), while a high pressure (near atmospheric) indicates high load (like wide-open throttle). The ECU uses this information to adjust engine parameters for optimal performance and efficiency.

Summary of Sensor Locations

Engine Component	Typical Sensors Located Here	Relevance to MAP Sensor
Crankshaft	Crankshaft Position Sensor	Measures rotational position/speed. Not for manifold pressure.
Inlet Manifold (Intake Manifold)	Manifold Absolute Pressure (MAP) Sensor, Intake Air Temperature (IAT) Sensor	Measures pressure in the intake system. Correct location for MAP sensor.
Camshaft	Camshaft Position Sensor	Measures rotational position/speed of camshaft relative to crankshaft. Not for manifold pressure.
Exhaust Manifold	Oxygen Sensor (O2 Sensor), Exhaust Gas Temperature (EGT) Sensor	Measures exhaust gas properties. Not for intake manifold pressure.

Conclusion on MAP Sensor Placement

To accurately measure the pressure within the intake system for engine control purposes, the manifold absolute pressure sensor is specifically designed and placed at the inlet manifold, also commonly referred to as the intake manifold.

Revision Table: Key Engine Sensors

Sensor Name	Common Location	Measured Parameter	Purpose
MAP Sensor (Manifold Absolute Pressure)	Intake Manifold	Absolute pressure in intake manifold	Engine load, fuel calculation
TPS (Throttle Position Sensor)	Throttle Body	Throttle plate angle	Driver input, engine load indication
CKP Sensor (Crankshaft Position)	Near Crankshaft Pulley or Flywheel	Crankshaft speed and position	Engine speed, ignition timing, injection timing
CMP Sensor (Camshaft Position)	Near Camshaft	Camshaft position relative to crankshaft	Cylinder identification (for sequential injection), timing checks
O2 Sensor (Oxygen Sensor)	Exhaust System (Exhaust Manifold/Pipe)	Oxygen content in exhaust gas	Measure air-fuel mixture richness/leanness

Additional Information on Engine Sensors and Manifolds

Understanding the different parts of an engine and their functions helps in knowing where sensors are placed and why. The manifolds are key components for managing airflow and exhaust gas flow.

- **Intake Manifold:** Distributes fresh air (and sometimes fuel) to the cylinders. Its design affects engine performance across different speed ranges. The pressure inside varies significantly with throttle position and engine speed, which is exactly what the MAP sensor measures.
- **Exhaust Manifold:** Collects burnt gases from the cylinders and directs them to the exhaust system. Conditions here (high temperature, pulsating pressure, exhaust gas composition) are monitored by different types of sensors.
- **Crankshaft and Camshaft:** These are rotating mechanical parts. Sensors monitoring them typically use magnetic pickups or Hall effect sensors to detect teeth or notches on wheels attached to these shafts, determining rotational speed and position.

The placement of each sensor is critical for it to accurately measure the parameter it is designed for, providing the ECU with the necessary data to control the engine efficiently.

140. Answer: b

Explanation:

Understanding the 12V Lead-Acid Battery Structure

A standard lead-acid battery is a type of rechargeable battery commonly used in vehicles and other applications. It is made up of several individual electrochemical cells connected together to achieve the desired voltage and current capacity.

Voltage of a Single Lead-Acid Cell

Each individual cell within a lead-acid battery produces a specific voltage. For a typical lead-acid cell, the voltage is approximately 2 volts (V). This is a fundamental characteristic of the chemistry involved in the cell.

Connecting Cells: Series vs. Parallel

To obtain higher voltages from multiple cells, they are connected in series. In a series connection, the positive terminal of one cell is connected to the negative terminal of the next cell, and so on. The total voltage of cells connected in series is the sum of the individual cell voltages.

To increase the current capacity (measured in Ampere-hours, Ah) while keeping the voltage the same, cells are connected in parallel. In a parallel connection, all positive terminals are connected together, and all negative terminals are connected together. The total current capacity is the sum of the individual cell capacities, but the voltage remains the same as that of a single cell (assuming identical cells).

Determining the Number of Cells for a 12V Battery

Given that a standard lead-acid cell produces approximately 2V, to build a 12V battery, we need to connect multiple cells in series to add up their voltages. Let N be the number of cells required.

The total voltage V_{total} from N cells in series, each with voltage V_{cell} , is given by the formula:

$$V_{total} = N \times V_{cell}$$

We want a total voltage of 12V, and each cell provides 2V. So, we can calculate the number of cells:

$$12V = N \times 2V/cell$$

$$N = \frac{12V}{2V/cell}$$

$$N = 6 \text{ cells}$$

Therefore, a 12V lead-acid battery consists of 6 cells.

Connection Type

Since we need to increase the voltage from 2V per cell to 12V total, the cells must be connected in series. A parallel connection of 2V cells would still result in a total voltage of only 2V.

Evaluating the Options

- **Option 1: three cells in parallel** – Three 2V cells in parallel would yield 2V (voltage doesn't add in parallel) with increased capacity. This is not 12V.
- **Option 2: six cells in a series** – Six 2V cells in series would yield $6 \times 2V = 12V$. This matches the requirement.
- **Option 3: six cells in parallel** – Six 2V cells in parallel would yield 2V (voltage doesn't add in parallel) with significantly increased capacity. This is not 12V.
- **Option 4: three cells in a series** – Three 2V cells in series would yield $3 \times 2V = 6V$. This is not 12V.

Based on the analysis, six cells connected in series are required to form a 12V lead-acid battery.

Summary of Cell Connections

Number of Cells	Connection Type	Resulting Voltage (assuming 2V/cell)
3	Parallel	2V
6	Series	12V
6	Parallel	2V
3	Series	6V

Revision Table: 12V Lead-Acid Battery

Component	Quantity	Connection Type	Purpose
Individual Lead-Acid Cells	6	Series	To sum the voltage of each cell to reach 12V total.

Additional Information on Lead-Acid Batteries

Lead-acid batteries are one of the oldest types of rechargeable batteries. They are widely used due to their low cost, reliability, and ability to provide high surge currents needed for engine starting in vehicles.

- **Internal Structure:** Each cell contains positive and negative plates (made of lead and lead dioxide) immersed in an electrolyte of sulfuric acid and water.
- **Working Principle:** During discharge, lead and lead dioxide react with sulfuric acid to form lead sulfate on the plates, releasing electrons and thus producing current. During charging, the process is reversed.
- **Battery Capacity (Ah):** This indicates how much current a battery can supply over a period. Connecting cells in parallel increases the total capacity.
- **Types:** Common types include Flooded Lead-Acid (wet cell), Sealed Lead-Acid (SLA), Gel Batteries, and Absorbed Glass Mat (AGM) batteries. While they differ in electrolyte handling, the basic cell

voltage (around 2V) remains consistent.

Understanding the construction, particularly the number and arrangement of cells, is key to understanding the voltage rating of lead-acid batteries like the common 12V automotive battery.

141. Answer: d

Explanation:

Understanding the Piston's Role in an Engine

A piston is a crucial component in many types of engines, including those found in cars, trucks, and motorcycles. It moves up and down inside a cylinder, driven by the expansion of gases during combustion. This reciprocating motion is then converted into rotational motion by the crankshaft, powering the vehicle. Because of its critical role, the material a piston is made of must withstand extreme conditions.

Essential Properties for Piston Materials

The harsh environment inside a combustion chamber places significant demands on the piston material. Key properties required include:

- **High Strength:** To withstand immense pressure forces during the power stroke.
- **High Temperature Resistance:** To operate reliably at temperatures reaching hundreds of degrees Celsius.
- **Good Thermal Conductivity:** To efficiently transfer heat away from the piston crown and prevent overheating.
- **Lightweight:** To minimize inertia, allowing the engine to rev higher and improving efficiency.
- **Wear Resistance:** To resist friction against the cylinder wall and piston rings.
- **Controlled Thermal Expansion:** The material should expand predictably with heat to maintain proper clearance within the cylinder.

Analyzing Potential Piston Materials

Let's consider the common materials listed as options and evaluate their suitability for pistons based on the required properties.

Aluminum Alloys for Pistons

Aluminum alloys are a very common material for pistons, especially in modern automotive gasoline engines and many diesel engines. They offer several key advantages:

- **Lightweight:** Aluminum is significantly lighter than steel or cast iron, reducing the overall weight of the reciprocating parts and improving engine performance and efficiency.

- **Excellent Thermal Conductivity:** Aluminum alloys conduct heat very well, helping to dissipate heat from the piston crown and reduce thermal stress.
- **Good Strength-to-Weight Ratio:** While not as strong as steel at high temperatures, alloys are developed to provide sufficient strength for typical engine pressures.
- **Cost-Effective:** Aluminum alloy pistons are generally more economical to manufacture in high volumes.

One challenge with aluminum is its relatively high thermal expansion compared to cast iron cylinder liners, which requires careful design considerations for piston clearance and sometimes specialized coatings or designs.

Cast Iron as a Piston Material

Cast iron was historically used for some pistons, particularly in larger, lower-speed engines. It offers good strength and wear resistance. However, its main disadvantages compared to aluminum alloys include:

- **Heavy Weight:** Cast iron is much heavier than aluminum, limiting engine speed and efficiency.
- **Lower Thermal Conductivity:** It does not transfer heat as effectively as aluminum.

Today, cast iron is more commonly used for engine blocks and cylinder liners, where its weight is less of a dynamic issue and its wear resistance is beneficial.

Steel and Carbon Steel for Pistons

Steel and carbon steel are very strong materials, capable of withstanding extremely high pressures and temperatures. They are sometimes used for pistons in very high-performance or heavy-duty diesel engines where ultimate strength is critical. However, they have drawbacks:

- **Heavy Weight:** Like cast iron, steel is much heavier than aluminum, negatively impacting engine dynamics.
- **Lower Thermal Conductivity:** Steel conducts heat less effectively than aluminum, potentially leading to higher piston temperatures unless compensated for by design.
- **Higher Thermal Expansion (relative to cast iron cylinders):** Similar to aluminum, steel's expansion needs careful management.

While strong, steel pistons are less common in typical passenger car engines compared to aluminum alloy pistons.

Comparing Piston Material Options

Here is a brief comparison of the materials often considered for pistons:

Material	Weight	Thermal Conductivity	Strength	Typical Use for Pistons
Aluminum Alloy	Low	High	Good (suitable for most engines)	Common in automotive gasoline/diesel engines
Cast Iron	High	Medium	Good	Less common (historically, or specific applications)
Steel / Carbon Steel	High	Medium-Low	Very High (suitable for heavy-duty/high-performance)	Heavy-duty engines, some high-performance

Conclusion on Piston Materials

Based on the requirements for a piston to be lightweight, strong, and have good thermal properties, aluminum alloys are the most widely used material for pistons in the majority of modern internal combustion engines, particularly in the automotive sector. While other materials like cast iron or steel might be used in specific or historical applications, aluminum alloy pistons represent the dominant technology due to their balanced properties and manufacturing advantages.

Revision Table: Key Piston Material Facts

Piston Material Type	Key Advantage(s)	Key Disadvantage(s)
Aluminum Alloy Piston	Lightweight, High Thermal Conductivity	Higher Thermal Expansion than iron/steel
Cast Iron Piston	Good Wear Resistance	Heavy, Lower Thermal Conductivity
Steel Piston	Very High Strength	Heavy, Lower Thermal Conductivity than aluminum

Additional Information: Piston Technology Advancements

Modern piston technology involves more than just the base material. Manufacturers often use various techniques to improve piston performance and durability, such as:

- **Piston Coatings:** Applying special coatings (like graphite, Teflon, or ceramic) to the piston skirt and crown to reduce friction, improve wear resistance, and manage heat.
- **Reinforced Piston Designs:** Using reinforced areas, different alloy compositions, or even composite materials in specific parts of the piston for added strength where needed, such as around the piston pin boss or the piston crown.
- **Cooling Channels:** Incorporating internal cooling galleries within the piston to allow engine oil to circulate and remove heat from the piston crown, especially in high-performance or diesel engines.

These advancements, often combined with the benefits of aluminum alloys, contribute to the efficiency and longevity of modern pistons and engines.

142. Answer: c

Explanation:

Understanding Turbocharger Components

A turbocharger is a device used in internal combustion engines to increase power output. It works by forcing more air into the combustion chamber. This is achieved by using the waste energy from the exhaust gases to spin a turbine, which in turn spins a compressor.

Turbocharger Turbine Connection Explained

A turbocharger consists of two main sections: the turbine and the compressor. These sections are connected by a shaft. Let's look at where the turbine side connects:

- **Turbine Side:** This part is located in the path of the engine's exhaust gases. The hot, high-pressure exhaust gas flows through the turbine housing and spins the turbine wheel. This rotational energy is used to power the compressor.
- The turbine wheel housing needs to be connected directly to the source of the exhaust gases.

Considering the flow of exhaust gases from the engine cylinders, they exit through the exhaust valves and collect in the exhaust manifold before being routed out of the vehicle.

Therefore, the turbine wheel housing, which captures the energy from the exhaust gases, is connected to the exhaust manifold.

Analyzing the Options

Let's examine why the other options are not where the turbine wheel housing is connected:

- **Flywheel side:** The flywheel is located at the back of the engine, connected to the crankshaft. It stores rotational energy and is part of the drivetrain system. It has no direct connection point for the turbocharger's turbine housing.
- **Front of the engine block:** Various components like the water pump, alternator, or accessory drives might be located at the front of the engine block, but not the exhaust manifold where the turbocharger's turbine housing attaches.
- **Intake manifold:** The intake manifold is on the opposite side of the engine's airflow path compared to the exhaust. It distributes fresh air (or air/fuel mixture) to the cylinders after it has been compressed by the turbocharger's compressor side. The compressor outlet connects to the intake manifold, not the turbine housing.

- **Exhaust manifold:** This component collects the exhaust gases from all the cylinders into a single pipe. It is the correct location for the turbocharger's turbine housing connection because the turbine uses the energy from these exhaust gases.

Based on how a turbocharger functions and the path of exhaust gases, the turbine wheel housing must be connected to the exhaust manifold to capture the energy from the engine's exhaust.

Turbocharger Component Connections

Component	Connects To	Function
Turbine Wheel Housing	Exhaust Manifold	Receives exhaust gas to spin turbine wheel.
Compressor Wheel Housing Outlet	Intake Manifold	Delivers compressed air to engine cylinders.

Revision Table: Turbocharger Connections

Review the key connections in a turbocharger system:

- Turbine Housing → Exhaust Manifold (Exhaust Gas Input)
- Turbine Housing Outlet → Exhaust System (Exhaust Gas Output)
- Compressor Housing Inlet → Air Filter/Intake Pipe (Fresh Air Input)
- Compressor Housing Outlet → Intake Manifold (Compressed Air Output)

Additional Information on Turbocharger Operation

Understanding the flow of gases is crucial to understanding turbochargers:

- Exhaust gases leave the engine cylinders and enter the exhaust manifold.
- From the exhaust manifold, these gases flow into the turbine housing, spinning the turbine wheel.
- The spinning turbine wheel rotates the shaft connected to the compressor wheel.
- The compressor wheel draws in fresh air from the intake system.
- The compressor wheel compresses this air, increasing its density and pressure.
- This compressed air is then pushed into the intake manifold and subsequently into the engine cylinders for combustion.
- This process allows more fuel to be burned, resulting in increased power compared to a naturally aspirated engine of the same size.

The connection of the turbine wheel housing to the exhaust manifold is fundamental to harnessing the energy from the exhaust gases to drive the turbocharger.

143. Answer: d

Explanation:

Understanding the Gudgeon Pin's Key Function in Engines

The question asks about the primary function of a component known as the gudgeon pin within an engine. This component plays a crucial role in the internal combustion engine mechanism.

Let's break down the possible connections mentioned in the options and identify the correct one based on standard engine design principles.

What is a Gudgeon Pin?

The gudgeon pin, also often called a piston pin or wrist pin, is a cylindrical shaft that connects the piston to the connecting rod. This connection allows the linear motion of the piston (moving up and down in the cylinder) to be converted into the rotational motion of the crankshaft via the connecting rod.

Analyzing the Options

- **Connect engine to gear box:** The engine is connected to the gearbox typically through the clutch (in manual transmissions) or a torque converter (in automatic transmissions). The gudgeon pin is an internal engine component and has no direct role in this connection.
- **Connect crankshaft to flywheel:** The crankshaft is connected to the flywheel, usually by bolts, at one end. The flywheel stores rotational energy and smooths out the engine's power delivery, and also serves as part of the clutch assembly or torque converter mount. The gudgeon pin is not involved in this connection.
- **Connect valve to camshaft:** Valves are connected to the camshaft through various mechanisms, such as tappets, pushrods, and rocker arms (depending on the engine design - OHV, OHC, DOHC). The camshaft's lobes push on these components to open and close the valves. The gudgeon pin is part of the piston/connecting rod assembly and unrelated to the valve train.
- **Connect piston to connecting rod:** This describes the exact location and function of the gudgeon pin. It passes through the piston bosses (holes in the piston skirt) and the small end bearing of the connecting rod, forming a pivot point that allows the connecting rod to swing as the piston moves up and down.

Based on the analysis of standard engine components and their functions, the gudgeon pin's role is to connect the piston to the connecting rod.

The Role of the Gudgeon Pin Connection

The gudgeon pin facilitates the transfer of force generated by the combustion process from the piston to the connecting rod. This force then pushes or pulls the connecting rod, causing the crankshaft to rotate. The connection is a pivot, allowing the connecting rod to change its angle relative to the piston as the crankshaft spins.

Importance of the Gudgeon Pin

The gudgeon pin is subjected to significant stress and force during engine operation. It must be strong enough to withstand the combustion pressure and the inertial forces of the piston and connecting rod. It

also requires proper lubrication to allow smooth movement at the connecting rod's small end.

Summary of Engine Component Connections

Component 1	Component 2	Connecting Part
Piston	Connecting Rod	Gudgeon Pin / Piston Pin
Connecting Rod	Crankshaft	Big End Bearing
Crankshaft	Flywheel	Bolts
Engine	Gearbox	Clutch / Torque Converter
Valve	Camshaft (indirectly)	Tappet, Pushrod, Rocker Arm etc.

Therefore, the function of the gudgeon pin is unequivocally to connect the piston to the connecting rod.

Revision Table: Gudgeon Pin Function

Key Facts about Gudgeon Pin			
Component Name	Other Names	Primary Function	Location
Gudgeon Pin	Piston Pin, Wrist Pin	Connects piston to connecting rod	Passes through piston bosses and connecting rod small end

Additional Information: Engine Components

Let's explore some other key components mentioned or related to the options to enhance our understanding of engine mechanics.

- **Piston:** A cylindrical component that moves up and down inside the cylinder. It receives the force from combustion and transmits it to the connecting rod.
- **Connecting Rod:** A rod that connects the piston to the crankshaft. It converts the reciprocating (up and down) motion of the piston into the rotary (circular) motion of the crankshaft. It has a small end (connected to the piston via the gudgeon pin) and a big end (connected to the crankshaft).
- **Crankshaft:** The main rotating shaft of the engine. It receives power from the connecting rods and drives the flywheel, and through it, the rest of the vehicle's drivetrain.
- **Flywheel:** A heavy wheel attached to the crankshaft. It smooths engine rotation and stores energy.
- **Camshaft:** A shaft with lobes that push open the engine's intake and exhaust valves at the correct time during the combustion cycle.
- **Valves:** Control the flow of air-fuel mixture into the cylinder (intake valve) and exhaust gases out of the cylinder (exhaust valve).

144. Answer: b

Explanation:

Understanding Power Production in a Single Cylinder Diesel Engine

A single-cylinder diesel engine, like most internal combustion engines, operates on a cycle that requires multiple revolutions of the crankshaft to produce power. Diesel engines typically use a four-stroke cycle to convert fuel energy into mechanical work.

The four strokes in a diesel engine cycle are:

1. Intake Stroke
2. Compression Stroke
3. Combustion (or Power) Stroke
4. Exhaust Stroke

Let's look at what happens during each stroke and how it relates to the crankshaft's rotation:

- **Intake Stroke:** The intake valve opens, and the piston moves downwards, drawing air into the cylinder. This stroke corresponds to 180 degrees of crankshaft rotation.
- **Compression Stroke:** The intake valve closes, and the piston moves upwards, compressing the air in the cylinder. This stroke also corresponds to 180 degrees of crankshaft rotation. At the end of this stroke, fuel is injected into the hot compressed air, which ignites the fuel.
- **Combustion (Power) Stroke:** The burning fuel creates high-pressure gases that push the piston downwards. This is the only stroke where the engine produces usable power (work). This stroke corresponds to another 180 degrees of crankshaft rotation.
- **Exhaust Stroke:** The exhaust valve opens, and the piston moves upwards, pushing the burnt gases out of the cylinder. This final stroke corresponds to the last 180 degrees of crankshaft rotation.

To complete all four strokes, the crankshaft rotates a total of $180^\circ + 180^\circ + 180^\circ + 180^\circ = 720^\circ$. Since one full revolution of the crankshaft is 360 degrees, 720 degrees is equal to two full revolutions ($720^\circ / 360^\circ = 2$).

Power is produced only during the Combustion (Power) stroke, which happens once every complete cycle. Since one complete cycle requires two revolutions of the crankshaft, a single-cylinder diesel engine produces power once every 2 revolutions.

Consider the sequence of events and crankshaft rotation:

Stroke	Crankshaft Rotation	Action	Power Produced?
Intake	0° to 180° (1st revolution)	Air drawn in	No
Compression	180° to 360° (1st revolution)	Air compressed	No
Combustion (Power)	360° to 540° (2nd revolution)	Fuel ignites, piston pushed down	Yes
Exhaust	540° to 720° (2nd revolution)	Burnt gases expelled	No

As shown in the table, the power stroke occurs during the second revolution. Therefore, power is generated during one stroke out of the four, and the complete cycle takes 2 revolutions.

Revision Table: Single Cylinder Diesel Engine Cycle

Cycle Event	Number of Strokes	Crankshaft Revolutions	Power Production
One Complete Cycle	4	2	Occurs once per cycle (during Power stroke)

Additional Information on Engine Cycles

While the four-stroke cycle is most common for diesel engines in vehicles and stationary applications, other engine types exist:

- **Two-Stroke Engines:** These engines complete the cycle in two strokes (one crankshaft revolution). They typically combine intake and compression, and combustion and exhaust, occurring simultaneously or overlapping significantly. Two-stroke engines produce a power stroke every revolution, but they are generally less fuel-efficient and produce more emissions than four-stroke engines.
- **Six-Stroke Engines:** Experimental designs exist that add extra strokes for purposes like increased expansion or air injection for better scavenging. These are not widely used in commercial applications.

The question specifically refers to a standard single-cylinder diesel engine, which implies the common four-stroke design used in most applications where fuel efficiency and emissions control are important.

145. Answer: d

Explanation:

Understanding the Vehicle Lubrication System Components

The lubrication system in a vehicle's engine is essential for its proper functioning and longevity. Its main purpose is to reduce friction between moving parts, cool the engine, clean internal components by carrying away debris, and create a seal between the piston rings and cylinder walls. This system involves several key parts working together to circulate engine oil.

Let's examine the components mentioned in the options to see which one is part of the vehicle lubrication system:

- **Oil Filter:** This component is definitely part of the lubrication system. Its job is to clean the engine oil, removing dirt, metal particles, and other contaminants before the oil is circulated back through the engine.
- **Oil Sump:** Also known as the oil pan, the oil sump is the reservoir at the bottom of the engine where the oil collects when the engine is not running. It stores the oil that the lubrication system circulates. So, it is a fundamental part of the system.
- **Oil Crank:** This term does not represent a standard component of the vehicle lubrication system. Perhaps it's a misunderstanding or a typo for 'crankshaft', which is an engine part that *gets* lubricated, but is not a component *of* the lubrication system itself like a pump or filter.
- **Oil Pump:** This is a critical component of the vehicle lubrication system. The oil pump is responsible for drawing oil from the oil sump and circulating it under pressure to all the parts of the engine that require lubrication, such as bearings, pistons, and valves. Without the oil pump, the oil would not reach these vital parts.

Considering the roles of these components within the vehicle's lubrication process, the oil pump is a fundamental and active part of the system that drives the circulation of oil.

Explanation of Components in Lubrication System

To clarify further, here's a brief look at how some common parts fit into the system:

- **Oil Sump:** Stores the oil.
- **Oil Strainer:** A coarse filter usually located in the sump, preventing large debris from entering the oil pump.
- **Oil Pump:** Pressurizes and circulates the oil throughout the engine.
- **Pressure Relief Valve:** Controls oil pressure, diverting excess oil back to the sump.
- **Oil Filter:** Filters fine contaminants from the oil before it reaches engine components.
- **Oil Passages/Galleries:** Channels within the engine block and cylinder head through which oil flows to bearings and other parts.
- **Oil Jets/Squirters:** Found in some engines, they spray oil onto the underside of pistons for cooling and lubrication.
- **Oil Pressure Gauge/Sensor:** Monitors the oil pressure in the system.

Based on this, the Oil Pump is clearly a core component responsible for the function of the vehicle lubrication system.

Revision Table: Key Lubrication System Parts

Component	Function in Lubrication System
Oil Sump	Stores engine oil
Oil Strainer	Filters large particles before pump
Oil Pump	Circulates oil under pressure
Oil Filter	Cleans oil of contaminants
Pressure Relief Valve	Regulates oil pressure
Oil Passages	Internal channels for oil flow

Additional Information on Vehicle Lubrication

Understanding the lubrication system is crucial for maintaining vehicle health. Proper lubrication reduces wear and tear, prevents overheating by carrying heat away from engine parts, and helps in sealing the combustion chamber.

Different types of oil pumps exist, commonly gear pumps or rotor pumps. The oil pressure needs to be within a specific range for effective lubrication, which is why the pressure relief valve is important. Regularly changing the oil and oil filter is the most basic and essential maintenance step for the lubrication system, ensuring the oil remains clean and effective.

146. Answer: c

Explanation:

Understanding Diesel Engine Starting Issues

Diesel engines require precise timing and high fuel pressure for combustion. When a diesel engine struggles to start, it indicates a problem within the system that prevents proper ignition.

Why Air in the Fuel Injection Pump Causes Poor Starting

The fuel injection pump in a diesel engine is a critical component responsible for pressurizing diesel fuel and delivering it to the cylinders at the exact right moment for combustion. Diesel engines rely on the heat generated by compressing air in the cylinder to ignite the fuel, unlike petrol engines which use a spark plug.

If there is air present in the fuel system, particularly in the fuel injection pump or lines, it can severely disrupt the process:

- Air is compressible, while liquid diesel fuel is not.

- The injection pump is designed to work with an incompressible fluid (diesel).
- When the pump tries to build high pressure with air present, the pressure pulse is weak or doesn't build correctly.
- This results in insufficient fuel pressure or improper fuel delivery to the injectors.
- Without adequate, high-pressure fuel injection, the fuel cannot atomize correctly and mix with the hot compressed air in the cylinder, preventing proper combustion and starting.
- This condition is often referred to as 'air-locking' the fuel system, and it must be 'bled' (air removed) for the system to function correctly again.

Therefore, air in the fuel injection pump is a direct and significant cause of poor starting in a diesel engine.

Analyzing Other Potential Issues

Let's consider why the other options are less likely causes of poor diesel engine starting:

- **Low brake fluid level:** Brake fluid is part of the hydraulic braking system. Its level affects braking performance but has no direct connection to the engine's ability to start.
- **Loose radiator fan belt:** The fan belt typically drives the radiator fan and possibly the alternator and water pump. While a loose belt can lead to issues like overheating (due to the fan/water pump not running) or battery not charging properly (if the alternator is affected while running), it does not prevent the initial cranking or fuel injection required for starting.
- **Low engine power:** Low engine power is usually a symptom that manifests *after* the engine has started and is running, especially under load. It indicates an issue affecting the engine's efficiency or output (like worn components, fuel delivery problems while running, or exhaust issues), but it is not typically a cause that prevents the engine from starting in the first place. However, severe underlying issues causing very low compression or massive fuel delivery problems while running could also cause poor starting, but 'low engine power' itself describes the running condition, not the starting issue's root cause.

Based on the function of a diesel engine and its components, air in the fuel injection pump directly impacts the ability to achieve combustion during starting.

Revision Table: Diesel Engine Starting Problems

Potential Cause	Impact on Starting	Explanation
Air in Fuel Injection Pump	Severe impact, likely cause of poor starting.	Prevents fuel pressurization and proper injection needed for combustion.
Low Brake Fluid Level	No impact on starting.	Related to braking system, not engine ignition.
Loose Radiator Fan Belt	No direct impact on starting.	Affects cooling and charging system while engine is running.
Low Engine Power	Usually a symptom, not a cause of poor starting.	Describes engine performance while running; root cause might also affect starting.

Additional Information: Diesel Fuel System

Understanding the key components of the diesel fuel system helps diagnose starting issues:

- **Fuel Tank:** Stores the diesel fuel.
- **Fuel Lines:** Carry fuel from the tank to the engine components.
- **Fuel Filter:** Removes contaminants from the fuel before it reaches the pump and injectors. A clogged filter can also cause fuel delivery issues.
- **Fuel Lift Pump (sometimes present):** Pumps fuel from the tank to the injection pump.
- **Fuel Injection Pump:** The heart of the system, increases fuel pressure dramatically and meters the correct amount of fuel to each cylinder at the right time.
- **Injectors:** Receive high-pressure fuel from the pump and spray (atomize) it into the combustion chamber.

Air can enter the system through loose connections, failed seals, running out of fuel, or improper bleeding after maintenance (like changing the fuel filter). Getting air out of the system is often necessary after certain repairs.

147. Answer: b

Explanation:

Understanding Electrical Conductors

An electrical conductor is a material that allows electric current to flow through it easily. This is because these materials have free electrons or ions that can move readily when a voltage is applied across them. Materials that do not allow electricity to flow easily are called electrical insulators.

Let's examine the given options to determine which one is an example of an electrical conductor.

Analyzing the Options

- **Porcelain:** Porcelain is a ceramic material. Ceramics generally have tightly bound electrons and do not have free charge carriers. Therefore, porcelain is typically a good electrical insulator and is used in applications like electrical insulators for power lines.
- **Copper:** Copper is a metal. Metals are known for having a "sea" of free electrons in their structure, which can move freely and carry electric current. Copper is widely used in electrical wiring because it is an excellent electrical conductor, second only to silver in terms of conductivity among common metals.
- **Plastic:** Plastics are polymers. Most common plastics have a molecular structure where electrons are tightly bound within covalent bonds and there are no free charge carriers. Thus, plastics are generally excellent electrical insulators and are used to cover electrical wires to prevent electric shock and short circuits.

- **Wood:** Wood is an organic material. Dry wood is a poor conductor of electricity and acts as an insulator. However, the conductivity of wood can increase significantly if it is wet, as the water containing dissolved ions can provide charge carriers. But in its typical dry state, wood is considered an insulator.

Based on the properties of these materials, copper stands out as the primary example of a good electrical conductor among the options provided.

Conductor vs. Insulator Comparison

Property	Electrical Conductor	Electrical Insulator
Electron Movement	Easy (Free electrons/ions)	Difficult (Tightly bound electrons)
Electric Current Flow	High conductivity (σ)	Low conductivity (σ)
Examples	Copper, Aluminum, Gold, Silver, Water (with impurities)	Plastic, Rubber, Glass, Porcelain, Dry Wood

The ability of a material to conduct electricity is quantified by its electrical conductivity (σ) or inversely by its electrical resistivity (ρ). Conductors have high conductivity (low resistivity), while insulators have low conductivity (high resistivity).

Conclusion on Electrical Conductors

Considering the definitions and examples of electrical conductors and insulators, copper is clearly an example of a material that allows electricity to flow easily, classifying it as an electrical conductor.

Revision Table: Electrical Properties

Key Electrical Properties Review

Term	Definition	Relevance
Electrical Conductor	Material allowing easy flow of electric current.	Used in wires, circuits.
Electrical Insulator	Material resisting flow of electric current.	Used for safety, preventing shorts.
Conductivity (σ)	Measure of a material's ability to conduct charge.	Higher for conductors. Unit: Siemens per meter (S/m).
Resistivity (ρ)	Measure of a material's opposition to current flow.	Higher for insulators. Unit: Ohm-meter ($\Omega \cdot m$).

Additional Information: Applications of Conductors and Insulators

Electrical conductors and insulators play crucial roles in countless technologies. Conductors like copper and aluminum are essential for transmitting electrical power over long distances and distributing it within buildings. They form the wires and traces on circuit boards that connect electronic components.

Insulators like plastic, rubber, and ceramic are vital for safety and functionality. They are used to:

- Cover wires to prevent electric shock.
- Separate components in circuits to prevent unwanted current paths (short circuits).
- Support high-voltage lines safely (e.g., porcelain insulators).
- Provide protective casings for electrical devices.

Understanding the difference between conductors and insulators is fundamental to studying electricity and electronics.

148. Answer: b

Explanation:

Understanding Washers for Locking Nuts

Washers serve various purposes in bolted joints, such as distributing load, providing spacing, reducing friction, and importantly, preventing loosening due to vibration or thermal expansion/contraction. When the primary goal is to lock a nut securely in place against rotation, specific types of washers are employed as locking devices. Let's examine the given options to determine which type of washer is typically used for this purpose.

Analyzing the Options for Nut Locking

- **Locking device:** This is a general term. Washers that prevent nuts from loosening are a type of locking device, but this option doesn't specify the particular type of washer.
- **Tab washer:** A tab washer is designed with one or more tabs that can be bent up against a face of the nut or a surrounding component (like the edge of a plate or a special slot). This physical interference prevents the nut from rotating. They are considered a positive locking mechanism.
- **Locking washer with lug:** This option is similar to a tab washer, likely referring to a washer with a specific lug or projection used for locking. This aligns closely with the function of a tab washer. However, "Tab washer" is a widely recognized term for this type of positive locking washer.
- **Spring washer:** Spring washers (like split lock washers, wave washers, or conical spring washers) work by providing a spring force that maintains tension in the joint. While they can help resist loosening under some conditions, they are not considered positive locking devices like tab washers. Their effectiveness can be limited under severe vibration or if the joint relaxes significantly.

Why Tab Washers are Used for Locking Nuts

Tab washers are specifically engineered to provide a positive lock. One or more tabs are initially flat. After the nut is tightened to the required torque, the relevant tab(s) are manually bent upwards against one of the flats of the nut (or bent down into a slot in the underlying component if the washer has an internal tab). This physical restraint prevents the nut from rotating relative to the washer and the bolted joint assembly. This makes them highly effective in applications where security against loosening is critical.

Washer Type	Primary Function for Nuts	Locking Mechanism	Considered Positive Locking?
Plain Washer	Load distribution, spacing	None	No
Spring Washer (e.g., Split Lock)	Maintain tension, slight resistance to loosening	Friction, spring force	No
Tab Washer	Prevent rotation (positive lock)	Physical interference (bent tab)	Yes

Based on their design and function, tab washers are a classic and effective type of washer used specifically for locking a nut by physically preventing its rotation.

Revision Table: Key Washer Types and Uses

Washer Type	Common Use Case
Flat/Plain Washer	Distributing load, protecting surfaces, spacing
Split Lock Washer	Adding tension, resisting loosening due to vibration (limited positive lock)
Internal Tooth Lock Washer	Locking against rotation by biting into surfaces (more effective on softer materials)
External Tooth Lock Washer	Locking against rotation by biting into surfaces (more effective on softer materials)
Wave Washer	Providing spring force over small deflection ranges
Conical Spring Washer (Belleville)	Providing spring force over larger deflection ranges, maintaining tension
Tab Washer	Positive locking of nut or bolt head using bent tabs

Additional Information on Locking Washers and Devices

While tab washers provide a strong positive lock, other methods and devices are also used to prevent nuts and bolts from loosening. These include:

- **Nylock Nuts (Nylon Insert Lock Nuts):** These have a nylon insert that deforms to grip the bolt threads, creating friction.
- **All Metal Lock Nuts:** These use distorted threads or a cone shape to create interference and friction.
- **Jam Nuts:** Using two nuts tightened against each other on the bolt thread.
- **Thread Locking Adhesives:** Liquid or semi-solid substances applied to the threads that cure to a hard plastic, filling gaps and bonding the threads together.
- **Castle Nuts and Cotter Pins:** A castle nut has slots cut into the top. A cotter pin or safety wire is passed through a hole in the bolt and through the slots in the nut to prevent rotation.

Each method has its advantages and disadvantages depending on the application, environment, vibration levels, and whether the joint needs to be frequently disassembled.

149. Answer: a

Explanation:

Understanding Engine Power Measurement After Reassembly

After an engine has been reassembled, perhaps after maintenance or repair, it is crucial to test its performance to ensure it is operating correctly and producing the expected power. Measuring engine power accurately requires a specific type of equipment designed for this purpose.

Methods for Testing Engine Performance

Several devices exist for testing various aspects of an engine's operation, but only one specifically measures its power output.

- **Dynamometer:** A dynamometer, often shortened to "dyno," is a device used to measure force, torque, or power. Engine dynamometers measure the torque and rotational speed (RPM) of the engine's output shaft. Engine power can then be calculated using the formula:

Power = Torque × Speed (with appropriate conversion factors depending on the units used).

There are different types of dynamometers, such as engine dynamometers (which measure power directly at the crankshaft or flywheel) and chassis dynamometers (which measure power at the wheels). Both are used for testing engine performance and power output after reassembly.

- **Hydrometer:** A hydrometer is an instrument used to measure the specific gravity (or relative density) of liquids. This is commonly used to test the concentration of solutions, such as the antifreeze concentration in coolant or the state of charge of a battery by measuring the density of its electrolyte. It has no function in measuring engine power.

- **Multimeter:** A multimeter is an electronic measuring instrument that combines several measurement functions, most commonly voltage, current, and resistance. It is used for diagnosing electrical issues in an engine or vehicle but cannot measure mechanical power output.
- **Test bench:** A test bench is a general term for a workbench or platform used to test components or systems. An engine might be mounted on a test bench *with* a dynamometer attached for testing. However, the test bench itself is the supporting structure, not the power-measuring device. The dynamometer is the instrument that performs the power measurement.

Analyzing the Options

Let's look at how each option relates to testing engine power after reassembly:

- **Dynamometer:** Directly measures or calculates engine power. This is the correct tool for the job.
- **Hydrometer:** Measures liquid density; irrelevant to engine power.
- **Multimeter:** Measures electrical properties; irrelevant to engine power.
- **Test bench:** A platform for testing; does not measure power itself.

Therefore, the correct instrument for testing engine power after reassembly is a dynamometer.

Conclusion on Engine Power Testing

To verify the performance and power output of an engine after it has been reassembled, a specialized tool that can measure mechanical power is required. The dynamometer is precisely this tool, measuring torque and speed to calculate the engine's power output.

Device	Primary Function	Used for Engine Power Test?
Dynamometer	Measures force, torque, speed, and calculates power	Yes
Hydrometer	Measures specific gravity of liquids	No
Multimeter	Measures electrical properties (voltage, current, resistance)	No
Test bench	Platform/setup for testing	Indirectly, as a support for testing equipment like a dynamometer

Revision Table: Engine Testing Tools

This table summarizes the primary uses of the tools mentioned in the context of engine work.

Tool	Use Case in Engine Service
Dynamometer	Testing engine power and torque output.
Hydrometer	Checking coolant mixture strength or battery electrolyte density.
Multimeter	Diagnosing electrical faults in wiring, sensors, ignition system, etc.
Test Bench	Providing a stable platform to mount and run an engine or component for testing.

Additional Information on Dynamometers and Engine Testing

Dynamometers are essential tools in automotive engineering, motorsports, and engine manufacturing. They allow engineers and technicians to simulate various load conditions and speeds to measure an engine's performance across its operating range. This data is critical for tuning, development, quality control, and verifying that an engine performs to specifications after reassembly or modification.

Types of dynamometers include:

- **Absorption Dynamometers:** These absorb the power produced by the engine and convert it into heat (e.g., eddy current dynos, water brake dynos) or electricity (e.g., AC motor/generator dynos).
- **Motoring Dynamometers:** These can also drive the engine (motor it) to measure friction losses or simulate downhill conditions. AC motor/generator dynos can often operate in both absorption and motoring modes.

Using a dynamometer provides objective data on engine power, torque, fuel efficiency, and emissions under controlled conditions, which is far more precise than subjective testing or simple measurements like compression checks.

Your Personal Exams Guide

150. Answer: b

Explanation:

Understanding Air Fuel Mixture Entry in Two-Stroke Engines

Two-stroke engines operate differently from four-stroke engines, especially in how they manage the intake, compression, power, and exhaust strokes within just two movements of the piston.

In a two-stroke engine, the crankshaft completes one revolution, and the piston completes two strokes (one up, one down) for every power cycle. The process of getting the fresh air-fuel mixture into the combustion chamber is one of the key differences compared to a four-stroke engine which uses valves driven by a camshaft.

How the Air Fuel Mixture Enters the Combustion Chamber

In a typical two-stroke engine design (like those used in many small motorcycles, scooters, or lawnmowers), the crankcase is sealed and plays a crucial role in the intake process. Here's a simplified explanation:

- When the piston moves upwards, it creates a vacuum in the crankcase.
- This vacuum draws the fresh air-fuel mixture from the carburetor (or fuel injection system) through an intake port into the crankcase.
- As the piston moves downwards after ignition and power stroke, it compresses the mixture that is now in the crankcase.
- Towards the end of the downward stroke, the piston uncovers ports in the cylinder wall. The exhaust port is uncovered first, allowing burnt gases to escape.
- Shortly after, the piston uncovers the **transfer ports**. The compressed fresh air-fuel mixture from the crankcase is then pushed upwards through these **transfer ports** into the combustion chamber, 'scavenging' or pushing out the remaining exhaust gases.

Therefore, the air-fuel mixture directly enters the combustion chamber via the **transfer ports**, not through an intake manifold leading directly to the combustion chamber like in a four-stroke engine or directly through the intake port which leads to the crankcase.

Analyzing the Options

- **intake port:** This port allows the air-fuel mixture to enter the crankcase from the carburetor/intake system, but not directly into the combustion chamber.
- **transfer port:** These ports connect the crankcase to the combustion chamber, allowing the mixture pressurized in the crankcase to enter the combustion chamber as the piston uncovers them. This is the correct passage for the mixture into the area where combustion occurs.
- **inlet manifold:** While an inlet manifold connects the carburetor/throttle body to the engine, in a two-stroke engine, it primarily leads to the intake port and the crankcase, not directly to the combustion chamber itself for final entry.
- **exhaust port:** This port is exclusively for allowing burnt gases to exit the combustion chamber and cylinder. It does not handle the entry of fresh air-fuel mixture.

Based on the operation of a two-stroke engine, the **transfer ports** are the passages through which the fresh air-fuel mixture moves from the crankcase into the combustion chamber for the next cycle of compression and combustion.

Revision Table: Two-Stroke Engine Ports

Port Type	Function in Two-Stroke Engine	Connects To
Intake Port	Allows fresh mixture entry into the crankcase.	Carburetor/Intake System & Crankcase
Transfer Port	Allows fresh mixture entry from crankcase into combustion chamber.	Crankcase & Combustion Chamber (via cylinder wall)
Exhaust Port	Allows burnt gases exit from combustion chamber.	Combustion Chamber (via cylinder wall) & Exhaust System

Additional Information on Two-Stroke Engine Cycles

Understanding the two strokes helps clarify the role of the transfer port:

1. Upward Stroke (Compression & Intake):

- Piston moves up.
- Top side: Compresses mixture in the combustion chamber (left from previous cycle's transfer).
- Bottom side (Crankcase): Creates vacuum, drawing fresh mixture via the intake port into the crankcase.

2. Downward Stroke (Power & Exhaust/Transfer):

- Piston moves down after ignition.
- Top side: Power stroke pushes piston down; towards end, uncovers exhaust port (exhaust gases exit) and then transfer ports.
- Bottom side (Crankcase): Compresses mixture drawn in during the upward stroke; this compressed mixture is then pushed via the transfer ports into the combustion chamber as they are uncovered by the piston.

This simultaneous action of exhaust and intake/transfer is known as 'scavenging'. The **transfer ports** are essential for moving the fresh charge from the crankcase into the cylinder head area for the next combustion event.

151. Answer: b

Explanation:

Understanding Coolants for Drilling Cast Iron

When performing machining operations like drilling, it's often necessary to use a coolant or lubricant. This helps manage the heat generated by friction and cutting, lubricate the cutting tool, and aid in removing chips or swarf from the cutting area. However, the choice of coolant depends heavily on the material being machined.

Why Coolant Choice Matters for Cast Iron Drilling

Cast iron is a unique material compared to many other metals like steel or aluminum. It is relatively brittle and tends to produce chips that break easily into small fragments, often creating a fine dust. Unlike ductile materials that produce long, stringy chips, cast iron chips are granular.

Using traditional liquid coolants or lubricants when drilling cast iron can cause several problems:

- The fine cast iron dust can mix with the liquid coolant, forming a thick, abrasive paste or sludge.
- This paste can clog the drill flutes, making chip removal difficult.

- The buildup can interfere with the cutting action and potentially damage the drill bit or the workpiece surface.
- It can also make the machining area messy and harder to clean.

For these reasons, a non-liquid method is often preferred for cooling and chip removal when drilling cast iron.

Analyzing Coolant Options for Cast Iron

Let's look at the provided options in the context of drilling cast iron:

- **Kerosene:** Kerosene is a liquid and a hydrocarbon. While sometimes used as a lubricant or coolant in specific machining applications, its use with cast iron would likely lead to the formation of the problematic sludge mentioned earlier. It also poses fire hazards.
- **Dry air:** Dry air is a gaseous coolant. When drilling cast iron, directing a stream of dry air into the cutting zone serves two main purposes: it cools the drill bit and the workpiece by convection, and importantly, it blows away the dry, powdery chips that are characteristic of cast iron machining. This prevents chip buildup and keeps the cutting area clear, which is highly effective for this material.
- **Lubricant oil:** Like other liquid oils, lubricant oil would mix with cast iron dust to form a paste, hindering chip evacuation and potentially causing cutting issues.
- **Machine oil:** Similar to lubricant oil, machine oil is a liquid that would create sludge with cast iron dust, making it unsuitable for drilling this material effectively and cleanly.

Optimal Coolant for Drilling Cast Iron

Based on the properties of cast iron and the effects of different coolants, dry air stands out as the most suitable option. It effectively removes the dry, brittle chips without creating messy or obstructive sludge.

Conclusion: Drilling Cast Iron Coolant Choice

When drilling cast iron, the primary concern shifts from lubricating the cut (as the chips break easily) to effectively removing the fine, dry chips. Liquid coolants tend to create a paste, hindering chip evacuation. Therefore, using dry air is the standard and most effective practice for drilling cast iron.

The correct answer is **Dry air**.

Revision Table: Drilling Cast Iron Coolants

Coolant Type	Effect on Cast Iron Drilling	Suitability
Liquid Coolants (Oils, Kerosene)	Mixes with fine dust, forms abrasive sludge, clogs drill flutes, hinders chip removal.	Not suitable for most applications.
Dry Air	Cools the cutting area, effectively blows away dry chips, prevents sludge formation, keeps area clean.	Highly suitable and preferred.

Additional Information: Machining Cast Iron

Machining cast iron is different from machining steel or aluminum due to its unique microstructure and chip formation characteristics. Key points include:

- **Chip Formation:** Cast iron produces discontinuous chips, which are small and brittle.
- **Tool Wear:** Due to abrasive particles in cast iron, tool wear can be significant. Carbide tools are often preferred.
- **Coolant Purpose:** For cast iron, the main purpose of a coolant (or air) is often chip evacuation and cooling rather than lubrication.
- **Alternatives to Dry Air:** In some specific applications, a minimal amount of mist coolant might be used, but dry machining (using dry air or no coolant at all) is very common for cast iron.

152. Answer: a

Explanation:

Understanding Injector Codes: What is IMA?

In modern vehicle fuel injection systems, particularly diesel engines, precise fuel delivery is crucial for optimal performance, fuel efficiency, and emissions control. When a fuel injector is replaced, the new injector often has slight variations in its flow characteristics compared to the original one. To compensate for these variations, a calibration code is typically associated with each individual injector.

This calibration code allows the vehicle's Engine Control Unit (ECU) to adjust the injection timing and fuel quantity specifically for that injector, ensuring all cylinders receive the correct amount of fuel. This compensation process is vital for smooth engine operation and meeting strict emission standards.

Common Terminology for Injector Calibration

The question asks for a common name for this "injector code." Among the options provided, one term is widely recognized in the automotive industry for this purpose.

Analyzing the Options

Let's look at the given options:

Option	Term	Relevance
1	IMA	Commonly stands for Injector Mass Adjustment or Injector Metering Adjustment . This term is widely used for the calibration code entered into the ECU when replacing modern fuel injectors, allowing the ECU to compensate for manufacturing variances.
2	MMA	Does not have a standard, widely recognized meaning related to injector codes in this context.
3	DMA	Could potentially refer to Direct Memory Access in computing, but is not related to fuel injector codes.
4	VMA	Does not have a standard, widely recognized meaning related to injector codes in this context.

Based on common automotive terminology, the calibration code associated with fuel injectors, especially in common rail diesel systems, is frequently referred to using the acronym **IMA**.

Why IMA is the Correct Term

IMA (Injector Mass Adjustment) codes are specific values printed on or supplied with new fuel injectors. These codes provide information about the injector's flow characteristics to the ECU. By inputting the correct IMA code for each injector during installation or replacement, technicians enable the ECU to make necessary adjustments to the fuel injection timing and pulse width, ensuring each cylinder receives the precisely calculated fuel amount.

This process is often called "IMA coding," "injector coding," or "injector calibration." It is a critical step to prevent issues like rough running, excessive smoke, poor fuel economy, and damage to emission control systems.

Conclusion

Therefore, the injector code is commonly known as IMA.

Revision Table: Key Injector Code Concepts

Term	Meaning	Purpose
Injector Code / Calibration Code	A unique value associated with an individual fuel injector.	Provides the ECU with specific flow characteristics of that injector.
IMA	Injector Mass Adjustment or Injector Metering Adjustment.	The common acronym for the injector calibration code used to compensate for manufacturing tolerances.
ECU	Engine Control Unit.	Uses the IMA code to make precise adjustments to fuel injection parameters for each injector.

Additional Information: Importance of Injector Coding (IMA)

Failure to correctly code (program the IMA value into the ECU) new or replacement injectors can lead to several problems:

- **Rough Idle:** Incorrect fuel delivery can cause cylinders to fire unevenly.
- **Increased Emissions:** Improper combustion due to incorrect fueling can lead to higher levels of pollutants.
- **Poor Fuel Economy:** The ECU might over-fuel or under-fuel, reducing efficiency.
- **Check Engine Light:** The system may detect incorrect fuel delivery and trigger diagnostic trouble codes (DTCs).
- **Reduced Performance:** The engine may not develop full power or torque.
- **Potential Damage:** In severe cases, prolonged incorrect fueling can potentially cause damage to engine components or emission control systems like Diesel Particulate Filters (DPFs).

For these reasons, correctly identifying and programming the IMA code is a standard procedure when working with fuel injection systems that utilize this technology.

153. Answer: a

Explanation:

Understanding Standard Electrode Lengths in Manufacturing

When working with processes like arc welding, electrodes are a crucial component. These electrodes are typically manufactured in specific, standardized lengths to ensure consistency, ease of use, and compatibility with various welding equipment and procedures. Using standard lengths helps in predicting material usage, managing inventory, and ensuring proper handling during the welding process.

Common Standard Electrode Lengths

In the manufacturing of electrodes, particularly those used for manual metal arc welding (MMAW), there are commonly recognized standard lengths. While various lengths might exist for specialized applications, two lengths are widely considered standard for general-purpose electrodes.

These standard lengths are:

- 350 mm
- 450 mm

These sizes are practical for manual handling and feeding the electrode during welding. Shorter electrodes would require more frequent changes, interrupting the welding process. Longer electrodes could become unwieldy or difficult to control, potentially affecting weld quality and welder fatigue.

Factors Influencing Electrode Lengths

The choice of standard lengths is influenced by several practical considerations:

- **Handling and Control:** The length must be manageable for a welder holding a stinger or electrode holder.
- **Welding Time:** A sufficient length allows for a reasonable amount of welding before the electrode needs to be replaced.
- **Manufacturing Efficiency:** Standardizing lengths streamlines the production process.
- **Packaging and Transport:** Standard sizes simplify packaging, shipping, and storage.

Comparing Electrode Length Options

Let's look at the given options in the context of standard electrode manufacturing lengths:

Option	Proposed Lengths	Alignment with Standard Lengths
1	350 mm and 450 mm	Commonly recognized standard lengths.
2	650 mm and 750 mm	Significantly longer than typical standard lengths for manual welding.
3	500 mm and 600 mm	While possible for some applications, less commonly cited as the primary standard lengths compared to 350/450 mm.
4	200 mm and 300 mm	Shorter lengths, might be used for specific purposes or smaller electrodes, but not the main standard lengths for general use.

Based on typical manufacturing standards for welding electrodes, the lengths 350 mm and 450 mm are widely produced and used, making them the standard lengths.

Conclusion on Standard Electrode Lengths

The manufacturing industry for electrodes adheres to certain standard dimensions for efficiency and usability. The most prevalent standard lengths for general-purpose electrodes are 350 mm and 450 mm. These dimensions strike a balance between providing sufficient material for welding runs and remaining easy for the welder to handle and control.

Revision Table: Key Electrode Manufacturing Standards

Aspect	Description
Purpose of Standard Lengths	Ensure consistency, ease of use, predictable material usage.
Primary Standard Lengths	350 mm, 450 mm
Process Using Standard Electrodes	Manual Metal Arc Welding (MMAW) is a common example.
Factors Influencing Length	Handling, welding time, manufacturing, logistics.

Additional Information: Electrode Types and Sizes

While 350 mm and 450 mm are standard lengths, electrodes also vary in diameter. The diameter depends on the welding current and the thickness of the material being welded. Common diameters can range from 1.5 mm up to 6 mm or more. The electrode's composition (flux coating type) also varies depending on the material being welded (e.g., mild steel, stainless steel, cast iron) and the desired weld properties.

Understanding both the standard lengths and the appropriate diameter and type is crucial for selecting the correct electrode for a specific welding job. Manufacturers provide specifications and guidelines for using their electrodes effectively, which often reference these standard dimensions.

Your Personal Exams Guide

154. Answer: c

Explanation:

Understanding Compound Rest Angle for ISO Metric Thread Cutting

Thread cutting on a lathe is a fundamental machining operation used to create screw threads on the external or internal surface of a workpiece. This process requires precise control over the tool's movement relative to the rotating workpiece. A crucial aspect of cutting threads, especially ISO Metric threads, is correctly setting the angle of the compound rest.

The Role of the Compound Rest

The compound rest is a component of the lathe that sits atop the cross slide. It can be swiveled to any angle in the horizontal plane. This allows the cutting tool, which is held in the tool post on the compound rest, to be fed into the workpiece at an angle. While the cross slide is used for diameter control and the

carriage for longitudinal movement, the compound rest is primarily used for feeding the tool into the work when turning short tapers or, as in this case, when cutting threads.

Compound Rest Angle for ISO Metric Threads

ISO Metric screw threads have a V-shaped profile with an included angle of 60° . When cutting these threads using a single-point tool, the traditional and most efficient method is to feed the tool predominantly into one flank of the thread. This helps in:

- Controlling the chip formed during cutting, directing it away from the cutting action.
- Reducing the load on the cutting tool by primarily cutting on one edge rather than forcing the tool directly into the root with both flanks cutting simultaneously.
- Producing a better surface finish on the thread flanks.

To achieve this single-flank feeding, the compound rest is typically swiveled to an angle equal to half of the included thread angle. For ISO Metric threads with a 60° included angle, this means the compound rest is set to 30° . The tool is then fed into the workpiece using the compound rest handle.

Thread System	Included Angle	Standard Compound Rest Angle
ISO Metric	60°	30°
Unified (UNC, UNF, etc.)	60°	29.5° or 30° (common practice is 29.5° or 30°)
Whitworth (BSW, BSF)	55°	27.5°

While other methods exist (like feeding straight in with the cross slide at 0° , which cuts on both flanks simultaneously), setting the compound rest at 30° is the standard practice for cutting ISO Metric threads, including coarse pitch ones, to ensure efficient cutting and good thread quality.

Analyzing the Options

- 45° : Setting the compound rest to 45° is not the standard practice for cutting ISO Metric threads. This angle is larger than half the thread angle and would not effectively feed the tool along one flank.
- 60° : Setting the compound rest to 60° is equal to the full thread angle. This would mean the tool feeds parallel to one of the thread flanks, which is not the method for cutting the V-profile itself.
- 30° : Setting the compound rest to 30° is exactly half the 60° included angle of an ISO Metric thread. This is the standard and recommended angle for feeding the tool predominantly into one flank, which is the most common and efficient method for cutting these threads.
- 15° : Setting the compound rest to 15° is less than half the thread angle. While it would still feed into one flank, it's not the optimal angle for standard ISO Metric thread cutting and wouldn't align the feed direction properly with the flank geometry compared to 30° .

Based on standard machining practices for cutting ISO Metric threads, the compound rest should be swiveled to 30° for feeding the cutting tool.

Revision Table: Key Concepts

Concept	Explanation
ISO Metric Thread Angle	60° included angle.
Compound Rest Role	Allows angled feeding of the cutting tool.
Standard Thread Cutting Method	Feed tool primarily into one flank.
Compound Rest Angle for Metric	Half of included angle, i.e., $60^\circ / 2 = 30^\circ$.

Additional Information: Thread Cutting on Lathes

When cutting threads on a lathe, several factors are critical:

- **Lead Screw Engagement:** The lead screw must be engaged using the half-nut lever to move the carriage longitudinally at the correct rate determined by the gearbox setting (for the desired thread pitch).
- **Tool Grinding:** The cutting tool must be ground precisely to the correct thread angle (60° for Metric) and must have the correct relief angles. The nose of the tool should have a slight radius for external threads or be sharp for internal threads, depending on the specific thread standard requirements (e.g., flat crest/root for ISO Metric).
- **Tool Alignment:** The tool must be set exactly on the centerline of the workpiece and perpendicular to the axis of rotation. A thread gauge is used to check the tool angle and alignment relative to the workpiece.
- **Depth of Cut:** Threads are cut incrementally by taking successive passes. The depth of cut for each pass is controlled by the compound rest (when set at an angle) or the cross slide.
- **Multiple Passes:** Cutting a full thread depth requires many light passes rather than a single deep cut, especially for coarse pitches, to manage heat, chip load, and tool wear.
- **Pitch vs. Lead:** For single-start threads, pitch and lead are the same. The lathe gearing is set based on the required pitch.

Understanding these elements, along with the correct compound rest angle, is essential for successfully cutting ISO Metric threads on a lathe.

155. Answer: b

Explanation:

Understanding Engine Cycles and Crankshaft Revolutions

The question asks about a type of engine where one working stroke occurs for every two revolutions of the crankshaft. This specific characteristic defines a particular type of internal combustion engine cycle.

An engine cycle involves a sequence of events that allows the engine to convert fuel into mechanical energy. The 'working stroke' or 'power stroke' is the part of the cycle where the ignited fuel expands and pushes the piston, which in turn rotates the crankshaft, producing usable power.

Four-Stroke Engine Operation

A four-stroke engine completes its cycle in four distinct strokes of the piston and two complete revolutions of the crankshaft. These strokes are:

1. **Intake Stroke:** The piston moves downwards, drawing the fuel-air mixture into the cylinder. The intake valve is open, and the exhaust valve is closed. The crankshaft completes half a revolution.
2. **Compression Stroke:** The piston moves upwards, compressing the fuel-air mixture. Both intake and exhaust valves are closed. The crankshaft completes the second half of the first revolution (one full revolution completed).
3. **Power (Working) Stroke:** The spark plug ignites the compressed fuel-air mixture, causing a rapid expansion of gases. This pushes the piston downwards forcefully. Both valves remain closed. This is the working stroke, and it occurs during the third stroke of the piston. The crankshaft completes the first half of the second revolution (one and a half revolutions completed).
4. **Exhaust Stroke:** The piston moves upwards, pushing the burnt gases out of the cylinder through the open exhaust valve. The intake valve is closed. The crankshaft completes the second half of the second revolution (two full revolutions completed).

As you can see, the working stroke (Power Stroke) occurs only once during these four strokes, which correspond to two complete revolutions of the crankshaft. Thus, in a four-stroke engine, there is one working stroke for every two crankshaft revolutions.

Two-Stroke Engine Operation Comparison

For comparison, a two-stroke engine completes its cycle in two strokes of the piston and one complete revolution of the crankshaft. It combines the intake and compression functions into the upward stroke and the power and exhaust functions into the downward stroke. A working stroke occurs once every revolution of the crankshaft.

Analyzing the Options

- **Two stroke:** One working stroke per revolution (not two).
- **Four stroke:** One working stroke per two revolutions. This matches the description in the question.
- **Three stroke:** Not a standard engine cycle classification.
- **One stroke:** Not a standard engine cycle classification.

Based on the analysis, the engine type described, with one working stroke for every two revolutions of the crankshaft, is the four-stroke engine.

Engine Cycle Comparison

Engine Type	Piston Strokes per Cycle	Crankshaft Revolutions per Cycle	Working Strokes per Cycle	Working Strokes per Crankshaft Revolution
Four Stroke	4	2	1	1 : 2
Two Stroke	2	1	1	1 : 1

Revision Table: Key Concepts

Engine Terminology Review

Term	Definition	Relevance to Engine Cycle
Stroke	Movement of the piston from one extreme (Top Dead Center or Bottom Dead Center) to the other.	Defines the phases of the engine cycle (Intake, Compression, Power, Exhaust).
Crankshaft Revolution	One complete rotation of the crankshaft (360°).	The number of revolutions corresponds to the completion of the engine cycle.
Working Stroke (Power Stroke)	The stroke where expanding gases push the piston down, producing power.	The energy-generating part of the cycle. Occurs once per cycle.

Additional Information on Engine Performance

The number of working strokes per crankshaft revolution directly impacts engine characteristics like power delivery and torque. While a two-stroke engine provides a working stroke every revolution, potentially offering higher power density for its size, a four-stroke engine offers better fuel efficiency, lower emissions, and smoother operation due to the dedicated exhaust stroke, making it suitable for most modern vehicles.

Understanding the relationship between piston strokes and crankshaft revolutions is fundamental to comprehending how internal combustion engines convert chemical energy into mechanical motion.

156. Answer: b

Explanation:

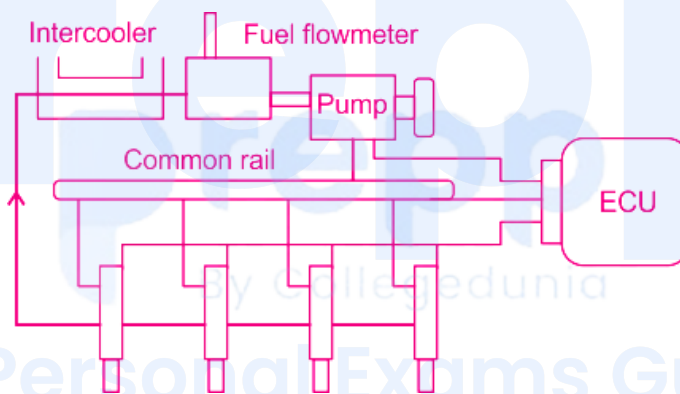
Explanation:

Common rail system:

- This type of system uses a high pressure fuel pump that is connected to a common fuel rail.
- Each cylinder's fuel injector is connected to the common fuel rail.

Common Rail Direct Injection (CRDI):

- The common rail fuel system consists of fuel tanks fuel pumps, common rail, pressure regulators, injectors, and sensors.
- The electrical fuel pump (low pressure) is placed inside the fuel tank.
- It develops pressure up to 6 bar and supplies to the high-pressure fuel pump (CRDI) through a fuel filter and water separator.
- The high-pressure fuel pump develops a pressure of 200 to 2000 bar and supplies to the common rail and common rail to fuel injectors inject fuel into the combustion chamber.
- Fuel injectors are operated by ECM through a solenoid valve.
- Common rail consists of a fuel pressure regulator rail pressure sensor and the fuel pressure regulator supplies the excess amount of fuel to the fuel tank (< 1 bar pressure).
- The common rail pressure sensor sends information to ECM/EDC, the existing pressure in the common rail will control the RPM of the fuel pump.
- Common rail will distribute the fuel to all the cylinders with equal pressure, then all cylinders will develop uniform power, which will reduce the vibration and noise of the engine.



157. Answer: b

Explanation:

Understanding Radiator Core Cleaning Methods

A vehicle's radiator is a crucial part of the cooling system, responsible for dissipating heat from the engine coolant. The radiator core, which consists of many tubes and fins, can accumulate deposits internally (like rust, scale, and sediment from coolant breakdown) and debris externally (dirt, insects, leaves).

Cleaning the radiator core is essential for maintaining efficient engine cooling and preventing overheating. There are different ways to approach cleaning, depending on whether the goal is to clear external blockages or internal deposits.

Analyzing the Options for Radiator Cleaning

Let's look at the provided options and how they relate to radiator core cleaning:

- **Option 1: From the outside to the fan side**

This method is typically used for cleaning the **external** fins of the radiator. Spraying water or compressed air from the outside (front of the car) towards the fan side helps remove debris like leaves, dirt, and insects that block airflow through the fins. This is important for airflow efficiency but doesn't address internal blockages.

- **Option 2: By reverse flushing**

Reverse flushing is a method used to clean the **inside** of the cooling system, including the radiator core. It involves circulating cleaning solution or water through the system in the opposite direction of normal coolant flow. This helps dislodge and push out accumulated rust, scale, and sediment that regular flushing might leave behind.

- **Option 3: From the top to the bottom of the core**

This describes the direction of normal coolant flow (from the top tank, down through the core tubes, to the bottom tank in most traditional down-flow radiators). Flushing in this direction is standard flushing, but reverse flushing is often more effective at removing stubborn deposits.

- **Option 4: From the top hose to the bottom hose**

This phrase also describes the normal flow path in the cooling system. Flushing or filling the system via the top hose and letting it exit from the bottom hose connection (or vice versa, depending on which is the inlet/outlet on the radiator) follows the typical direction. Again, reverse flushing specifically reverses this direction.

Why Reverse Flushing is Effective for Internal Cleaning

Internal deposits build up over time due to corrosion and degradation of coolant. These deposits can restrict the flow of coolant through the narrow tubes of the radiator core, reducing its ability to transfer heat. Reverse flushing works by:

1. Reversing the normal flow path.
2. Often using higher pressure or volume in the reverse direction.
3. Dislodging particles that may be stuck against the direction of normal flow.

This process is particularly effective at breaking loose and flushing out sediment and scale from the radiator tubes and other parts of the cooling system.

While external cleaning is important for airflow, the question asks how the radiator core itself should be cleaned, and reverse flushing is a specific technique for cleaning the **internal** passages to remove accumulated deposits.

Conclusion

Based on the methods described, reverse flushing is a recognized and effective technique specifically aimed at cleaning the internal passages of the radiator core and the cooling system by reversing the flow direction to remove stubborn sediment and scale.

Revision Table: Radiator Cleaning

Cleaning Method	Target	Direction	Purpose
External Cleaning	Radiator Fins (outside)	Outside → Fan Side	Remove debris, improve airflow
Standard Flushing	Cooling System (inside)	Normal Flow Direction	Replace coolant, remove loose contaminants
Reverse Flushing	Cooling System & Radiator Core (inside)	Opposite of Normal Flow	Remove stubborn internal deposits (rust, scale, sediment)

Additional Information: Cooling System Maintenance

Proper cooling system maintenance is vital for engine longevity and performance. Here are some related points:

- **Coolant Type:** Always use the correct type of coolant specified for your vehicle, as mixing types can cause corrosion and deposit formation.
- **Flushing Frequency:** The cooling system should be flushed and refilled according to the vehicle manufacturer's recommendations, typically every 30,000 to 100,000 miles or every 2 to 5 years, depending on the coolant type.
- **Signs of a Clogged Radiator:** Symptoms can include engine overheating, reduced heater performance, discolored coolant, or finding sediment when draining the system.
- **Safety:** Always work on a cool engine. Coolant is toxic, so handle it with care and dispose of used coolant properly.

158. Answer: a

Explanation:

Understanding Chromium Stainless Steel Alloy Applications

Chromium stainless steel alloys are a type of stainless steel known primarily for their corrosion resistance, attributed to the presence of chromium. They also offer good strength and wear resistance, properties that are crucial for various engineering applications, particularly in demanding environments like those found within an internal combustion engine.

Common Materials for Engine Components

Different components within an engine are subjected to varying stresses, temperatures, and corrosive environments, requiring specific materials. Let's look at common materials for the components listed in the options:

- **Crankshaft:** Typically made from forged steel (like carbon steel or alloy steel containing chromium, molybdenum, or manganese) or nodular cast iron. These materials provide high strength, stiffness, and fatigue resistance to withstand the massive rotational forces and bending moments.
- **Piston:** Commonly made from aluminum alloys due to their light weight and good thermal conductivity, although some high-performance or diesel engines might use cast iron or steel.
- **Camshaft:** Often made from cast iron or forged steel. The lobes, which endure high contact stresses and wear against tappets, are often hardened.
- **Inlet and exhaust valve:** Valves operate at high temperatures and are exposed to corrosive combustion gases. Inlet valves are typically made from chromium-nickel alloys or similar heat-resistant steels. Exhaust valves operate at even higher temperatures and are often made from more advanced alloys like nickel-based alloys or high-chromium stainless steels (such as austenitic stainless steels or specific heat-resistant martensitic stainless steels) for enhanced heat and corrosion resistance.

Analyzing the Options and Chromium Stainless Steel suitability

The question asks which component is made from Chromium stainless steel alloy. While many engine parts use alloys containing chromium, 'Chromium stainless steel' specifically refers to steel with a minimum of 10.5% chromium, forming a passive layer that resists corrosion. Let's consider the suitability:

- **Piston:** While some pistons might use steel, aluminum alloys are dominant. Chromium stainless steel isn't the primary material due to density and thermal expansion considerations compared to aluminum.
- **Camshaft:** Cast iron and forged steel are common. While chromium is present in alloy steels used for camshafts, describing it specifically as "Chromium stainless steel alloy" isn't the most typical material classification for mass-produced camshafts.
- **Inlet and exhaust valve:** As mentioned, exhaust valves, in particular, require high heat and corrosion resistance. High-chromium stainless steels, especially austenitic types, are indeed used for valves due to their excellent performance at elevated temperatures and resistance to exhaust gases. This makes valves a strong candidate.
- **Crankshaft:** Crankshafts require immense strength, fatigue resistance, and wear resistance at bearing surfaces. Traditional forged carbon or alloy steels are common. However, certain high-strength or corrosion-resistant applications might consider stainless steels. While less common than conventional steels or cast iron for standard crankshafts, specific types of stainless steel alloys (potentially certain martensitic or precipitation-hardening stainless steels offering high strength) could theoretically be used for specialized or high-performance crankshafts, particularly where corrosion resistance is a significant factor (though internal engine environments are usually lubricated). Given the options, and considering the provided correct answer text points to Crank shaft, we must assume a specific application or type of Chromium stainless steel alloy is suitable for a

crankshaft. It's important to note that standard crankshaft materials are typically chosen for properties like fatigue strength and toughness, which are traditionally met by heat-treated alloy steels or nodular cast iron. If Chromium stainless steel is used, it implies specific design considerations benefiting from its unique properties.

Conclusion

Based on the typical material usage, both valves and potentially certain specialized crankshafts or other engine components could utilize chromium stainless steel alloys. However, aligning with the provided component indicated, Chromium stainless steel alloy is stated as being used for making the Crank shaft.

Component	Common Materials	Chromium Stainless Steel Suitability (General)
Crankshaft	Forged Steel, Nodular Cast Iron	Possible for specialized applications needing high strength/corrosion resistance, but not typical.
Piston	Aluminum Alloys, Cast Iron, Steel	Unlikely to be the primary choice.
Camshaft	Cast Iron, Forged Steel	Chromium alloys used, but "stainless steel" isn't the typical descriptor.
Inlet and exhaust valve	Heat-resistant Steels, Nickel Alloys, High-Chromium Stainless Steels	High suitability, especially for exhaust valves.

Despite valves being a more common application for high-chromium stainless steels due to heat and corrosion resistance, the question implies Chromium stainless steel alloy is used for making the Crank shaft among the given options.

Revision Table: Key Concepts

Term	Brief Description
Chromium Stainless Steel	Steel alloy with >10.5% Chromium for corrosion resistance.
Crankshaft	Engine component converting linear piston motion to rotational motion.
Valve	Controls fuel/air mixture entry (inlet) and exhaust gas exit (exhaust) from cylinders.

Additional Information: Material Selection for Engine Components

Selecting materials for engine components is a complex process involving trade-offs between properties, cost, manufacturability, and application requirements. Key properties considered include:

- **Strength:** Ability to withstand load without deformation.
- **Fatigue Strength:** Ability to withstand repeated cycles of stress.
- **Hardness:** Resistance to wear and indentation.
- **Toughness:** Ability to absorb energy before fracturing.
- **Corrosion Resistance:** Ability to resist degradation from chemical reactions.
- **Heat Resistance:** Ability to maintain properties at high temperatures.
- **Thermal Conductivity:** Rate at which heat is transferred through the material.
- **Density:** Mass per unit volume (impacts weight).

While conventional forged steels and cast iron are predominant for crankshafts due to their excellent mechanical properties and cost-effectiveness, specialized crankshaft designs or specific environmental factors could potentially justify the use of certain high-strength stainless steel alloys, leveraging their corrosion resistance alongside necessary mechanical properties obtained through appropriate heat treatments.

159. Answer: b

Explanation:

Understanding DC Motor Applications in Cooling Systems

Cooling systems are essential in various applications, from vehicles to electronic devices, to dissipate heat and maintain optimal operating temperatures. These systems often utilize different components, and motors play a crucial role in driving some of these parts.

Components in a Cooling System and Motor Use

Let's examine the given options and their typical functions within a cooling system context:

- **Hose pipes:** These are conduits that transport the coolant fluid between different parts of the system, such as the radiator, engine block, and heater core. They do not require a motor for their function; they simply carry the fluid moved by a pump.
- **Fan:** Fans are used to move air, which helps in heat transfer. In cooling systems, fans often push or pull air through a radiator or heatsink to increase the rate of cooling. Electric cooling fans, commonly found in automotive cooling systems (especially for auxiliary cooling or when the engine-driven fan isn't sufficient) and in electronics, are powered by electric motors. A small DC motor is perfectly suited for driving such a fan, particularly in smaller systems or for variable speed control.
- **Pump:** Pumps are responsible for circulating the coolant fluid throughout the system. While pumps are indeed driven by motors, the size and type of motor can vary. In many automotive systems, the main coolant pump is mechanically driven by the engine belt. However, electric coolant pumps also exist and are driven by electric motors, often larger than the small DC motors used for fans, depending on the system's flow requirements. In smaller electronic cooling loops, tiny pumps are used, potentially driven by small DC motors, but fans are a more widespread application of small DC motors for air cooling.

- **Belt:** Belts are used to transfer mechanical power from one component to another, typically from a power source like the engine crankshaft to accessories like the alternator, power steering pump, or mechanical water pump. Belts are driven by a motor (the engine) but are not devices that themselves use a motor for their primary function in the cooling system context.

Identifying the Device Using a Small DC Motor

Considering the common applications of small DC motors in cooling systems:

- Electronic cooling systems often use small DC fans to cool components like CPUs, GPUs, or power supplies.
- Automotive cooling systems use electric fans (often driven by DC motors) to pull air through the radiator, especially when the vehicle is stationary or moving slowly. These auxiliary electric fans are distinct from engine-driven mechanical fans.

Among the given options, the fan is the device that most commonly uses a small DC motor to perform its function within a cooling system. The motor spins the fan blades to move air, facilitating heat removal.

Therefore, in a cooling system, a small DC motor is typically used in the fan.

Revision Table: Cooling System Components

Component	Primary Function	Typical Motor Use
Hose pipes	Transport coolant fluid	No motor required
Fan	Move air for heat dissipation	Often uses a small DC motor (electric fan)
Pump	Circulate coolant fluid	Uses a motor (electric pump) or mechanically driven (mechanical pump)
Belt	Transfer mechanical power	Does not use a motor itself in this context; is driven by one (the engine)

Additional Information: Types of Cooling System Motors

While small DC motors are common for fans, cooling systems can utilize various types of motors depending on the application and power requirements:

- **Brushed DC Motors:** Simple, cost-effective, often used in basic cooling fans.
- **Brushless DC (BLDC) Motors:** More efficient, longer lifespan, used in variable-speed fans and pumps in more advanced systems.
- **AC Motors:** Less common in standard automotive 12V systems but can be found in industrial cooling or some electric vehicle (EV) cooling components if an AC power source is readily available or generated.

The choice of motor depends on factors like voltage, power needed, required lifespan, efficiency goals, and cost constraints of the cooling system design.

160. Answer: d

Explanation:

Understanding Passenger Car Emissions Measurement Units

The question asks about the standard unit used to measure emissions from passenger cars. Emissions are pollutants released into the environment, and measuring them is crucial for environmental regulations and monitoring.

Standard Unit for Passenger Car Emissions

For passenger cars and other road vehicles, emissions are typically measured based on the distance traveled. This makes sense because the amount of pollution generated is directly related to how far the vehicle drives.

The standard unit expresses the mass of a specific pollutant emitted per unit of distance. The most common unit for this measurement in regulations and reporting is grams per kilometer.

- **Grams (g):** This is a unit of mass, representing the amount of pollutant.
- **Kilometer (km):** This is a unit of distance, representing how far the car has traveled.

Therefore, the unit **g/km** represents the mass of pollutant emitted for every kilometer the vehicle travels.

Analyzing the Options for Emission Units

Let's look at the provided options and see which one fits the standard measurement for passenger cars:

- **kg/km:** Kilograms per kilometer. While dimensionally correct (mass per distance), kilograms are usually too large a unit for typical passenger car emissions, which are measured in grams.
- **g/m:** Grams per meter. This is also mass per distance, but meters are a much shorter distance unit than kilometers. While possible, g/km is the internationally recognized standard for vehicle emissions.
- **g/kwh:** Grams per kilowatt-hour. This unit measures emissions based on energy output. It is commonly used for stationary engines, power plants, or sometimes for heavy-duty engines, but not typically for measuring the total emissions of a passenger car over a distance.
- **g/km:** Grams per kilometer. This unit measures the mass of pollutant per unit of distance traveled, which is the standard way to express passenger car emissions in regulations worldwide.

Based on the analysis of common practices and regulations for passenger cars, the unit grams per kilometer (g/km) is the correct measure for their emissions.

Why g/km is Used for Car Emissions

Using g/km allows for direct comparison of the environmental impact of different vehicles based on their fuel efficiency and emission control systems over a standard distance. Regulations often set limits in g/km for various pollutants like CO (Carbon Monoxide), NOx (Nitrogen Oxides), HC (Hydrocarbons), and Particulate Matter (PM).

Unit	Meaning	Common Use Case
kg/km	Mass per distance	Less common for passenger cars, possibly for very heavy pollutants or vehicles
g/m	Mass per distance	Dimensionally correct but not standard for vehicle regulations
g/kWh	Mass per energy output	Stationary engines, power plants, heavy-duty engines
g/km	Mass per distance	Standard for passenger car emissions

Therefore, the unit used for measuring emissions in passenger cars is g/km.

Revision Table: Passenger Car Emissions

Concept	Description	Standard Unit
Passenger Car Emissions	Pollutants released by cars into the air (e.g., CO, NOx, PM)	g/km (grams per kilometer)
Measurement Basis	Distance traveled by the vehicle	Per kilometer (km)
Why g/km?	Allows comparison based on distance, aligns with regulations	Mass (g) per distance (km)

Additional Information: Vehicle Emission Standards

Emission standards for passenger cars vary globally (e.g., Euro standards in Europe, EPA standards in the US). These standards set limits on the maximum permissible emissions of specific pollutants, measured in g/km.

- **Key Pollutants Measured:**
 - Carbon Monoxide (CO)
 - Nitrogen Oxides (NOx)
 - Hydrocarbons (HC) or Total Hydrocarbons (THC)
 - Non-Methane Hydrocarbons (NMHC)
 - Particulate Matter (PM)
 - Carbon Dioxide (CO₂) – often measured in g/km as well, though it's a greenhouse gas, not a traditional pollutant, and its measurement reflects fuel efficiency.

- **Test Cycles:** Emissions are measured during specific driving test cycles that simulate real-world driving conditions (e.g., urban, highway).
- **Importance:** Measuring and regulating passenger car emissions is crucial for improving air quality and public health in urban areas.

161. Answer: b

Explanation:

Understanding Chisels for Squaring Corners

Chisels are essential hand tools used for cutting, shaping, and removing material. They consist of a blade with a cutting edge at the end, and a handle. Different types of chisels are designed for specific tasks based on the shape and grind of their cutting edge.

Types of Chisels and Their Uses

Let's look at the common types of chisels mentioned in the options and their primary applications in metalworking or similar tasks:

- **Flat Chisel:** This is one of the most common types. It has a flat, straight cutting edge and is used for cutting flat surfaces, chipping large areas of metal, and removing excess material.
- **Diamond Point Chisel:** The cutting end of this chisel is ground to a square shape, with the cutting edge being the corner of the square. This unique shape makes it ideal for cleaning out square corners and V-shaped grooves. It can effectively reach into sharp angles that other chisels cannot.
- **Cross-cut Chisel:** Also known as a cape chisel, this chisel has a narrow cutting edge. It is primarily used for cutting grooves, channels, or keyways. The narrow blade allows it to cut along a specific line or channel.
- **Half Round Nose Chisel:** This chisel has a rounded cutting edge. It is used for cutting curved grooves, channels, or chipping out concave surfaces.

Why Diamond Point Chisel Squares Corners

The question specifically asks about squaring materials at the corners. Squaring a corner means making the angle precisely 90 degrees and ensuring the material is clean and sharp at that junction. As described, the **diamond point chisel** has a cutting edge that comes to a point, effectively a corner itself. When this point is applied to an inside corner, it can reach the very apex of the corner and remove material precisely, allowing the corner to be cleaned up and made square. Flat chisels or cross-cut chisels are not designed to work effectively right into a sharp internal corner.

Consider the shape of the cutting edge:

- A flat chisel has a wide, straight edge, good for flat areas.
- A cross-cut chisel has a narrow, straight edge, good for grooves.

- A half-round nose chisel has a curved edge, good for rounded grooves.
- A **diamond point chisel** has a point (like the corner of a square) as its cutting edge, which fits perfectly into a corner to make it square.

Therefore, the diamond point chisel is the correct tool for the task of squaring materials at the corners.

Comparison of Chisel Types and Uses

Chisel Type	Cutting Edge Shape	Primary Use
Flat Chisel	Straight, wide	Chipping flat surfaces, removing bulk material
Diamond Point Chisel	Pointed (like a square corner)	Squaring corners, cleaning V-grooves
Cross-cut Chisel (Cape Chisel)	Narrow, straight	Cutting grooves, channels, keyways
Half Round Nose Chisel	Rounded	Cutting curved grooves, chipping concave surfaces

Conclusion on Squaring Corners

Based on the specific design and purpose of each chisel type, the **diamond point chisel** is uniquely suited for the task of squaring materials at the corners due to its pointed cutting edge.

Revision Table: Key Chisel Uses

Quick Review of Chisel Types and Applications

Chisel Type	Best For
Flat	Flat surface chipping
Diamond Point	Squaring corners
Cross-cut	Cutting grooves
Half Round Nose	Cutting curved grooves

Additional Information on Chisels

Understanding more about chisels can help in selecting the right tool for various tasks:

- **Material:** Chisels are typically made from hardened steel to maintain a sharp edge. The handle can be made of wood, metal, or plastic.

- **Sharpening:** Chisels need regular sharpening to remain effective. The sharpening process depends on the shape of the cutting edge.
- **Safety:** Always use a hammer or mallet appropriate for the chisel handle. Wear safety glasses to protect your eyes from flying debris. Ensure the workpiece is securely held.
- **Cold vs. Wood Chisels:** The chisels discussed here are primarily cold chisels used for metal, stone, or brick. Wood chisels have a different bevel angle and are designed specifically for working with wood. The question implies a context where squaring material corners is needed, which could be metalworking or masonry, making cold chisels the relevant type.

Choosing the correct **type of chisel**, like the **diamond point chisel** for **squaring corners**, is crucial for efficiency and achieving a clean, accurate result in material shaping tasks.

162. Answer: d

Explanation:

Understanding Keys on an Automobile Shaft

Keys are essential mechanical components used to connect a rotating machine element (like a gear, pulley, or coupling) to a shaft. Their primary function is to prevent relative rotation between the shaft and the connected element, ensuring that torque is transmitted effectively from one to the other.

Different types of keys are designed for various applications based on the load, speed, and type of connection required. Let's look at the types of keys mentioned in the options and consider their typical uses.

Types of Keys and Their Applications

- **Saddle Key:** This type of key is tapered and fits in a keyway in the hub but has no keyway in the shaft. It relies on friction between the key and the shaft to transmit torque. Saddle keys are suitable only for light loads as they don't provide a positive lock.
- **Taper Key:** A taper key is rectangular or square in cross-section and is tapered along its length. It fits into keyways in both the shaft and the hub. The taper helps in seating the key firmly, providing a strong connection suitable for moderate loads. A gib head key is a type of taper key with a gib head for easy removal.
- **Gib Head Key:** As mentioned, this is a specific type of taper key. It has a head on one end, which makes it easier to drive in or remove the key from the keyway. Like other taper keys, it is used for transmitting moderate to heavy loads.
- **Woodruff Key:** This key has a semicircular shape and fits into a semicircular keyway milled into the shaft. It fits into a rectangular keyway in the hub. The semicircular shape allows the key to tilt and align itself with a tapered keyway in the hub, which can be advantageous, especially on tapered shafts. Woodruff keys are often used in machine tools, automobiles, and other applications where self-alignment is beneficial and moderate loads are involved. Their design helps prevent the concentration of stress in the shaft keyway compared to traditional rectangular keys.

Why Woodruff Keys are Used on Automobile Shafts

Automobile shafts, such as those in steering mechanisms or gearboxes, often require a connection that can handle moderate torque and potentially cope with minor misalignments or tapered sections. The Woodruff key is well-suited for these applications due to:

- Its ability to align itself, which is useful on tapered shafts.
- The keyway design on the shaft, which can be less stress-concentrating than a standard square or rectangular keyway under certain conditions.
- Providing a secure connection for transmitting torque in automotive systems.

While other keys like taper keys are used in various automotive applications, the Woodruff key is specifically noted for its use on certain automobile shafts because of its unique shape and self-aligning feature.

Comparison of Key Types

Key Type	Shape	Shaft Keyway	Hub Keyway	Load Capacity	Typical Application
Saddle Key	Tapered	None (Relies on friction)	Yes	Light	Light duty applications
Taper Key	Rectangular/Square, Tapered	Yes	Yes	Moderate to Heavy	General machine applications
Gib Head Key	Taper Key with Head	Yes	Yes	Moderate to Heavy	Applications requiring easy removal
Woodruff Key	Semicircular	Semicircular	Rectangular	Moderate	Automobile shafts, machine tools (often with tapered shafts)

Based on the characteristics and typical applications, the Woodruff key is a common choice for use on automobile shafts.

Revision Table: Key Types for Shafts

Let's quickly review the key points about different key types used for connecting components to shafts.

- Keys prevent relative rotation and transmit torque.
- Saddle keys use friction, suitable for light loads.
- Taper keys and Gib head keys provide a positive lock for moderate to heavy loads.
- Woodruff keys are semicircular, self-aligning, and often used on tapered shafts in applications like automobile shafts.

Additional Information: Key Design Considerations

When selecting a key for a specific application like an automobile shaft, engineers consider factors such as:

- The amount of torque to be transmitted.
- The speed of rotation.
- The potential for shock or vibration.
- The ease of assembly and disassembly.
- The required lifespan of the connection.
- The cost of manufacturing the key and the keyways.

The shape and dimensions of the keyway also play a crucial role in the strength and reliability of the keyed joint. Proper fitting of the key is essential to avoid stress concentrations and potential failure.

163. Answer: a

Explanation:

Understanding Ti-VCT: Twin Independent Variable Camshaft Timing

The question asks for the full form of the automotive technology acronym Ti-VCT. Ti-VCT is a type of variable valve timing system used in internal combustion engines. Let's break down what the acronym stands for.

Decoding Ti-VCT: What Each Part Means

The acronym Ti-VCT stands for:

- **Ti:** Twin Independent
- **V:** Variable
- **C:** Camshaft
- **T:** Timing

So, the full form is Twin Independent Variable Camshaft Timing.

What is Variable Camshaft Timing (VCT)?

Variable Camshaft Timing, or VCT, is a technology that allows the timing of the engine's intake and exhaust valves to be adjusted while the engine is running. Normally, the camshaft is fixed, and valve timing is set for a specific engine speed and load. VCT allows the engine control unit (ECU) to advance or retard the opening and closing of the valves based on driving conditions. This can improve engine performance, fuel efficiency, and reduce emissions.

What does 'Twin Independent' mean in Ti-VCT?

The "Twin Independent" part of Ti-VCT signifies that this system can independently adjust the timing of both the intake camshaft and the exhaust camshaft. Many earlier or simpler VCT systems might only adjust one camshaft (either intake or exhaust) or adjust both together in a fixed relationship. 'Twin Independent' means the intake timing can be adjusted separately from the exhaust timing, offering greater flexibility and control over the engine's breathing process across different engine speeds and loads.

Analyzing the Options

Let's look at the provided options and see why only one is correct:

- **Option 1: Twin Independent Variable Camshaft Timing** – This matches our breakdown of the acronym Ti-VCT.
- **Option 2: Twin Interdependent Variable Crankshaft Timing** – This option uses "Interdependent" instead of "Independent" and refers to "Crankshaft Timing" instead of "Camshaft Timing". Valve timing is primarily controlled by the camshaft relative to the crankshaft, but the adjustment is on the camshaft. Also, 'Interdependent' suggests the timing is linked, contradicting the 'Independent' nature of Ti-VCT.
- **Option 3: Twin Interdependent Variable Controller Timing** – This option also uses "Interdependent" and introduces "Controller Timing," which is not the correct term related to valve timing.
- **Option 4: Twin Interdependent Variable Countershaft Timing** – Similar to options 2 and 3, this uses "Interdependent" and refers to "Countershaft Timing," which is not the correct component related to valve timing in this context.

Based on the meaning of Ti-VCT technology, the correct full form is Twin Independent Variable Camshaft Timing.

Conclusion

The full form of Ti-VCT is Twin Independent Variable Camshaft Timing. This advanced variable valve timing system provides independent control over both intake and exhaust camshafts to optimize engine operation under various conditions.

Revision Table: Key Terms

Term	Explanation
Ti-VCT	Twin Independent Variable Camshaft Timing
VCT	Variable Camshaft Timing (a general term for systems adjusting valve timing)
Camshaft	A rotating shaft with lobes that push against valves, controlling their opening and closing.
Valve Timing	The precise moments when engine valves open and close relative to the crankshaft position.

Additional Information on Variable Valve Timing Systems

Variable Valve Timing (VVT) systems have evolved over time. Here are some key aspects:

- **Purpose:** To make the engine more efficient and powerful across a wider range of RPMs (engine speeds). At low RPMs, different timing is needed compared to high RPMs for optimal performance.
- **Benefits:** Improved fuel economy, increased power and torque, reduced emissions (especially NOx), smoother idling, and better engine response.
- **Mechanism:** Most VVT systems, like Ti-VCT, work by changing the phase angle of the camshaft relative to the crankshaft. This is often done using a hydraulic actuator controlled by the engine's ECU.
- **Types:** Some systems only vary intake timing, some only exhaust, some vary both together (like older VCT), and advanced systems like Ti-VCT vary both independently. Some systems can also vary valve lift and duration (how far the valve opens and for how long).

Ti-VCT is a specific implementation of VVT technology used by various manufacturers, including Ford.

164. Answer: b

Explanation:

Understanding Drill Holding Devices

Drilling machines use various accessories to hold the drill bits securely during operation. The type of holding device used depends primarily on the shank type and size of the drill bit.

The question asks specifically about the device used for holding **straight shank drills** with a diameter less than 12 mm.

Analyzing the Options for Drill Holding

Let's look at the common devices mentioned in the options and their typical uses:

- **Sleeve:** A sleeve is a tapered tube used to adapt a drill with a smaller Morse taper shank to fit into a machine spindle or socket with a larger Morse taper bore. It is used for **tapered shank drills**.
- **Drill chuck:** A drill chuck is a jawed mechanism that grips the cylindrical surface of a drill shank. It is the standard method for holding **straight shank drills** of various sizes, particularly those with smaller diameters.
- **Drill drift:** A drill drift is a wedge-shaped tool used to remove tapered shank drills, reamers, or sleeves from the tapered bore of a machine spindle or socket. It is not a holding device.
- **Socket:** A socket is similar to a sleeve but is typically inserted directly into the machine spindle bore. It has a tapered internal bore (usually Morse taper) to receive a drill or sleeve with a corresponding tapered shank. It is used for **tapered shank drills**.

Identifying the Correct Holding Device for Straight Shank Drills

Based on the analysis, the primary device used for holding **straight shank drills** is the **drill chuck**. While chucks can hold straight shank drills of various sizes, they are particularly common for diameters below approximately 12 mm (or 1/2 inch), where straight shanks are most prevalent.

Larger drills often come with tapered shanks and are held directly in sockets or sleeves, which fit into the machine spindle's tapered bore.

Therefore, for a **straight shank drill** with a diameter less than 12 mm, the appropriate holding device is the **drill chuck**.

Summary of Drill Holding Devices

Device	Shank Type Held	Typical Use
Drill Chuck	Straight Shank	Holds straight shank drills, especially smaller sizes (< 12 mm)
Socket	Tapered Shank (and Sleeves)	Holds tapered shank drills directly in spindle or receives sleeves
Sleeve	Tapered Shank	Adapts a smaller tapered shank drill to a larger tapered bore
Drill Drift	N/A (Removal Tool)	Removes tapered shank tools/sleeves from bores

Conclusion

The tool used to hold **straight shank drill** bits having a diameter less than 12 mm is typically a **drill chuck**.

Drill Chuck Revision Table

Feature	Description
Purpose	To securely grip and hold cylindrical tools, primarily drill bits with straight shanks.
Mechanism	Typically uses three jaws that move inward simultaneously to grip the shank.
Types	Keyed chucks (tightened with a key) and keyless chucks (tightened by hand).
Connection to Machine	Can have a tapered bore to fit directly onto a tapered spindle nose or a threaded bore to screw onto a threaded spindle.
Applications	Drilling machines (drill presses, hand drills, milling machines), lathes (for tailstock drilling).

Additional Information on Drill Shanks and Holders

Understanding the different types of drill shanks is crucial for selecting the correct holding device. The two main types are:

- **Straight Shank:** Cylindrical shank, held by friction grip in a chuck. Common for smaller drill sizes.
- **Tapered Shank (e.g., Morse Taper):** Conical shank that fits into a corresponding tapered bore (socket or spindle). Held by the taper's friction and often includes a tang for positive drive and easier removal. Common for larger drill sizes.

The choice between a straight shank and a tapered shank often depends on the torque transmission requirement and the size of the drill bit. Straight shanks are adequate for smaller diameters and lower torque, while tapered shanks provide a more secure and robust connection for larger diameters and higher torque applications.

When using a drill chuck, it is essential to ensure the drill bit is properly centered and tightened firmly to prevent slippage, which can damage the shank and the chuck jaws.

165. Answer: d

Explanation:

Understanding Pulleys for Speed Variation

The question asks us to identify the type of pulley that is specifically used to achieve different speeds in a power transmission system. Let's examine the functions of the given pulley types.

Analysis of Pulley Types

- **Solid Pulley:** A solid pulley is typically cast in one piece. It is used for transmitting power between shafts with a fixed speed ratio, determined by the diameters of the driving and driven pulleys. It does not allow for easy changes in speed.
- **Split Pulley:** A split pulley is made in two halves that are bolted together around the shaft. This design makes installation easier without removing the shaft. Like a solid pulley, it is generally used for a fixed speed ratio.
- **Jockey Pulley:** A jockey pulley (also known as an idler pulley) is primarily used to maintain tension in a belt or to increase the arc of contact of the belt on the main pulleys. It does not directly change the speed ratio between the driving and driven shafts.
- **Stepped Pulley:** A stepped pulley, also known as a cone pulley, has multiple steps or different diameters machined onto the same casting or assembly. By shifting the belt from one step to another on both the driving and driven stepped pulleys, the speed ratio between the shafts can be changed. For example, if the belt connects a small diameter on the driving pulley to a large diameter on the driven pulley, the driven shaft will rotate slower. Conversely, connecting a large diameter on the driving pulley to a small diameter on the driven pulley will make the driven shaft rotate faster. This allows for obtaining different speeds from a single driving speed.

Based on the functions described, the stepped pulley is the type specifically designed and used to obtain different speeds.

Explanation of Stepped Pulley Function

A stepped pulley system typically involves two stepped pulleys mounted on parallel shafts, connected by a belt. Each pulley has a set of varying diameters (steps). The speed ratio between the shafts is determined by the ratio of the diameters of the steps on which the belt is currently running:

$$\text{Speed Ratio} = \frac{\text{Speed of Driven Shaft}}{\text{Speed of Driving Shaft}} = \frac{\text{Diameter of Driving Pulley Step}}{\text{Diameter of Driven Pulley Step}}$$

By moving the belt to different pairs of steps, the diameter ratio changes, resulting in a different speed ratio and thus a different output speed for the driven shaft. This mechanism is commonly found in machinery like drill presses and lathes to adjust cutting speeds.

Therefore, the pulley used to obtain different speeds is the stepped pulley.

Pulley Type	Primary Function	Used for Different Speeds?
Solid Pulley	Power transmission (fixed ratio)	No
Split Pulley	Power transmission (fixed ratio), Easy installation	No
Jockey Pulley	Belt tensioning, Increase contact angle	No (Indirect effect on slip, but not primary function)
Stepped Pulley	Power transmission (variable ratio)	Yes

Revision Table: Pulley Types and Uses

Pulley Type	Description	Application Example
Solid Pulley	Single piece casting	General power transmission
Split Pulley	Made in two halves	Retrofitting onto existing shafts
Jockey Pulley	Used with tensioner mechanism	Improving belt drive efficiency, reducing vibration
Stepped Pulley	Multiple diameters (steps)	Drill presses, Lathes, Milling machines

Additional Information on Variable Speed Drives

While stepped pulleys provide a mechanical way to change speed, other methods for achieving variable speeds include:

- **Variable Frequency Drives (VFDs):** Electronic devices used with AC motors to change motor speed by adjusting the frequency of the power supply.
- **Variable Speed Belts and Pulleys (e.g., V-belt with variable width pulleys):** Systems where the effective diameter of the pulleys can be changed continuously while the system is running, offering more flexibility than stepped pulleys.
- **Gearboxes:** Systems of gears that provide discrete or continuous speed changes.

Stepped pulleys are a simple, robust, and cost-effective method for providing multiple distinct speed options.

166. Answer: d

Explanation:

Understanding Why an Engine Might Not Start

An engine requires several things to start and run: fuel, air, spark (in petrol engines), and proper timing. If any of these essential components or processes are missing or significantly faulty, the engine will likely fail to start.

Analyzing the Potential Reasons for a No-Start Condition

Let's look at each option provided and determine its potential effect on an engine's ability to start.

- **Punctured tires:** Punctured tires affect the vehicle's ability to move safely, but they have no direct impact on the engine starting mechanism. The engine system is separate from the wheels and tires.
- **Timing defect:** Engine timing refers to the precise synchronization of the crankshaft and camshaft. This controls when the valves open and close and when the spark plug fires (or fuel is injected). A significant timing defect can indeed prevent an engine from starting because the combustion process happens at the wrong time, or not at all. However, other issues like lack of fuel are also common no-start causes.
- **Clogged air cleaner:** The air cleaner filters the air entering the engine. A clogged air cleaner restricts airflow. While a severely clogged air cleaner can make an engine run poorly, lose power, or even be difficult to start, it typically doesn't cause a complete failure to start unless the restriction is extreme. The engine still gets *some* air, and with fuel and spark, it might at least try to turn over or sputter.
- **No fuel in the tank:** Fuel is absolutely essential for combustion. If there is no fuel in the tank, the fuel pump cannot deliver fuel to the engine's combustion chambers. Without fuel mixing with air and being ignited, the combustion process cannot occur, and the engine will not start, no matter how well the other systems are working. This is a fundamental requirement.

Considering the options, having no fuel in the tank is a very direct and guaranteed reason for an engine not starting. While a severe timing defect can also prevent starting, running out of fuel is a common and definite cause.

Impact of Issues on Engine Starting

Issue	Likelihood of Preventing Start	Primary Effect
Punctured tires	Very Low (None)	Prevents driving, safety hazard
Timing defect	High	Incorrect valve/spark timing
Clogged air cleaner	Moderate (if severe)	Restricted air intake, poor performance
No fuel in the tank	Very High (Guaranteed)	No fuel supply for combustion

Revision Table: Engine Starting Causes

Here's a quick summary of how the different issues relate to engine starting problems:

- Running out of fuel means the engine lacks a key component (fuel) needed for combustion, directly preventing starting.
- Incorrect engine timing means the combustion attempt happens at the wrong time, which can also prevent starting.
- A clogged air filter restricts air, making starting difficult or causing poor running, but usually doesn't stop it completely unless extremely severe.
- Punctured tires affect mobility, not the starting process itself.

Additional Information on Engine No-Start Issues

Beyond the options listed, other common reasons an engine might not start include:

- **Dead Battery:** The starter motor needs power from the battery to turn the engine over.
- **Faulty Starter Motor:** If the starter motor itself is broken, it cannot crank the engine.
- **Ignition System Problems:** Issues with spark plugs, ignition coils, or the distributor (in older cars) can prevent the spark needed for combustion.
- **Fuel Pump Failure:** Even with fuel in the tank, a faulty fuel pump cannot deliver it to the engine.
- **Immobilizer System Issues:** Modern cars have security systems that can prevent the engine from starting if they detect an issue.

Diagnosing a no-start condition often involves checking these key areas: fuel, air, spark, and mechanical timing/compression.

167. Answer: c

Explanation:

Understanding Diesel Engine Components

This question asks us to identify the component that is not typically found in a diesel engine among the given options. Let's look at each option to understand its role in engine operation.

Analysis of Components

- **Injector:** In a diesel engine, fuel is injected directly into the combustion chamber (or a pre-combustion chamber). The injector atomizes the fuel into fine droplets, which is essential for proper mixing with air and efficient combustion. Therefore, an injector is a crucial part of a diesel engine.
- **Fuel injection pump:** Diesel engines use a fuel injection pump to deliver fuel to the injectors at very high pressure and at the correct timing. This high pressure is necessary to inject fuel against the high compression pressure in the cylinder. The fuel injection pump is a fundamental component of the diesel fuel system.
- **Spark plug:** Spark plugs are used to ignite the air-fuel mixture in petrol (gasoline) engines. They create an electric spark at the precise moment needed for combustion. Diesel engines, however, rely on compression ignition. The air in the cylinder is compressed to such a high pressure that its temperature rises significantly. When fuel is injected into this hot, compressed air, it auto-ignites without the need for a spark. Thus, a spark plug is not required in a typical diesel engine.
- **Fuel filter:** The fuel filter is an important component in any fuel system, including diesel engines. It removes contaminants and particles from the diesel fuel before it reaches the injection pump and injectors. Clean fuel is vital to protect the precision components of the fuel injection system from wear and damage. So, a fuel filter is definitely part of a diesel engine system.

Comparison: Diesel vs. Petrol Engines

The key difference relevant to this question lies in the method of ignition:

Feature	Diesel Engine	Petrol (Gasoline) Engine
Ignition Method	Compression ignition (auto-ignition of fuel upon injection into hot compressed air)	Spark ignition (ignition initiated by an electric spark from a spark plug)
Fuel Injection	Fuel injected directly into cylinder near the end of the compression stroke	Air-fuel mixture drawn into cylinder during intake stroke (or fuel injected into intake port/cylinder)
Components	Injectors, Fuel injection pump, Fuel filter, Glow plugs (for cold starts), etc.	Spark plugs, Carburetor or Fuel injectors, Ignition coil, Fuel filter, etc.

Based on this comparison and the analysis of each option, the component that is not a part of a diesel engine is the spark plug, as diesel engines use compression ignition, not spark ignition.

Conclusion

Reviewing the components, the injector, fuel injection pump, and fuel filter are all essential parts of a diesel engine's fuel and injection system. The spark plug, however, is a component used specifically for igniting the fuel-air mixture in petrol engines, which operate on the spark ignition principle. Diesel engines operate on the compression ignition principle and therefore do not use spark plugs.

The component that is not part of a diesel engine is the spark plug.

Revision Table: Key Diesel Engine Components

Component	Function in Diesel Engine
Injector	Sprays fuel into the combustion chamber.
Fuel Injection Pump	Delivers high-pressure fuel to injectors.
Fuel Filter	Cleans fuel before injection.
Spark Plug	Not used. Used for spark ignition in petrol engines.
Glow Plug	Used for heating the combustion chamber for easier cold starting (common in many diesel engines).

Additional Information on Diesel Engine Operation

Diesel engines work based on the Diesel cycle. Here's a simplified overview:

- Intake Stroke:** Only air is drawn into the cylinder.
- Compression Stroke:** The air is compressed to a very high pressure (typically 18:1 to 25:1 ratio), causing its temperature to rise significantly.
- Power Stroke:** Near the end of the compression stroke, fuel is injected into the hot compressed air. The high temperature of the air causes the fuel to auto-ignite. The resulting combustion creates high pressure, pushing the piston down and generating power.
- Exhaust Stroke:** The burnt gases are expelled from the cylinder.

This reliance on high compression and auto-ignition is the defining characteristic of a diesel engine, distinguishing it from a petrol engine's need for a spark plug to initiate combustion.

168. Answer: a

Explanation:

Understanding Metal Properties and Lightness

When we talk about a metal being "lightest," we are typically referring to its density. Density is a fundamental physical property of a substance that measures its mass per unit volume. A material with lower density is considered "lighter" than a material with higher density, assuming equal volumes are compared. The formula for density (ρ) is given by:

$$\rho = \frac{m}{V}$$

where m is the mass and V is the volume.

To determine which among the given metals is the lightest, we need to compare their densities. The options provided are Aluminium, Tin, Lead, and Copper. Let's look at the approximate densities of these metals:

Approximate Densities of Given Metals

Metal	Chemical Symbol	Approximate Density (g/cm^3)
Aluminium	Al	2.7
Tin	Sn	7.3
Lead	Pb	11.3
Copper	Cu	9.0

Comparing Densities to Find the Lightest Metal

Based on the density values listed in the table:

- Aluminium has a density of approximately 2.7 g/cm^3 .
- Tin has a density of approximately 7.3 g/cm^3 .
- Lead has a density of approximately 11.3 g/cm^3 .
- Copper has a density of approximately 9.0 g/cm^3 .

Comparing these values, we can see that Aluminium has the lowest density (2.7 g/cm^3) among the four metals. A lower density means that for the same volume, Aluminium would have less mass compared to Tin, Lead, or Copper.

Therefore, Aluminium is considered the lightest metal among the given options. Its low density makes it useful in applications where weight is a critical factor, such as in aerospace and automotive industries.

Conclusion on Lightest Metal

By comparing the densities of Aluminium, Tin, Lead, and Copper, it is clear that Aluminium has the lowest density. This indicates that Aluminium is the lightest metal among the choices provided.

Let's re-examine the options:

- Aluminium (Density $\approx 2.7 \text{ g/cm}^3$)

- Tin (Density $\approx 7.3 \text{ g/cm}^3$)
- Lead (Density $\approx 11.3 \text{ g/cm}^3$)
- Copper (Density $\approx 9.0 \text{ g/cm}^3$)

The smallest value among these densities is 2.7 g/cm^3 , which corresponds to Aluminium.

Revision Table: Metal Properties Summary

Summary of Metal Properties Related to Density

Metal	Approximate Density (g/cm^3)	Relative Lightness
Aluminium	2.7	Lightest
Tin	7.3	Heavier than Aluminium
Copper	9.0	Heavier than Aluminium and Tin
Lead	11.3	Heaviest among these options

Additional Information on Metal Density

While density is often used to describe how "light" or "heavy" a metal is, it's important to remember that it is an intensive property, meaning it doesn't depend on the amount of the substance. Atomic weight, on the other hand, is the average mass of atoms of an element. Although there is often a correlation between atomic weight and density (heavier atoms packed tightly tend to result in higher density), it's not a direct proportional relationship due to differences in crystal structure and atomic packing efficiency. For instance, Gold (Au) has a much higher atomic weight than Lead (Pb), and also a higher density ($\approx 19.3 \text{ g/cm}^3$). However, comparing Lead and Copper, Lead has a higher atomic weight but also a higher density, illustrating the general trend. But the most accurate way to compare the "lightness" of equal volumes of different metals is by comparing their densities.

169. Answer: b

Explanation:

Understanding Woodruff Keys

Woodruff keys are a type of machine key used to secure a rotating machine element, such as a gear or pulley, onto a shaft. Their main purpose, like other keys, is to prevent relative rotation between the shaft and the element mounted on it, transmitting torque.

Distinctive Shape and Installation of Woodruff Keys

What sets a **woodruff key** apart is its shape. Unlike standard rectangular or square keys, a **woodruff key** is segment-shaped, essentially a semicircle. The key fits into a semicircular slot (keyseat) machined into the shaft and a rectangular keyway in the hub of the mating part (like a gear or pulley).

This unique geometry offers specific advantages during assembly:

- The semicircular slot on the shaft allows the key to tilt or pivot slightly.
- This pivoting ability helps the key align itself with the keyway in the mating part (the hub) as the hub is slid onto the shaft.
- This characteristic is known as **self-aligning**.
- The **ease of fitting** comes from this self-aligning nature, making assembly quicker and less prone to binding, especially when there might be slight axial or angular misalignment between the shaft keyseat and the hub keyway.

Analyzing the Options

Let's evaluate the given options in the context of why **woodruff keys** are chosen:

- **Option 1: Because they are less expensive** – Cost can be a factor in design, but it's not the primary technical reason for selecting a **woodruff key** over other types. Their specific advantage lies in their function and installation.
- **Option 2: Because they are easy to fit and self-align** – As explained above, the semicircular shape and the way they fit into the shaft make **woodruff keys** very easy to assemble and inherently **self-aligning**. This is a significant advantage in manufacturing and maintenance.
- **Option 3: Because they can take shear stress** – All machine keys are designed to transmit torque by resisting shear stress. Taking shear stress is the fundamental function of a key, not a unique reason to choose a **woodruff key** specifically over other key types that also handle shear stress.
- **Option 4: Because they can take heavy load** – In fact, the deep slot required for a **woodruff key** in the shaft can potentially weaken the shaft more than the shallower slot for a standard rectangular key. **Woodruff keys** are generally used for lighter loads or in applications where assembly ease is prioritized over maximum torque transmission or shaft strength.

Considering the defining characteristics and primary advantages, the ease of fitting and self-alignment are the main reasons for using **woodruff keys**.

Why Ease of Fit Matters for Woodruff Keys

The **ease of fitting** and **self-aligning** of **woodruff keys** simplify the assembly process. In mass production or repair scenarios, reducing assembly time and complexity is valuable. This feature is particularly beneficial when tolerances might accumulate, or when dealing with longer keyways.

Key Type	Shape	Shaft Keyseat	Hub Keyway	Primary Advantage	Load Capacity (Generally)
Parallel Key	Rectangular/Square	Rectangular slot	Rectangular slot	Simple, good for torque transmission	Moderate to Heavy
Taper Key	Tapered rectangular	Tapered slot	Tapered slot	Provides tight fit, prevents axial movement	Moderate to Heavy
Woodruff Key	Semicircular segment	Semicircular slot	Rectangular slot	Easy to fit, self-aligning	Light to Moderate

Revision Table: Comparing Key Types

Feature	Woodruff Key	Parallel Key
Shape	Semicircular segment	Rectangular or Square
Shaft Slot (Keyseat)	Semicircular, deeper cut	Rectangular, shallower cut
Hub Slot (Keyway)	Rectangular	Rectangular
Installation Ease	High (Self-aligning)	Requires precise alignment
Shaft Weakening	More significant due to deep cut	Less significant
Load Suitability	Lighter to Moderate	Moderate to Heavy

Additional Information: Applications of Woodruff Keys

Woodruff keys are commonly found in applications where assembly convenience is important and where the loads are not excessively high. Some examples include:

- Machine tool spindles
- Automotive steering mechanisms
- Some types of agricultural machinery
- Where short keyways are used

Despite the advantage of **self-alignment** and **ease of fitting**, designers must consider the potential weakening of the shaft due to the deep semicircular cut when selecting a key type for a specific application, especially under high stress or fatigue conditions.

170. Answer: a

Explanation:

Understanding Heat Treatment Processes and Surface Scale

Heat treatment processes are crucial for altering the microstructure and properties of metals, such as increasing hardness, toughness, or ductility. However, exposing metals to high temperatures, especially in the presence of oxygen (like in air), can lead to oxidation, commonly known as scaling. Scale is a layer of metal oxide formed on the surface, which can affect the surface finish and dimensional accuracy of the component. The question asks which of the given heat treatment processes typically results in a scale-free surface.

What is Scale in Heat Treatment?

Scale refers to the oxides formed on the surface of metal components when they are heated to high temperatures in an oxidizing atmosphere. This occurs because the metal atoms react with oxygen from the environment.

Analyzing the Heat Treatment Options

Let's look at each heat treatment process listed and consider its tendency to cause scaling:

- **Induction Hardening:** This process uses electromagnetic induction to rapidly heat the surface of a metal component. Heating is localized and very quick, often followed by immediate quenching. Because the time at high temperature is minimal and rapid cooling occurs, there is very little opportunity for significant oxidation or scale formation, especially when compared to furnace-based processes that involve longer soaking times at high temperatures in air. Some induction hardening setups also use a controlled atmosphere or submerged quenching to further minimize scale.
- **Flame Hardening:** Similar to induction hardening, flame hardening heats the surface rapidly, but it uses an oxy-acetylene or similar flame directly applied to the surface. While heating is fast, the process is performed in air, and the hot surface is directly exposed to the atmosphere, which promotes oxidation and scale formation.
- **Case Hardening:** This is a group of processes (like carburizing, nitriding, carbonitriding) aimed at increasing the hardness of the surface layer (case) of a component while keeping the core soft. Carburizing and carbonitriding are typically performed at high temperatures ($\approx 850^{\circ}\text{C}$ to 1050°C) in atmospheres containing carbon. While the furnace atmosphere is controlled, subsequent steps like quenching and tempering (often in air) can still lead to scale formation. Nitriding is performed at lower temperatures ($\approx 500^{\circ}\text{C}$ to 570°C) in a nitrogen-rich atmosphere (e.g., ammonia), which significantly reduces *oxide* scale compared to high-temperature processes in air, but it still forms a distinct nitride layer on the surface.
- **Nitriding:** As mentioned under Case Hardening, nitriding uses a nitrogen atmosphere at moderate temperatures. It forms a hard nitride layer. While it avoids the high-temperature oxidation common in air hardening, the surface is not strictly "scale-free" in the sense of a clean metallic surface, as it has a chemically formed nitride layer.

Comparing these processes, induction hardening, due to its rapid heating, localized nature, and quick quenching, is the process among the options that is best known for minimizing oxide scale and yielding a

relatively scale-free surface finish.

Conclusion

Based on the analysis of how scale forms and the characteristics of each heat treatment process, induction hardening is the process most likely to result in a scale-free surface on the component because of the limited time the metal spends at high temperature in an oxidizing environment.

Revision Table: Comparing Heat Treatment Processes

Process	Heating Method	Atmosphere	Typical Temperature Range	Tendency for Oxide Scale	Surface Finish
Induction Hardening	Electromagnetic Induction	Air or Controlled	Very High (surface)	Low (due to rapid heating/cooling)	Scale-free to low scale
Flame Hardening	Flame	Air	Very High (surface)	High	Scaled
Case Hardening (Carburizing)	Furnace/Bath	Carbon-rich	850–1050°C	Moderate to High (depends on process/quenching/tempering)	Can be scaled
Nitriding	Furnace	Nitrogen-rich (Ammonia)	500–570°C	Very Low (oxide scale), forms nitride layer	Relatively clean, nitride layer present

Additional Information on Surface Finish and Scaling

Achieving a good surface finish and minimizing scale are important considerations in many engineering applications. Scaling can lead to:

- Loss of material.
- Reduced dimensional accuracy.
- Poor surface appearance.
- Difficulty in subsequent machining or finishing operations.

Processes that minimize time at high temperatures or are conducted in controlled, non-oxidizing atmospheres (like vacuum, inert gas, or specific reducing atmospheres) are preferred when a scale-free surface is required. Induction hardening falls into this category due to its speed and localization. Other methods to prevent scaling include applying protective coatings or using salt baths.

171. Answer: d

Explanation:

Understanding Crankshaft Manufacturing Processes

The question asks about the primary manufacturing process used to make a crankshaft. A crankshaft is a critical component in an internal combustion engine, responsible for converting the linear motion of the pistons into rotational motion, which powers the vehicle or machinery.

Due to the high stresses and dynamic loads a crankshaft experiences, it must be made from materials with high strength, toughness, and fatigue resistance. The manufacturing process chosen significantly impacts these properties.

Let's look at the given options:

- **Casting:** Casting involves pouring molten metal into a mold. While suitable for complex shapes and often more cost-effective for high volume, cast crankshafts generally have lower strength and fatigue resistance compared to wrought processes like forging. Some modern engines do use cast iron or cast steel crankshafts, but forging is traditionally preferred for high-performance or heavy-duty applications.
- **Turning:** Turning is a machining process where a workpiece rotates against a cutting tool. It is used for shaping cylindrical parts and is commonly used as a *finishing* process for crankshafts (to create journals and surfaces accurately) but is not the primary process for forming the overall shape and achieving the required material properties from raw material.
- **Pressing:** Pressing, or stamping, is a process typically used for forming sheet metal into shapes. While presses are used in forging operations, pressing by itself usually refers to sheet metal work and is not suitable for creating a solid, complex, high-strength component like a crankshaft.
- **Forging:** Forging is a manufacturing process involving shaping metal by plastic deformation under compressive forces, typically applied with a hammer or press. Forging refines the grain structure of the metal, aligning it along the direction of stress, which significantly increases the strength, toughness, and fatigue life of the component. This makes forging an ideal process for manufacturing components subjected to high dynamic loads, such as crankshafts, connecting rods, and gears. Forged crankshafts are known for their durability and reliability.

Considering the functional requirements of a crankshaft and the properties imparted by different manufacturing methods, forging is the process most commonly used to manufacture high-strength, durable crankshafts.

Comparing Crankshaft Manufacturing Methods

Process	Description	Typical Crankshaft Use	Key Benefit for Crankshafts	Limitation for Crankshafts
Forging	Shaping metal using compressive force (hammer/press).	High-performance, heavy-duty engines.	Improved strength, toughness, fatigue resistance due to favorable grain structure.	Higher initial tooling cost, requires subsequent machining.
Casting	Pouring molten metal into a mold.	Standard passenger car engines.	Can create complex shapes easily, lower cost for high volume.	Lower strength and fatigue resistance compared to forging.
Turning	Machining process removing material from a rotating workpiece.	Finishing journals and surfaces after forging or casting.	Achieves precise dimensions and surface finish.	Not suitable for initial bulk shaping or achieving internal material properties like forging.
Pressing	Shaping material (usually sheet) with a press.	Not used for primary crankshaft manufacturing.	Suitable for sheet metal components.	Not suitable for solid, complex, high-stress components like crankshafts.

Conclusion on Crankshaft Manufacturing

Based on the analysis of the common manufacturing processes and the specific demands placed on a crankshaft, forging stands out as the preferred method for producing crankshafts that require high strength and excellent fatigue life. While casting is also used, forging provides superior mechanical properties essential for many engine applications.

Revision Table: Manufacturing Processes

Process Type	Typical Application Examples	Key Characteristic
Casting	Engine blocks, cylinder heads, simple machine parts.	Forms complex shapes by pouring liquid material into a mold.
Turning	Shafts, bolts, cylindrical components.	Removes material from a rotating part to create cylindrical forms.
Pressing	Car body panels, simple brackets, deep-drawn parts.	Shapes sheet metal using dies and presses.
Forging	Crankshafts, connecting rods, gears, tools, turbine blades.	Shapes solid metal by hammering or pressing, improving strength and grain structure.

Additional Information on Crankshafts and Forging

Crankshafts are typically made from carbon steel or alloy steel. The specific grade of steel is chosen based on the required strength and durability. After forging, crankshafts often undergo heat treatments, such as hardening and tempering, to further enhance their mechanical properties. Machining processes, including turning, milling, and grinding, are then used to achieve the final precise dimensions and surface finishes required for the journals and other critical features.

The forging process for crankshafts can be done using various methods, including open-die forging and closed-die forging. Closed-die forging is more common for mass production as it allows for higher precision and consistency in shape.

Understanding the manufacturing process helps appreciate why components like crankshafts are robust and reliable under demanding operating conditions.

172. Answer: a

Explanation:

Understanding the Twist Drill Point Angle

A **twist drill** is a type of drill bit commonly used in drilling machines. It is designed to cut holes in various materials. A key feature of a twist drill is its **point angle**, which is the angle formed at the tip of the drill by the cutting lips.

The **point angle** significantly affects how the drill performs, including its ability to penetrate the material, center itself, and remove chips. Different materials require different point angles for efficient drilling.

Standard Twist Drill Point Angle Explained

For general-purpose drilling in common materials like mild steel, cast iron, and plastics, a standard **point angle** is widely used. This standard angle provides a good balance between cutting efficiency, chip removal, and drill strength for these common applications.

The widely accepted standard **point angle** for a general-purpose **twist drill** is 118° . This angle is sharp enough to penetrate effectively but also robust enough to resist damage during drilling.

Point Angle Options and Their Uses

Let's consider the provided options and their typical uses:

- 118° : This is the standard **point angle** for general-purpose drilling in common materials like mild steel, cast iron, and non-ferrous metals.
- 135° : This is a flatter **point angle** often used for harder materials like stainless steel or for applications requiring less walking (where the drill tip slides before cutting). It has a stronger tip but requires more

thrust force.

- 60°: This is a very sharp **point angle**, sometimes used for drilling plastics or soft materials, or for specific applications like spotting or center drilling.
- 90°: This angle is not typical for standard **twist drill** points but might be found on specialty tools like countersinks.

Based on the question asking for the **point angle** of a **twist drill** without specifying a material, the standard, general-purpose angle is implied.

Point Angle	Typical Material Use
60° – 90°	Plastics, very soft materials, spotting/center drilling
118°	Standard for mild steel, cast iron, brass, aluminum
135°	Hard steel, stainless steel, materials requiring less walking

Therefore, the most appropriate answer representing the standard **point angle** of a **twist drill** is 118°.

Revision Table – Twist Drill Basics

Feature	Description
Body	The main part of the drill, with flutes.
Flutes	Spiral grooves that remove chips from the hole.
Cutting Lips	The sharpened edges at the point that do the cutting.
Point Angle	The angle formed by the cutting lips at the tip.
Shank	The part held by the drill chuck.

Additional Information – Factors Affecting Twist Drill Performance

Several factors influence how well a **twist drill** works:

- **Point Angle:** As discussed, affects penetration, centering, and material suitability.
- **Helix Angle:** The angle of the flutes. A larger angle is better for soft materials (better chip evacuation), while a smaller angle is better for hard materials (stronger cutting edge).
- **Lip Clearance Angle:** The angle behind the cutting lips that prevents the drill from rubbing against the workpiece. Too little clearance causes rubbing; too much weakens the cutting edge.
- **Web Thickness:** The thickness of the metal connecting the flutes at the center. A thicker web provides more strength but makes drilling more difficult.
- **Material:** The material the drill bit is made from (e.g., High-Speed Steel (HSS), Cobalt HSS, Carbide) affects its hardness, wear resistance, and suitability for different workpiece materials.

Understanding these features helps in selecting the correct drill for a specific task and material, ensuring efficient and accurate drilling.

173. Answer: c

Explanation:

Understanding Broken Bolts in Tapped Holes

When a bolt or screw breaks off inside a tapped hole, it creates a difficult situation. The remaining piece is often flush with or below the surface, making it impossible to grip with standard tools. This broken piece obstructs the hole, preventing a new fastener from being installed. Removing it without damaging the threads of the tapped hole is crucial.

Methods for Removing Broken Bolts

Several methods can be attempted to remove a broken bolt from a tapped hole. These methods vary in effectiveness and risk to the surrounding material and threads. Let's examine some common approaches and the tools used.

Using a Tap Extractor

A tap extractor (or screw extractor/bolt extractor) is a specialized tool designed specifically for removing broken threaded fasteners like bolts or screws from holes. These tools work by gripping the inside of the broken bolt and allowing torque to be applied in the reverse direction of tightening (counter-clockwise for standard right-hand threads) to unscrew the broken piece.

There are different types of extractors, but a common type for bolts uses flutes or splines that grip the internal wall of the broken bolt or requires drilling a pilot hole into the broken bolt first before inserting a spiral flute extractor. The principle is to apply outward pressure and rotational force simultaneously.

Tool Type	How it Works	Suitability for Broken Bolts
Tap Extractor / Bolt Extractor	Grips internal surface (or drilled hole) of broken bolt and rotates it counter-clockwise.	Highly suitable, designed for this purpose.
Heat Treatment	Applying heat can expand or contract metal, sometimes loosening seized parts.	Can be risky, might affect material properties or damage surrounding area/threads. Not a direct removal tool for the broken piece itself.
Pin Punching Hammer	Used to drive pins or other objects.	Not suitable for removing a broken bolt from a tapped hole; could further wedge the bolt or damage threads/hole.
Annealing	Heat treatment process to soften metal.	Softening the bolt might make drilling easier, but it doesn't directly remove the broken piece and requires further steps (like drilling and using an extractor). Annealing the *tapped hole* itself isn't a method for removing the bolt.

Why Other Methods Are Not Ideal

- **With heat treatment:** While sometimes used to loosen seized fasteners, applying heat directly to a broken bolt in a tapped hole can be tricky. Overheating could damage the threads in the tapped hole or alter the material properties of the surrounding component. It's not a direct removal tool but rather a potential aid in loosening.
- **With a pin punching hammer:** A punch is used to drive objects. Using it on a broken bolt stuck in a tapped hole will likely just drive the bolt piece deeper or at an angle, potentially damaging the threads and making removal even harder. It's not the correct tool for extraction.
- **By annealing the tapped drill hole:** Annealing softens metal by heating it. Annealing the *tapped hole* (the part with the threads) is generally undesirable as it could weaken the threads. While annealing the *broken bolt* itself might make it easier to drill, it doesn't remove the bolt directly and requires subsequent steps like drilling and using an extractor. Annealing the hole is incorrect for removal.

Based on the function and design of the tools and methods, a tap extractor or bolt extractor is the most appropriate tool specifically designed for removing a broken bolt from a tapped hole effectively and with minimal risk of damage to the threads.

Revision Table: Broken Bolt Removal

Problem	Recommended Tool	Outcome
Broken bolt in tapped hole	Tap Extractor / Bolt Extractor	Removes broken bolt, preserves threads.

Additional Information: Using Extractors Safely

When using a tap or bolt extractor:

- Always ensure the tool is properly sized for the broken bolt.
- If required, drill a pilot hole carefully down the center of the broken bolt using the correct drill bit size.
- Insert the extractor straight and tap it gently to seat it firmly.
- Apply steady, firm, counter-clockwise pressure. Avoid sudden jerks.
- Using penetrating oil can help loosen the broken bolt before attempting extraction.
- If the bolt is very tight or seized, alternative methods like drilling it out completely may be necessary, often followed by thread repair using a tap or helicoil.

174. Answer: c

Explanation:

Understanding Cylinder Liners in Compressed Ignition Engines

Compressed ignition engines, commonly known as diesel engines, rely on compressing air to a high temperature and then injecting fuel, which ignites due to the heat. A crucial component in the engine's combustion chamber is the cylinder liner.

What are Cylinder Liners?

A cylinder liner is essentially a cylindrical tube that is fitted inside the cylinder block. It forms the inner wall of the engine cylinder, where the piston moves up and down during the engine cycle. Liners are typically made from wear-resistant materials like cast iron alloys.

Why Cylinder Liners are Used

Cylinder liners serve several important functions in compressed ignition engines:

- They provide a wear surface for the piston rings. As the piston moves, there is friction between the piston rings and the cylinder wall. The liner provides a hard, durable surface to minimize wear on the cylinder block itself.
- They can be replaced. If the cylinder wall wears out or gets damaged (e.g., scratched), it is much easier and cheaper to replace the liner than to replace the entire cylinder block.
- They help in heat transfer. The liner helps transfer heat from the combustion process to the coolant flowing through the cylinder block.
- Some types of liners (wet liners) also form a seal with the coolant jacket.

Location of Cylinder Liners

In compressed ignition engines, cylinder liners are typically provided **inside the cylinder block**. The cylinder block is the main structural component of the engine, containing the bores where the pistons operate. The liners are inserted into these bores. There are generally two types:

- **Dry Liners:** These are simple sleeves pressed into the cylinder block bore. They do not come into direct contact with the engine coolant; heat is transferred through the liner wall to the cylinder block and then to the coolant.
- **Wet Liners:** These liners are in direct contact with the engine coolant on their outer surface. They provide better cooling but require seals (like O-rings) at their top and bottom ends to prevent coolant leakage.

Regardless of whether they are wet or dry, the fundamental location is within the bore of the cylinder block.

Analyzing the Options

Let's look at the provided options:

- Option 1: inside the cylinder head - The cylinder head is at the top of the engine and contains valves, spark plugs/injectors, and combustion chambers, but not the main cylindrical wall where the piston slides.
- Option 2: outside the cylinder block - This is incorrect. The cylinder block is the outer structure; the liners go inside it.
- Option 3: inside the cylinder block - This correctly describes the location where cylinder liners are fitted to form the inner wall of the cylinder bore.
- Option 4: on the cylinder head - Similar to option 1, the liner is not located on the cylinder head.

Therefore, the correct location for cylinder liners in compressed ignition engines is inside the cylinder block.

Revision Table: Key Engine Components

Component	Primary Function	Typical Location
Cylinder Block	Main structure, houses cylinders, crankshaft	Base of the engine assembly (above oil pan)
Cylinder Liner	Forms inner wall of cylinder, wear surface	Inside the cylinder block bore
Piston	Moves up/down, transfers force to crankshaft	Inside the cylinder/liner
Cylinder Head	Seals top of cylinder, houses valves/injectors	Top of the cylinder block
Crankshaft	Converts linear piston motion to rotational motion	Housed in the lower part of the cylinder block

Additional Information: Types of Engine Cylinders

Engines can be designed with or without separate cylinder liners:

- **Integral Cylinders:** In some engines (often smaller or gasoline engines), the cylinder bores are machined directly into the cylinder block material. If these bores wear out, the block may need to be rebored to an oversized dimension, or the entire block might need replacement.
- **Engines with Liners:** As discussed, these engines use separate liners inserted into the block. This design is very common in heavy-duty compressed ignition engines due to the high stresses and wear encountered, making maintenance and repair easier and more cost-effective. The use of liners allows the cylinder block to be made from a different material than the wear surface, optimizing properties for different requirements (strength for the block, wear resistance for the liner).

The use of cylinder liners inside the cylinder block is a standard practice in robust engine design for durability and ease of service.

175. Answer: a

Explanation:

Understanding VVTL-i: Full Form and Function

The question asks for the full form of VVTL-i, a technology used in some automobile engines.

VVTL-i stands for a specific type of variable valve timing system. This system is designed to optimize engine performance and efficiency across different engine speeds by adjusting when the engine's intake and exhaust valves open and close, and also how much they open (the valve lift).

Full Form of VVTL-i

The widely recognized and correct full form for VVTL-i is **Variable Valve Timing Lift Intelligence system**.

Let's break down the components of the name:

- **Variable Valve Timing (VVT):** This part refers to the ability of the system to change the timing of the valve events (when they open and close) based on engine operating conditions. Traditional engines have fixed valve timing.
- **Lift (L):** This adds the capability to also change the extent to which the valves open, known as valve lift. Different valve lifts can be beneficial at different engine speeds – for example, higher lift might be better for power at high RPM, while lower lift might be better for efficiency at low RPM.
- **Intelligence (i):** This typically indicates that the system is electronically controlled by the engine's computer (ECU) to make these adjustments in real-time based on various sensor inputs.

Analysis of Options

Let's look at the given options and compare them to the correct full form:

- **Option 1: Variable Valve Timing Lift Intelligence system** – This matches the standard full form of VVTL-i.

- **Option 2: Valve Variation Time Line In** – This does not correspond to any known automotive engine technology acronym. The terms are incorrect.
- **Option 3: Variable Valve Total Line In system** – This also does not represent the full form of VVTL-i. The words "Total Line In system" are inaccurate.
- **Option 4: Variable Valve Time and Lift** – While this option includes "Variable Valve Time and Lift", it misses the crucial "Intelligence system" part, which is integral to the full name and indicates the electronic control aspect.

Based on the analysis, only Option 1 provides the accurate and complete full form for VVTL-i.

Revision Table: Key Terms

Term	Explanation
VVTL-i	Variable Valve Timing Lift Intelligence system
VVT	Variable Valve Timing (adjusts timing only)
Valve Timing	When engine valves open and close relative to piston position.
Valve Lift	How far the engine valves open.
ECU	Engine Control Unit (the computer that controls engine functions like VVTL-i).

Additional Information: How VVTL-i Works

VVTL-i is an advanced engine technology built upon basic Variable Valve Timing (VVT). A standard VVT system might only change the timing duration (how long the valve is open) or overlap (when intake and exhaust valves are open simultaneously). VVTL-i adds the capability to change the valve lift. This is often achieved using different cam profiles (shapes on the camshaft). At low engine speeds, a less aggressive cam profile might be used for better fuel efficiency and smoother idle. As the engine speed increases, the system can switch to a more aggressive cam profile, increasing valve lift and duration, allowing more air into the engine for increased power output.

This ability to change both timing and lift provides a significant advantage in optimizing engine performance across a wide range of operating conditions, improving both fuel economy and power compared to engines with fixed valve operation or only VVT.