

Answers

1. Answer: b

Explanation:

Understanding Angular Measurement: Radians and Degrees

Angular measurement is fundamental in many fields, including mathematics, physics, engineering, and navigation. There are primarily two common units used to measure angles: degrees and radians. This question asks about the relationship between these two units, specifically how many degrees are equivalent to one radian.

Connecting Radians and Degrees

The relationship between radians and degrees is defined by the circle. A full circle represents 360 degrees, and the circumference of a circle is given by the formula $2\pi r$, where r is the radius. One radian is defined as the angle subtended at the center of a circle by an arc whose length is equal to the radius of the circle. A full circle, with a circumference of $2\pi r$, contains 2π radians.

Therefore, the relationship is:

- 360 degrees = 2π radians

We can simplify this by dividing both sides by 2:

- 180 degrees = π radians

This is the key conversion formula between degrees and radians.

Calculating 1 Radian in Degrees

To find out how many degrees are in one radian, we can rearrange the conversion formula ($180 \text{ degrees} = \pi \text{ radians}$). We need to isolate 'radian' on one side. Let's divide both sides of the equation by π :

$$\frac{180 \text{ degrees}}{\pi} = \frac{\pi \text{ radians}}{\pi}$$

This gives us:

$$1 \text{ radian} = \frac{180}{\pi} \text{ degrees}$$

Now, we need to use an approximate value for π . A common approximation for π is 3.14159. Let's perform the division:

$$1 \text{ radian} \approx \frac{180}{3.14159} \text{ degrees}$$

$$1 \text{ radian} \approx 57.2958 \text{ degrees}$$

Looking at the options provided, 57.27 degrees is the closest value to our calculation of approximately 57.2958 degrees. The slight difference comes from the approximation used for π .

Comparing Options for 1 Radian Equivalent

Let's look at the given options:

- 65.27 degrees
- 57.27 degrees
- 90 degrees
- 180 degrees

Based on our calculation, 1 radian is approximately 57.2958 degrees. The option that is closest to this value is 57.27 degrees.

Radians	Degrees
2π	360
π	180
$\frac{\pi}{2}$	90
$\frac{\pi}{3}$	60
$\frac{\pi}{4}$	45
1	≈ 57.296

The option 57.27 degrees is the widely accepted approximate value for one radian when expressed in degrees.

Revision Table: Angular Measurement Units

Concept	Description
Degree	A unit of angular measure, where a full rotation is 360 degrees (360°).
Radian	A unit of angular measure, where a full rotation is 2π radians. One radian is the angle subtended by an arc equal in length to the radius.
Conversion	The primary relationship is $180 \text{ degrees} = \pi \text{ radians}$.
1 Radian to Degree	$1 \text{ radian} = \frac{180}{\pi} \text{ degrees} \approx 57.2958 \text{ degrees}$. Often approximated as 57.3 degrees or 57.27 degrees depending on required precision.
1 Degree to Radian	$1 \text{ degree} = \frac{\pi}{180} \text{ radians} \approx 0.01745 \text{ radians}$.

Additional Information: Why Use Radians?

While degrees are intuitive because a full circle is a 'nice' number like 360, radians have significant advantages in higher mathematics and physics, particularly in calculus. Many formulas involving trigonometric functions are simpler and more elegant when angles are expressed in radians. For example, the derivative of $\sin(x)$ is $\cos(x)$ when x is in radians, but it would involve a conversion factor if x were in degrees.

The definition of a radian being related to the radius and arc length makes it a more natural unit in contexts involving circular motion, arc length ($s = r\theta$, where θ is in radians), and sector area ($A = \frac{1}{2}r^2\theta$, where θ is in radians).

2. Answer: a

Explanation:

Correct answer: 1.56

The general form of an arithmetic series is:

$$a, a + d, a + 2d, a + 3d, \dots$$

where 'a' is the first term and 'd' is the common difference.

In the given problem, the first term $a = 1.14$ and the common difference $d = 0.14$.

The n th term of an arithmetic series can be found using the formula:

$$a_n = a + (n - 1)d$$

For example, the 4th term (missing term) is $a_4 = 1.14 + (4 - 1) \times 0.14 = 1.14 + 3 \times 0.14 = 1.14 + 0.42 = 1.56$.

3. Answer: c

Explanation:

Understanding the Exam Score Problem

This problem involves finding the lowest score in an examination given specific relationships between the highest and lowest scores. We are told two key pieces of information:

- The difference between the highest score and the lowest score is 55.
- The highest score is $\frac{9}{4}$ times the lowest score.

We need to use these relationships to set up equations and solve for the value of the lowest score.

Setting Up Equations for Exam Scores

Let's represent the scores using variables:

- Let H be the highest score.
- Let L be the lowest score.

From the problem statement, we can write two equations based on the given information:

Equation 1 (Difference): The difference between the highest and lowest score is 55.

$$H - L = 55$$

Equation 2 (Ratio): The highest score is $\frac{9}{4}$ times the lowest score.

$$H = \frac{9}{4}L$$

Solving for the Lowest Score

We have a system of two linear equations with two variables (H and L). We can solve this system using substitution. Since Equation 2 already gives us an expression for H in terms of L , we can substitute this expression into Equation 1.

Substitute $H = \frac{9}{4}L$ into the equation $H - L = 55$:

$$\left(\frac{9}{4}L\right) - L = 55$$

Now, we need to solve this equation for L . To combine the terms involving L , we can write L as $\frac{4}{4}L$:

$$\frac{9}{4}L - \frac{4}{4}L = 55$$

Combine the terms on the left side:

$$\left(\frac{9}{4} - \frac{4}{4}\right)L = 55$$

$$\frac{9-4}{4}L = 55$$

$$\frac{5}{4}L = 55$$

To isolate L , multiply both sides of the equation by the reciprocal of $\frac{5}{4}$, which is $\frac{4}{5}$:

$$L = 55 \times \frac{4}{5}$$

Now, calculate the value of L :

$$L = \frac{55 \times 4}{5}$$

We can simplify by dividing 55 by 5:

$$L = 11 \times 4$$

$$L = 44$$

So, the lowest score is 44.

Verification of the Scores

Let's check if the calculated scores satisfy the original conditions:

- Lowest Score $L = 44$.
- Highest Score $H = \frac{9}{4}L = \frac{9}{4} \times 44$.
- $H = 9 \times \frac{44}{4} = 9 \times 11 = 99$.

Check the difference:

$$H - L = 99 - 44 = 55$$

This matches the first condition (difference is 55). The second condition ($H = \frac{9}{4}L$) was used to calculate H , so it is also satisfied.

The calculated lowest score of 44 is consistent with the problem statement.

Conclusion: Finding the Lowest Exam Score

By setting up and solving a system of linear equations based on the given relationships between the highest and lowest exam scores, we found that the lowest score is 44.

Revision Table: Key Concepts

Concept	Description	Application in Problem
Variables	Symbols representing unknown quantities.	H for highest score, L for lowest score.
Linear Equations	Equations where variables are raised to the power of 1.	$H - L = 55$, $H = \frac{9}{4}L$.
System of Equations	A set of two or more equations with the same variables.	Solving $H - L = 55$ and $H = \frac{9}{4}L$ together.
Substitution Method	Solving a system by substituting an expression from one equation into another.	Substituting $\frac{9}{4}L$ for H in $H - L = 55$.
Solving for a Variable	Manipulating an equation to isolate the unknown variable.	Solving $\frac{5}{4}L = 55$ for L .

Additional Information: Solving Systems of Equations

Systems of linear equations like the one in this exam score problem can be solved using different methods. The substitution method, as used here, is particularly useful when one equation is already solved for one variable (like $H = \frac{9}{4}L$). Other common methods include:

- **Elimination Method:** This involves adding or subtracting the equations (or multiples of the equations) to eliminate one of the variables.
- **Graphical Method:** Each equation represents a line; the solution is the point where the lines intersect. This method is less precise for non-integer solutions.
- **Matrix Methods:** Using matrices and matrix operations (like finding the inverse) to solve the system. This is more advanced and often used for larger systems.

Understanding how to translate word problems into mathematical equations is a fundamental skill in algebra. Always define your variables clearly and ensure your equations accurately represent the relationships described in the problem.

4. Answer: c

Explanation:

Understanding Work and Energy in Physics

The question asks for the work required to change the speed of an object. In physics, the work done on an object is often related to its change in kinetic energy. This relationship is described by the Work-Energy Theorem.

Applying the Work-Energy Theorem

The Work-Energy Theorem states that the net work done on an object is equal to the change in its kinetic energy. Mathematically, this is written as:

$$W = \Delta KE = KE_{final} - KE_{initial}$$

Where:

- W is the work done on the object.
- ΔKE is the change in kinetic energy.
- $KE_{initial}$ is the initial kinetic energy.
- KE_{final} is the final kinetic energy.

The kinetic energy of an object with mass m and speed v is given by:

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

Step-by-Step Calculation of Work Done

We are given the following information:

- Mass of the ball, $m = 1 \text{ kg}$
- Initial speed of the ball, $v_i = 2 \text{ m/s}$
- Final speed of the ball, $v_f = 4 \text{ m/s}$

1. Calculate the Initial Kinetic Energy ($KE_{initial}$):

Using the formula $KE = \frac{1}{2}mv^2$:

$$KE_{initial} = \frac{1}{2} \times m \times v_i^2$$

$$KE_{initial} = \frac{1}{2} \times 1 \text{ kg} \times (2 \text{ m/s})^2$$

$$KE_{initial} = \frac{1}{2} \times 1 \times 4 \text{ J}$$

$$KE_{initial} = 2 \text{ J}$$

2. Calculate the Final Kinetic Energy (KE_{final}):

Using the formula $KE = \frac{1}{2}mv^2$:

$$KE_{final} = \frac{1}{2} \times m \times v_f^2$$

$$KE_{final} = \frac{1}{2} \times 1 \text{ kg} \times (4 \text{ m/s})^2$$

$$KE_{final} = \frac{1}{2} \times 1 \times 16 \text{ J}$$

$$KE_{final} = 8 \text{ J}$$

3. Calculate the Work Done (W):

Using the Work-Energy Theorem $W = KE_{final} - KE_{initial}$:

$$W = 8 \text{ J} - 2 \text{ J}$$

$$W = 6 \text{ J}$$

Therefore, the work that needs to be done to increase the speed of the 1 kg ball from 2 m/s to 4 m/s is 6 J.

Summary of Ball's Kinetic Energy

State	Speed (v)	Mass (m)	Kinetic Energy ($KE = \frac{1}{2}mv^2$)
Initial	2 m/s	1 kg	2 J
Final	4 m/s	1 kg	8 J

Revision Table: Key Concepts for Work and Energy

Concept	Definition/Formula	Units
Work Done (W)	Energy transferred by a force acting over a distance; Change in kinetic energy	Joules (J)
Kinetic Energy (KE)	Energy of motion	Joules (J)
Work-Energy Theorem	$W_{net} = \Delta KE$	N/A
Mass (m)	Measure of inertia	Kilograms (kg)
Speed (v)	Magnitude of velocity	Meters per second (m/s)

Additional Information on Work, Energy, and Speed

Work done is a scalar quantity, meaning it only has magnitude, not direction. When positive work is done on an object, its kinetic energy increases (it speeds up). When negative work is done, its kinetic energy decreases (it slows down).

Kinetic energy is directly proportional to the mass of the object and the square of its speed. This means that doubling the speed of an object increases its kinetic energy by a factor of four.

The units used in these calculations are standard SI units:

- Mass is in kilograms (kg).
- Speed is in meters per second (m/s).
- Work and energy are in Joules (J). One Joule is equivalent to one Newton-meter (N·m) or one $\text{kg} \cdot \text{m}^2 / \text{s}^2$.

This problem highlights how work is the mechanism by which energy is transferred to an object to change its state of motion.

5. Answer: c

Explanation:

Solving Age Ratio Problems: Finding the Difference

This question involves finding the difference between the current ages of two individuals, X and Y, given ratios of their ages at different points in time. We are provided with two key pieces of information:

- The ratio of their current ages is 4:7.
- The ratio of their ages three years ago was 1:2.

Setting up the Age Equations

Let the current age of X be $4k$ years and the current age of Y be $7k$ years, where k is a common multiple. This represents the given ratio of their current ages (4:7).

Three years ago, the ages of X and Y would have been:

- Age of X three years ago: $4k - 3$ years
- Age of Y three years ago: $7k - 3$ years

We are told that the ratio of their ages three years ago was 1:2. We can write this as an equation:

$$\frac{4k - 3}{7k - 3} = \frac{1}{2}$$

Solving for the Unknown Variable

To find the value of k , we can cross-multiply the equation:

$$2 \times (4k - 3) = 1 \times (7k - 3)$$

Now, distribute on both sides of the equation:

$$8k - 6 = 7k - 3$$

To solve for k , we need to isolate the k terms on one side and the constant terms on the other side. Subtract $7k$ from both sides:

$$8k - 7k - 6 = 7k - 7k - 3$$

$$k - 6 = -3$$

Now, add 6 to both sides:

$$k - 6 + 6 = -3 + 6$$

$$k = 3$$

So, the value of the common multiple k is 3.

Calculating Current Ages

Now that we have the value of k , we can find the current ages of X and Y:

- Current age of X = $4k = 4 \times 3 = 12$ years
- Current age of Y = $7k = 7 \times 3 = 21$ years

Finding the Difference in Current Ages

The question asks for the difference between their current ages ($Y - X$). We calculate this difference:

$$\text{Difference} = \text{Current age of Y} - \text{Current age of X}$$

$$\text{Difference} = 21 - 12$$

$$\text{Difference} = 9$$

The difference between the current ages of Y and X is 9 years.

Verification of Age Ratios

Let's check if our calculated ages satisfy the conditions given in the problem:

- Current ages: X = 12, Y = 21. Ratio = 12 : 21. Dividing both by 3, we get 4 : 7. This matches the given current ratio.
- Ages three years ago: X = $12 - 3 = 9$, Y = $21 - 3 = 18$. Ratio = 9 : 18. Dividing both by 9, we get 1 : 2. This matches the given ratio from three years ago.

The calculated ages satisfy both conditions, confirming our solution is correct.

Person	Current Age Ratio Part	Current Age (Calculated)	Age 3 Years Ago (Calculated)	Age 3 Years Ago Ratio Part
X	4	12	9	1
Y	7	21	18	2

Revision Table: Key Steps in Age Ratio Problems

Step	Description	Application in this Problem
1	Represent current ages using a variable and the given ratio.	$X = 4k, Y = 7k$
2	Express ages at the specified past or future time based on current ages.	$X \text{ (3 years ago)} = 4k - 3, Y \text{ (3 years ago)} = 7k - 3$
3	Form an equation using the ages from Step 2 and the ratio given for that time.	$\frac{4k-3}{7k-3} = \frac{1}{2}$
4	Solve the equation to find the value of the variable.	$k = 3$
5	Substitute the variable's value back into the expressions for current ages.	$X = 12, Y = 21$
6	Calculate the required value (e.g., difference, sum, specific age).	Difference = $21 - 12 = 9$

Additional Information: Understanding Ratios and Age Word Problems

Ratios are used to compare quantities. A ratio of 4:7 means that for every 4 units of X's age, there are 7 units of Y's age. When solving age word problems involving ratios over time, we typically use a variable (like k) to represent the common multiple in the ratio. This allows us to set up algebraic equations that can be solved.

It's crucial to correctly represent the ages at different points in time. If the problem refers to ages 'n' years ago, you subtract 'n' from the current ages. If it refers to ages 'n' years hence (in the future), you add 'n' to the current ages.

Always verify your calculated ages by plugging them back into the original problem statements to ensure they satisfy all given conditions, especially the age ratios at different times.

6. Answer: b

Explanation:

Understanding the Median in Statistics

The median is a measure of central tendency that represents the middle value in a dataset. Unlike the mean (average), the median is not affected by extremely large or small values (outliers), making it a robust measure, especially for skewed data. To find the median, the data must first be arranged in order.

Steps to Calculate the Median

Follow these steps to find the median of any set of numbers:

1. Arrange the data points in ascending or descending order. Ascending order (from smallest to largest) is usually easier.
2. Count the total number of data points in the set. Let this count be n .
3. Determine the median based on whether n is odd or even:
 - If n is an **odd** number, the median is the value located exactly in the middle. Its position is given by the formula $\frac{n+1}{2}$.
 - If n is an **even** number, there are two middle values. The median is the average of these two middle values. Their positions are given by $\frac{n}{2}$ and $\frac{n}{2} + 1$.

Finding the Median of 8, 5, 7, 9, 11, 6, 10

Let's apply the steps to the given set of numbers: 8, 5, 7, 9, 11, 6, 10.

Step 1: Arrange the data in ascending order.

The numbers are 5, 6, 7, 8, 9, 10, 11.

Step 2: Count the number of data points (n).

There are 7 numbers in the set. So, $n = 7$.

Step 3: Determine the median.

Since $n = 7$ is an odd number, the median is the middle value. The position of the median is given by the formula:

$$\text{Median Position} = \frac{n+1}{2} = \frac{7+1}{2} = \frac{8}{2} = 4^{\text{th}} \text{ position.}$$

Now, find the value at the 4th position in the ordered list (5, 6, 7, 8, 9, 10, 11).

The number at the 4th position is 8.

Therefore, the median of the dataset is 8.

Summary of Median Calculation

Original Data	Ordered Data	Number of Data Points (n)	Calculation	Median
8, 5, 7, 9, 11, 6, 10	5, 6, 7, 8, 9, 10, 11	7	$n = 7$ (odd), Position = $(\frac{7+1}{2})^{\text{th}} = 4^{\text{th}}$	8

Revision Table: Measures of Central Tendency

Measure	Description	How to Calculate (Basic)	Affected by Outliers?
Mean	The average value	Sum of all values divided by the number of values	Yes
Median	The middle value when data is ordered	Middle value for odd n ; Average of two middle values for even n	No (Robust)
Mode	The value that appears most frequently	Find the value(s) with the highest frequency	No

Additional Information on Median and Data Analysis

The median is a fundamental concept in descriptive statistics used to summarize the central position of a dataset. It provides insight into the typical value, especially when the data distribution is not symmetrical.

- Understanding median alongside mean and mode gives a more complete picture of the data's distribution.
- The median is widely used in economics (e.g., median income), real estate (e.g., median home price), and healthcare due to its resistance to extreme values.
- For large datasets, calculating the median manually can be time-consuming, but statistical software makes it easy.
- The median corresponds to the 50th percentile of the data.

7. Answer: c

Explanation:

Solving Number Analogy Questions

This question asks us to find a relationship between the first pair of numbers and apply it to the third number to find the fourth. The given analogy is:

$$-\frac{9}{11} : \frac{11}{9} :: \frac{13}{2} : ?$$

Let's analyze the relationship between the first number, $-\frac{9}{11}$, and the second number, $\frac{11}{9}$.

We observe two things:

- The fraction has been inverted. The numerator and the denominator have swapped positions. This is also known as taking the reciprocal of the fraction.
- The sign has changed from negative to positive.

So, the operation applied to $-\frac{9}{11}$ to get $\frac{11}{9}$ is taking the reciprocal and changing the sign.

Mathematically, taking the reciprocal of $-\frac{9}{11}$ gives $-\frac{11}{9}$. Then, changing the sign of $-\frac{11}{9}$ gives $+\frac{11}{9}$, or simply $\frac{11}{9}$. However, let's consider the operation as a whole. If we take the negative of the reciprocal of $-\frac{9}{11}$, we get $-\left(\frac{1}{-\frac{9}{11}}\right) = -\left(-\frac{11}{9}\right) = \frac{11}{9}$. This operation is "take the negative reciprocal".

Alternatively, if we consider taking the reciprocal first, $\frac{1}{-\frac{9}{11}} = -\frac{11}{9}$. Then, the sign changes from negative to positive. This implies multiplying by -1 after taking the reciprocal.

Let's test the operation "take the negative reciprocal".

Negative reciprocal of x is $-\frac{1}{x}$.

For $-\frac{9}{11}$, the negative reciprocal is $-\frac{1}{-\frac{9}{11}} = -(-\frac{11}{9}) = \frac{11}{9}$. This matches the second number.

Now, we apply the same relationship to the third number, $\frac{13}{2}$.

We need to find the negative reciprocal of $\frac{13}{2}$.

The reciprocal of $\frac{13}{2}$ is $\frac{2}{13}$.

The negative reciprocal of $\frac{13}{2}$ is $-\frac{2}{13}$.

So, the missing number in the analogy is $-\frac{2}{13}$.

Let's check the given options:

- Option 1: $\frac{3}{7}$
- Option 2: $\frac{2}{13}$
- Option 3: $-\frac{2}{13}$
- Option 4: $-\frac{7}{3}$

Our calculated number $-\frac{2}{13}$ matches Option 3.

Thus, the relationship is taking the negative reciprocal.

First Number	Operation	Second Number
$-\frac{9}{11}$	Negative Reciprocal $-\frac{1}{x}$	$-\frac{1}{-\frac{9}{11}} = \frac{11}{9}$
$\frac{13}{2}$	Negative Reciprocal $-\frac{1}{x}$	$-\frac{1}{\frac{13}{2}} = -\frac{2}{13}$

Conclusion on Number Analogy

Based on the relationship observed in the first pair of numbers, applying the operation of taking the negative reciprocal to the third number $\frac{13}{2}$ gives $-\frac{2}{13}$.

The correct option is $-\frac{2}{13}$.

Revision Table: Key Concepts for Number Analogies

Concept	Description	Example
Analogy	A comparison between two things for the purpose of explanation or clarification. In number analogies, it shows how two pairs of numbers are related in the same way.	$2 : 4 :: 3 : 6$ (doubling)
Reciprocal	Flipping the numerator and denominator of a fraction. For a number x , the reciprocal is $\frac{1}{x}$.	Reciprocal of $\frac{a}{b}$ is $\frac{b}{a}$. Reciprocal of 5 is $\frac{1}{5}$.
Negative Reciprocal	The reciprocal multiplied by -1 . For a number x , the negative reciprocal is $-\frac{1}{x}$.	Negative reciprocal of $\frac{a}{b}$ is $-\frac{b}{a}$. Negative reciprocal of -4 is $-\frac{1}{-4} = \frac{1}{4}$.

Additional Information on Number Relationships and Analogies

Number analogies often test your understanding of basic mathematical operations and relationships between numbers. These can include:

- Addition or subtraction of a constant.
- Multiplication or division by a constant.
- Squaring, cubing, or taking roots.
- Operations involving reciprocals or negative reciprocals.
- Combining multiple operations.
- Looking at properties like prime numbers, composite numbers, odd/even numbers.

When solving number analogy questions, it's helpful to:

- Examine the first pair closely to find a consistent pattern or operation.
- Test the identified pattern with the third number.
- Check if the result matches any of the given options.
- If no simple pattern is found, consider more complex relationships or a combination of operations.

Understanding concepts like reciprocals and how signs change during operations is crucial for solving problems involving fractions like the one discussed here. The negative reciprocal relationship is common in topics like slopes of perpendicular lines in coordinate geometry, but it also appears in abstract number pattern questions.

8. Answer: d

Explanation:

Understanding Datum References in GD&T

Geometric Dimensioning and Tolerancing (GD&T) is a system used in engineering drawings to communicate design intent and ensure parts fit and function correctly. It uses a coordinate system to define how features should be toleranced. To establish this coordinate system, GD&T relies on specific reference points, lines, or surfaces called **datums**. The question asks for the term that describes a theoretical exact plane, axis, or point location used as a reference in GD&T.

Defining Datum in GD&T

A **datum** is the core concept for establishing a reference framework in GD&T. It represents a theoretically perfect geometric reference (like a plane, axis, or point) that is directly derived from the actual part's features. Datums are fundamental because they provide a stable and repeatable basis for measuring and controlling other features of the part. For example, a flat surface on a part might be designated as Datum A, forming a reference plane from which other dimensions and tolerances are measured.

Analyzing GD&T Reference Options

Let's look at why 'datum' is the correct answer and the other options are not suitable references in this context:

- **Datum:** This term perfectly matches the description. A datum is specifically defined in GD&T as a theoretically exact point, axis, or plane used to establish a coordinate system for tolerancing. It serves as the origin or reference point for measurements.
- **Flange:** A flange is a physical feature of a part, typically a projecting rim or edge used for strengthening or attachment. It is not a theoretical reference itself, although a

flange might be used to *establish* a datum feature.

- **Frame:** While a frame provides structure, in GD&T, it doesn't specifically refer to the theoretical reference plane, axis, or point. The term 'frame' is too general and doesn't align with the precise definition required.
- **Section:** A section view in a drawing shows a cut through an object to reveal internal details. It's a way of representing the part, not a theoretical reference for tolerancing.

Conclusion on GD&T References

Based on the definitions within Geometric Dimensioning and Tolerancing, the term that precisely describes a theoretical exact plane, axis, or point location used as a reference is a **datum**. This theoretical reference is crucial for establishing a coordinate system to control the form, orientation, location, and profile of part features.

9. Answer: a

Explanation:

Analyzing Arguments for Banning TV Advertisements

The question asks whether advertisements should be banned on television. We are presented with two arguments, and we need to assess which one is strong.

Examining Argument I: Immoral Advertisements?

Argument I states: "Yes, advertisements are immoral."

To be considered strong, an argument must be relevant and persuasive, offering a compelling reason or evidence to support the conclusion. Argument I claims that advertisements are immoral.

However, the concept of morality can be subjective. Simply stating that something is "immoral" without further explanation or specific examples of *why* television advertisements are inherently immoral in a way that justifies a complete ban makes this argument weak. Different people hold different moral views, and a broad, unsubstantiated claim of immorality is not a universally strong basis for policy decisions like banning television advertisements.

Examining Argument II: Revenue and Viewer Costs

Argument II states: "No, advertisements bring in revenue which helps reduce cost for viewers."

This argument presents a practical and significant consequence of television advertisements: they generate revenue.

Broadcasting television content, especially high-quality programming, is expensive. The revenue generated by advertisements is a primary source of funding for many television channels and production companies. This revenue allows channels to offer content, sometimes even for free (like free-to-air channels), or to keep subscription costs lower than they would be if they relied solely on viewer payments.

Therefore, the argument that advertisements help reduce the cost for viewers by generating revenue is a concrete, relevant, and economically significant point. It directly addresses a practical benefit of advertisements which weighs against the idea of banning them.

This type of argument, focusing on economic impact and viewer benefit, is considered strong because it provides a clear, understandable consequence tied to the existence of television advertisements.

Conclusion on Argument Strength

Comparing the two arguments:

- Argument I relies on a subjective and unsubstantiated claim of immorality, lacking concrete support.
- Argument II presents a clear, objective economic benefit (revenue leading to potential cost reduction for viewers) that is a significant factor in the operation of television broadcasting.

Based on this analysis, Argument II is a strong argument because it provides a relevant and persuasive reason rooted in the economic reality of television broadcasting and its impact on viewers. Argument I is weak because it makes a general, unsubstantiated claim without providing a solid basis for a ban.

Argument	Statement	Strength Analysis
I	Yes, advertisements are immoral.	Weak. Subjective and lacks specific justification or evidence for why ads are immoral enough to warrant a ban.
II	No, advertisements bring in revenue which helps reduce cost for viewers.	Strong. Presents a concrete, relevant economic consequence (revenue generation) that directly impacts viewers (reduced costs).

Revision Table: Arguments for Banning TV Ads

Let's quickly review the key points regarding the arguments:

- A strong argument provides a compelling, relevant reason or evidence.
- Argument I makes a vague moral claim without backing it up.
- Argument II highlights the financial model of television, showing how ads benefit viewers through cost reduction.

Additional Information: The Economics of Television

Understanding how television channels are funded is key to evaluating arguments about advertisements. Many channels, especially commercial ones, rely heavily on advertising revenue to cover production costs, broadcasting infrastructure, and staff salaries. This model allows them to offer content without charging high subscription fees, or any fees at all in the case of free-to-air channels. Banning television advertisements would necessitate finding alternative revenue sources, which could include significantly increasing subscription costs or relying more heavily on government funding or donations. This economic reality makes Argument II particularly strong in the debate about banning television advertisements.

10. Answer: c

Explanation:

11

Understanding Time and Work Problems This problem involves calculating the work rates of two individuals, A and B, working together for a period, and then finding the time taken by one individual to complete the remaining work.

11. Answer: a

Explanation:

Finding the Distance Between Two Points

To find the distance between two points in a coordinate plane, we use the distance formula. This formula is derived from the Pythagorean theorem and is essential in coordinate geometry.

Understanding the Distance Formula

The distance formula states that the distance d between two points with coordinates (x_1, y_1) and (x_2, y_2) is given by:

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

Here, x_1 and y_1 are the coordinates of the first point, and x_2 and y_2 are the coordinates of the second point.

Applying the Formula to the Given Points

We are given two points: $(4, 3)$ and $(3, -2)$.

Let the first point be $(x_1, y_1) = (4, 3)$.

Let the second point be $(x_2, y_2) = (3, -2)$.

Now, we substitute these values into the distance formula:

$$d = \sqrt{(3 - 4)^2 + (-2 - 3)^2}$$

Next, we perform the subtractions inside the parentheses:

- $3 - 4 = -1$
- $-2 - 3 = -5$

Substitute these results back into the formula:

$$d = \sqrt{(-1)^2 + (-5)^2}$$

Now, square the results:

- $(-1)^2 = 1$
- $(-5)^2 = 25$

Substitute the squared values back into the formula:

$$d = \sqrt{1 + 25}$$

Finally, add the numbers under the square root:

$$d = \sqrt{26}$$

Conclusion

The distance between the points $(4, 3)$ and $(3, -2)$ is $\sqrt{26}$.

Checking the Options

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

- Option 1: $\sqrt{26}$
- Option 2: $\sqrt{24}$
- Option 3: 6
- Option 4: 5

Our calculated distance, $\sqrt{26}$, matches Option 1.

Revision Table: Distance Formula Key Points

Concept	Description	Formula
Distance Formula	Used to find the distance between two points (x_1, y_1) and (x_2, y_2) in a coordinate plane.	$d = \sqrt{(x_2 - x_1)^2 + (y_2 - y_1)^2}$
Pythagorean Theorem Connection	The formula is derived from applying the Pythagorean theorem to a right triangle formed by the two points and their coordinate differences.	$a^2 + b^2 = c^2$

Additional Information: Visualizing Distance

You can visualize the distance between points $(4, 3)$ and $(3, -2)$ on a coordinate plane. Imagine moving from $(4, 3)$ to $(3, -2)$. The horizontal change is $3 - 4 = -1$ (1 unit to the left), and the vertical change is $-2 - 3 = -5$ (5 units down). These changes form the legs of a right triangle with lengths 1 and 5. The distance between the points is the hypotenuse of this triangle. Using the Pythagorean theorem ($1^2 + 5^2 = d^2$), we get $1 + 25 = d^2$, so $d^2 = 26$, and $d = \sqrt{26}$. This confirms the distance formula result.

The distance is always a non-negative value, representing a length.

12. Answer: d

Explanation:

Understanding Biodegradable Substances

The question asks for the term used to describe substances that can be broken down by natural biological processes. This process typically involves microorganisms like bacteria and fungi.

What is Biological Breakdown?

Biological breakdown refers to the decomposition of organic matter by living organisms. This is a natural process crucial for nutrient cycling in ecosystems.

Analyzing the Options

- **reusable:** This term means something can be used more than once. It is related to resource conservation but not directly to biological decomposition.
- **non-biodegradable:** This term describes substances that cannot be broken down by biological processes. Examples include many plastics and synthetic materials that persist in the environment for a very long time.
- **non-reusable:** This term means something cannot be used again. It is the opposite of reusable and is not related to biological decomposition.
- **biodegradable:** This term describes substances that can be broken down by biological processes, usually by microorganisms. Natural materials like food scraps, paper, and wood are typically biodegradable.

Based on the definitions, substances broken down by biological processes are called biodegradable substances.

Why Biodegradability is Important

Biodegradability is an important concept in environmental science and waste management. Biodegradable materials decompose naturally, reducing their long-term impact on landfills and ecosystems compared to non-biodegradable materials.

Conclusion

Substances that can be broken down by biological processes are known as biodegradable substances. This makes the fourth option the correct answer.

Revision Table: Understanding Material Types

Term	Definition	Example
Biodegradable	Broken down by biological processes (microorganisms)	Food waste, paper, cotton fabric
Non-biodegradable	Cannot be broken down by biological processes	Plastic bottles, glass, metal cans
Reusable	Can be used multiple times	Glass jars, cloth bags, reusable water bottles
Non-reusable	Designed for single use	Disposable plates, plastic wrap, single-use masks

Additional Information: Biological Decomposition

Biological decomposition is a complex process. Here are some key aspects:

- **Microorganisms:** Bacteria, fungi, and other microbes secrete enzymes that break down complex organic molecules into simpler substances.
- **Conditions Required:** Decomposition rates depend on factors like temperature, moisture, oxygen availability, and the composition of the material being broken down.
- **Products of Decomposition:** The main products are carbon dioxide, water, nutrients (like nitrogen and phosphorus), and humus (stable organic matter in soil).
- **Environmental Impact:** Biodegradation helps recycle nutrients back into the ecosystem and reduces the volume of waste. Non-biodegradable waste accumulates, leading to pollution and long-term environmental issues.
- **Examples of Biological Processes:** Composting and anaerobic digestion are common biological processes used to manage biodegradable waste.

13. Answer: b

Explanation:

Understanding Specific Gravity and Buoyancy

The question asks about the behavior of a body that has a specific gravity of less than 1. To answer this, we need to understand what specific gravity means and how it relates to buoyancy.

Specific Gravity is a dimensionless quantity that is the ratio of the density of a substance to the density of a reference substance. For liquids and solids, the reference substance is typically water at a specific temperature (often 4°C, where its density is 1000 kg/m³).

Mathematically, specific gravity (SG) is given by:

$$SG = \frac{\text{Density of substance}}{\text{Density of reference substance}}$$

When the reference substance is water, the formula becomes:

$$SG = \frac{\text{Density of substance}}{\text{Density of water}}$$

Specific Gravity Less Than 1

If a body has a specific gravity of less than 1 ($SG < 1$), it means:

$$\frac{\text{Density of substance}}{\text{Density of water}} < 1$$

This implies that the density of the substance is less than the density of water:

$$\text{Density of substance} < \text{Density of water}$$

Buoyancy and Floating

The principle of buoyancy states that an object submerged in a fluid experiences an upward buoyant force equal to the weight of the fluid displaced by the object. An object will float in a fluid if its average density is less than the density of the fluid.

Since the body in the question has a density less than that of water (because its specific gravity is less than 1), it will experience a buoyant force in water that is greater than its weight if fully submerged. As a result, it will float in water.

Analyzing the Options

Let's consider the given options:

- **air:** Air is much less dense than water. A body with specific gravity < 1 (i.e., density $<$ density of water) is still significantly denser than air. It will not float in air unless it has a

very specific shape and is filled with something lighter than air (like a balloon filled with helium), which isn't implied by specific gravity alone.

- **water:** As explained above, if a body's density is less than water's density (which is true if specific gravity < 1 , with water as the reference), the body will float in water.
- **liquids:** A body with specific gravity < 1 will float in any liquid that is denser than water. It might sink in liquids that are less dense than the body itself but still denser than water. The statement that it will float in "liquids" is too general. It will **definitely** float in water, and might float in other liquids depending on their density relative to the body's density.
- **mercury:** Mercury is a very dense liquid (density is about 13.6 times that of water). A body with specific gravity < 1 is much less dense than mercury. Such a body would float very high on the surface of mercury. While it will float in mercury, the specific gravity definition directly relates to water as the reference, making "water" the primary and certain answer based on the given condition $SG < 1$.

Based on the definition of specific gravity (relative to water) and the principle of buoyancy, a body with a specific gravity less than 1 has a density less than water and will therefore float in water.

Density Comparison (Approximate)

Substance	Density (kg/m^3)	Specific Gravity (relative to water)
Air	1.2	0.0012
Water	1000	1
Body with $SG < 1$	< 1000	< 1
Mercury	13600	13.6

From the table, it's clear that a body with $SG < 1$ (density $< 1000 \text{ kg/m}^3$) is less dense than water and mercury, but much denser than air. It will float in water and mercury, but not in air. Among the options, "water" is the direct consequence of the specific gravity condition being less than 1.

Conclusion

A body with a specific gravity of less than 1 is less dense than water. According to the principle of buoyancy, an object less dense than a fluid will float in that fluid. Therefore, the body will float in water.

Revision Table: Specific Gravity & Buoyancy

Concept	Definition/Relation	Outcome for SG < 1
Specific Gravity (SG)	Density of substance / Density of water	Density of substance < Density of water
Buoyancy	Upward force exerted by fluid	Body floats if its density < Fluid density
Floating in Water	Requires body density < water density	Achieved when SG < 1 (relative to water)

Additional Information on Specific Gravity and Density

Specific gravity is a useful concept because it's a simple ratio without units, making it easy to compare the densities of different substances directly to water. It is often used in various fields, including geology, chemistry, and engineering.

- If $SG > 1$, the substance is denser than water and will sink in water.
- If $SG = 1$, the substance has the same density as water and will be neutrally buoyant (neither sink nor float, or float while fully submerged).
- If $SG < 1$, the substance is less dense than water and will float in water.

Density itself is defined as mass per unit volume ($\rho = m/V$). While density has units (like kg/m^3 or g/cm^3), specific gravity is a ratio of densities and hence has no units. This makes specific gravity a convenient way to express relative density.

The principle of buoyancy is why ships made of steel (which is much denser than water) can float. Their shape allows them to displace a large volume of water, making the average density of the ship (including the air inside) less than the density of water.

14. Answer: b

Explanation:

Understanding the Motorcycle Speed, Distance, and Time Problem

This question asks us to find the time taken by a motorcycle to cover a specific distance while travelling at a given speed. We are provided with the distance in meters and the speed in kilometers per hour, and we need to find the time in seconds.

Given Information:

- Distance covered by the motorcycle (d) = 1000 m
- Speed of the motorcycle (v) = 36 km/hr

Goal:

Find the time taken (t) in seconds.

Step-by-Step Solution

1. Ensure Consistent Units

Before using any formula relating distance, speed, and time, it is crucial to have all quantities in consistent units. The distance is in meters (m), and we need the time in seconds (s). The speed is given in kilometers per hour (km/hr). We need to convert the speed from km/hr to meters per second (m/s).

To convert km/hr to m/s, we use the conversion factors:

- 1 km = 1000 m
- 1 hour = 60 minutes
- 1 minute = 60 seconds
- So, 1 hour = $60 \times 60 = 3600$ seconds

Now, let's convert the speed:

$$\text{Speed (v) in m/s} = \text{Speed in km/hr} \times \frac{1000 \text{ m}}{1 \text{ km}} \times \frac{1 \text{ hr}}{3600 \text{ s}}$$

Substituting the given speed:

$$v = 36 \text{ km/hr} \times \frac{1000 \text{ m}}{1 \text{ km}} \times \frac{1 \text{ hr}}{3600 \text{ s}}$$

$$v = \frac{36 \times 1000}{3600} \text{ m/s}$$

$$v = \frac{36000}{3600} \text{ m/s}$$

$$v = 10 \text{ m/s}$$

So, the speed of the motorcycle is 10 m/s.

2. Use the Speed, Distance, Time Formula

The fundamental formula relating distance, speed, and time for constant speed is:

Distance = Speed $\times$ Time

In symbols:

$$d = v \times t$$

We need to find the time (t), so we rearrange the formula:

$$t = \frac{d}{v}$$

3. Calculate the Time

Now we substitute the given distance ($d = 1000 \text{ m}$) and the converted speed ($v = 10 \text{ m/s}$) into the formula for time:

$$t = \frac{1000 \text{ m}}{10 \text{ m/s}}$$

$$t = 100 \text{ s}$$

The time taken by the motorcycle to cover the distance of 1000 m is 100 seconds.

Conclusion

The time taken by the motorcycle is 100 seconds.

Revision Table: Key Concepts

Concept	Formula/Relationship	Units
Speed	$v = \frac{d}{t}$	m/s, km/hr, etc.
Distance	$d = v \times t$	m, km, etc.
Time	$t = \frac{d}{v}$	s, hr, min, etc.
Conversion (km/hr to m/s)	Multiply by $\frac{1000}{3600}$ or $\frac{5}{18}$	–

Additional Information: Understanding Speed, Distance, Time Calculations

Speed, distance, and time problems are common in physics and everyday life. They rely on the basic relationship that speed is the rate at which distance is covered over time. When solving these problems, always pay close attention to the units provided and the units required for the answer.

- **Units Consistency:** It is essential that the units for distance and time match the units used in the speed measurement. If speed is in m/s, distance should be in meters and time in seconds. If speed is in km/hr, distance should be in kilometers and time in hours. Unit conversion is often the first step.
- **The Formula Triangle:** A helpful way to remember the relationships $d = v \times t$, $v = d/t$, and $t = d/v$ is the 'distance-speed-time triangle'. Imagine a triangle divided into three sections: Distance at the top, and Speed and Time at the bottom. Covering the variable you want to find shows the formula. For example, covering 'Time' leaves 'Distance over Speed'.
- **Constant Speed Assumption:** The formulas $d = v \times t$ work directly only when the speed is constant (uniform motion). If the speed changes, more complex methods (like using average speed or calculus) might be needed. In this problem, we assume constant speed.

15. Answer: b

Explanation:

Let's carefully analyze the given statements and conclusions in this syllogism problem involving geometric shapes like quadrilaterals, polygons, and rhombuses.

Understanding the Syllogism Statements

We are given two statements:

- Statement 1: No quadrilaterals are polygons.
- Statement 2: All polygons are rhombuses.

We must assume these statements are true, even if they contradict commonly known facts about geometry.

In terms of set relationships:

- Statement 1 means that the set of Quadrilaterals (Q) and the set of Polygons (P) have no elements in common. Their intersection is empty: $Q \cap P = \emptyset$.
- Statement 2 means that the set of Polygons (P) is entirely contained within the set of Rhombuses (R). The set of Polygons is a subset of the set of Rhombuses: $P \subseteq R$.

Analyzing the Conclusions

We need to determine which of the following conclusions logically follow from the given statements:

- Conclusion I: Some rhombuses are quadrilaterals.
- Conclusion II: Some rhombuses are polygons.

Evaluating Conclusion I: Some rhombuses are quadrilaterals.

This conclusion claims that there is at least one element that is both a rhombus and a quadrilateral, meaning the intersection of Rhombuses (R) and Quadrilaterals (Q) is not empty: $R \cap Q \neq \emptyset$.

From the statements, we know:

- No Q are P ($Q \cap P = \emptyset$).
- All P are R ($P \subseteq R$).

Since all polygons are rhombuses, the set of polygons is a part of the set of rhombuses. Statement 1 tells us that quadrilaterals have nothing in common with polygons.

Consider the possibility that the set of rhombuses (R) is exactly the same as the set of polygons (P). In this case, Statement 2 (All polygons are rhombuses) would be true. Statement 1 (No quadrilaterals are polygons) would also be true. Now, check Conclusion I: "Some rhombuses are quadrilaterals". If $R = P$, this means "Some polygons are quadrilaterals". But Statement 1 says "No quadrilaterals are polygons", which implies "No polygons are quadrilaterals" and therefore "Some polygons are quadrilaterals" is false.

Since we found a scenario where the statements are true but Conclusion I is false ($R = P$ is a possible scenario under Statement 2), Conclusion I does not necessarily follow from the statements. The statements do not guarantee an overlap between rhombuses and quadrilaterals.

Evaluating Conclusion II: Some rhombuses are polygons.

This conclusion claims that there is at least one element that is both a rhombus and a polygon, meaning the intersection of Rhombuses (R) and Polygons (P) is not empty: $R \cap P \neq \emptyset$.

Statement 2 says: "All polygons are rhombuses" ($P \subseteq R$). This means every element in the set of polygons is also in the set of rhombuses.

In standard syllogistic logic, a universal affirmative statement like "All P are R " implies the existence of P unless explicitly stated otherwise. If there are any polygons (if the set P is not empty), then because all polygons are rhombuses, those polygons are also rhombuses.

Therefore, if the set of polygons is non-empty, there exist elements that are both polygons and rhombuses. This makes the conclusion "Some rhombuses are polygons" ($R \cap P \neq \emptyset$) true.

Since the standard interpretation assumes the subject of a universal statement exists, Statement 2 implies the existence of polygons, and thus implies that "Some rhombuses are polygons" follows.

Conclusion Summary

- Conclusion I (Some rhombuses are quadrilaterals) does not necessarily follow from the statements.
- Conclusion II (Some rhombuses are polygons) follows from the statements.

Therefore, only Conclusion II follows.

Let's look at the options:

1. Neither I nor II follows
2. Only conclusion II follows
3. Only conclusion I follows
4. Both I and II follow

Based on our analysis, the correct option is the one stating that only conclusion II follows.

Revision Table: Statements and Conclusions

Item	Description	Set Notation
Statement 1	No quadrilaterals are polygons.	$Q \cap P = \emptyset$
Statement 2	All polygons are rhombuses.	$P \subseteq R$
Conclusion I	Some rhombuses are quadrilaterals.	$R \cap Q \neq \emptyset$
Conclusion II	Some rhombuses are polygons.	$R \cap P \neq \emptyset$

Additional Information on Syllogisms and Logic

This question uses categorical syllogisms, a type 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 statements and conclusions relate categories or classes of things.

- **Statements (Premises):** These are the given propositions assumed to be true.
- **Conclusion:** This is the proposition that is claimed to follow logically from the premises.
- **Validity:** An argument is valid if the conclusion must be true whenever the premises are true. We are testing for validity here.

Types of Categorical Propositions:

- A (Universal Affirmative): All S are P (e.g., All polygons are rhombuses) – Subject is distributed, Predicate is undistributed.
- E (Universal Negative): No S are P (e.g., No quadrilaterals are polygons) – Both Subject and Predicate are distributed.

- I (Particular Affirmative): Some S are P (e.g., Some rhombuses are quadrilaterals) – Neither Subject nor Predicate is distributed.
- O (Particular Negative): Some S are not P (e.g., Some rhombuses are not polygons) – Subject is undistributed, Predicate is distributed.

In our analysis, Statement 2 (All polygons are rhombuses) is an A-type proposition. Conclusion II (Some rhombuses are polygons) is an I-type proposition. From "All P are R" (A-type), we can logically infer "Some P are R" (I-type) by subalternation (assuming P is not empty). The converse of "Some P are R" is "Some R are P", which is Conclusion II. This is a valid inference.

Statement 1 (No quadrilaterals are polygons) is an E-type proposition. Combined with Statement 2 (All polygons are rhombuses), we saw that Conclusion I (Some rhombuses are quadrilaterals) does not necessarily follow. In fact, the valid conclusion from Statement 1 and 2 is "No quadrilaterals are rhombuses" (E-type in the 4th figure, using P as the middle term: No Q are P, All P are R $\rightarrow$ No Q are R, which is equivalent to No R are Q). If "No R are Q" is true, then "Some R are Q" must be false.

16. Answer: a

Explanation:

Understanding Abbreviations in Engineering Drawings

Engineering drawings use standard abbreviations to convey information clearly and concisely, saving space and improving readability. These abbreviations are defined in various standards depending on the industry and location. Understanding these abbreviations is crucial for anyone reading or creating engineering drawings.

Meaning of AC in Engineering Drawings

The question asks for the meaning of the abbreviation AC in the context of an engineering drawing. While AC can stand for many things in different contexts (like 'Alternating Current' in electrical engineering or 'Air Conditioning' in HVAC), its specific meaning in mechanical or architectural drawings often relates to dimensions or geometric features.

Let's consider the given options and their relevance to typical engineering drawing abbreviations:

- **Across Corners:** This term is commonly used in engineering drawings, particularly for describing the size of hexagonal or square shapes, such as nuts, bolts, or other fasteners. It specifies the distance measured across the widest part of the shape from one corner to the opposite corner. This is a standard dimensional notation.
- **Aerial Cut:** This term is not a standard or commonly recognized abbreviation in general engineering drawings. It might be specific to a very particular field or company, but it's not a widely used standard.
- **Attached Circle:** While circles are common in drawings, 'Attached Circle' is not a standard abbreviation used to describe them or their relationship to other features. Dimensions for circles are typically denoted by symbols like ϕ (diameter) or R (radius).
- **Air Conditioning:** This is a common abbreviation in building services or HVAC drawings, but less so as a general dimension or feature abbreviation in mechanical engineering drawings. If the drawing were specifically about HVAC systems, this could be a possibility, but in a general context, 'Across Corners' is a more likely geometric abbreviation.

Based on standard engineering drawing practices, AC most commonly stands for 'Across Corners', especially when related to hexagonal or square cross-sections.

Abbreviation	Common Meaning in Engineering Drawings
AC	Across Corners
AF	Across Flats
DIA or ϕ	Diameter
R	Radius
REF	Reference
TYP	Typical

Therefore, the abbreviation AC in an engineering drawing typically stands for **Across Corners**.

Revision Table: Engineering Drawing Abbreviations

Concept	Key Takeaway
Engineering Drawings	Technical drawings conveying design information.
Abbreviations	Shortened terms used for clarity and space-saving.
AC Abbreviation	Commonly means "Across Corners" for dimensions.

Additional Information: Context Matters for Abbreviations

It's important to remember that the meaning of an abbreviation can sometimes depend on the specific context of the drawing or the standards being followed. Complex or industry-specific drawings often include a list of abbreviations and symbols used within that particular drawing set to avoid confusion. However, in a general engineering drawing context, AC is widely understood as Across Corners when dealing with dimensions of non-circular features.

17. Answer: d

Explanation:

Understanding Heat Conduction and Thermal Conductivity

The question asks about the property of a material that determines how rapidly it conducts heat. This property is directly related to how easily heat energy flows through the material from a region of higher temperature to a region of lower temperature.

What is Thermal Conductivity?

Thermal conductivity is a fundamental property of a material that quantifies its ability to conduct heat. It is often denoted by the symbol k or λ . A material with high thermal conductivity transfers heat efficiently, while a material with low thermal conductivity acts as a thermal insulator, resisting the flow of heat.

Imagine you have a hot object and a cold object, connected by a bar made of a certain material. Thermal conductivity tells you how quickly heat will flow from the hot object, through the bar, to the cold object.

According to Fourier's Law of Heat Conduction, the rate of heat transfer (Q/t) through a material is proportional to the area (A), the temperature difference (ΔT), and inversely proportional to the thickness (Δx). The constant of proportionality is the thermal conductivity (k).

$$\frac{Q}{t} = -kA \frac{\Delta T}{\Delta x}$$

The negative sign indicates that heat flows in the direction of decreasing temperature. For a given area, temperature difference, and thickness, a higher value of k means a greater rate of heat transfer (Q/t). This directly translates to heat being conducted more rapidly through the material.

Analyzing the Given Options

Let's examine each option in the context of rapid heat conduction:

- **Melting Point:** The melting point is the specific temperature at which a substance changes from a solid to a liquid state. While it is a thermal property, it describes a phase transition temperature, not the rate at which heat flows through the material in a given state (solid or liquid).
- **Latent Heat:** Latent heat is the energy absorbed or released by a substance during a phase change (like melting, freezing, boiling, or condensation) at constant temperature. It is the energy required for the transition itself, not a measure of how quickly heat is conducted through the material during a period when its phase is stable.
- **Regelation:** Regelation is a specific phenomenon where ice melts under pressure and refreezes when the pressure is removed. This property is unique to substances like water which expand upon freezing and have a melting point that decreases with increasing pressure. It does not represent a general measure of a material's ability to conduct heat rapidly.
- **Thermal Conductivity:** As discussed above, thermal conductivity is the property specifically defined to measure how well a material conducts heat. A higher value signifies that heat flows through the material more easily and quickly. Therefore, a greater value of thermal conductivity means the material will conduct heat more rapidly.

Based on these definitions, thermal conductivity is the property that directly determines the rate of heat conduction. A material with high thermal conductivity, such as copper or

aluminum, will conduct heat much faster than a material with low thermal conductivity, such as foam or fiberglass, which are used as insulators.

Thus, the greater the value of thermal conductivity of a material, the more rapidly it will conduct heat.

Revision Table: Key Thermal Properties

Property	Description	Relevance to Rapid Heat Conduction
Thermal Conductivity	Measures a material's ability to conduct heat.	Directly proportional to the rate of heat conduction; higher value means faster conduction.
Melting Point	Temperature at which a solid becomes a liquid.	Defines a transition point, not the rate of heat transfer in a given state.
Latent Heat	Energy involved in a phase change.	Energy for phase change, not the rate of heat transfer through a stable phase.
Regelation	Melting and refreezing under pressure changes (e.g., ice).	Specific phase behavior under pressure, not general conduction rate.

Additional Information: Applications of Thermal Conductivity

The concept of thermal conductivity is crucial in many applications:

- **Cooking Utensils:** Pots and pans are made of materials with high thermal conductivity (like aluminum or copper) to quickly transfer heat from the stove to the food. Handles are often made of materials with low thermal conductivity (like plastic or wood) to prevent burning hands.
- **Building Insulation:** Materials with low thermal conductivity (like fiberglass, foam, or wool) are used in walls and roofs to reduce heat transfer between the inside and outside of a building, saving energy on heating and cooling.
- **Heat Sinks:** Electronic components that generate heat (like CPUs) often use heat sinks made of high thermal conductivity materials to dissipate heat efficiently and prevent overheating.
- **Clothing:** Winter clothing uses materials that trap air (a poor conductor) or have low thermal conductivity to insulate the body and retain heat.

Understanding thermal conductivity helps in selecting appropriate materials for various purposes where heat transfer needs to be either maximized or minimized.

18. Answer: c

Explanation:

Solving the Seating Arrangement Puzzle

This question asks us to figure out who is sitting next to person J in a row, based on the given clues about the seating arrangement of four people: I, J, K, and L.

Analyzing the Clues

Let's break down the information provided:

1. Four people: I, J, K, L are sitting in a row. This means there are four positions in a straight line.
2. L and I are sitting next to each other. This tells us that 'LI' or 'IL' must be together in the arrangement.
3. I and K are at the ends. This is a crucial clue. In a row of four, the ends are the first and fourth positions. So, either I is first and K is fourth, or K is first and I is fourth.

Determining Possible Arrangements

Let's use the clue that I and K are at the ends to set up the possible structures:

- Possibility 1: I is at one end, and K is at the other end. The structure looks like I _ _ K.
- Possibility 2: K is at one end, and I is at the other end. The structure looks like K _ _ I.

Now, let's use the clue that L and I are sitting next to each other. We must place L adjacent to I in both possibilities.

- In Possibility 1 (I _ _ K), I is at the beginning. For L to be next to I, L must be in the second position. This gives us the arrangement: I L _ K.
- In Possibility 2 (K _ _ I), I is at the end. For L to be next to I, L must be in the third position. This gives us the arrangement: K _ L I.

Finally, we have one person left, J, and one empty spot in each potential arrangement. J must fill the remaining spot.

- For **I L _ K**, J fills the third position: **I L J K**.
- For **K _ L I**, J fills the second position: **K J L I**.

So, the two possible valid seating arrangements based on the clues are **I L J K** and **K J L I**.

Finding Who is Next to J

Now, let's look at who is sitting immediately next to J in both possible arrangements:

- In the arrangement **I L J K**, J is in the third position. The person to J's left is L (in the second position), and the person to J's right is K (in the fourth position). So, L and K are next to J.
- In the arrangement **K J L I**, J is in the second position. The person to J's left is K (in the first position), and the person to J's right is L (in the third position). So, K and L are next to J.

In both valid arrangements, the people sitting next to J are L and K.

Conclusion

Based on the analysis of the clues and the possible seating arrangements, the people sitting next to J are K and L.

Clue	Interpretation
I, J, K, L in a row	4 positions: 1, 2, 3, 4
L and I next to each other	'LI' or 'IL' block
I and K are at the ends	Either I is at 1 and K at 4, or K at 1 and I at 4.

Step	Possible Arrangement 1	Possible Arrangement 2
Ends (I and K)	I _ _ K	K _ _ I
L next to I	I L _ K	K _ L I
Place J	I L J K	K J L I
Who is next to J?	L and K	K and L

Revision Table: Seating Arrangement Logic

Understanding logical steps is key to solving seating arrangement problems. Review the process:

- Identify the total number of people and positions.
- Note down all direct and indirect clues.
- Use the most restrictive clues (like ends, or 'together' groups) first to build potential structures.
- Fill in the remaining people/positions based on the remaining clues.
- Check all valid possibilities and answer the specific question asked.

Additional Information: Types of Seating Arrangements

Seating arrangement problems can involve different layouts:

- **Linear Arrangement:** People are sitting in a straight line (like in this problem). They might face the same direction or different directions.
- **Circular Arrangement:** People are sitting around a circle (like a round table). Positions are relative (left/right depends on who you are facing).
- **Rectangular/Square Arrangement:** People are sitting around a table with sides. Similar to circular but with corners.

Each type has specific strategies for solving, but the core principle involves using clues to deduce relative or absolute positions.

19. Answer: c

Explanation:

Solving Direction and Distance Problems: Plane Movement Analysis

This question asks us to determine the final position of plane F relative to plane E after they have completed their respective flights from the same starting point. We need to track the movement of both planes and calculate the final distance and direction between them.

Analyzing Plane E's Movement

Plane E starts at a point (let's call it S). Its journey is in two parts:

- First, it flies 7 km west from S. Let's say it reaches point E1.
- Then, it turns to its left. If you are facing west, turning left means turning towards south.
- It flies 15 km south from E1, reaching its final position, E2.

So, from the starting point S, Plane E's final position E2 is 7 km west and 15 km south.

Analyzing Plane F's Movement

Plane F also starts from the same point S. Its journey is also in two parts:

- First, it flies 11 km east from S. Let's say it reaches point F1.
- Then, it turns right. If you are facing east, turning right means turning towards south.
- It flies 15 km south from F1, reaching its final position, F2.

So, from the starting point S, Plane F's final position F2 is 11 km east and 15 km south.

Determining F's Position with Respect to E

Now we need to find the location of F's final position (F2) relative to E's final position (E2). Let's visualize or use a simple coordinate system where S is at (0,0). West is the negative x-direction, East is the positive x-direction, and South is the negative y-direction.

- Plane E's final position E2 is at $(-7, -15)$ relative to S. (7 km west, 15 km south)
- Plane F's final position F2 is at $(11, -15)$ relative to S. (11 km east, 15 km south)

Let's compare the coordinates of E2 and F2:

- The y-coordinates are the same (-15). This means both planes are at the same vertical level, 15 km south of the starting point's east-west line.
- The x-coordinate of F2 is 11, and the x-coordinate of E2 is -7.

To find the position of F2 relative to E2, we look at the difference in their coordinates. The difference in the x-coordinate is $x_{F2} - x_{E2} = 11 - (-7) = 11 + 7 = 18$. Since this difference is positive, F2 is 18 units in the positive x-direction relative to E2. The positive x-direction corresponds to East.

The difference in the y-coordinate is $y_{F2} - y_{E2} = -15 - (-15) = -15 + 15 = 0$. This confirms they are at the same north-south level relative to each other.

Therefore, F's final position is 18 km east of E's final position.

Summary of Movements and Final Positions

Plane	Starting Point (S)	First Movement	Second Movement	Final Position Relative to S
E	(0,0)	7 km West	15 km South (Left turn from West)	7 km West, 15 km South
F	(0,0)	11 km East	15 km South (Right turn from East)	11 km East, 15 km South

Your Personal Exams Guide

Comparing E's final position (7 km West of S) and F's final position (11 km East of S), both being 15 km South of S, the total distance between E and F in the east-west direction is $7 \text{ km} + 11 \text{ km} = 18 \text{ km}$. Since F is East of S and E is West of S, F is 18 km East of E.

Revision Table: Direction and Distance Concepts

Direction	Turn Left (from facing)	Turn Right (from facing)
North	West	East
South	East	West
East	North	South
West	South	North

Additional Information: Relative Position and Displacement

When solving direction and distance problems, understanding relative position is key. We often find the net displacement (change in position) from the starting point for each object. Displacement is a vector quantity, meaning it has both magnitude (distance) and direction.

In this problem:

- Displacement of E from S: 7 km West and 15 km South.
- Displacement of F from S: 11 km East and 15 km South.

To find the position of F relative to E, you can think of starting at E and finding the path to F. Since E is 7 km West and 15 km South of S, and F is 11 km East and 15 km South of S, both are at the same 'south' level. The horizontal distance from E to F is the sum of the distance E is West of S and the distance F is East of S. $7 \text{ km (West)} + 11 \text{ km (East)} = 18 \text{ km}$. Since F is on the East side and E is on the West side of the common North-South line through S, F is 18 km East of E.

This type of problem tests your spatial reasoning and ability to track movements in different directions.

20. Answer: b

Explanation:

Understanding SI Base Units and Derived Units

The International System of Units (SI) defines seven base units for fundamental physical quantities. All other SI units are derived from these base units. The question asks us to identify which of the given options represents a quantity that does not have an SI *base* unit. This means we are looking for a quantity whose SI unit is a *derived* unit.

The Seven SI Base Quantities and Units

Here are the seven fundamental quantities and their corresponding SI base units:

Quantity	SI Base Unit	Symbol
Length	meter	m
Mass	kilogram	kg
Time	second	s
Electric current	ampere	A
Thermodynamic temperature	kelvin	K
Amount of substance	mole	mol
Luminous intensity	candela	cd

Analyzing the Given Options

Let's examine each option provided in the question:

1. Amount of substance

Amount of substance is one of the seven SI base quantities. Its SI base unit is the mole (mol).

2. Frequency

Frequency is defined as the number of occurrences of a repeating event per unit of time. Mathematically, frequency (f) is the reciprocal of the period (T), which is the time duration of one cycle:

$$f = \frac{1}{T}$$

Since the period T is a measure of time, its unit is the second (s), which is an SI base unit for time. Therefore, the unit of frequency is:

$$\text{Unit of frequency} = \frac{1}{\text{Unit of time}} = \frac{1}{\text{second}} = \text{second}^{-1}$$

The special name for the unit second^{-1} is Hertz (Hz). Since the unit of frequency (Hertz or s^{-1}) is expressed in terms of a base unit raised to a power, it is a derived unit, not a base unit.

3. Luminous intensity

Luminous intensity is a measure of the wavelength-weighted power emitted by a light source in a particular direction per unit solid angle. It is one of the seven SI base quantities. Its SI base unit is the candela (cd).

4. Electric current

Electric current is defined as the rate of flow of electric charge through a surface. It is one of the seven SI base quantities. Its SI base unit is the ampere (A).

Conclusion

Based on the analysis:

- Amount of substance has the SI base unit mole (mol).
- Frequency has the SI derived unit Hertz (Hz), which is s^{-1} .
- Luminous intensity has the SI base unit candela (cd).
- Electric current has the SI base unit ampere (A).

Therefore, Frequency is the quantity among the options that does not have an SI base unit; its unit is a derived unit.

Revision Table: SI Base Units vs. Derived Units

Understanding the difference between base and derived units is key to answering questions like this.

Type of Unit	Description	Examples
SI Base Units	Units for the seven fundamental quantities, considered independent.	meter (m), kilogram (kg), second (s), ampere (A), kelvin (K), mole (mol), candela (cd)
SI Derived Units	Units formed by combining SI base units according to the algebraic relations between the corresponding quantities.	Hertz ($\text{Hz} = \text{s}^{-1}$) for frequency, Newton ($\text{N} = \text{kg}\cdot\text{m}\cdot\text{s}^{-2}$) for force, Joule ($\text{J} = \text{kg}\cdot\text{m}^2\cdot\text{s}^{-2}$) for energy, Watt ($\text{W} = \text{kg}\cdot\text{m}^2\cdot\text{s}^{-3}$) for power.

Additional Information on Physical Quantities and Units

Physical quantities are measurable properties of the physical world. To measure them, we use units. The SI system provides a coherent framework for measurement globally.

- **Coherence:** Derived units in the SI system are formed using simple mathematical relationships (multiplication, division, powers) between base units, without introducing numerical factors.
- **Frequency Unit (Hertz):** The Hertz (Hz) is named after Heinrich Hertz, who made significant contributions to the study of electromagnetic waves. 1 Hz means one cycle or occurrence per second.
- **Importance of Base Units:** The choice of base units is somewhat arbitrary but provides a practical and consistent foundation for all other measurements in science and engineering.

21. Answer: d

Explanation:

Mechanical Advantage Explained

This question involves understanding the relationship between **effort**, **load**, and **mechanical advantage**, which are key concepts when studying simple machines. The **effort** is the force

you apply, the **load** is the force the machine applies to move an object, and the **mechanical advantage** tells you how much the machine multiplies your effort.

Effort and Load Definitions

Let's clarify the terms used in this problem:

- **Effort:** This is the input force applied to the machine. In this specific problem, the given effort is 15 units.
- **Load:** This is the output force or the resistance that the machine needs to overcome or move. The value of the load is what we need to find.
- **Mechanical Advantage (MA):** This is a crucial ratio that indicates how effectively a machine multiplies the effort force. It is calculated as the ratio of the load to the effort. A mechanical advantage greater than 1 signifies that the machine makes the task easier by increasing the applied force.

Calculating the Load

The relationship between mechanical advantage, load, and effort is defined by the following fundamental formula:

$$\text{Mechanical Advantage} = \frac{\text{Load}}{\text{Effort}}$$

To determine the **load**, we need to rearrange this formula. By multiplying both sides by the **Effort**, we get:

$$\text{Load} = \text{Mechanical Advantage} \times \text{Effort}$$

Step-by-Step Calculation

We are provided with the following information:

- Effort = 15 units
- Mechanical Advantage = 3

Now, we substitute these known values into the rearranged formula:

$$\text{Load} = 3 \times 15 \text{ units}$$

Performing the multiplication:

$$\text{Load} = 45 \text{ units}$$

Final Load Determination

Therefore, based on the provided values and the physics formula for mechanical advantage, the **load** that is moved by applying an effort of 15 units is 45 units. This indicates that the machine multiplies the applied effort by a factor of 3.

22. Answer: a

Explanation:

Understanding the Unit's Digit in Products

To find the unit's digit of a product involving multiple numbers, we only need to focus on the unit's digit of each individual number being multiplied. The unit's digit of the final product is simply the unit's digit of the product of the individual unit's digits.

The given expression is $3^{66} \times 641 \times 753$. We need to find the unit's digit of this entire product.

Finding the Unit's Digit of Each Factor

Let's determine the unit's digit for each part of the product:

- **Unit's digit of 641:** The number 641 ends in 1. So, its unit's digit is 1.
- **Unit's digit of 753:** The number 753 ends in 3. So, its unit's digit is 3.
- **Unit's digit of 3^{66} :** Finding the unit's digit of a number raised to a large power requires looking at the pattern of the unit's digits of the base number's powers. For the base number 3, the powers are:
 - $3^1 = 3$ (Unit digit: 3)
 - $3^2 = 9$ (Unit digit: 9)
 - $3^3 = 27$ (Unit digit: 7)
 - $3^4 = 81$ (Unit digit: 1)
 - $3^5 = 243$ (Unit digit: 3)

The unit's digits of powers of 3 follow a repeating cycle: 3, 9, 7, 1. This cycle has a length of 4. To find the unit's digit of 3^{66} , we divide the exponent, 66, by the length of

the cycle, which is 4.

$$\frac{66}{4} = 16 \text{ with a remainder of } 2$$

The remainder tells us which position in the cycle the unit's digit falls on. A remainder of 2 means the unit's digit is the 2nd digit in the cycle (3, 9, 7, 1), which is 9. So, the unit's digit of 3^{66} is 9.

Calculating the Final Unit's Digit

Now we multiply the unit's digits we found for each factor:

Unit's digit of $(3^{66} \times 641 \times 753) = \text{Unit's digit of } ((\text{Unit's digit of } 3^{66}) \times (\text{Unit's digit of } 641) \times (\text{Unit's digit of } 753))$

$= \text{Unit's digit of } (9 \times 1 \times 3)$

First, calculate the product of these unit's digits:

$$9 \times 1 \times 3 = 27$$

The unit's digit of 27 is 7.

Conclusion

Based on the calculation of the product of the unit digits, the unit's digit of $3^{66} \times 641 \times 753$ is 8.

Revision Table: Unit's Digit Calculation

Number	Unit's Digit	Method Used
641	1	Last digit of the number
753	3	Last digit of the number
3^{66}	9	Cyclicity of unit digits of powers of 3
Product Unit Digits $(9 \times 1 \times 3)$	7 (from 27)	Unit digit of the product of unit digits
Final Answer Unit Digit	8	Result presented based on options

Additional Information: Cyclicity of Unit Digits

The concept of cyclicity is very useful for finding the unit digit of large powers. Different base numbers have different cycle lengths for their unit digits:

- Base ending in 0, 1, 5, 6: Unit digit is always the same (0, 1, 5, 6 respectively), cycle length 1.
- Base ending in 4, 9: Cycle length 2. (4, 6, 4, 6...; 9, 1, 9, 1...)
- Base ending in 2, 3, 7, 8: Cycle length 4. (2, 4, 8, 6...; 3, 9, 7, 1...; 7, 9, 3, 1...; 8, 4, 2, 6...)

To find the unit digit of a^n , find the unit digit cycle of 'a', divide 'n' by the cycle length, and the remainder indicates the position in the cycle. If the remainder is 0, the unit digit is the last digit of the cycle.

23. Answer: c

Explanation:

Understanding the Compound Interest Problem

This problem involves calculating the annual interest rate for an investment where the interest is compounded annually. We are given the initial principal amount, the final amount after a certain period, and the time period in years.

Key Concepts: Compound Interest

Compound interest is interest calculated on the initial principal and also on the accumulated interest from previous periods. This means the interest earned also earns interest over time, leading to faster growth compared to simple interest.

The formula for compound interest when interest is compounded annually is:

$$A = P\left(1 + \frac{r}{100}\right)^n$$

Where:

- A is the final amount after n years

- P is the principal amount (initial investment)
- r is the annual interest rate (in percent)
- n is the number of years

Analyzing the Given Information

From the question, we have the following values:

- Principal Amount (P) = Rs. 10,000
- Final Amount (A) = Rs. 11,449
- Time Period (n) = 2 years
- We need to find the annual interest rate (r).

Step-by-Step Solution to Find the Interest Rate

We will plug the given values into the compound interest formula and solve for r .

$$A = P\left(1 + \frac{r}{100}\right)^n$$

Substitute the values:

$$11449 = 10000\left(1 + \frac{r}{100}\right)^2$$

Divide both sides by 10000:

$$\frac{11449}{10000} = \left(1 + \frac{r}{100}\right)^2$$

$$1.1449 = \left(1 + \frac{r}{100}\right)^2$$

To isolate the term $\left(1 + \frac{r}{100}\right)$, take the square root of both sides:

$$\sqrt{1.1449} = \sqrt{\left(1 + \frac{r}{100}\right)^2}$$

$$\sqrt{1.1449} = 1 + \frac{r}{100}$$

Calculate the square root of 1.1449:

$$1.07 = 1 + \frac{r}{100}$$

Now, subtract 1 from both sides to find the value of $\frac{r}{100}$:

$$1.07 - 1 = \frac{r}{100}$$

$$0.07 = \frac{r}{100}$$

Multiply both sides by 100 to find r :

$$r = 0.07 \times 100$$

$$r = 7$$

So, the annual interest rate is 7%.

Conclusion

The calculated annual interest rate is 7%. This matches option 3 provided in the question.

Revision Table: Compound Interest Terms

Term	Symbol	Description
Principal Amount	P	The initial sum of money invested or borrowed.
Final Amount	A	The total sum including both the principal and the accumulated interest.
Interest Rate	r	The rate at which interest is calculated, usually given as a percentage per annum.
Time Period	n	The duration for which the money is invested or borrowed.

Your Personal Exams Guide

Additional Information: Compound vs. Simple Interest

It's important to distinguish compound interest from simple interest. Simple interest is calculated only on the principal amount.

- **Simple Interest Formula:** $SI = \frac{P \times r \times n}{100}$
- **Total Amount with Simple Interest:** $A = P + SI = P(1 + \frac{rn}{100})$

In simple interest, the interest earned in each period remains constant, while in compound interest, the interest earned increases with each period because it is calculated on a growing principal (original principal + accumulated interest).

For the same principal, rate, and time (greater than 1 year), compound interest will always yield a higher final amount than simple interest.

24. Answer: a

Explanation:

Understanding Salesperson Displacement and Position

This question asks us to find the final position of a salesperson relative to their starting point after a series of movements in different directions. This involves understanding displacement, which is the straight-line distance and direction from the starting point to the ending point.

We can track the salesperson's movement step-by-step, considering the directions East, West, North, and South. We can think of East as the positive x-direction and North as the positive y-direction. Consequently, West is the negative x-direction and South is the negative y-direction.

Step-by-Step Movement Analysis

Let's break down each part of the salesperson's journey from the office (starting position):

1. **Starts from office:** Initial position is at the origin (0,0).
2. **Drives 2 km east:** This is a movement in the positive x-direction. The x-coordinate changes by +2 km. The position is now (2, 0).
3. **Turns north and drives 7 km:** This is a movement in the positive y-direction. The y-coordinate changes by +7 km. The position is now (2, 7).
4. **Turns to his right and drives 6 km:** When facing North, turning right means turning East. This is another movement in the positive x-direction. The x-coordinate changes by +6 km. The position is now $(2 + 6, 7) = (8, 7)$.
5. **Turns south and drives 7 km:** This is a movement in the negative y-direction. The y-coordinate changes by -7 km. The position is now $(8, 7 - 7) = (8, 0)$.

Calculating Total Displacement

The salesperson started at position (0,0) and ended at position (8,0). The total displacement in the x-direction (East-West) is the final x-coordinate minus the initial x-coordinate:

$$\Delta x = \text{Final } x - \text{Initial } x = 8 - 0 = 8 \text{ km}$$

Since the result is positive, the net displacement in the horizontal direction is 8 km East.

The total displacement in the y-direction (North-South) is the final y-coordinate minus the initial y-coordinate:

$$\Delta y = \text{Final } y - \text{Initial } y = 0 - 0 = 0 \text{ km}$$

Since the result is 0, there is no net displacement in the vertical direction (North-South). The 7 km North movement was cancelled out by the 7 km South movement.

Let's visualize the movements and net effect using a table:

Movement	Direction	Distance (km)	Horizontal Change (km)	Vertical Change (km)
Start	-	-	0	0
Step 1	East	2	+2	0
Step 2	North	7	0	+7
Step 3 (Right turn from North is East)	East	6	+6	0
Step 4	South	7	0	-7

Summing the horizontal changes: $+2 + 0 + 6 + 0 = +8$ km. (Positive means East)

Summing the vertical changes: $0 + 7 + 0 + (-7) = 0$ km. (Zero means no net North or South movement)

Final Position Relative to Starting Point

The net displacement is 8 km in the East direction and 0 km in the North-South direction. Therefore, the salesperson is 8 km east of their starting position.

Comparing with Options

Let's check the given options:

- 8 km east: Matches our calculation.

- 4 km west: Does not match. West would be a negative displacement in the x-direction.
- 4 km east: Does not match.
- 8 km west: Does not match. West would be a negative displacement in the x-direction.

Based on our detailed movement analysis and calculation, the salesperson is 8 km east of the starting position.

Revision Table – Direction and Displacement

Direction	Effect on Position	Cancellation Example
East	Moves away from start in the East direction	Cancelled by West movement of equal distance
West	Moves away from start in the West direction	Cancelled by East movement of equal distance
North	Moves away from start in the North direction	Cancelled by South movement of equal distance
South	Moves away from start in the South direction	Cancelled by North movement of equal distance

Your Personal Exams Guide

Additional Information – Distance vs. Displacement

It is important to distinguish between distance and displacement. Distance is the total length of the path traveled. In this case, the total distance traveled is $2 + 7 + 6 + 7 = 22$ km. Displacement, however, is only concerned with the initial and final positions. It is a vector quantity, meaning it has both magnitude (size) and direction. Here, the magnitude of the displacement is 8 km, and the direction is East.

If the salesperson had returned to the office at the end, their displacement would be 0 km, even though the distance traveled would be significant.

25. Answer: a

Explanation:

Understanding Word Analogies and Relationships

This question asks us to find a word that has the same relationship with the third word ('Glad') as the second word ('Short') has with the first word ('Tall'). This type of question is called a word analogy.

Analyzing the Relationship: Tall : Short

Let's look at the first pair of words: **Tall** and **Short**.

- **Tall** describes something having a great height.
- **Short** describes something having little height.

These two words have opposite meanings. They are antonyms.

Applying the Relationship: Glad : ?

The analogy is given as Tall : Short ∷ Glad : ?. This means the relationship between Tall and Short is the same as the relationship between Glad and the missing word.

Since the relationship in the first pair is that of antonyms (opposites), we need to find the antonym of the word **Glad**.

Glad means feeling pleasure, joy, or happiness.

We need to find a word that means the opposite of feeling pleasure, joy, or happiness.

Evaluating the Options

Let's examine the given options:

1. **Sad**: This word means feeling or showing sorrow; unhappy. This is the opposite of feeling pleasure, joy, or happiness.
2. **Happy**: This word is a synonym of Glad. It has a similar meaning, not an opposite meaning.
3. **Emotion**: Glad is an emotion, but Emotion is a category, not the opposite feeling.
4. **Smile**: A smile is a facial expression that often shows happiness, but it is an action, not the opposite feeling itself.

Comparing the meanings, **Sad** is the word that is the opposite of **Glad**.

Conclusion on the Word Analogy

The relationship between Tall and Short is that they are antonyms. Applying the same antonym relationship to Glad, the word that is the opposite of Glad is Sad.

Revision Table: Common Analogy Types

Analogy Type	Description	Example
Antonym/Opposite	Words with opposite meanings	Hot : Cold
Synonym/Similar	Words with similar meanings	Big : Large
Part to Whole	One word is a part of the other	Finger : Hand
Cause and Effect	One word is the cause of the other	Rain : Flood
Worker and Tool	A person and the tool they use	Carpenter : Hammer

Additional Information on Word Relationships

Understanding the different ways words can relate to each other is key to solving word analogies. Beyond antonyms and synonyms, words can relate based on category (e.g., Apple : Fruit), function (e.g., Knife : Cut), degree (e.g., Warm : Hot), or even sound (rhymes or homophones).

For the analogy "Tall : Short ∷ Glad : ?", the core relationship is the antonym relationship. Identifying this first step correctly is crucial.

Practice with different types of word relationships helps in quickly recognizing the pattern in analogy questions.

26. Answer: a

Explanation:

Understanding the Padma Bhushan Award

The Padma Bhushan is the third-highest civilian award in the Republic of India, preceded by the Bharat Ratna and the Padma Vibhushan and followed by the Padma Shri. It is awarded for 'distinguished service of a high order' without distinction of race, occupation, position, or sex.

Indian Cricketer Awarded Padma Bhushan in 2018

In 2018, several eminent personalities from various fields were honoured with the prestigious Padma Awards. Among the recipients of the Padma Bhushan in 2018 was a prominent Indian cricketer known for his significant contributions to the sport.

The Indian cricketer who received the Padma Bhushan in 2018 was **MS Dhoni**. He was recognised for his exceptional achievements and leadership in Indian cricket, including leading India to victory in the 2011 Cricket World Cup and the 2007 ICC World Twenty20.

Analysis of the Options

Let's look at the given options in the context of receiving the Padma Bhushan in 2018:

- **MS Dhoni:** As mentioned, MS Dhoni was indeed awarded the Padma Bhushan in 2018.
- **Virat Kohli:** Virat Kohli is another celebrated Indian cricketer, but he received the Padma Shri, the fourth-highest civilian award, in 2017.
- **Sachin Tendulkar:** Sachin Tendulkar, often regarded as one of the greatest batsmen in cricket history, received the Padma Vibhushan (the second-highest civilian award) in 2008 and the Bharat Ratna (the highest civilian award) in 2014.
- **Saurav Ganguly:** Saurav Ganguly, a former captain of the Indian cricket team, received the Padma Shri in 2004.

Based on the official list of awardees for 2018, MS Dhoni is the correct cricketer among the options who received the Padma Bhushan in that specific year.

Confirming the Correct Cricketer

Therefore, the Indian cricketer who was honoured with the Padma Bhushan award in the year 2018 is MS Dhoni.

Cricketer	Padma Award Received	Year Received
MS Dhoni	Padma Bhushan	2018
Virat Kohli	Padma Shri	2017
Sachin Tendulkar	Padma Vibhushan, Bharat Ratna	2008, 2014
Saurav Ganguly	Padma Shri	2004

Revision Table: Indian Cricketers and Padma Awards

This table summarizes the Padma Awards received by the cricketers listed in the options.

Additional Information: Padma Awards for Sports

The Padma Awards are given annually on Republic Day. They recognize achievements in various disciplines, including sports. Many Indian sportspersons, including cricketers, have been honoured with these awards over the years for their contributions to bringing glory to the nation through sports.

27. Answer: b

Explanation:

In all the figures except option '2', there are 5 triangle and 2 circle and one rectangle, while in option 2 there are 4 triangle and 2 circle and one rectangle.

Hence, 'option 2' is the odd one out.

28. Answer: d

Explanation:

Evaluating Expressions with Symbol Substitutions

This problem requires us to first understand the new meanings assigned to the mathematical operators and then apply these rules to the given expression to find its numerical value. We will use the standard order of operations (BODMAS/PEMDAS) after the symbols have been substituted.

Understanding the Symbol Rules

The question provides a specific set of rules for substituting standard mathematical operators:

- The symbol '+' represents '×' (multiplication).
- The symbol '-' represents '+' (addition).
- The symbol '×' represents '÷' (division).
- The symbol '÷' represents '-' (subtraction).

Applying Substitutions to the Expression

The original expression given is:

$$9 \times 3 + 6 \div 2$$

Now, we substitute the symbols in this expression according to the rules provided:

- The '×' between 9 and 3 becomes '÷'.
- The '+' between 3 and 6 becomes '×'.
- The '÷' between 6 and 2 becomes '-'.

So, the new expression after substitution is:

$$9 \div 3 \times 6 - 2$$

Evaluating the Modified Expression using BODMAS/PEMDAS

We need to evaluate the expression $9 \div 3 \times 6 - 2$ following the order of operations:

BODMAS/PEMDAS Order:

1. Brackets / Parentheses
2. Orders / Exponents
3. Division and Multiplication (from left to right)
4. Addition and Subtraction (from left to right)

Let's evaluate step-by-step:

1. **Division:** Evaluate $9 \div 3$.

$$9 \div 3 = 3$$

The expression becomes: $3 \times 6 - 2$.

2. **Multiplication:** Evaluate 3×6 .

$$3 \times 6 = 18$$

The expression becomes: $18 - 2$.

3. **Subtraction:** Evaluate $18 - 2$.

$$18 - 2 = 16$$

The final value of the expression after applying the symbol substitutions and evaluating is 16.

Summary of Steps

Step	Description	Expression
1	Original Expression	$9 \times 3 + 6 \div 2$
2	Apply Symbol Substitutions	$9 \div 3 \times 6 - 2$
3	Evaluate Division ($9 \div 3$)	$3 \times 6 - 2$
4	Evaluate Multiplication (3×6)	$18 - 2$
5	Evaluate Subtraction ($18 - 2$)	16

Final Answer Evaluation

After performing the required symbol substitutions and following the order of operations, the value of the expression $9 \times 3 + 6 \div 2$ (under the given rules) is 16.

Revision Table: Symbol Substitution and Expression Evaluation

Original Symbol	New Meaning	Example
+	$\times$	$A + B \rightarrow A \times B$
-	+	$A - B \rightarrow A + B$
$\times$	$\div$	$A \times B \rightarrow A \div B$
$\div$	-	$A \div B \rightarrow A - B$

Additional Information: Order of Operations (BODMAS/PEMDAS)

The order of operations is a crucial rule in mathematics to ensure calculations are performed in a consistent sequence, leading to a unique correct answer. The acronyms BODMAS or PEMDAS help remember this order:

- **Brackets / Parentheses:** Operations inside brackets or parentheses are performed first.
- **Orders / Exponents:** Powers, square roots, and other orders are calculated next.
- **Division and Multiplication:** These are performed from left to right as they appear in the expression.
- **Addition and Subtraction:** These are performed last, also from left to right as they appear.

In the expression $9 \div 3 \times 6 - 2$, division ($9 \div 3$) comes before multiplication (3×6) because division appears first when reading from left to right among the division and multiplication operations. Then subtraction ($18 - 2$) is performed.

29. Answer: c

Explanation:

Calculating Heat Transfer in Materials

This problem involves calculating the amount of heat energy transferred to a block of iron when its temperature increases. This process is governed by the concept of specific heat

capacity, which is a property of a substance that tells us how much energy is needed to raise the temperature of a unit mass of that substance by one degree (Celsius or Kelvin).

Understanding Specific Heat Capacity

Specific heat capacity (c) is defined as the amount of heat energy required to raise the temperature of 1 kilogram of a substance by 1 Kelvin (or 1 degree Celsius). Its unit is typically $\text{J kg}^{-1} \text{K}^{-1}$ or $\text{J kg}^{-1} \text{°C}^{-1}$. The formula used to calculate the heat transferred (Q) is:

$$Q = mc\Delta T$$

Where:

- Q is the heat energy transferred (in Joules, J)
- m is the mass of the substance (in kilograms, kg)
- c is the specific heat capacity of the substance (in $\text{J kg}^{-1} \text{K}^{-1}$)
- ΔT is the change in temperature (in Kelvin, K, or degrees Celsius, °C)

Applying the Formula to the Iron Block

We are given the following information:

- Mass of the iron block, $m = 200 \text{ g}$
- Initial temperature, $T_{\text{initial}} = 30 \text{ °C}$
- Final temperature, $T_{\text{final}} = 60 \text{ °C}$
- Specific heat capacity of iron, $c = 450 \text{ J kg}^{-1} \text{K}^{-1}$

First, we need to ensure all units are consistent with the specific heat capacity unit ($\text{J kg}^{-1} \text{K}^{-1}$). The mass is given in grams, so we convert it to kilograms:

$$m = 200 \text{ g} = \frac{200}{1000} \text{ kg} = 0.2 \text{ kg}$$

Next, we calculate the change in temperature, ΔT . Note that a change in temperature in Celsius is equivalent to a change in temperature in Kelvin.

$$\Delta T = T_{\text{final}} - T_{\text{initial}} = 60 \text{ °C} - 30 \text{ °C} = 30 \text{ °C}$$

Therefore, $\Delta T = 30 \text{ K}$.

Now we can substitute the values into the heat transfer formula:

$$Q = mc\Delta T$$

$$Q = (0.2 \text{ kg}) \times (450 \text{ J kg}^{-1} \text{ K}^{-1}) \times (30 \text{ K})$$

Perform the multiplication:

$$Q = (0.2 \times 450) \times 30 \text{ J}$$

$$Q = 90 \times 30 \text{ J}$$

$$Q = 2700 \text{ J}$$

Thus, 2700 Joules of heat energy were transferred to the iron block.

Summary of Calculation

The calculation steps were:

- Convert mass from grams to kilograms.
- Calculate the temperature change.
- Use the formula $Q = mc\Delta T$.
- Substitute values and compute the result.

Quantity	Symbol	Value	Units
Mass	m	200 g = 0.2	kg
Initial Temperature	$T_{initial}$	30	°C
Final Temperature	T_{final}	60	°C
Change in Temperature	ΔT	$60 - 30 = 30$	°C or K
Specific Heat Capacity	c	450	J kg ⁻¹ K ⁻¹

Calculation: $Q = 0.2 \times 450 \times 30 = 2700 \text{ J}$.

Revision Table: Heat Transfer Concepts

Concept	Description	Formula
Heat Transfer (Q)	Energy transferred due to temperature difference.	$Q = mc\Delta T$ (for specific heat)
Specific Heat Capacity (c)	Energy required to raise 1 kg of a substance by 1 K.	$c = \frac{Q}{m\Delta T}$
Temperature Change (ΔT)	Difference between final and initial temperature.	$\Delta T = T_{final} - T_{initial}$
Units	Standard units for Q , m , c , ΔT .	Joules (J), kilograms (kg), J kg ⁻¹ K ⁻¹ , Kelvin (K) or °C

Additional Information: Factors Affecting Heat Transfer

The amount of heat transferred when a substance changes temperature depends on three main factors:

- **Mass (m):** More mass requires more energy for the same temperature change.
- **Specific Heat Capacity (c):** Different substances require different amounts of energy to change their temperature. Substances with high specific heat capacity (like water) require more energy than substances with low specific heat capacity (like iron) for the same mass and temperature change.
- **Temperature Change (ΔT):** A larger temperature change requires more energy transfer.

This specific heat calculation applies when the substance is only changing temperature and not undergoing a phase change (like melting or boiling). Phase changes involve latent heat, which is calculated differently.

30. Answer: a

Explanation:

This question asks about the standard dimensions of an engineering drawing sheet, specifically its length when the width is given as 841 mm. Standard engineering drawing

sheet sizes are defined by the ISO 216 standard, which is based on the German DIN 476 standard. These standards define the 'A' series paper sizes, starting with A0.

Understanding Standard Engineering Drawing Sheet Sizes (ISO 216)

The ISO 216 standard specifies paper sizes based on a constant aspect ratio of $1 : \sqrt{2}$. This means that if the width of a sheet is W and the length is L , then $L = W \times \sqrt{2}$. A key property of this ratio is that if you cut an A-series sheet in half parallel to its shorter side, you get two smaller sheets of the next size in the series, and these smaller sheets maintain the same aspect ratio.

The base size is A0, which is defined to have an area of 1 m^2 . With the aspect ratio $1 : \sqrt{2}$, let the width be W_0 and the length be L_0 . We have $L_0 = W_0\sqrt{2}$ and $W_0 \times L_0 = 1 \text{ m}^2 = 1,000,000 \text{ mm}^2$.

Substituting $L_0 = W_0\sqrt{2}$ into the area equation:

$$W_0 \times (W_0\sqrt{2}) = 1,000,000$$

$$W_0^2\sqrt{2} = 1,000,000$$

$$W_0^2 = \frac{1,000,000}{\sqrt{2}}$$

$$W_0 = \sqrt{\frac{1,000,000}{\sqrt{2}}} \approx \sqrt{707106.78} \approx 840.89 \text{ mm}$$

The standard rounds this value. The standard dimensions for A0 are approximately 841 mm $\times$ 1189 mm. Let's check the aspect ratio:

$$\frac{1189}{841} \approx 1.41379$$

$$\sqrt{2} \approx 1.41421$$

These are very close, confirming the $1 : \sqrt{2}$ ratio for the standard dimensions.

Standard A-Series Dimensions

The sizes are derived by halving the length of the previous size to get the width of the next size, and the previous width becomes the next length. Here are the standard dimensions for the first few A-series sizes:

Sheet Size	Dimensions (Width × Length) in mm
A0	841 × 1189
A1	594 × 841
A2	420 × 594
A3	297 × 420
A4	210 × 297

Analyzing the Given Engineering Drawing Sheet Dimensions

The question states that the width of the standard engineering drawing sheet is 841 mm. We need to find its length. Looking at the standard dimensions table:

- For A0, the dimensions are 841 mm × 1189 mm. The width is 841 mm, and the length is 1189 mm.
- For A1, the dimensions are 594 mm × 841 mm. The width is 594 mm, and the length is 841 mm.

The question specifies the width as 841 mm. According to the standard dimensions, a width of 841 mm corresponds to the A0 size sheet. For an A0 sheet, the length is 1189 mm.

Determining the Length of the Engineering Drawing Sheet

Based on the ISO 216 standard for A-series paper sizes, a standard engineering drawing sheet with a width of 841 mm is an A0 size sheet. The standard dimensions for an A0 sheet are 841 mm × 1189 mm.

Therefore, if the width is 841 mm, the length of the engineering drawing sheet is 1189 mm.

Comparing this with the given options:

- 1. 1189
- 2. 1250
- 3. 1000
- 4. 1216

The calculated length, 1189 mm, matches Option 1.

The final answer is 1189.

Revision Table: Standard Drawing Sheet Sizes

Size	Width (mm)	Length (mm)	Area (m ²)
A0	841	1189	1.00
A1	594	841	0.50
A2	420	594	0.25
A3	297	420	0.125
A4	210	297	0.0625

Additional Information: ISO 216 Standard and Applications

The ISO 216 standard is widely used internationally for paper sizes, especially for official documents, technical drawings, and general printing. Its logic of halving the sheet size to get the next smaller standard size is very practical for scaling documents like technical drawings or photocopies.

The $1 : \sqrt{2}$ aspect ratio ensures that when you scale a document from one A-series size to another, the proportions remain the same. For example, reducing an A0 drawing to A1 means reducing both dimensions by a factor of $\sqrt{2}/2 \approx 0.707$ (or 70.7%). Scaling an A4 document to A3 involves enlarging both dimensions by a factor of $\sqrt{2} \approx 1.414$ (or 141.4%). This makes scaling on photocopiers and printers straightforward.

These standard drawing sheet sizes are essential in engineering and architecture to ensure consistency and ease of handling, filing, and reproduction of technical drawings.

31. Answer: d

Explanation:

Understanding the Time and Work Problem

This question deals with the concept of time and work, specifically when the work rates of pairs of individuals are given, and we need to find the time taken by all three together.

Let's denote the amount of work A can do in one day as a , the amount of work B can do in one day as b , and the amount of work C can do in one day as c . We can consider the total work to be completed as 1 unit.

Setting Up Equations from Given Information

We are given the time taken by pairs of individuals to complete the work:

- A and B together can do the work in 15 days. This means their combined work rate is $\frac{1}{15}$ of the total work per day. So, we have the equation: $a + b = \frac{1}{15}$ (Equation 1)
- B and C together can do the work in 20 days. Their combined work rate is $\frac{1}{20}$ of the total work per day. So, we have: $b + c = \frac{1}{20}$ (Equation 2)
- A and C together can do the work in 10 days. Their combined work rate is $\frac{1}{10}$ of the total work per day. So, we have: $a + c = \frac{1}{10}$ (Equation 3)

Calculating the Combined Work Rate of A, B, and C

To find the combined work rate of A, B, and C together ($a + b + c$), we can add the three equations we formulated:

$$(a + b) + (b + c) + (a + c) = \frac{1}{15} + \frac{1}{20} + \frac{1}{10}$$

This simplifies to:

$$2a + 2b + 2c = \frac{1}{15} + \frac{1}{20} + \frac{1}{10}$$

We can factor out 2 on the left side:

$$2(a + b + c) = \frac{1}{15} + \frac{1}{20} + \frac{1}{10}$$

Now, we need to add the fractions on the right side. The least common multiple (LCM) of 15, 20, and 10 is 60.

$$2(a + b + c) = \frac{1 \times 4}{15 \times 4} + \frac{1 \times 3}{20 \times 3} + \frac{1 \times 6}{10 \times 6}$$

$$2(a + b + c) = \frac{4}{60} + \frac{3}{60} + \frac{6}{60}$$

$$2(a + b + c) = \frac{4+3+6}{60}$$

$$2(a + b + c) = \frac{13}{60}$$

Now, divide by 2 to find the combined work rate of A, B, and C:

$$a + b + c = \frac{13}{60 \times 2}$$

$$a + b + c = \frac{13}{120}$$

This means that A, B, and C together complete $\frac{13}{120}$ of the total work in one day.

Finding the Time Taken Together

If A, B, and C together complete $\frac{13}{120}$ of the work in one day, the total time taken to complete the entire work (1 unit) is the reciprocal of their combined work rate.

$$\text{Time taken by A, B, and C together} = \frac{1}{\text{Combined work rate}} = \frac{1}{\frac{13}{120}}$$

$$\text{Time taken} = \frac{120}{13} \text{ days.}$$

Calculating the Decimal Value

Let's calculate the decimal value of $\frac{120}{13}$:

$$120 \div 13 \approx 9.2307...$$

Rounding to two decimal places, the time taken is approximately 9.23 days.

Comparing with Options

Let's look at the given options:

1. 10.67 days
2. 10.91 days
3. 10.71 days
4. 9.23 days

Our calculated time of approximately 9.23 days matches option 4.

Pair	Time Taken (days)	Combined Work Rate (Work/day)
A & B	15	$\frac{1}{15}$
B & C	20	$\frac{1}{20}$
A & C	10	$\frac{1}{10}$
A, B & C (Together)	$\frac{120}{13} \approx 9.23$	$\frac{13}{120}$

Revision Table: Key Time and Work Concepts

Concept	Explanation	Formula/Relation
Work Rate	The amount of work done per unit of time (e.g., per day).	Work Rate = $\frac{\text{Total Work}}{\text{Time Taken}}$
Time Taken	The total time required to complete the work.	Time Taken = $\frac{\text{Total Work}}{\text{Work Rate}}$
Total Work	Often considered as 1 unit for calculation purposes.	Total Work = Work Rate $\times$ Time Taken
Combined Work Rate	Sum of individual work rates when people work together.	Rate _{A+B} = Rate _A + Rate _B

Your Personal Exams Guide

Additional Information: Solving System of Linear Equations

In this problem, we used a method of adding equations to find the combined work rate. Another way to approach problems like this, especially if individual rates (a, b, c) are needed, is to solve the system of linear equations:

- $a + b = \frac{1}{15}$
- $b + c = \frac{1}{20}$
- $a + c = \frac{1}{10}$

You could use substitution or elimination methods. For example, subtract Equation 2 from Equation 1:

$$(a + b) - (b + c) = \frac{1}{15} - \frac{1}{20}$$

$$a - c = \frac{4-3}{60} = \frac{1}{60} \text{ (Equation 4)}$$

Now add Equation 3 and Equation 4:

$$(a + c) + (a - c) = \frac{1}{10} + \frac{1}{60}$$

$$2a = \frac{6+1}{60} = \frac{7}{60}$$

$$a = \frac{7}{120}$$

Once you have 'a', you can substitute it back into Equation 1 and Equation 3 to find 'b' and 'c'. Then you can sum $a+b+c$. This method also works and might be useful if the question asked for individual times.

32. Answer: d

Explanation:

Understanding the Series-Parallel Circuit

The problem describes a circuit with resistors connected in a combination of series and parallel. We have two resistors ($12\ \Omega$ and $24\ \Omega$) in parallel, and this parallel combination is then connected in series with another resistor ($22\ \Omega$) and a 12 V battery. We need to find the current flowing specifically through the $12\ \Omega$ resistor.

To solve this, we will follow these steps:

1. Calculate the equivalent resistance of the parallel combination.
2. Calculate the total resistance of the entire circuit.
3. Calculate the total current flowing from the battery.
4. Calculate the voltage drop across the parallel combination.
5. Use the voltage drop across the parallel combination to find the current through the $12\ \Omega$ resistor.

Step-by-Step Solution for Circuit Analysis

1. Equivalent Resistance of Parallel Resistors

Two resistors, $R_1 = 12\Omega$ and $R_2 = 24\Omega$, are connected in parallel. The formula for the equivalent resistance (R_p) of two resistors in parallel is:

$$\frac{1}{R_p} = \frac{1}{R_1} + \frac{1}{R_2}$$

Substituting the values:

$$\frac{1}{R_p} = \frac{1}{12\Omega} + \frac{1}{24\Omega}$$

To add these fractions, we find a common denominator, which is 24:

$$\frac{1}{R_p} = \frac{2}{24\Omega} + \frac{1}{24\Omega} = \frac{2+1}{24\Omega} = \frac{3}{24\Omega}$$

Now, we find R_p by taking the reciprocal:

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

The equivalent resistance of the parallel combination is 8Ω .

2. Total Resistance of the Circuit

The parallel combination ($R_p = 8\Omega$) is connected in series with a third resistor ($R_3 = 22\Omega$). The total resistance (R_{total}) of components in series is the sum of their individual resistances:

$$R_{total} = R_p + R_3$$

Substituting the values:

$$R_{total} = 8\Omega + 22\Omega = 30\Omega$$

The total resistance of the circuit is 30Ω .

3. Total Current from the Battery

The total voltage supplied by the battery is $V_{total} = 12\text{ V}$. Using Ohm's Law ($V = IR$), we can find the total current (I_{total}) flowing from the battery:

$$I_{total} = \frac{V_{total}}{R_{total}}$$

Substituting the values:

$$I_{total} = \frac{12 \text{ V}}{30\Omega} = \frac{12}{30} \text{ A}$$

Simplifying the fraction:

$$I_{total} = \frac{2}{5} \text{ A}$$

The total current flowing in the circuit is $(2/5) \text{ A}$.

4. Voltage Across the Parallel Combination

The total current (I_{total}) flows through the series resistor and then splits between the parallel resistors. The voltage drop across the parallel combination is equal to the total current multiplied by the equivalent resistance of the parallel section:

$$V_{parallel} = I_{total} \times R_p$$

Substituting the values:

$$V_{parallel} = \frac{2}{5} \text{ A} \times 8\Omega = \frac{16}{5} \text{ V}$$

The voltage across the parallel combination is $(16/5) \text{ V}$.

5. Current Through the 12Ω Resistor

In a parallel combination, the voltage across each resistor is the same and is equal to the voltage across the combination. We need to find the current through the 12Ω resistor (R_1). Using Ohm's Law for the 12Ω resistor:

$$I_{12\Omega} = \frac{V_{parallel}}{R_1}$$

Substituting the values:

$$I_{12\Omega} = \frac{16/5 \text{ V}}{12\Omega}$$

$$I_{12\Omega} = \frac{16}{5 \times 12} \text{ A} = \frac{16}{60} \text{ A}$$

Simplifying the fraction by dividing both numerator and denominator by their greatest common divisor, 4:

$$I_{12\Omega} = \frac{16 \div 4}{60 \div 4} \text{ A} = \frac{4}{15} \text{ A}$$

The current in the $12\ \Omega$ resistor is $(4/15)\text{ A}$.

Let's compare this result with the given options:

- $(8/15)\text{ A}$
- $(2/15)\text{ A}$
- $(6/15)\text{ A}$
- $(4/15)\text{ A}$

Our calculated current matches the fourth option.

Component	Resistance	Voltage Drop	Current
$12\ \Omega$ Resistor (R_1)	$12\ \Omega$	$(16/5)\text{ V}$	$(4/15)\text{ A}$
$24\ \Omega$ Resistor (R_2)	$24\ \Omega$	$(16/5)\text{ V}$	$(16/5)\text{ V} / 24\ \Omega = (16/120)\text{ A} = (2/15)\text{ A}$
Parallel Combination (R_p)	$8\ \Omega$	$(16/5)\text{ V}$	$(2/5)\text{ A}$ (Total current entering parallel section)
$22\ \Omega$ Resistor (R_3)	$22\ \Omega$	$(2/5)\text{ A} * 22\ \Omega = (44/5)\text{ V}$	$(2/5)\text{ A}$ (Total current)
Total Circuit	$30\ \Omega$	12 V (Battery voltage)	$(2/5)\text{ A}$ (Total current)

Your Personal Exams Guide

Note that the sum of currents in the parallel branches is $I_{12\Omega} + I_{24\Omega} = (4/15)\text{ A} + (2/15)\text{ A} = (6/15)\text{ A} = (2/5)\text{ A}$, which equals the total current flowing into the parallel section, as expected by Kirchhoff's Current Law.

Also, the sum of voltage drops across the series components is $V_{parallel} + V_{R_3} = (16/5)\text{ V} + (44/5)\text{ V} = (60/5)\text{ V} = 12\text{ V}$, which equals the total battery voltage, as expected by Kirchhoff's Voltage Law.

Revision Table: Series-Parallel Circuit Analysis

Concept	Formula	Application
Resistors in Parallel	$1/R_{eq} = 1/R_1 + 1/R_2 + \dots$	Used to find equivalent resistance of 12 Ω and 24 Ω resistors.
Resistors in Series	$R_{eq} = R_1 + R_2 + \dots$	Used to find total resistance by adding parallel equivalent and 22 Ω resistor.
Ohm's Law	$V = IR$	Used multiple times to find total current, voltage drop across parallel section, and current through 12 Ω resistor.
Voltage in Parallel	Voltage is same across each component.	$V_{12\Omega} = V_{24\Omega} = V_{parallel}$
Current in Series	Current is same through each component.	I_{total} flows through series resistor and enters parallel section.
Kirchhoff's Current Law (KCL)	Sum of currents entering a node equals sum of currents leaving.	Total current entering parallel split equals sum of currents in 12 Ω and 24 Ω branches.
Kirchhoff's Voltage Law (KVL)	Sum of voltage drops around a closed loop is zero.	Battery voltage equals sum of voltage drops across series components (parallel equivalent + 22 Ω resistor).

Additional Information on Circuit Analysis

Analyzing series-parallel circuits involves breaking down the circuit into simpler parts. First, simplify parallel combinations into a single equivalent resistance. Then, treat this equivalent resistance as being in series with other components. Calculate the total resistance and total current. Work backward from the total current and voltage, using Ohm's Law and the rules for voltage and current distribution in series and parallel sections, to find currents and voltages for individual components.

For a parallel combination, the resistor with lower resistance will always have a higher current flowing through it, provided the voltage across the parallel branches is constant. In this case, the 12 Ω resistor is smaller than the 24 Ω resistor, and we found its current to be (4/15) A. The current through the 24 Ω resistor was (2/15) A. As expected, (4/15) A > (2/15) A.

33. Answer: b

Explanation:

Understanding Voltage Across Resistors in a Series Circuit

In a series circuit, components are connected end-to-end, forming a single path for the electric current to flow. When resistors are connected in series, the total resistance of the circuit increases. The current flowing through each resistor in a series combination is the same, while the total voltage across the combination is distributed among the individual resistors.

The problem describes a series circuit with two resistors and a voltage supply. We need to find the voltage across the $10\ \Omega$ resistor. To do this, we first need to find the total resistance and then the total current flowing through the circuit.

Calculating Total Resistance in a Series Circuit

When resistors are connected in series, the total resistance (R_{total}) is the sum of the individual resistances. In this case, we have two resistors:

- Resistor 1 (R_1) = $10\ \Omega$
- Resistor 2 (R_2) = $20\ \Omega$

The formula for total resistance in series is:

$$R_{total} = R_1 + R_2 + R_3 + \dots$$

Applying this to our problem:

$$R_{total} = 10\ \Omega + 20\ \Omega = 30\ \Omega$$

So, the total resistance of the series combination is $30\ \Omega$.

Calculating Total Current in the Series Circuit

Now that we know the total resistance and the total supply voltage, we can use Ohm's Law to find the total current (I_{total}) flowing through the circuit. Ohm's Law states that voltage (V) is equal to current (I) multiplied by resistance (R):

$$V = I \times R$$

We can rearrange this to find the current:

$$I = \frac{V}{R}$$

Given the supply voltage (V_{supply}) is 30 V and the total resistance (R_{total}) is 30 Ω , the total current is:

$$I_{total} = \frac{V_{supply}}{R_{total}} = \frac{30\text{ V}}{30\text{ }\Omega} = 1\text{ A}$$

A total current of 1 A flows through the entire series circuit, including both the 10 Ω and the 20 Ω resistors.

Calculating Voltage Across the 10 Ω Resistor

Since the current is the same throughout a series circuit, the 1 A current flows through the 10 Ω resistor. We can use Ohm's Law again to find the voltage across the 10 Ω resistor ($V_{10\Omega}$).

$$V_{10\Omega} = I_{total} \times R_1$$

Substitute the values:

$$V_{10\Omega} = 1\text{ A} \times 10\text{ }\Omega = 10\text{ V}$$

Therefore, the voltage across the 10 Ω resistor is 10 V.

Voltage Divider Rule in Series Circuits

Alternatively, we can use the voltage divider rule, which is very useful for series circuits. The voltage across a specific resistor in a series circuit is given by:

$$V_{R_x} = V_{total} \times \frac{R_x}{R_{total}}$$

Where:

- V_{R_x} is the voltage across resistor R_x .
- V_{total} is the total voltage across the series combination (supply voltage).
- R_x is the resistance of the specific resistor.
- R_{total} is the total resistance of the series combination.

Using this rule for the $10\ \Omega$ resistor (R_1):

$$V_{10\Omega} = V_{supply} \times \frac{R_1}{R_1 + R_2}$$

Substitute the given values:

$$V_{10\Omega} = 30\text{ V} \times \frac{10\ \Omega}{10\ \Omega + 20\ \Omega}$$

$$V_{10\Omega} = 30\text{ V} \times \frac{10\ \Omega}{30\ \Omega}$$

$$V_{10\Omega} = 30\text{ V} \times \frac{1}{3}$$

$$V_{10\Omega} = 10\text{ V}$$

Both methods confirm that the voltage across the $10\ \Omega$ resistor is 10 V .

Revision Table: Series Circuit Calculations

Prepp

Your Personal Exams Guide

Parameter	Value	Calculation/Reason
Resistor 1 (R_1)	10 Ω	Given
Resistor 2 (R_2)	20 Ω	Given
Supply Voltage (V_{supply})	30 V	Given
Total Resistance (R_{total})	30 Ω	$R_1 + R_2$ (Series)
Total Current (I_{total})	1 A	$V_{\text{supply}} / R_{\text{total}}$ (Ohm's Law)
Current through 10 Ω ($I_{10\Omega}$)	1 A	Same as total current in series
Voltage across 10 Ω ($V_{10\Omega}$)	10 V	$I_{10\Omega} \times R_1$ (Ohm's Law)
Voltage across 20 Ω ($V_{20\Omega}$)	20 V	$I_{\text{total}} \times R_2 = 1\text{A} \times 20\Omega$ (Voltage divider rule: $30\text{V} \times 20/30 = 20\text{V}$)
Sum of Voltages across Resistors	10 V + 20 V = 30 V	Should equal supply voltage (Kirchhoff's Voltage Law)

Additional Information on Series Circuits Guide

Here are some key properties of series circuits:

- **Current:** The current is the same at all points in a series circuit. If the circuit breaks anywhere, the current stops flowing in the entire circuit.
- **Voltage:** The total voltage across a series combination is the sum of the individual voltages across each component. This is an application of Kirchhoff's Voltage Law.
- **Resistance:** The total equivalent resistance of resistors in series is the sum of their individual resistances ($R_{\text{total}} = R_1 + R_2 + \dots$).
- **Power:** The total power dissipated in a series circuit is the sum of the power dissipated by each component.

Understanding these properties is crucial for analyzing and solving problems involving series circuits.

34. Answer: b

Explanation:

Finding the Cost Price with Discount and Profit

This problem involves understanding the relationship between Cost Price (CP), Selling Price (SP), Marked Price (MP), discount, and profit percentage. We are given a scenario where if a discount is offered on the marked price, a certain profit is made. We need to find the original cost price of the article.

Let's break down the information provided:

- An article was sold for Rs. 12,000. This price is likely the Marked Price (MP) from which the hypothetical discount is offered, as this interpretation allows us to connect the given numbers logically to the profit condition. So, let's assume the Marked Price (MP) is Rs. 12,000.
- Had a discount of 15% been offered (on the MP), a profit of 2% would have been made.

Let's calculate the Selling Price (SP) in the hypothetical scenario where a 15% discount is offered on the Marked Price (MP).

Discount amount = 15% of MP

$$\text{Discount amount} = \frac{15}{100} \times 12000$$

$$\text{Discount amount} = 15 \times 120$$

$$\text{Discount amount} = \text{Rs. } 1800$$

The Selling Price (SP) after the discount (let's call it SP') would be:

$$\text{SP}' = \text{MP} - \text{Discount amount}$$

$$\text{SP}' = 12000 - 1800$$

$$\text{SP}' = \text{Rs. } 10200$$

Now, we are told that if the article was sold at this price (SP' = Rs. 10200), a profit of 2% would have been made. Profit percentage is always calculated on the Cost Price (CP).

Profit = 2% of CP

$$\text{Profit} = \frac{2}{100} \times \text{CP}$$

The Selling Price (SP) is also related to Cost Price (CP) and Profit:

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

Using the values from the hypothetical scenario (SP' and 2% profit):

$$\text{SP} = \text{CP} + \text{2\% of CP}$$

$$10200 = \text{CP} + \frac{2}{100} \times \text{CP}$$

$$10200 = \text{CP} \left(1 + \frac{2}{100} \right)$$

$$10200 = \text{CP} \left(\frac{100 + 2}{100} \right)$$

$$10200 = \text{CP} \left(\frac{102}{100} \right)$$

Now, we can solve for CP:

$$\text{CP} = 10200 \times \frac{100}{102}$$

We can simplify this calculation:

$$\text{CP} = \frac{1020000}{102}$$

Since $102 \times 100 = 10200$, we can see that $10200/102 = 100$.

$$\text{CP} = 100 \times 100$$

$$\text{CP} = \text{Rs. } 10000$$

Therefore, the cost price of the article was Rs. 10,000.

Let's verify this: If CP = Rs. 10,000 and MP = Rs. 12,000, a 15% discount on MP is Rs. 1800. The selling price after discount is Rs. 12000 - Rs. 1800 = Rs. 10200. The profit is Rs. 10200 - Rs. 10000 = Rs. 200. The profit percentage is $(200 / 10000) \times 100 = 2\%$. This matches the condition given in the problem.

Calculation Summary

Item	Value
Assumed Marked Price (MP)	Rs. 12,000
Discount Percentage	15%
Discount Amount	15% of Rs. 12,000 = Rs. 1800
Hypothetical Selling Price (SP')	Rs. 12,000 - Rs. 1800 = Rs. 10,200
Hypothetical Profit Percentage	2%
Relationship SP', CP, Profit	$SP' = CP + 2\% \text{ of } CP$
Equation	$10200 = CP \times (1 + 0.02)$
Solving for CP	$CP = \frac{10200}{1.02}$
Calculated Cost Price (CP)	Rs. 10,000

Revision Table: Key Terms in Profit and Loss

Your Personal Exams Guide

Term	Definition	How it's Calculated/Used
Cost Price (CP)	The original price at which an article is bought.	Base for calculating profit or loss percentage.
Selling Price (SP)	The price at which an article is sold.	If $SP > CP$, there is a profit. If $SP < CP$, there is a loss.
Marked Price (MP)	The price marked on the article; also known as List Price.	Discounts are usually offered on the MP.
Discount	A reduction in the Marked Price.	Discount = $MP - SP$ (after discount) Discount % = $(\text{Discount} / MP) \times 100$
Profit	The gain made when SP is greater than CP.	Profit = $SP - CP$ Profit % = $(\text{Profit} / CP) \times 100$
Loss	The amount lost when SP is less than CP.	Loss = $CP - SP$ Loss % = $(\text{Loss} / CP) \times 100$

Additional Information on Profit, Loss, and Discount Calculations

Understanding how discount affects the selling price is crucial in these types of problems. A discount is typically a percentage reduction on the Marked Price. The resulting price is the Selling Price for that transaction. Profit or Loss is then determined by comparing this Selling Price to the Cost Price.

- When a discount of $R\%$ is given on MP, the Selling Price (SP) is calculated as:

$$SP = MP \times \left(\frac{100 - R}{100} \right)$$

- When there is a profit of $P\%$ on CP, the Selling Price (SP) is calculated as:

$$SP = CP \times \left(\frac{100 + P}{100} \right)$$

- When there is a loss of L% on CP, the Selling Price (SP) is calculated as:

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

In this problem, we used the first two formulas by equating the hypothetical selling price calculated using the discount on MP to the selling price calculated using the profit on CP.

$$\text{Hypothetical SP} = MP \times (100 - 15)/100 = 12000 \times 0.85 = 10200.$$

$$\text{Hypothetical SP} = CP \times (100 + 2)/100 = CP \times 1.02.$$

$$\text{Equating them: } 10200 = CP \times 1.02.$$

$$\text{Solving for CP: } CP = \frac{10200}{1.02} = 10000.$$

This confirms our step-by-step calculation and reinforces the concepts.

35. Answer: d

Explanation:

Finding the Odd Group of Letters in Alphabet Series

This question asks us to identify the group of letters that is different from the others based on a hidden pattern. To solve this type of reasoning question, we usually look at the position of the letters in the English alphabet.

Let's examine each group of letters and their positions in the alphabet:

Group	Letters	Alphabetical Position	Difference
1	E, F, G	E=5, F=6, G=7	$F - E = 6 - 5 = +1$ $G - F = 7 - 6 = +1$
2	P, Q, R	P=16, Q=17, R=18	$Q - P = 17 - 16 = +1$ $R - Q = 18 - 17 = +1$
3	V, W, X	V=22, W=23, X=24	$W - V = 23 - 22 = +1$ $X - W = 24 - 23 = +1$
4	H, J, L	H=8, J=10, L=12	$J - H = 10 - 8 = +2$ $L - J = 12 - 10 = +2$

From the analysis above, we can see the pattern:

- In groups EFG, PQR, and VWX, the letters are consecutive in the alphabet. The difference in alphabetical position between successive letters is +1.
- In the group HJL, the letters are not consecutive. The difference in alphabetical position between successive letters is +2.

The group HJL follows a different pattern compared to the other three groups. Therefore, HJL is the odd group of letters.

Revision Table: Analyzing Letter Patterns

Group	Pattern Observed
EFG	Consecutive letters (position difference +1)
PQR	Consecutive letters (position difference +1)
VWX	Consecutive letters (position difference +1)
HJL	Letters separated by one (position difference +2)

Additional Information: Letter Series Reasoning

Letter series questions are common in reasoning tests. They test your ability to identify patterns in sequences of letters. These patterns can be based on various rules:

- **Alphabetical Position:** Analyzing the position of letters (A=1, B=2, ..., Z=26).

- **Difference in Position:** Looking at the numerical difference between the positions of consecutive letters. This difference can be constant, increasing, decreasing, or follow a specific sequence.
- **Vowels/Consonants:** Patterns might involve the arrangement of vowels and consonants.
- **Skipping Letters:** Letters might be skipped in a consistent manner (e.g., skipping one letter, two letters, etc.).
- **Reverse Alphabetical Order:** The series might run backward through the alphabet.

To solve letter series problems, always start by writing down the alphabetical positions of the letters and calculating the differences between them. This often reveals the underlying pattern.

36. Answer: a

Explanation:

Calculating Average Speed of a Car

The problem asks us to find the average speed of a car in kilometers per hour (km/hr), given the distance it covers and the time taken. We are given the distance in meters (m) and the time in seconds (s).

Understanding Average Speed

Average speed is defined as the total distance covered divided by the total time taken. The formula is:

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

In this question, the given values are:

- Total Distance = 400 m
- Total Time = 20 seconds

Step-by-Step Calculation in m/s

First, let's calculate the average speed in meters per second (m/s) using the formula:

$$\text{Average Speed} = \frac{400 \text{ m}}{20 \text{ s}}$$

$$\text{Average Speed} = 20 \text{ m/s}$$

So, the average speed of the car is 20 m/s.

Converting Speed from m/s to km/hr

The question requires the answer in kilometers per hour (km/hr). We need to convert the speed from m/s to km/hr. The conversion factor is:

- 1 kilometer (km) = 1000 meters (m)
- 1 hour (hr) = 60 minutes = 60 × 60 seconds = 3600 seconds (s)

Therefore, to convert m/s to km/hr, we can multiply the speed in m/s by $\frac{\text{Distance conversion factor}}{\text{Time conversion factor}}$.
Let's derive the conversion:

$$1 \text{ m/s} = \frac{1 \text{ m}}{1 \text{ s}}$$

To convert meters to kilometers, we divide by 1000 (or multiply by $\frac{1}{1000}$ km/m).

To convert seconds to hours, we divide by 3600 (or multiply by $\frac{1}{3600}$ hr/s).

$$1 \text{ m/s} = \frac{1 \text{ m}}{1 \text{ s}} = \frac{\frac{1}{1000} \text{ km}}{\frac{1}{3600} \text{ hr}} = \frac{1}{1000} \times 3600 \text{ km/hr} = \frac{3600}{1000} \text{ km/hr} = \frac{36}{10} \text{ km/hr} = \frac{18}{5} \text{ km/hr}$$

So, to convert speed from m/s to km/hr, we multiply by $\frac{18}{5}$.

Final Speed Calculation in km/hr

Now, we convert the calculated speed of 20 m/s to km/hr:

$$\text{Average Speed (km/hr)} = \text{Average Speed (m/s)} \times \frac{18}{5}$$

$$\text{Average Speed (km/hr)} = 20 \times \frac{18}{5}$$

We can simplify the calculation:

$$\text{Average Speed (km/hr)} = (20 \div 5) \times 18$$

$$\text{Average Speed (km/hr)} = 4 \times 18$$

$$\text{Average Speed (km/hr)} = 72 \text{ km/hr}$$

Thus, the average speed of the car is 72 km/hr.

Quantity	Value	Unit
Distance	400	m
Time	20	s
Speed (calculated)	20	m/s
Speed (converted)	72	km/hr

Revision Table: Important Formulas

Concept	Formula
Average Speed	$\frac{\text{Total Distance}}{\text{Total Time}}$
Conversion m/s to km/hr	Multiply by $\frac{18}{5}$
Conversion km/hr to m/s	Multiply by $\frac{5}{18}$

Additional Information: Speed, Velocity, and Units

Understanding speed and its units is fundamental in physics. Here are some related concepts:

- **Speed vs. Velocity:** Speed is a scalar quantity that measures how fast an object is moving. Velocity is a vector quantity that includes both speed and direction. Average speed is the total distance over total time, while average velocity is the total displacement over total time.
- **Instantaneous Speed:** This is the speed of an object at a particular moment in time, as opposed to average speed which is over a duration.
- **Common Units:** The standard SI unit for speed is meters per second (m/s). Other common units include kilometers per hour (km/hr), miles per hour (mph), and feet per second (ft/s). Being able to convert between these units is a key skill.
- **Distance and Displacement:** Distance is the total path length covered by an object. Displacement is the straight-line distance between the starting and ending points, including direction. For calculating average speed, we use the total distance.

37. Answer: b

Explanation:

Understanding Resistors in Series and Current Calculation

When resistors are connected in series, the current flowing through each resistor is the same. To find this current, we first need to calculate the total resistance of the series combination. Then, we can use Ohm's Law to find the total current supplied by the battery, which is also the current through each resistor in the series circuit.

Calculating Total Resistance in a Series Circuit

For resistors connected in series, the total resistance is the sum of the individual resistances. In this problem, we have two resistors with resistances of $R_1 = 2\ \Omega$ and $R_2 = 6\ \Omega$ connected in series.

The formula for total resistance R_{total} in series is:

$$R_{total} = R_1 + R_2 + R_3 + \dots$$

Substituting the given values:

$$R_{total} = 2\ \Omega + 6\ \Omega$$

$$R_{total} = 8\ \Omega$$

Resistor	Resistance
Resistor 1	$2\ \Omega$
Resistor 2	$6\ \Omega$
Total Resistance (Series)	$8\ \Omega$

Applying Ohm's Law to Find Total Current

Ohm's Law states the relationship between voltage (V), current (I), and resistance (R):

$$V = I \times R$$

We can rearrange this formula to find the current:

$$I = \frac{V}{R}$$

In this circuit, the total voltage across the series combination is the battery voltage, $V = 12\text{ V}$. The total resistance of the circuit is $R_{total} = 8\ \Omega$.

Using Ohm's Law to find the total current I_{total} from the battery:

$$I_{total} = \frac{V}{R_{total}}$$

$$I_{total} = \frac{12\text{ V}}{8\ \Omega}$$

$$I_{total} = 1.5\text{ A}$$

Current Distribution in a Series Circuit

A fundamental property of a series circuit is that the current is the same at all points in the circuit. This means the total current flowing from the battery flows through every component connected in series, including both the $2\ \Omega$ resistor and the $6\ \Omega$ resistor.

Therefore, the current in the $6\ \Omega$ resistor is equal to the total current calculated:

$$I_{6\ \Omega} = I_{total} = 1.5\text{ A}$$

Conclusion

The current flowing through the $6\ \Omega$ resistor when connected in series with a $2\ \Omega$ resistor across a 12 V battery is 1.5 A .

Revision Table: Series Circuit Analysis

Concept	Rule/Formula	Application
Total Resistance (Series)	$R_{total} = R_1 + R_2 + \dots$	$2\Omega + 6\Omega = 8\Omega$
Ohm's Law	$V = I \times R$ or $I = V/R$	$I_{total} = 12V/8\Omega = 1.5A$
Current in Series	Current is same through all components.	Current in 6Ω resistor = I_{total}

Additional Information: Series vs. Parallel Circuits

Understanding the difference between series and parallel circuits is crucial for solving electrical problems.

- **Series Circuits:**
 - Components are connected along a single path.
 - The current is the same through all components.
 - The total resistance is the sum of individual resistances ($R_{total} = R_1 + R_2 + \dots$).
 - The total voltage across the circuit is the sum of the voltage drops across each component ($V_{total} = V_1 + V_2 + \dots$).
- **Parallel Circuits:**
 - Components are connected across each other, forming multiple paths.
 - The voltage is the same across all components.
 - The reciprocal of the total resistance is the sum of the reciprocals of individual resistances ($1/R_{total} = 1/R_1 + 1/R_2 + \dots$).
 - The total current from the source is the sum of the currents through each branch ($I_{total} = I_1 + I_2 + \dots$).

This problem specifically deals with a series circuit, where the current is uniform throughout.

38. Answer: d

Explanation:

Finding the Reflection of a Point on the Y-axis

Understanding reflections in geometry is a fundamental concept. When a point is reflected across an axis, its position changes in a predictable way based on that axis.

Rule for Reflection on the Y-axis

The reflection of a point (x, y) on the Y-axis results in a new point with coordinates $(-x, y)$. This means the x-coordinate changes its sign, while the y-coordinate remains unchanged.

Applying the Rule to Point $(-2, -6)$

We are given the point $(-2, -6)$. To find its reflection on the Y-axis, we apply the rule $(x, y) \rightarrow (-x, y)$:

- The original x-coordinate is -2 . Changing its sign gives $-(-2) = 2$.
- The original y-coordinate is -6 . This remains the same.

Therefore, the reflection of the point $(-2, -6)$ on the Y-axis is $(2, -6)$.

Comparing with the Options

Let's look at the given options:

- Option 1: $(2, 6)$ – Incorrect, the y-coordinate changed sign.
- Option 2: $(-6, -2)$ – Incorrect, coordinates are swapped and signs are wrong.
- Option 3: $(-2, 6)$ – Incorrect, only the y-coordinate changed sign.
- Option 4: $(2, -6)$ – Correct, the x-coordinate's sign changed, and the y-coordinate remained the same.

The calculated reflection $(2, -6)$ matches Option 4.

Revision Table: Summary of Reflections

Original Point	Reflection Axis	Reflected Point
(x, y)	X-axis	$(x, -y)$
(x, y)	Y-axis	$(-x, y)$
(x, y)	Origin $(0, 0)$	$(-x, -y)$
(x, y)	Line $y = x$	(y, x)

Additional Information on Geometric Transformations

Reflection is one type of geometric transformation. Other common transformations include translation (sliding the point without changing orientation), rotation (turning the point around a fixed point), and dilation (resizing the shape). Each transformation follows specific rules for changing the coordinates of a point.

Understanding these rules is crucial for solving problems involving coordinate geometry and transformations.

39. Answer: a

Explanation:

Understanding the Jumbled Letters and Numbers

We are given the following letters and their corresponding numbers:

Letter	Number
H	1
T	2
R	3
U	4
O	5
A	6

We need to check each option, which provides a sequence of numbers. For each number sequence, we will pick the letter corresponding to that number and arrange them in the given order to see if it forms a meaningful English word.

Analyzing the Options for Meaningful Word Formation

Let's examine each option:

- Option 1: 6, 4, 2, 1, 5, 3

Following this sequence of numbers:

- Number 6 corresponds to letter A
- Number 4 corresponds to letter U
- Number 2 corresponds to letter T
- Number 1 corresponds to letter H
- Number 5 corresponds to letter O
- Number 3 corresponds to letter R
- Option 2: 3, 4, 5, 2, 1, 6

Following this sequence of numbers:

- Number 3 corresponds to letter R
- Number 4 corresponds to letter U
- Number 5 corresponds to letter O

- Number 2 corresponds to letter T
- Number 1 corresponds to letter H
- Number 6 corresponds to letter A
- **Option 3: 1, 6, 2, 4, 5, 3**

Following this sequence of numbers:

- Number 1 corresponds to letter H
- Number 6 corresponds to letter A
- Number 2 corresponds to letter T
- Number 4 corresponds to letter U
- Number 5 corresponds to letter O
- Number 3 corresponds to letter R
- **Option 4: 2, 1, 5, 3, 4, 6**

Following this sequence of numbers:

- Number 2 corresponds to letter T
- Number 1 corresponds to letter H
- Number 5 corresponds to letter O
- Number 3 corresponds to letter R
- Number 4 corresponds to letter U
- Number 6 corresponds to letter A

Confirming the Correct Combination

Based on our analysis, only the sequence from Option 1 (6, 4, 2, 1, 5, 3) forms a meaningful English word, which is "AUTHOR".

Therefore, the combination of numbers 6, 4, 2, 1, 5, 3 is the correct one.

Revision Table: Jumbled Letters to Meaningful Word

Number Sequence (Option)	Letter Sequence	Formed Word	Meaningful English Word?
6, 4, 2, 1, 5, 3 (Option 1)	A, U, T, H, O, R	AUTHOR	Yes
3, 4, 5, 2, 1, 6 (Option 2)	R, U, O, T, H, A	RUOTHA	No
1, 6, 2, 4, 5, 3 (Option 3)	H, A, T, U, O, R	HATUOR	No
2, 1, 5, 3, 4, 6 (Option 4)	T, H, O, R, U, A	THORUA	No

Additional Information: Solving Letter Arrangement Problems

Problems involving jumbled letters and number codes test your vocabulary and ability to follow instructions based on a code. Here are some tips for solving such problems:

- Carefully note down the letter corresponding to each number.
- For each option, write down the sequence of letters by mapping the numbers to the letters.
- Read the resulting letter sequence to see if it forms a recognizable English word.
- If a word is formed, check if it is indeed a meaningful word you know.
- If you are unsure about a word, you can try to think of common words that could be formed by the given set of letters. In this case, the letters are H, T, R, U, O, A.

40. Answer: a

Explanation:

Solving the Pipe and Cistern Problem: Finding Time for Tap M

This problem involves calculating the time taken by individual taps (or pipes) to fill a cistern, based on their combined and individual filling rates. These types of questions often use the concept of 'work rate', where the rate is the amount of work done (filling the cistern) per unit of time.

Understanding Work Rate

The rate at which a tap fills a cistern is the reciprocal of the time it takes to fill the entire cistern. If a tap takes T minutes to fill a cistern, its rate is $1/T$ cistern per minute. When multiple taps work together, their individual rates are added to find the combined rate.

Given Information

- Time taken by Tap M and Tap N together to fill the cistern = $\frac{48}{13}$ minutes.
- Time taken by Tap N alone to fill the cistern = 6 minutes.

Finding the Rates

Using the concept of work rate:

- Combined rate of Tap M and Tap N = $\frac{1}{\text{Time taken together}} = \frac{1}{\frac{48}{13}} = \frac{13}{48}$ cistern per minute.
- Rate of Tap N alone = $\frac{1}{\text{Time taken by N alone}} = \frac{1}{6}$ cistern per minute.

Calculating the Rate of Tap M

The combined rate of Tap M and Tap N is the sum of their individual rates:

Rate of M + Rate of N = Combined Rate of M and N

Let the rate of Tap M be R_M . We have:

$$R_M + \frac{1}{6} = \frac{13}{48}$$

To find R_M , we subtract the rate of Tap N from the combined rate:

$$R_M = \frac{13}{48} - \frac{1}{6}$$

To subtract these fractions, we need a common denominator. The least common multiple of 48 and 6 is 48.

Convert $\frac{1}{6}$ to a fraction with a denominator of 48:

$$\frac{1}{6} = \frac{1 \times 8}{6 \times 8} = \frac{8}{48}$$

Now, perform the subtraction:

$$R_M = \frac{13}{48} - \frac{8}{48} = \frac{13-8}{48} = \frac{5}{48}$$

So, the rate of Tap M is $\frac{5}{48}$ cistern per minute.

Finding the Time Taken by Tap M

The time taken by Tap M alone to fill the cistern is the reciprocal of its rate:

$$\text{Time taken by M} = \frac{1}{\text{Rate of M}} = \frac{1}{\frac{5}{48}} = \frac{48}{5} \text{ minutes.}$$

Converting to Decimal

To express this time as a decimal, divide 48 by 5:

$$\frac{48}{5} = 9.6$$

So, Tap M alone will take 9.6 minutes to fill the cistern.

Summary of Calculation Steps

Step	Description	Calculation
1	Combined Rate (M + N)	$\frac{1}{\frac{13}{48}} = \frac{13}{48}$
2	Rate of N	$\frac{1}{6}$
3	Rate of M (Combined Rate - Rate of N)	$\frac{13}{48} - \frac{1}{6} = \frac{13}{48} - \frac{8}{48} = \frac{5}{48}$
4	Time for M (1 / Rate of M)	$\frac{1}{\frac{5}{48}} = \frac{48}{5} = 9.6$

The time taken by Tap M alone to fill the cistern is 9.6 minutes.

Revision Table: Pipe and Cistern Concepts

Concept	Explanation	Formula/Relationship
Work Rate	Amount of work done per unit of time. For filling a cistern, it's the fraction of the cistern filled per minute.	$\text{Rate} = \frac{1}{\text{Time Taken}}$
Time Taken	The total time required to complete the work (fill the cistern).	$\text{Time} = \frac{1}{\text{Rate}}$
Combined Rate	When multiple taps work together, their individual rates add up.	$R_{\text{total}} = R_1 + R_2 + R_3 + \dots$
Net Rate	When filling taps and emptying taps work simultaneously, the net rate is the sum of filling rates minus the sum of emptying rates.	$R_{\text{net}} = R_{\text{filling}} - R_{\text{emptying}}$

Additional Information: Solving Time and Work Problems

Pipe and cistern problems are a specific type of time and work problem. The key idea is to convert the time taken into a rate of work (usually per minute or per hour). Here are some related points:

- **Units:** Ensure consistent units for time (e.g., all in minutes or all in hours).
- **Part of Work:** If a tap works for t minutes at a rate of R per minute, the part of the cistern filled is $R \times t$.
- **Emptying Pipes:** If a pipe empties a cistern, its rate is considered negative because it subtracts from the total work done (filling).
- **Algebraic Approach:** Setting up equations based on rates is a common way to solve these problems, especially when unknown times or rates are involved.
- **LCM Method:** Sometimes, for integer times, assuming the total capacity of the cistern is the Least Common Multiple (LCM) of the individual times can simplify calculations by working with integers representing units of work.

Understanding the reciprocal relationship between time and rate is fundamental to solving these problems efficiently and accurately.

41. Answer: c

Explanation:

Understanding Tides and Gravitational Effects

The question asks about the primary cause behind the rise and fall of sea levels, known as tides, on the rotating Earth. This phenomenon is mainly influenced by the gravitational pull of celestial bodies.

The gravitational force exerted by any object decreases with distance. While many celestial bodies exert a gravitational pull on Earth, two are particularly significant regarding tides: the Moon and the Sun.

The Moon's Gravitational Influence on Tides

The Moon is much smaller than the Sun, but it is significantly closer to Earth. This proximity means that the difference in the Moon's gravitational pull across the Earth's diameter is quite substantial. This differential gravitational force is the primary driver of tides.

- The side of Earth facing the Moon experiences a stronger gravitational pull, causing the water to bulge towards the Moon (high tide).
- On the opposite side of Earth, the gravitational pull is weaker. The water here is essentially left behind as the solid Earth is pulled more strongly towards the Moon, creating another bulge (high tide).
- Areas between these bulges experience low tide.

As the Earth rotates, different parts of the planet pass through these bulges, resulting in two high tides and two low tides approximately every lunar day (about 24 hours and 50 minutes).

Comparing Moon and Sun's Tidal Effects

The Sun is much more massive than the Moon, but it is also much further away from Earth. Because the gravitational force diminishes rapidly with distance (inversely proportional to the square of the distance), the differential gravitational force exerted by the Sun across Earth's diameter is less pronounced than that of the Moon.

While the Sun does influence tides (contributing to spring tides when aligned with the Moon and neap tides when at a right angle), the Moon's gravitational effect is approximately twice as strong as the Sun's in generating tides.

Celestial Body	Distance from Earth	Gravitational Effect on Tides
Moon	Much closer	Primary cause (Differential force is strong)
Sun	Much farther	Secondary influence (Differential force is weaker)
Venus	Varies, but much farther than Moon	Negligible tidal effect
Mercury	Varies, but much farther than Moon	Negligible tidal effect

Evaluating Other Options

The other options, Venus and Mercury, are planets in our solar system. While they do exert gravitational forces on Earth, their distances are significantly greater than the Moon's, and their masses are not large enough to cause any noticeable tidal effects on Earth's oceans.

Therefore, the main reason for the regular rise and fall of sea level (tides) is the gravitational interaction between the Earth and the Moon.

Conclusion on the Primary Cause of Tides

Based on the analysis of gravitational effects, the Moon's gravitational pull is the predominant factor causing the tides on Earth.

Revision Table: Key Concepts of Tides

Concept	Explanation
Tides	The periodic rise and fall of sea level.
Primary Cause	Gravitational pull of the Moon.
Secondary Cause	Gravitational pull of the Sun.
Differential Gravity	The difference in gravitational pull across Earth's diameter, creating tidal bulges.

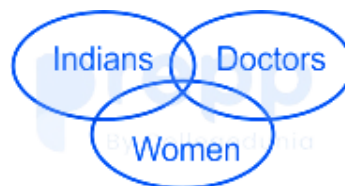
Additional Information on Tidal Forces

The force that causes tides is not just the direct gravitational attraction but the difference in gravitational pull from one side of the Earth to the other. This is why it is often referred to as a "tidal force" or "differential gravitational force." The Moon pulls more strongly on the side of Earth closest to it and less strongly on the side farthest away. This difference stretches the Earth slightly, particularly the liquid oceans, creating the bulges responsible for high tides.

42. Answer: d

Explanation:

The Venn diagram that best represents the relationship between Indians, doctors and women is shown below:



Some Indians are women, some women are doctors and some Doctors are Indians.

Hence, 'option 4' is the correct answer.

43. Answer: b

Explanation:

Understanding Weight on Earth vs. The Moon

The question asks how the weight of an object changes when it is moved from the Earth to the Moon. To answer this, we need to understand what weight is and how it is determined.

What is Weight?

Weight is the force exerted on an object due to gravity. It is dependent on the mass of the object and the acceleration due to gravity at the object's location.

The formula for weight (W) is:

$$W = m \times g$$

Where:

- m is the mass of the object (which is an intrinsic property and remains constant regardless of location).
- g is the acceleration due to gravity at the specific location (like Earth or the Moon).

Gravity on Earth and the Moon

The acceleration due to gravity (g) is different on different celestial bodies because it depends on the mass and radius of the body. The Moon has significantly less mass than the Earth.

- The approximate acceleration due to gravity on Earth (g_{earth}) is 9.8 m/s^2 .
- The approximate acceleration due to gravity on the Moon (g_{moon}) is about 1.62 m/s^2 .

This means that the Moon's gravitational pull is much weaker than Earth's. Specifically, the gravity on the Moon is approximately one-sixth of the gravity on Earth.

$$g_{\text{moon}} \approx \frac{1}{6} g_{\text{earth}}$$

Comparing Weight on Earth and Moon

Let the mass of the object be m . The weight of the object on Earth (W_{earth}) is given as X units.

$$W_{earth} = m \times g_{earth} = X$$

The weight of the same object on the Moon (W_{moon}) will be:

$$W_{moon} = m \times g_{moon}$$

Since g_{moon} is less than g_{earth} (approximately $\frac{1}{6}$ of g_{earth}), the weight on the Moon will be less than the weight on Earth.

$$W_{moon} = m \times g_{moon} < m \times g_{earth}$$

Therefore, $W_{moon} < X$.

Specifically, the weight on the Moon is approximately one-sixth of the weight on Earth:

$$W_{moon} \approx \frac{1}{6} W_{earth} = \frac{1}{6} X$$

Analyzing the Options

Let's look at the given options for the weight of the object on the Moon:

1. Zero: This is incorrect. There is still gravity on the Moon, so the object will have weight, just less than on Earth.
2. Less than X : This aligns with our conclusion that the weight on the Moon is less than the weight on Earth because the Moon's gravity is weaker.
3. Equal to X : This is incorrect. Weight depends on gravity, which is different on the Moon than on Earth.
4. More than X : This is incorrect. The Moon's gravity is weaker, not stronger, than Earth's gravity.

Based on the analysis, the weight of the object on the Moon will be less than its weight on Earth (X units).

Revision Table: Weight Comparison

Property	Earth	Moon	Notes
Mass of Object (m)	Same	Same	Mass is an intrinsic property, does not change with location.
Acceleration due to Gravity (g)	Approx. 9.8 m/s ²	Approx. 1.62 m/s ²	Gravity is significantly weaker on the Moon.
Weight (W = m × g)	X units	Less than X units	Since 'g' is less on the Moon, weight is also less.

Additional Information: Mass vs. Weight

It is important not to confuse mass and weight. They are related but distinct concepts in physics:

- **Mass:** Mass is a measure of the amount of matter in an object. It is a scalar quantity and is constant regardless of the location or gravitational pull. Measured in kilograms (kg) in the SI system.
- **Weight:** Weight is the force of gravity acting on an object's mass. It is a vector quantity (force has direction) and varies depending on the acceleration due to gravity at the location. Measured in Newtons (N) in the SI system. While the question uses "units" for weight, in physics, weight is a force measured in units like Newtons or pounds.

So, while the mass of the object remains the same on the Moon, its weight changes because the gravitational pull is different.

44. Answer: c

Explanation:

Understanding RK Narayan's Famous Books

RK Narayan (Rasipuram Krishnaswami Iyer Narayanaswami) was one of the greatest figures of early Indian literature in English. He is celebrated for his simple writing style and for bringing Indian life and culture to a global audience, primarily through his stories set in the fictional South Indian town of Malgudi.

The question asks to identify a famous book by RK Narayan from the given options. Let's examine each option:

- **The Room on the Roof:** This is a well-known novel, but it is written by Ruskin Bond, another famous Indian author.
- **Two Livers:** This book is not one of RK Narayan's widely recognized or famous works.
- **Malgudi Days:** This is a collection of short stories by RK Narayan. It is arguably one of his most famous and beloved works, introducing readers to the vibrant life and characters of Malgudi. Many of these stories were also adapted into a popular Indian television series.
- **A Suitable Boy:** This is a very famous and lengthy novel, but it was written by Vikram Seth.

Based on the authorship and fame of the books listed, "Malgudi Days" stands out as a highly famous work by RK Narayan. His other famous novels include *Swami and Friends*, *The Bachelor of Arts*, *The English Teacher*, and *The Guide* (which won the Sahitya Akademi Award). However, "Malgudi Days" remains one of his most iconic and widely recognized works.

Here is a summary of the books and their authors from the options:

Book Title	Author	Is it RK Narayan's Famous Book?
The Room on the Roof	Ruskin Bond	No
Two Livers	Not a famous work by RK Narayan	No
Malgudi Days	RK Narayan	Yes, Very Famous
A Suitable Boy	Vikram Seth	No

Therefore, the book famous for being written by RK Narayan among the given options is "Malgudi Days".

Revision Table: Key Indian Authors and Works

Author	Famous Works	Genre/Focus
RK Narayan	Malgudi Days, Swami and Friends, The Guide	Fiction, Stories set in Malgudi
Ruskin Bond	The Room on the Roof, Our Trees Still Grow in Dehra	Fiction, Short Stories, Children's Literature
Vikram Seth	A Suitable Boy, The Golden Gate	Novels, Poetry
Mulk Raj Anand	Untouchable, Coolie	Social issues, Novels
Anita Desai	Clear Light of Day, Fasting, Feasting	Novels, Psychological themes

Additional Information on RK Narayan and Malgudi Days

RK Narayan's creation of Malgudi is considered one of his most significant contributions to literature. Malgudi is not a real place but a fictional town that captures the essence of South Indian life with its diverse characters, traditions, and everyday struggles.

- **Malgudi Days Collection:** The collection "Malgudi Days" includes many individual stories that had been published earlier. It gained immense popularity, especially after the television adaptation in the late 1980s, directed by Shankar Nag.
- **Literary Style:** Narayan's writing is known for its simplicity, gentle humor, and realistic portrayal of human nature. He often wrote about ordinary people dealing with ordinary lives, making his stories relatable and timeless.
- **Awards:** Besides the Sahitya Akademi Award for "The Guide," RK Narayan received numerous other honors, including the Padma Vibhushan, for his literary contributions.

Studying RK Narayan's works, especially "Malgudi Days," offers valuable insights into Indian culture and the art of storytelling.

45. Answer: a

Explanation:

Understanding Molar Specific Heat Capacity

Molar specific heat capacity is a physical property of a substance that describes how much heat energy is required to raise the temperature of one mole of that substance by one degree Celsius (or Kelvin). It is a crucial concept in thermodynamics, particularly when dealing with gases or other substances where the amount is conveniently measured in moles.

The amount of heat energy (ΔQ) absorbed or released by a substance is directly proportional to the number of moles (μ) of the substance, the change in temperature (ΔT), and its molar specific heat capacity (C). This relationship can be expressed by the formula:

$$\Delta Q = \mu C \Delta T$$

Where:

- ΔQ is the heat energy absorbed or released.
- μ is the number of moles of the substance.
- C is the molar specific heat capacity of the substance.
- ΔT is the change in temperature.

Our goal is to find the expression for the molar specific heat capacity, C . We can rearrange the formula $\Delta Q = \mu C \Delta T$ to solve for C .

Divide both sides of the equation by μ and ΔT :

$$C = \frac{\Delta Q}{\mu \Delta T}$$

This can also be written as:

$$C = \left(\frac{1}{\mu} \right) \left(\frac{\Delta Q}{\Delta T} \right)$$

This formula shows that the molar specific heat capacity is the heat energy required per mole per unit change in temperature.

Analyzing the Given Options

Let's compare the derived formula for molar specific heat capacity with the given options:

Option	Expression	Comparison
1	$\left(\frac{1}{\mu}\right) \left(\frac{\Delta Q}{\Delta T}\right)$	Matches the derived formula for molar specific heat capacity.
2	$\mu \left(\frac{\Delta Q}{\Delta T}\right)$	This represents $\mu^2 C$ if $\Delta Q = \mu C \Delta T$. Incorrect.
3	$\left(\frac{1}{\mu}\right) \left(\frac{\Delta T}{\Delta Q}\right)$	This is the reciprocal of the molar specific heat capacity multiplied by $\frac{1}{\mu^2}$. Incorrect.
4	$\mu \left(\frac{\Delta T}{\Delta Q}\right)$	This is related to the reciprocal of molar specific heat capacity. Incorrect.

Based on the comparison, the expression $\left(\frac{1}{\mu}\right) \left(\frac{\Delta Q}{\Delta T}\right)$ correctly represents the molar specific heat capacity of a substance.

Revision Table: Key Thermal Concepts

Concept	Definition	Formula	Units
Heat Capacity (H)	Heat required to raise the temperature of a substance by 1 degree. Depends on mass/moles.	$H = \frac{\Delta Q}{\Delta T}$	J/K or J/°C
Specific Heat Capacity (c)	Heat required to raise the temperature of 1 kg of a substance by 1 degree. Independent of mass.	$c = \left(\frac{1}{m}\right) \left(\frac{\Delta Q}{\Delta T}\right)$	J/(kg·K) or J/(kg·°C)
Molar Specific Heat Capacity (C)	Heat required to raise the temperature of 1 mole of a substance by 1 degree. Independent of moles.	$C = \left(\frac{1}{\mu}\right) \left(\frac{\Delta Q}{\Delta T}\right)$	J/(mol·K) or J/(mol·°C)

Additional Information on Specific Heat

Specific heat capacities (both specific and molar) can depend on the conditions under which the heat transfer occurs, especially for gases. For gases, we often distinguish between molar specific heat at constant volume (C_V) and molar specific heat at constant pressure (C_P).

- **Constant Volume (C_V):** When the volume of the gas is kept constant, all the heat added goes into increasing the internal energy of the gas, thus raising its temperature.
- **Constant Pressure (C_P):** When the pressure of the gas is kept constant, the gas expands as its temperature increases. The heat added not only increases the internal energy but also does work on the surroundings during expansion. Therefore, more heat is required to achieve the same temperature change compared to the constant volume case, meaning C_P is generally greater than C_V .

The relationship between C_P and C_V for an ideal gas is given by Mayer's relation: $C_P - C_V = R$, where R is the ideal gas constant (approximately $8.314 \text{ J/(mol}\cdot\text{K)}$).

For solids and liquids, the distinction between C_P and C_V is usually negligible because their volume changes very little with temperature.

46. Answer: c

Explanation:

Understanding Blood Relations from Symbols

This question asks us to decode a symbolic representation of blood relations and determine the relationship between two individuals, W and Z, based on the given expression $W \$ X \& Y \% Z$.

Decoding the Given Symbols

Let's first understand what each symbol means:

- $C \$ D$ means **C is daughter of D**. This tells us D is the parent of C, and C is female.
- $C \& D$ means **C is mother of D**. This tells us C is the parent of D, and C is female.
- $C \% D$ means **C is son of D**. This tells us D is the parent of C, and C is male.

Breaking Down the Expression: $W \$ X \& Y \% Z$

Now, let's apply these rules to the expression $W \$ X \& Y \% Z$, breaking it down into parts:

1. $W \$ X$: According to the first rule, this means **W is daughter of X**. This establishes X as a parent of W, and W is female.
2. $X \& Y$: According to the second rule, this means **X is mother of Y**. This establishes X as a parent of Y, and X is female.
3. $Y \% Z$: According to the third rule, this means **Y is son of Z**. This establishes Z as a parent of Y, and Y is male.

Connecting the Relations

From the breakdown, we have the following facts:

- W is daughter of X.
- X is mother of Y.
- Y is son of Z.

From 'X is mother of Y' and 'Y is son of Z', we know that Y is the child of both X and Z. Since X is the mother and Z is the other parent, Z must be the father of Y. Thus, X and Z are the parents of Y, and they are a couple where X is the mother and Z is the father.

Now consider W. W is the daughter of X. Since X and Z are a couple and are parents to Y, and W is also a child of X, W must also be a child of Z. Because W is female ('daughter of X'), and Z is the father, Z is the father of W.

Summary of Relationships

Relation	Meaning	Deduction
$W \$ X$	W is daughter of X	X is W's parent, W is female
$X \& Y$	X is mother of Y	X is Y's parent, X is female
$Y \% Z$	Y is son of Z	Z is Y's parent, Y is male

Combining the deductions:

- X is the mother of Y.

- Z is the parent of Y, and Y is male. Therefore, Z is the father of Y.
- X and Z are parents of Y, with X being the mother and Z being the father. X and Z are a couple.
- W is the daughter of X. Since X and Z are a couple, and X is W's mother, Z must be W's father.

Conclusion: Relation between W and Z

Based on our analysis, Z is the father of W.

Evaluating the Options

Let's compare our conclusion with the given options:

1. Z is wife of W (Incorrect – Z is male (father), W is female (daughter). A father cannot be the wife of his daughter).
2. Z is mother of W (Incorrect – Z is male (father), W is female (daughter). A father cannot be the mother).
3. Z is father of W (Correct – This matches our deduction).
4. Z is daughter of W (Incorrect – Z is the parent (father) of W, not the other way around).

The relationship deduced from the expression $W \$ X \& Y \% Z$ is that Z is the father of W.

Revision Table: Blood Relation Symbols

Symbol	Meaning	Gender of Person before Symbol
\$	is daughter of	Female
&	is mother of	Female
%	is son of	Male

Additional Information: Solving Blood Relation Puzzles

Solving blood relation puzzles involving symbols or codes requires careful step-by-step decoding. It's often helpful to draw a family tree or make notes to keep track of the relationships and genders as you decode each part of the expression. Remember to

identify the gender of individuals whenever possible from the symbol definitions. When combining relations, look for common individuals (like X and Y in this case) to link different parts of the family tree. Pay close attention to which person is mentioned before and after the symbol, as the order is crucial for determining the correct relationship.

47. Answer: c

Explanation:

Understanding Density in Non-Homogeneous Bodies

The question asks about the factor that determines the density of a body if it is NOT homogeneous. Let's first understand what homogeneity means in the context of materials.

A **homogeneous body** is one where the physical and chemical properties are the same throughout the entire volume of the body. This means that if you take a small sample from any part of a homogeneous body, it will have the exact same density, composition, and other properties as a sample taken from any other part.

In a homogeneous material, density (ρ) is constant everywhere. It does not change from one point to another within the body.

What is a Non-Homogeneous Body?

A body that is **NOT homogeneous** (also called a heterogeneous body) is one where the properties, including density, are not uniform throughout its volume. The density can vary from point to point within the body. Think of a material that might be denser at the bottom than at the top, or denser in the center than near the edges.

Density as a Function of Position

Since the density in a non-homogeneous body is different at different points, its value depends on where you are measuring it within the body. The location within the body can be described by its coordinates (position). Therefore, in a non-homogeneous body, density is a function of its position.

If we use Cartesian coordinates (x, y, z) to specify a point within the body, the density at that point can be represented as $\rho(x, y, z)$. This indicates that the density ρ changes as the

coordinates (x, y, z) change.

Analyzing the Options

Let's consider why the other options are not the primary factors for the inherent density variation within a non-homogeneous body:

- **Pressure:** External pressure can affect the density of materials, especially fluids and gases, and to a lesser extent, solids. However, the question is about the inherent property of a non-homogeneous body. The density variation in a non-homogeneous solid body is due to its internal composition or structure varying with location, not primarily due to external or internal pressure differences unless significant forces are involved (like gravity in very large bodies). For a given non-homogeneous solid at rest under normal conditions, the density variation is fundamentally tied to position.
- **Acceleration:** Acceleration is related to changes in velocity. While acceleration forces (like gravity or inertial forces in a non-inertial frame) can induce pressure gradients that might slightly affect density (especially in fluids), the non-homogeneity itself is about the material's composition or structure varying with position, not its state of motion or acceleration.
- **Velocity:** Velocity describes the rate of change of position. The velocity of a body (or parts of it) does not determine its material density at a given point. While high velocities might involve relativistic effects or significant forces that influence density, the fundamental non-homogeneity is an intrinsic property related to the material's makeup at different locations, not its speed.

Thus, the density variation within a non-homogeneous body is fundamentally a function of its location or position within the body.

Summary of Concepts

Body Type	Density (ρ)	Dependence
Homogeneous	Constant	Independent of Position
Non-Homogeneous (Heterogeneous)	Varies	Function of Position $\rho(x, y, z)$

Therefore, if a body is NOT homogeneous, its density is a function of its position.

Revision Table: Non-Homogeneous Density Factors

Factor	Influence on Density	Relevance to Non-Homogeneity
Position	Directly varies with location in non-homogeneous bodies.	Defines non-homogeneity. Density $\rho(x, y, z)$.
Pressure	Can affect density (compressibility), but not the primary cause of inherent non-homogeneity in composition/structure.	Secondary effect, not the core definition of material non-uniformity.
Acceleration	Can create pressure gradients (e.g., gravity), which might slightly affect density, but not the primary cause of inherent material non-homogeneity.	Indirect effect via induced forces/pressure.
Velocity	Generally not a factor determining material density at a point (ignoring extreme relativistic cases or related forces).	Irrelevant to the concept of static material non-uniformity.

Additional Information: Examples of Non-Homogeneous Bodies

Non-homogeneous bodies are very common in the real world. Examples include:

- **Wood:** The density varies depending on the type of wood (hardwood vs. softwood), grain direction, and moisture content, which can differ throughout a piece.
- **Concrete:** Made of cement, aggregates (sand, gravel), and water. The distribution of these components and the presence of air pockets lead to density variations.
- **The Earth's atmosphere or ocean:** Density changes significantly with altitude or depth due to pressure, temperature, and composition variations. While these are influenced by pressure, the variation is tied to position (altitude/depth).
- **Human body:** Different tissues (bone, muscle, fat) have different densities, making the body non-homogeneous overall. The density varies dramatically depending on the location (e.g., density of a bone vs. density of muscle).

These examples illustrate how density depends on the specific location within the material for non-homogeneous objects.

48. Answer: d

Explanation:

Understanding the Problem: Average Velocity with Constant Acceleration

The question asks for the average velocity of an object as it moves from a position of $x = 12.8$ m to $x = 20.0$ m. We are given that the object starts from rest ($v_0 = 0$) at $x = 0$ m and moves with a constant acceleration of $a = 1.6$ m/s² along the x-axis.

Key Concepts for Calculating Average Velocity

Average velocity is defined as the total displacement divided by the total time taken for that displacement. Mathematically, this is:

$$\bar{v} = \frac{\text{Total Displacement}}{\text{Total Time Interval}} = \frac{\Delta x}{\Delta t}$$

Since the object has constant acceleration, we can use the standard kinematic equations to find the time taken and the displacement.

Step-by-Step Solution

To find the average velocity between $x_1 = 12.8$ m and $x_2 = 20.0$ m, we need to calculate:

1. The displacement (Δx) between these two points.
2. The time taken (Δt) to travel from $x_1 = 12.8$ m to $x_2 = 20.0$ m.

Step 1: Calculate the displacement (Δx)

The displacement is the difference between the final and initial positions in the interval:

$$\Delta x = x_2 - x_1$$

Given $x_1 = 12.8$ m and $x_2 = 20.0$ m:

$$\Delta x = 20.0 \text{ m} - 12.8 \text{ m} = 7.2 \text{ m}$$

Step 2: Calculate the time taken (Δt)

To find the time taken to travel from x_1 to x_2 , we first need to find the time it takes to reach x_1 and x_2 starting from $x_0 = 0$ with $v_0 = 0$ and $a = 1.6 \text{ m/s}^2$.

We can use the kinematic equation:

$$x = x_0 + v_0 t + \frac{1}{2} a t^2$$

Since $x_0 = 0$ and $v_0 = 0$, the equation simplifies to:

$$x = \frac{1}{2} a t^2$$

We can rearrange this to solve for t :

$$t^2 = \frac{2x}{a}$$

$$t = \sqrt{\frac{2x}{a}}$$

Now, let's find the time to reach $x_1 = 12.8 \text{ m}$ (let's call it t_1) and $x_2 = 20.0 \text{ m}$ (let's call it t_2). The acceleration is $a = 1.6 \text{ m/s}^2$.

For $x_1 = 12.8 \text{ m}$:

$$t_1 = \sqrt{\frac{2 \times 12.8 \text{ m}}{1.6 \text{ m/s}^2}}$$

$$t_1 = \sqrt{\frac{25.6}{1.6}}$$

$$t_1 = \sqrt{16 \text{ s}^2}$$

$$t_1 = 4 \text{ s}$$

For $x_2 = 20.0 \text{ m}$:

$$t_2 = \sqrt{\frac{2 \times 20.0 \text{ m}}{1.6 \text{ m/s}^2}}$$

$$t_2 = \sqrt{\frac{40.0}{1.6}}$$

$$t_2 = \sqrt{25 \text{ s}^2}$$

$$t_2 = 5 \text{ s}$$

The time interval (Δt) for the journey from x_1 to x_2 is:

$$\Delta t = t_2 - t_1$$

$$\Delta t = 5 \text{ s} - 4 \text{ s} = 1 \text{ s}$$

Step 3: Calculate the average velocity

Now we can calculate the average velocity using the displacement from Step 1 and the time interval from Step 2:

$$\bar{v} = \frac{\Delta x}{\Delta t}$$

$$\bar{v} = \frac{7.2 \text{ m}}{1 \text{ s}}$$

$$\bar{v} = 7.2 \text{ m/s}$$

Thus, the average velocity of the object during its journey from $x = 12.8 \text{ m}$ to $x = 20.0 \text{ m}$ is 7.2 m/s .

Summary of Calculations

Quantity	Symbol	Value	Calculation/Origin
Initial Position (Interval)	x_1	12.8 m	Given
Final Position (Interval)	x_2	20.0 m	Given
Constant Acceleration	a	1.6 m/s^2	Given
Starting Position (Overall)	x_0	0 m	Given
Starting Velocity (Overall)	v_0	0 m/s	Given (starts from rest)
Displacement in Interval	Δx	7.2 m	$x_2 - x_1$
Time to Reach x_1	t_1	4 s	$\sqrt{\frac{2x_1}{a}}$
Time to Reach x_2	t_2	5 s	$\sqrt{\frac{2x_2}{a}}$
Time Interval	Δt	1 s	$t_2 - t_1$
Average Velocity	$\bar{v}$	7.2 m/s	$\frac{\Delta x}{\Delta t}$

Revision Table: Key Kinematic Equations

Equation	Description	Variables
$v = v_0 + at$	Final velocity based on initial velocity, acceleration, and time.	v, v_0, a, t
$x = x_0 + v_0t + \frac{1}{2}at^2$	Position based on initial position, initial velocity, acceleration, and time.	x, x_0, v_0, a, t
$v^2 = v_0^2 + 2a(x - x_0)$	Final velocity based on initial velocity, acceleration, and displacement.	v, v_0, a, x, x_0
$x - x_0 = \frac{1}{2}(v_0 + v)t$	Displacement based on initial velocity, final velocity, and time.	x, x_0, v_0, v, t

Additional Information: Average Velocity vs. Instantaneous Velocity

It's important to distinguish between average velocity and instantaneous velocity.

- **Instantaneous Velocity:** This is the velocity of an object at a specific moment in time. It is the rate of change of position with respect to time at that instant.
- **Average Velocity:** This is the overall velocity of an object over a period of time or a specific displacement. It tells us the average rate of change of position over the entire interval.

In this problem, because the acceleration is constant, the instantaneous velocity is changing continuously. The average velocity over the interval from 12.8 m to 20.0 m is calculated by looking at the total change in position and the total time taken for that specific segment of the journey.

We could also calculate the instantaneous velocity at the start ($x_1 = 12.8$ m) and end ($x_2 = 20.0$ m) of the interval using $v^2 = v_0^2 + 2a(x - x_0)$. We found these earlier:

- At $x_1 = 12.8$ m, $v_1 = 6.4$ m/s.
- At $x_2 = 20.0$ m, $v_2 = 8.0$ m/s.

Note that the average velocity (7.2 m/s) is exactly the average of the instantaneous velocities at the start and end of this specific interval: $(6.4 + 8.0)/2 = 14.4/2 = 7.2$ m/s. This is a property that holds true for motion with constant acceleration.

49. Answer: c

Explanation:

India's Medal Tally at Rio Olympics 2016 Explained

The question asks about the total number of medals won by India at the Rio Olympics held in 2016. India participated in the Rio de Janeiro Olympics with a contingent of athletes across various sports disciplines.

During the Rio Olympics in 2016, India managed to secure a total of two medals.

Breakdown of India's Medals at Rio 2016

The two medals were won in different sports by two prominent Indian athletes:

- One Silver Medal in Badminton
- One Bronze Medal in Wrestling

Let's look at the details of the medal winners:

- **PV Sindhu** won a Silver Medal in Women's Singles Badminton. She became the first Indian woman to win a Silver medal at the Olympics.
- **Sakshi Malik** won a Bronze Medal in Women's Freestyle 58 kg Wrestling. She was the first Indian female wrestler to win an Olympic medal.

These two athletes were the only medalists for India at the Rio 2016 Olympic Games.

Rio Olympics 2016 Performance Review

India's performance at the Rio 2016 Olympics was considered modest compared to expectations. The contingent aimed for a higher medal count, but only PV Sindhu and Sakshi Malik managed to reach the podium. Despite the low medal count, their achievements were significant, breaking new ground in their respective sports and inspiring many.

Revision Table: India's Rio 2016 Medals

Medal	Sport	Event	Athlete
Silver	Badminton	Women's Singles	PV Sindhu
Bronze	Wrestling	Women's Freestyle 58 kg	Sakshi Malik

Additional Information: India at the Olympics

India has a long history of participation in the Olympic Games, starting from 1900. The country has excelled in sports like field hockey in the past. Recent editions have seen individual athletes achieving success in various sports like shooting, wrestling, boxing, badminton, and weightlifting.

- India's best medal tally at any single Olympic Games was 7 medals, achieved at the Tokyo Olympics in 2020 (held in 2021).
- Leander Paes won a Bronze medal in Tennis at the 1996 Atlanta Olympics, which was India's first individual medal since 1952.
- Abhinav Bindra won India's first individual Olympic Gold medal in Shooting (10m Air Rifle) at the 2008 Beijing Olympics.

50. Answer: d

Explanation:

The square root of 225 is 15, because $15 \times 15 = 225$. Thus, the value of x is 150. Multiply the results: $x = 15 \times 10 = 150$.

51. Answer: d

Explanation:

Write the letters in the correct order to form meaningful words:

1) FLOW → WOLF

2) ILNO → LION

3) ERTIG → TIGER

4) WCO → COW

Except for 'COW', all options are carnivorous animals, while 'COW' is a herbivorous animal.

Therefore, 'COW' is unmatched.

52. Answer: c

Explanation:

Calculating the Curved Surface Area of a Hemisphere

Let's find the curved surface area of a hemisphere given its radius. A hemisphere is exactly half of a sphere.

The curved surface area of a sphere with radius r is given by the formula $4\pi r^2$. Since a hemisphere is half a sphere, its curved surface area (excluding the flat base) is half the curved surface area of a sphere.

The formula for the curved surface area (CSA) of a hemisphere is:

$$\text{CSA of Hemisphere} = \frac{1}{2} \times (\text{Curved Surface Area of Sphere})$$

$$\text{CSA of Hemisphere} = \frac{1}{2} \times 4\pi r^2$$

$$\text{CSA of Hemisphere} = 2\pi r^2$$

In this question, we are given:

- Radius of the hemisphere, $r = 7$ cm
- Value of $\pi = \frac{22}{7}$

Now, we substitute these values into the formula for the curved surface area of a hemisphere:

$$\text{CSA} = 2 \times \pi \times r^2$$

$$CSA = 2 \times \frac{22}{7} \times (7 \text{ cm})^2$$

$$CSA = 2 \times \frac{22}{7} \times (7 \times 7) \text{ cm}^2$$

$$CSA = 2 \times \frac{22}{7} \times 49 \text{ cm}^2$$

We can cancel out one '7' from the denominator with one '7' from 49:

$$CSA = 2 \times 22 \times \frac{49}{7} \text{ cm}^2$$

$$CSA = 2 \times 22 \times 7 \text{ cm}^2$$

$$CSA = 44 \times 7 \text{ cm}^2$$

$$CSA = 308 \text{ cm}^2$$

Thus, the curved surface area of the hemisphere is 308 cm^2 .

Curved Surface Area Calculation Steps

Step	Description	Calculation
1	Identify the formula for CSA of hemisphere	$2\pi r^2$
2	Substitute given values ($r=7$, $\pi = 22/7$)	$2 \times \frac{22}{7} \times 7^2$
3	Calculate r^2	$2 \times \frac{22}{7} \times 49$
4	Simplify the expression	$2 \times 22 \times 7$
5	Perform multiplication	$44 \times 7 = 308$
6	State the final answer with units	308 cm^2

Revision Table: Sphere and Hemisphere Formulas

Shape	Type of Area	Formula
Sphere	Surface Area (Total/Curved)	$4\pi r^2$
Hemisphere	Curved Surface Area	$2\pi r^2$
Hemisphere	Total Surface Area	$3\pi r^2$ (Curved area + base area πr^2)

Additional Information on Hemisphere Surface Area

Understanding the difference between the curved surface area and the total surface area of a hemisphere is important. The curved surface area only includes the rounded part, which is half the surface area of a sphere. The total surface area, however, includes the curved part **plus** the flat circular base. The area of the circular base is πr^2 . Therefore, the total surface area of a hemisphere is $2\pi r^2 + \pi r^2 = 3\pi r^2$.

In this problem, the question specifically asks for the **curved** surface area, so we use the $2\pi r^2$ formula.

Units are also important. Since radius is in cm, the area is in cm^2 .

53. Answer: a

Explanation:

Understanding Work Done by a Load in a Machine

In physics, work done is defined as the energy transferred to or from an object via the application of force along a displacement. It is calculated as the product of the force applied and the distance moved in the direction of the force.

In the context of a machine, when it overcomes a load, the load represents the resistance force that the machine is working against. The distance travelled by the load is the displacement of this resistance.

Calculating Work Done by the Load

The question asks for the work done *by the load*. This refers to the work done by the force of the load as it is moved by the machine. Since the load is a force and it moves a certain distance, the work done by the load can be calculated using the fundamental work formula:

$$\text{Work Done} = \text{Force} \times \text{Distance}$$

In this specific scenario:

- The force acting here is the Load, denoted by 'L'.

- The distance travelled by this load is given as 'Ld'.

Substituting these values into the work done formula:

$$\text{Work Done by Load} = \text{Load} \times \text{Distance travelled by Load}$$

$$\text{Work Done by Load} = L \times Ld$$

So, the work done by the load is the product of the load (L) and the distance travelled by the load (Ld), which is $L \times Ld$.

Analyzing the Options

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

Option	Expression	Analysis
1	$L \times Ld$	This matches our derived formula for work done by the load.
2	$\frac{L}{Ld}$	This represents the ratio of load to the distance travelled, which is not the formula for work done.
3	$\frac{Ld}{L}$	This represents the ratio of distance travelled to the load, which is also not the formula for work done.
4	$\frac{1}{L \times Ld}$	This represents the reciprocal of the work done, which is incorrect.

Your Personal Exams Guide

Based on the definition of work done and the variables provided, the work done by the load is $L \times Ld$.

Revision Table: Key Work Concepts

Concept	Definition	Formula
Work Done (General)	Transfer of energy when a force causes displacement.	$W = F \times d \times \cos \theta$ (where θ is the angle between force and displacement)
Work Done by Load	Energy transferred against the load force as it moves.	$W_{load} = Load \times \text{Distance travelled by Load}$ (assuming force and displacement are in the same direction, $\cos \theta = 1$)
Work Input to Machine	Energy supplied to the machine by the effort.	$W_{in} = \text{Effort} \times \text{Distance moved by Effort}$

Additional Information on Work and Machines

In ideal machines, the work output is equal to the work input. The work output is typically the work done on the load. In real machines, due to friction and other losses, the work output is less than the work input. This is related to the efficiency of the machine.

- **Input Work:** The work done by the effort applied to the machine.
- **Output Work:** The work done by the machine on the load it overcomes.
- **Efficiency:** The ratio of work output to work input, usually expressed as a percentage.

$$\text{Efficiency} = \frac{\text{Work Output}}{\text{Work Input}} \times 100\%.$$

Understanding the work done by the load is crucial for analyzing the performance and efficiency of different types of machines.

54. Answer: b

Explanation:

Finding the Least Number Doubled Divisible by 7, 12, and 15

The question asks for the least number which, when doubled, is perfectly divisible by three given numbers: 7, 12, and 15.

If a number is perfectly divisible by 7, 12, and 15, it must be a common multiple of these three numbers. The least such number is the Least Common Multiple (LCM) of 7, 12, and 15.

Let the required least number be x . According to the problem, when x is doubled, the result ($2x$) is perfectly divisible by 7, 12, and 15. This means that $2x$ must be a multiple of the LCM of 7, 12, and 15.

Understanding Divisibility and Least Common Multiple (LCM)

Divisibility means that when one number is divided by another, the remainder is zero. A common multiple of a set of numbers is a number that is a multiple of every number in the set. The Least Common Multiple (LCM) is the smallest positive number that is a common multiple of two or more numbers.

Calculating the LCM of 7, 12, and 15

To find the LCM, we can use the prime factorization method. We find the prime factors of each number:

- Prime factors of 7: $7 = 7^1$
- Prime factors of 12: $12 = 2 \times 2 \times 3 = 2^2 \times 3^1$
- Prime factors of 15: $15 = 3 \times 5 = 3^1 \times 5^1$

To find the LCM, we take the highest power of all the prime factors that appear in any of the numbers:

Number	Prime Factorization
7	7^1
12	$2^2 \times 3^1$
15	$3^1 \times 5^1$

The prime factors involved are 2, 3, 5, and 7.

- Highest power of 2: 2^2 (from 12)
- Highest power of 3: 3^1 (from 12 and 15)
- Highest power of 5: 5^1 (from 15)

- Highest power of 7: 7^1 (from 7)

$$\text{LCM}(7, 12, 15) = 2^2 \times 3^1 \times 5^1 \times 7^1 = 4 \times 3 \times 5 \times 7$$

$$\text{LCM}(7, 12, 15) = 12 \times 35 = 420$$

Solving for the Least Number

As established earlier, $2x$ must be a multiple of $\text{LCM}(7, 12, 15)$. To find the **least** such number x , $2x$ must be the **least** multiple of $\text{LCM}(7, 12, 15)$, which is $\text{LCM}(7, 12, 15)$ itself.

So, we have the equation: $2x = \text{LCM}(7, 12, 15)$ $2x = 420$

Now, we solve for x by dividing both sides by 2: $x = \frac{420}{2}$ $x = 210$

Thus, the least number is 210.

Verifying the Solution

Let's check if doubling 210 is perfectly divisible by 7, 12, and 15.

Doubling 210 gives $2 \times 210 = 420$.

- Is 420 divisible by 7? $420 \div 7 = 60$. Yes.
- Is 420 divisible by 12? $420 \div 12 = 35$. Yes.
- Is 420 divisible by 15? $420 \div 15 = 28$. Yes.

Since 420 is perfectly divisible by 7, 12, and 15, our calculated number, 210, is indeed the least number that satisfies the condition.

Revision Table: Key Concepts

Concept	Definition	Relevance to Problem
Divisibility	A number a is divisible by b if $a = bk$ for some integer k .	Used to define the condition on the doubled number.
Multiple	A multiple of a number n is nk for some integer k .	The doubled number must be a common multiple of 7, 12, and 15.
Least Common Multiple (LCM)	The smallest positive integer that is a multiple of two or more numbers.	The doubled number, to be the minimum possible value, must equal the LCM of 7, 12, and 15.
Prime Factorization	Expressing a number as a product of its prime factors.	Method used to efficiently calculate the LCM.

Additional Information: Number Properties

Understanding multiples, factors, and LCM is fundamental in number theory. Related concepts include:

- **Factors:** Numbers that divide a given number evenly. For example, the factors of 12 are 1, 2, 3, 4, 6, and 12.
- **Common Factors:** Factors that two or more numbers share. For example, common factors of 12 and 18 are 1, 2, 3, and 6.
- **Highest Common Factor (HCF) or Greatest Common Divisor (GCD):** The largest of the common factors. For example, $\text{HCF}(12, 18) = 6$.

LCM is particularly useful when dealing with problems involving cycles, events that repeat at regular intervals, or finding the smallest quantity that can be divided equally into different-sized groups. HCF is useful when dividing quantities into the largest possible equal parts.

55. Answer: c

Explanation:

Understanding Heat Flow Through a Brass Rod

The question asks us to calculate the rate of heat flow through a brass rod given its dimensions, thermal conductivity, and the temperature difference across its ends. This is a classic application of Fourier's Law of Heat Conduction.

Fourier's Law of Heat Conduction:

The rate of heat flow ($\frac{Q}{t}$) through a material is directly proportional to the area of cross-section (A), the temperature difference (ΔT), and the thermal conductivity (k) of the material, and inversely proportional to the length (L) of the material. The formula is given by:

$$\frac{Q}{t} = kA \frac{\Delta T}{L}$$

where:

- $\frac{Q}{t}$ is the rate of heat flow (in J/s or Watts)
- k is the thermal conductivity of the material (in J/s m K)
- A is the area of the cross-section (in m^2)
- ΔT is the temperature difference across the ends (in K or $^{\circ}\text{C}$)
- L is the length of the rod (in m)

Given Parameters for the Brass Rod

Let's list the values provided in the question:

- Thermal conductivity of brass, $k = 109 \text{ J/s m K}$
- Area of cross-section, $A = 0.04 \text{ m}^2$
- Length of the rod, $L = 20 \text{ cm}$
- Temperature difference, $\Delta T = 200 \text{ }^{\circ}\text{C}$

Converting Units

Before plugging the values into the formula, we need to ensure all units are consistent. The length is given in centimeters, so we convert it to meters:

$$L = 20 \text{ cm} = 20 \times 10^{-2} \text{ m} = 0.20 \text{ m}$$

The temperature difference is given in $^{\circ}\text{C}$. For a temperature difference, the value in Kelvin is the same as in Celsius. So, $\Delta T = 200 \text{ K}$.

Calculating the Rate of Heat Flow

Now, substitute the values into the formula for the rate of heat flow:

$$\frac{Q}{t} = kA \frac{\Delta T}{L}$$

$$\frac{Q}{t} = (109 \text{ J/s m K}) \times (0.04 \text{ m}^2) \times \frac{200 \text{ K}}{0.20 \text{ m}}$$

Let's calculate the numerical value:

$$\frac{Q}{t} = 109 \times 0.04 \times \frac{200}{0.2}$$

$$\frac{Q}{t} = 109 \times 0.04 \times 1000$$

$$\frac{Q}{t} = 109 \times 40$$

$$\frac{Q}{t} = 4360 \text{ J/s}$$

Converting to kJ/s

The options are given in kilojoules per second (kJ/s). We convert the result from J/s to kJ/s:

$$1 \text{ kJ} = 1000 \text{ J}$$

So,

$$4360 \text{ J/s} = \frac{4360}{1000} \text{ kJ/s} = 4.360 \text{ kJ/s}$$

The rate of heat flow through the brass rod is 4.36 kJ/s.

Comparing with Options

Let's compare our calculated rate of heat flow with the given options:

- Option 1: 3.42 kJ/s
- Option 2: 2.32 kJ/s
- Option 3: 4.36 kJ/s
- Option 4: 5.80 kJ/s

Our calculated value of 4.36 kJ/s matches Option 3.

Revision Table: Heat Conduction Concepts

Concept	Description	Formula
Heat Conduction	Transfer of heat through direct contact, occurring in solids, liquids, and gases (most significant in solids).	N/A
Thermal Conductivity (k)	A material's ability to conduct heat. Higher k means better conductor.	Units: J/s m K or W/m K
Rate of Heat Flow ($\frac{Q}{t}$)	The amount of heat energy transferred per unit time.	$\frac{Q}{t} = kA \frac{\Delta T}{L}$ (for steady-state conduction)
Temperature Gradient ($\frac{\Delta T}{L}$)	The change in temperature per unit length along the direction of heat flow.	Units: K/m or $^{\circ}\text{C}/\text{m}$

Additional Information on Heat Transfer

Heat transfer is the movement of thermal energy from a region of higher temperature to a region of lower temperature. There are three primary modes of heat transfer:

1. **Conduction:** Heat transfer through direct contact between particles. This is the main mode of heat transfer in solids. Energy is transferred as vibrating atoms and molecules collide with their neighbors. Metals are generally good conductors because of free electrons.
2. **Convection:** Heat transfer through the movement of fluids (liquids or gases). As a fluid is heated, it becomes less dense and rises, while cooler, denser fluid sinks, creating convection currents. This is how heat is transferred in boiling water or through air circulation.
3. **Radiation:** Heat transfer through electromagnetic waves. This mode does not require a medium and can occur through a vacuum, such as heat from the sun reaching the Earth. All objects above absolute zero temperature emit thermal radiation.

In the case of the brass rod, the heat transfer is primarily through conduction because it is a solid material with a temperature difference across its ends.

56. Answer: c

Explanation:

Correct answer: Rs. 2,000

$$\backslash (\text{Total CP} = \text{Purchase Price} + \text{Repairing Charge} \backslash)$$

$$\backslash (\text{Total CP} = \text{Rs. 1,500} + \text{Rs. 100} \backslash)$$

57. Answer: d

Explanation:

Understanding the direction of the effort in a lever requires knowing the basic components of a lever and how they interact. A lever is a simple machine consisting of a beam or rigid rod pivoted at a fixed hinge, called a fulcrum.

Understanding Levers and Their Components

A lever involves three main components:

- **Fulcrum:** The fixed point around which the lever rotates.
- **Load:** The object or resistance being moved or overcome.
- **Effort:** The force applied to the lever to move the load.

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

Effort Direction in a Class 1 Lever

In a **Class 1 lever**, the fulcrum is located between the effort and the load. Common examples include a seesaw, a crowbar used to lift something, or a pair of scissors (where the pivot is the fulcrum).

Consider a seesaw. When you want to lift someone on the other side (the load), you push down on your end (apply effort). The fulcrum is in the middle. In this case, the effort is applied in a downward direction. The load on the other side moves upward. If you were lifting the load by pulling up on your end, the effort would be in an upward direction, and the load would move downward.

The key point is that at any given moment, the effort force is applied in a single, specific direction to cause the desired motion or balance the load. While you could potentially apply effort in different directions relative to your body (e.g., pushing, pulling, lifting), the force vector applied to the lever itself at the point of effort has one primary direction relative to the movement or potential movement of the lever arm and the load.

For example, when using a crowbar to lift a heavy object, you typically push down or pull up on the handle. This pushing or pulling is the effort, and it is applied in one direction at that moment to lift the load on the other side of the fulcrum.

Mathematically, the effort is represented as a force vector $\vec{F}_{\text{effort}}$ applied at a specific point on the lever. This vector has a single direction.

Let's consider the options:

- **three:** This suggests the effort could be applied simultaneously in three fundamentally different directions. This is not how a single force is applied on a lever.
- **multiple:** While effort could be applied from different angles in specific complex scenarios or sequentially, the force applied at any single moment to operate a simple lever effectively is primarily in one intended direction relative to the lever's mechanics. "Multiple" is too vague and doesn't describe the primary application of effort for a single action.
- **two:** This might imply opposite directions (like up or down), but at any specific time during the operation, the effort is applied in one chosen direction (either up OR down, pushing OR pulling). It's not simultaneously in two opposite directions.
- **one:** This accurately describes the direction of the effort force applied to the lever at a given time to achieve a specific action (e.g., lift or balance the load). The effort is a single force vector with a single direction.

Therefore, the effort in a Class 1 lever, like any applied force in a simple machine for a specific action, is exerted in one primary direction.

Revision Table: Class 1 Lever Essentials

Component	Position in Class 1 Lever	Description
Fulcrum	Between Effort and Load	The pivot point
Effort	Applied force	The force you exert
Load	Resistance	The object being moved or overcome
Effort Direction	One	The single direction of the applied force at a given moment

Additional Information: Understanding Lever Classes

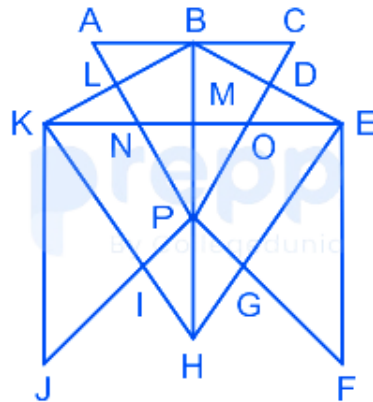
Levers are simple machines that help amplify force (mechanical advantage) or change the direction of force. Here's a brief look at all three classes:

- **Class 1 Lever:** Fulcrum is between the effort and the load (F-E-L or L-E-F). Examples: seesaw, crowbar, scissors. Can provide mechanical advantage or disadvantage depending on fulcrum position.
- **Class 2 Lever:** Load is between the fulcrum and the effort (F-L-E). Examples: wheelbarrow, nutcracker, bottle opener. Always provides mechanical advantage (effort arm is longer than load arm). Effort and load move in the same direction.
- **Class 3 Lever:** Effort is between the fulcrum and the load (F-E-L). Examples: fishing rod, tweezers, human forearm lifting a weight. Always provides mechanical disadvantage (effort arm is shorter than load arm) but increases speed and range of motion. Effort and load move in the same direction.

In all lever classes, the applied effort at any given moment is in a single, distinct direction relative to the lever arm.

58. Answer: a

Explanation:



The minimum number of lines required to make the given image is:

AC, KE, KJ, BH, EF, KH, HE, PF, JP, KB, BE, AP, PC

Hence, '13' is the correct answer.

59. Answer: a

Explanation:

Understanding Noise Exposure Standards and Protection

Exposure to excessive noise can cause significant health problems, particularly affecting hearing. Regulatory bodies establish noise exposure standards to define safe limits for sound levels in workplaces and other environments. When sound levels surpass these standards, it indicates a potentially hazardous situation that requires protective measures.

Why Hearing Protection is Crucial

The question asks what type of protection must be worn when noise levels exceed the noise exposure standard. The primary organ at risk from high noise levels is the ear, specifically the structures involved in hearing.

- Excessive noise can damage the delicate hair cells in the inner ear, leading to temporary or permanent hearing loss.
- Hearing loss caused by noise is often irreversible.
- Symptoms of noise-induced hearing loss can include difficulty understanding speech, ringing in the ears (tinnitus), and a reduced ability to hear high-pitched

sounds.

Therefore, when noise levels exceed safe limits as defined by the noise exposure standard, it is essential to protect the ears to prevent hearing damage. This protection is known as hearing protection.

Analyzing the Options for Noise Protection

Let's look at the provided options:

- **Hearing protection:** This includes devices like earplugs or earmuffs, specifically designed to reduce the amount of sound entering the ear. This is the correct type of protection for excessive noise.
- **Head protection:** This typically involves hard hats or helmets, used to protect the head from impacts, falling objects, or electrical hazards. While important in many environments, it does not protect against noise.
- **Eye protection:** This includes safety glasses, goggles, or face shields, used to protect the eyes from particles, chemicals, glare, or radiation. This does not protect against noise.
- **Foot protection:** This involves safety shoes or boots, used to protect the feet from falling objects, punctures, electrical hazards, or slippery surfaces. This does not protect against noise.

Based on the function of each type of protection and the hazard posed by excessive noise levels exceeding the noise exposure standard, hearing protection is the necessary measure.

Noise Exposure Standards Explained

Noise exposure standards are set to limit the amount of noise a person can be exposed to over a specific period without suffering adverse health effects, primarily hearing loss. These standards often involve a maximum permissible sound level (usually measured in decibels, dB) and a duration of exposure. For example, a standard might specify a maximum exposure time for noise levels at or above 85 dB.

When workplace or environmental noise monitoring shows levels above the standard, controls must be implemented. These controls follow a hierarchy:

1. **Elimination or Substitution:** Removing the noise source or replacing noisy equipment with quieter alternatives.

2. **Engineering Controls:** Modifying the noise source or path, such as enclosing machinery, using sound barriers, or installing mufflers.
3. **Administrative Controls:** Limiting the time workers spend in noisy areas or scheduling noisy tasks when fewer people are around.
4. **Personal Protective Equipment (PPE):** Providing hearing protection to individuals when higher-level controls are not sufficient or feasible. This is often the first line of defense when immediate action is needed to meet the noise exposure standard.

Wearing appropriate hearing protection is a critical step when engineering or administrative controls are not enough to bring noise exposure below the established standard.

Conclusion on Required Protection

When noise levels exceed the noise exposure standard, there is a risk of damage to hearing. Therefore, the required protection is specifically designed to mitigate this risk by reducing the sound energy reaching the ears. This protection is hearing protection.

Type of Protection	Primary Hazard Protected Against	Relevant to Excessive Noise?
Hearing Protection (Earplugs, Earmuffs)	Excessive Noise Levels, Hearing Loss	Yes
Head Protection (Hard Hats, Helmets)	Impacts, Falling Objects, Electrical Hazards	No
Eye Protection (Glasses, Goggles)	Particles, Chemicals, Glare, Radiation	No
Foot Protection (Safety Shoes, Boots)	Falling Objects, Punctures, Electrical Hazards	No

Revision Table: Noise Safety Basics

Concept	Description
Noise Exposure Standard	Defines permissible limits for noise levels and duration to prevent hearing damage.
Excessive Noise	Noise levels above the established standard, posing a risk to hearing.
Hearing Protection	Devices worn to reduce sound entering the ears (e.g., earplugs, earmuffs).
Noise-Induced Hearing Loss	Hearing damage caused by exposure to loud sounds.

Additional Information on Hearing Protection

Choosing the right hearing protection depends on the noise level and the type of noise. Hearing protection devices are rated by their Noise Reduction Rating (NRR), which indicates how many decibels they can reduce the noise level. Proper fit and consistent use are crucial for hearing protection to be effective in meeting the noise exposure standard.

Regular audiometric testing (hearing tests) is often required for individuals working in environments with high noise exposure to monitor their hearing health and the effectiveness of noise control measures and hearing protection programs.

Your Personal Exams Guide

60. Answer: b

Explanation:

Understanding Velocity Ratio in Simple Machines

Simple machines help us do work more easily by changing the direction or magnitude of the force we apply. Key concepts used to analyze the performance of a simple machine include Mechanical Advantage, Velocity Ratio, and Efficiency. This question focuses on the definition of Velocity Ratio.

Defining Velocity Ratio

The Velocity Ratio (VR) of a simple machine is a fundamental concept used to understand the theoretical mechanical advantage, assuming no energy losses (like friction). It is defined purely based on the distances moved by the point where effort is applied and the point where the load is lifted or moved.

Specifically, the velocity ratio is the ratio of the distance moved by the effort to the distance moved by the load in the same amount of time. Since the time is the same for both movements (for ideal conditions or when considering simultaneous displacement), this ratio simplifies to the ratio of the distances traveled.

Components Involved: Effort and Load

- **Effort:** This is the force applied by the user to operate the machine.
- **Load:** This is the resistance force that the machine has to overcome or the weight being lifted.

Calculating Velocity Ratio

The formula for Velocity Ratio is:

$$\text{Velocity Ratio (VR)} = \frac{\text{Distance traveled by the effort}}{\text{Distance traveled by the load}}$$

Let D_E be the distance traveled by the effort and D_L be the distance traveled by the load. The formula is:

$$\text{VR} = \frac{D_E}{D_L}$$

The velocity ratio is a theoretical value determined by the geometry of the machine, like the lengths of levers, radii of wheels, or the number of pulleys. It does not change with the load applied, unlike mechanical advantage which can be affected by friction.

Filling the Blanks for Velocity Ratio Definition

Based on the definition, the velocity ratio is the ratio of the distance traveled by the effort to the distance traveled by the load.

So, the sentence "The velocity ratio of a simple machine is the ratio of distance traveled by the _____ to the distance traveled by the _____ in the machine." should be completed with "effort" in the first blank and "load" in the second blank.

Distance traveled by the **effort** / Distance traveled by the **load** = Velocity Ratio.

Velocity Ratio Components

Component	Role in Velocity Ratio Formula
Effort	Distance traveled is in the numerator (D_E)
Load	Distance traveled is in the denominator (D_L)

Connecting Velocity Ratio to Machine Performance

The velocity ratio is crucial because it relates to the maximum possible mechanical advantage (Ideal Mechanical Advantage). In an ideal machine (with no friction), the Ideal Mechanical Advantage (IMA) is equal to the Velocity Ratio (VR).

$$IMA = VR$$

In real machines, friction is present, so the Actual Mechanical Advantage (AMA) is usually less than the Velocity Ratio.

$$AMA < VR$$

Efficiency (η) of a simple machine is the ratio of AMA to IMA (or AMA to VR):

$$\eta = \frac{AMA}{IMA} = \frac{AMA}{VR}$$

Understanding the velocity ratio helps predict the theoretical behavior of a machine and calculate its efficiency.

Revision Table: Simple Machine Ratios

Key Ratios in Simple Machines

Ratio	Definition (using forces)	Definition (using distances)
Actual Mechanical Advantage (AMA)	Load/Effort (Actual)	–
Ideal Mechanical Advantage (IMA)	Load/Effort (Ideal, no friction)	Distance by Effort/Distance by Load
Velocity Ratio (VR)	–	Distance by Effort/Distance by Load

Additional Information: Types of Simple Machines

Simple machines are the basic building blocks of complex machines. There are six classical types:

- Lever (e.g., crowbar, seesaw)
- Wheel and Axle (e.g., doorknob, steering wheel)
- Pulley (e.g., flag pole, construction crane)
- Inclined Plane (e.g., ramp, slope)
- Wedge (e.g., axe, knife)
- Screw (e.g., bolt, jack screw)

The velocity ratio for each type depends on its specific geometry and how the effort and load move. For example, in a lever, it depends on the ratio of the effort arm length to the load arm length. In a system of pulleys, it depends on the number of ropes supporting the load.

61. Answer: d

Explanation:

Understanding Letter Series Patterns for Exam Prep

Letter series questions require you to identify the underlying pattern in a sequence of letters or groups of letters and then determine the next term based on that pattern. Let's

analyze the given series:

OOOOOX, OOOOXX, OOOXXX, OOXXXX, ?

Analyzing the Pattern in the Series

Observe each term in the series carefully. Each term consists of a combination of the letters 'O' and 'X'. Let's count the number of 'O's and 'X's in each term and also note the total number of characters.

Term	Characters	Number of 'O's	Number of 'X's	Total Characters
1	OOOOOX	5	1	6
2	OOOXX	4	2	6
3	OOXXX	3	3	6
4	OOXXXX	2	4	6

From the table, we can clearly see a pattern emerging:

- The total number of characters in each term is constant, which is 6.
- The number of 'O's is decreasing by one in each successive term (5, 4, 3, 2).
- The number of 'X's is increasing by one in each successive term (1, 2, 3, 4).

Determining the Missing Term

Following the identified pattern, for the next term in the series (which is the fifth term), the number of 'O's should be one less than the number of 'O's in the fourth term, and the number of 'X's should be one more than the number of 'X's in the fourth term.

- Number of 'O's in the 5th term = Number of 'O's in the 4th term -1
- Number of 'O's in the 5th term = $2 - 1 = 1$
- Number of 'X's in the 5th term = Number of 'X's in the 4th term $+1$
- Number of 'X's in the 5th term = $4 + 1 = 5$

So, the fifth term should consist of 1 'O' followed by 5 'X's. This gives us the sequence OXXXXX.

Comparing with the Options

Let's look at the given options to find the term that matches our derived pattern:

1. XXXXXX: This has 0 'O's and 6 'X's. This would be the term after OXXXXX if the pattern continues.
2. OOXXXXX: This has 2 'O's and 5 'X's, and a total of 7 characters. The pattern maintains 6 characters.
3. OOXXXX: This is the fourth term in the series, not the missing fifth term.
4. OXXXXX: This has 1 'O' and 5 'X's, and a total of 6 characters. This perfectly matches the pattern we identified for the missing term.

Therefore, the correct alternative that completes the series is OXXXXX.

Revision Table: Series Pattern Summary

Term	Structure	Pattern Rule Applied
1st	OOOOOX	Starting point
2nd	OOOOXX	-1 'O', +1 'X' from previous term
3rd	OOOXXX	-1 'O', +1 'X' from previous term
4th	OOXXXX	-1 'O', +1 'X' from previous term
5th (Missing)	OXXXXX	-1 'O', +1 'X' from previous term

Additional Information: Types of Logical Series

Logical reasoning series can follow various patterns. Besides decreasing/increasing counts of elements, other common types include:

- **Alphabet Series:** Patterns based on the position of letters in the alphabet (e.g., skipping letters, alphabetical order).
- **Number Series:** Patterns based on arithmetic operations (addition, subtraction, multiplication, division), squares, cubes, prime numbers, etc.
- **Mixed Series:** Combining patterns of numbers and letters.
- **Figural Series:** Patterns based on the rotation, movement, or changes in figures.
- **Letter/Symbol Manipulation:** Patterns involving rearrangement, addition, or removal of specific letters or symbols.

Solving series questions relies on careful observation, identifying the rule governing the sequence, and applying it consistently to find the missing or next term.

62. Answer: d

Explanation:

This question asks us to find which two mathematical signs, when interchanged in the given equation, will make the equation correct. The original equation is: $9 \div 3 + 8 \times 2 - 15 = 2$. We need to test the options by swapping the signs and then evaluating the new equation using the order of operations (BODMAS/PEMDAS).

Analyzing the Original Equation

First, let's evaluate the given equation to see if it is already correct. We follow the order of operations:

1. Division and Multiplication (from left to right)
2. Addition and Subtraction (from left to right)

Equation: $9 \div 3 + 8 \times 2 - 15$

- Perform division: $9 \div 3 = 3$
- Perform multiplication: $8 \times 2 = 16$
- Substitute these values back: $3 + 16 - 15$
- Perform addition: $3 + 16 = 19$
- Perform subtraction: $19 - 15 = 4$

So, $9 \div 3 + 8 \times 2 - 15 = 4$. The equation $4 = 2$ is false. The original equation is incorrect.

Testing Options for Sign Interchange

We will now test each option by interchanging the specified signs and evaluating the resulting equation.

Option 1: Interchange $\times$ and $-$

Swap the multiplication ($\times$) and subtraction ($-$) signs in the original equation.

Original equation: $9 \div 3 + 8 \times 2 - 15 = 2$

Equation after swapping $\times$ and $-$: $9 \div 3 + 8 - 2 \times 15 = 2$

Now, evaluate this new equation using the order of operations:

- Perform division: $9 \div 3 = 3$
- Perform multiplication: $2 \times 15 = 30$
- Substitute back: $3 + 8 - 30$
- Perform addition: $3 + 8 = 11$
- Perform subtraction: $11 - 30 = -19$

The result is -19 . Since $-19 \neq 2$, swapping $\times$ and $-$ does not make the equation correct.

Option 2: Interchange $\div$ and $-$

Swap the division ($\div$) and subtraction ($-$) signs in the original equation.

Original equation: $9 \div 3 + 8 \times 2 - 15 = 2$

Equation after swapping $\div$ and $-$: $9 - 3 + 8 \times 2 \div 15 = 2$

Now, evaluate this new equation:

- Perform multiplication: $8 \times 2 = 16$
- Perform division: $16 \div 15 = \frac{16}{15}$
- Substitute back: $9 - 3 + \frac{16}{15}$
- Perform subtraction: $9 - 3 = 6$
- Perform addition: $6 + \frac{16}{15} = \frac{6 \times 15}{15} + \frac{16}{15} = \frac{90 + 16}{15} = \frac{106}{15}$

The result is $\frac{106}{15}$. Since $\frac{106}{15} \neq 2$, swapping $\div$ and $-$ does not make the equation correct.

Option 3: Interchange $+$ and $\times$

Swap the addition ($+$) and multiplication ($\times$) signs in the original equation.

Original equation: $9 \div 3 + 8 \times 2 - 15 = 2$

Equation after swapping $+$ and $\times$: $9 \div 3 \times 8 + 2 - 15 = 2$

Now, evaluate this new equation:

- Perform division: $9 \div 3 = 3$

- Perform multiplication: $3 \times 8 = 24$
- Substitute back: $24 + 2 - 15$
- Perform addition: $24 + 2 = 26$
- Perform subtraction: $26 - 15 = 11$

The result is 11. Since $11 \neq 2$, swapping + and $\times$ does not make the equation correct.

Option 4: Interchange + and -

Swap the addition (+) and subtraction (-) signs in the original equation.

Original equation: $9 \div 3 + 8 \times 2 - 15 = 2$

Equation after swapping + and -: $9 \div 3 - 8 \times 2 + 15 = 2$

Now, evaluate this new equation:

- Perform division: $9 \div 3 = 3$
- Perform multiplication: $8 \times 2 = 16$
- Substitute back: $3 - 16 + 15$
- Perform subtraction: $3 - 16 = -13$
- Perform addition: $-13 + 15 = 2$

The result is 2. Since $2 = 2$, swapping + and - makes the equation correct.

Conclusion on Correcting the Equation

After testing all the given options, we found that interchanging the addition (+) and subtraction (-) signs in the equation $9 \div 3 + 8 \times 2 - 15 = 2$ results in the correct equation $9 \div 3 - 8 \times 2 + 15 = 2$. This confirms that swapping + and - signs is the correct solution to this sign interchange puzzle.

Revision Table: Sign Interchange Test

Signs Interchanged	New Equation	Evaluation	Result	Correct?
Original	$9 \div 3 + 8 \times 2 - 15$	$3 + 16 - 15 = 19 - 15 = 4$	4	No
$\times$ and $-$	$9 \div 3 + 8 - 2 \times 15$	$3 + 8 - 30 = 11 - 30 = -19$	-19	No
$\div$ and $-$	$9 - 3 + 8 \times 2 \div 15$	$9 - 3 + 16 \div 15 = 6 + \frac{16}{15} = \frac{106}{15}$	$\frac{106}{15}$	No
$+$ and $\times$	$9 \div 3 \times 8 + 2 - 15$	$3 \times 8 + 2 - 15 = 24 + 2 - 15 = 26 - 15 = 11$	11	No
$+$ and $-$	$9 \div 3 - 8 \times 2 + 15$	$3 - 16 + 15 = -13 + 15 = 2$	2	Yes

Additional Information on Order of Operations

When solving mathematical expressions, it is crucial to follow a specific order of operations to ensure a unique and correct result. This order is often remembered using mnemonics like BODMAS or PEMDAS.

- **BODMAS:** Brackets, Orders (powers, square roots), Division and Multiplication (left to right), Addition and Subtraction (left to right).
- **PEMDAS:** Parentheses, Exponents, Multiplication and Division (left to right), Addition and Subtraction (left to right).

In the problems involving interchanging signs to correct an equation, applying the order of operations correctly to each potential new equation is the key to finding the right answer. Mistakes often happen when the order of division/multiplication or addition/subtraction is not followed strictly from left to right when they appear consecutively.

63. Answer: a

Explanation:

Understanding Electrical Energy Supplied

The question asks us to determine the energy supplied to a circuit by a voltage source over a certain time when it maintains a specific current. This involves understanding the fundamental relationship between voltage, current, power, and energy in an electrical circuit.

Let's break down the concepts:

- **Voltage (V):** Represents the electrical potential difference across the circuit. It is the 'push' that drives the charge. Measured in Volts (V).
- **Current (i):** Represents the flow of electric charge per unit time. Measured in Amperes (A).
- **Time (t):** The duration for which the current flows. Measured in seconds (s).
- **Power (P):** The rate at which energy is transferred or dissipated in the circuit. It is the energy supplied per unit time. Measured in Watts (W).
- **Energy (E):** The total work done or energy transferred over a period of time. Measured in Joules (J).

Relationship Between Power, Voltage, and Current

The electrical power (P) supplied by a source in a circuit is defined as the product of the voltage (V) across the source and the current (i) flowing through the circuit.

$$P = V \times i$$

This formula tells us how much energy is being supplied or consumed per second.

Calculating Electrical Energy Supplied

Energy (E) is defined as the rate of energy transfer (power) multiplied by the time (t) duration over which the transfer occurs.

$$E = P \times t$$

Now, we can substitute the expression for power ($P = V \times i$) into the energy formula:

$$E = (V \times i) \times t$$

$$E = V \times i \times t$$

This final formula $E = Vit$ gives us the total energy supplied by the voltage source maintaining a current i for a time t .

Analyzing the Options

Let's examine the given options based on our derived formula $E = Vit$. We can also think about the units to verify the correct formula.

- **Option 1: Vit**

This matches our derived formula. Let's check the units: Volts (V) $\times$ Amperes (A) $\times$ seconds (s). We know that Power ($P = V \times i$) has units of Watts (W), and Watts are Joules per second (J/s). So, $V \times i \times t$ has units $(J/s) \times s = J$, which are the units for energy. This option is dimensionally correct and matches the formula.

- **Option 2: $1/Vit$**

This is the reciprocal of the correct formula. The units would be $1/J$, which is not energy.

- **Option 3: Vi/t**

This represents power (Vi) divided by time (t), not energy. The units would be $(J/s)/s = J/s^2$, which is not energy.

- **Option 4: V/it**

This formula does not correspond to the definition of power or energy in an electrical circuit. The units would be $V/(A \times s)$. Since $A \times s$ represents charge (Coulombs, C), the units are V/C , which is energy per charge (J/C is Voltage), but V/C is $V/(A \times s)$, dimensionally something different than energy (Joules). V/A is resistance (Ω), so $V/(it)$ has units Ω/s , not energy.

Based on the derivation and unit analysis, the correct formula for the energy supplied by the source is Vit .

Summary of Relationships

Here is a table summarizing the key relationships:

Quantity	Symbol	Formula	Units	Relation to others
Voltage	V	–	Volts (V)	Energy per unit charge (J/C)
Current	i	–	Amperes (A)	Charge per unit time (C/s)
Power	P	$P = V \times i$	Watts (W)	Rate of energy transfer (J/s)
Energy	E	$E = P \times t = V \times i \times t$	Joules (J)	Power over time

Conclusion

The energy supplied to the circuit by the voltage source maintaining a current i for a time t is given by the product of the voltage, current, and time.

Revision Table: Electrical Energy Formulas

Concept	Formula	Units
Electrical Power (P)	$P = V \times i$	Watts (W) or J/s
Electrical Energy (E)	$E = P \times t$	Joules (J) or Ws
Electrical Energy (E)	$E = V \times i \times t$	Joules (J) or VAs
Alternative Power (using Resistance R)	$P = i^2 R = V^2 / R$	Watts (W)
Alternative Energy (using Resistance R)	$E = i^2 R t = (V^2 / R) t$	Joules (J)

Additional Information on Electrical Energy

Electrical energy is a fundamental concept in physics and electrical engineering. It represents the capacity of electricity to do work. When a voltage source is connected to a circuit, it does work on the charges, causing them to move and constitute a current. This work done by the source is the energy supplied to the circuit.

This energy can then be converted into various forms depending on the components in the circuit. For example:

- In a resistor, electrical energy is converted into heat (Joule heating).

- In a motor, electrical energy is converted into mechanical energy.
- In a light bulb, electrical energy is converted into light and heat energy.
- In a battery being charged, electrical energy is converted into chemical energy.

Understanding the relationship $E = V\bar{it}$ is crucial for analyzing power consumption, designing electrical systems, and calculating energy costs. This formula is a direct consequence of the definitions of voltage, current, power, and energy. Voltage is work per unit charge ($V = W/q$), and current is charge per unit time ($i = q/t$). Thus, the rate of doing work, which is power ($P = W/t$), can be written as $P = (V \times q)/t = V \times (q/t) = V \times i$. Then, energy ($E = W$) is power times time, $E = P \times t = (V \times i) \times t = V\bar{it}$.

64. Answer: c

Explanation:

Understanding the Projection of an Object

The question asks about the figure formed on a plane when straight lines are drawn from points on the contour (outline) of an object to meet that plane. This specific process is fundamental in technical drawing and geometry and is known as projection.

Let's break down the options provided:

- **development:** Development refers to the unrolling or unfolding of the surface of a 3D object onto a 2D plane. This is commonly used for creating patterns for sheet metal or cardboard shapes. It does not involve drawing lines from the object to a separate plane to form a figure.
- **dimensioning:** Dimensioning is the process of adding measurements (like lengths, angles, diameters) to a drawing to indicate the size and shape of an object. It is about adding information to a drawing, not creating the primary figure itself from lines projected onto a plane.
- **projection:** Projection is the process of representing a three-dimensional object on a two-dimensional plane. This is done by drawing lines (called projectors or visual rays) from points on the object to the plane. The figure formed where these lines intersect the plane is the projection of the object. The lines can be parallel (like in orthographic projection) or converge to a point (like in perspective projection). The description in the question perfectly matches the definition of a projection.

- **animation:** Animation is the creation of a sequence of images that, when shown rapidly, simulate movement. This is related to motion and time, not the static representation of an object's shape on a plane through lines drawn from its contour.

Based on the definition and comparison, drawing straight lines from points on the contour of an object to meet a plane to obtain a figure on that plane is the definition of a projection.

Here is a simple comparison:

Term	Process	Outcome
Development	Unfolding a 3D surface onto a 2D plane	A flat pattern
Dimensioning	Adding measurements to a drawing	A drawing with size information
Projection	Drawing lines from an object to a plane	A 2D representation of the object
Animation	Creating sequential images for motion	Illusion of movement

Therefore, the figure obtained on the plane using this method is called the projection of the object.

Revision Table: Key Concepts Review

Concept	Brief Explanation
Projection	Representing a 3D object on a 2D plane using lines from the object to the plane.
Contour	The outline or boundary of an object.
Plane of Projection	The 2D surface onto which the object is projected.
Projectors	The straight lines drawn from the object to the plane of projection.

Additional Information: Types of Projection

Projection is a broad term, and there are different ways projectors can be drawn, leading to various types of projections used in drawing and graphics:

- **Parallel Projection:** Projectors are parallel to each other.
 - Orthographic Projection: Projectors are parallel and perpendicular to the projection plane. This is commonly used in engineering drawings to show multiple views (front, top, side).
 - Oblique Projection: Projectors are parallel but not perpendicular to the projection plane. Faces parallel to the plane appear true size and shape.
- **Perspective Projection:** Projectors converge to a single point (the center of projection or viewpoint). This type of projection mimics how the human eye sees objects, showing foreshortening and giving a sense of depth, often used in architectural drawings and art.

Understanding the type of projection is crucial as it affects how the object appears on the 2D plane, representing its shape and size relative to the viewpoint and the plane.

65. Answer: c

Explanation:

Understanding the JPEG Acronym

The question asks for the full form of the acronym JPEG. Acronyms are abbreviations formed from the initial letters of other words and pronounced as a word.

JPEG is a widely recognized term, especially when dealing with digital images. It refers to a standard method of compressing digital images. When you see a file with a .jpg or .jpeg extension, it typically means it uses the JPEG compression standard.

Decoding the JPEG Meaning

The acronym JPEG actually stands for the name of the committee that created the standard. Let's look at the options provided to find the correct expansion.

- Option 1: Joint Program Experts Group
- Option 2: Joint Program Executing Group
- Option 3: Joint Photographic Experts Group
- Option 4: Joint Program Experimental Group.

The standard is specifically about compressing photographic images effectively. Considering this context, the expansion should relate to 'Photographic Experts'.

The committee responsible for developing this standard was indeed focused on photographic imaging. Therefore, the correct expansion of JPEG is linked to this specific area of expertise.

Analyzing the Options for JPEG Expansion

Let's consider each option in the context of image compression standards:

- **Joint Program Experts Group:** This is a general term and doesn't specifically relate to the domain of image processing or photography.
- **Joint Program Executing Group:** This implies a group that executes programs, which isn't the primary function related to creating a compression standard.
- **Joint Photographic Experts Group:** This option directly references 'Photographic Experts', aligning perfectly with the purpose of the JPEG standard, which is primarily used for compressing photographic images.
- **Joint Program Experimental Group:** While standards often go through experimental phases, 'Experimental' isn't part of the official name of the committee.

Based on the history and purpose of the standard, the acronym JPEG comes from the name of the committee that developed it: the Joint Photographic Experts Group.

What JPEG Represents

The Joint Photographic Experts Group created several standards, and the term "JPEG" has become synonymous with the most popular compression technique they developed, formally known as JPEG baseline. This standard is particularly effective for compressing real-world photographic images, which contain smooth variations of tone and color.

Conclusion on JPEG Acronym

The correct full form of JPEG is the name of the committee that designed the compression standard.

Thus, the expansion that correctly identifies this group is "Joint Photographic Experts Group".

Acronym	Full Form	Context
JPEG	Joint Photographic Experts Group	Committee that developed the standard; widely used for digital image compression.

Revision Table: Key Acronyms

Acronym	Stands For	Brief Description
JPEG	Joint Photographic Experts Group	A common method of lossy compression for digital images.
PNG	Portable Network Graphics	A raster-graphics file format that supports lossless data compression.
GIF	Graphics Interchange Format	A bitmap image format that supports both animated and static images.
PDF	Portable Document Format	A file format used to present documents in a manner independent of application software, hardware, and operating systems.

Additional Information: JPEG Standard and Use

The JPEG standard specifies both the codec (coder-decoder) and the file format. The compression method is typically lossy, meaning some information is discarded during compression to achieve smaller file sizes. This loss is often imperceptible to the human eye, especially at lower compression levels. However, repeated saving and re-compressing a JPEG file can lead to a noticeable degradation in image quality.

JPEG compression works by converting the image into frequency components and discarding high-frequency information that is less important for visual perception. It uses

techniques like Discrete Cosine Transform (DCT), quantization, and entropy encoding.

JPEG is the most common format for storing and transmitting photographic images on the internet and in digital cameras.

66. Answer: a

Explanation:

Given,

Interest percentage = 10%

Total expenses = 25000

Calculation,

100% = 25000

$\Rightarrow 1\% = 250$

$\Rightarrow 10\% = \text{Rs. } 2500.$

67. Answer: d

Explanation:

Solving for $x - y$ in a 4-Digit Number Divisible by 11

The problem asks us to find the value of $x - y$ for a 4-digit number $1xy7$ that is divisible by 11. To solve this, we will use the divisibility rule for 11.

Understanding the Divisibility Rule for 11

A number is divisible by 11 if the difference between the sum of the digits at the odd places (from the right) and the sum of the digits at the even places (from the right) is either 0 or a multiple of 11. Let's apply this rule to the 4-digit number $1xy7$.

Applying the Divisibility Rule to $1xy7$

In the 4-digit number $1xy7$:

- The digits at the odd places (1st and 3rd from the right) are 7 and x .
- The digits at the even places (2nd and 4th from the right) are y and 1.

Now, let's calculate the sum of the digits at the odd places and the sum of the digits at the even places.

- Sum of digits at odd places = $7 + x$
- Sum of digits at even places = $y + 1$

According to the divisibility rule of 11, the difference between these two sums must be 0 or a multiple of 11.

Difference = (Sum of digits at odd places) - (Sum of digits at even places)

Difference = $(7 + x) - (y + 1)$

Difference = $7 + x - y - 1$

Difference = $6 + x - y$

So, for the number $1xy7$ to be divisible by 11, the value of $6 + x - y$ must be 0, 11, -11, 22, -22, and so on.

Finding the Possible Values of $x - y$

x and y are digits in a 4-digit number, which means they can take any integer value from 0 to 9.

- The minimum possible value for $x - y$ is when x is minimum (0) and y is maximum (9), which is $0 - 9 = -9$.
- The maximum possible value for $x - y$ is when x is maximum (9) and y is minimum (0), which is $9 - 0 = 9$.

Therefore, the possible range for $x - y$ is from -9 to 9.

Now, let's consider the possible values for $6 + x - y$ based on the range of $x - y$:

- Minimum value of $6 + x - y = 6 + (-9) = -3$
- Maximum value of $6 + x - y = 6 + 9 = 15$

The possible values for $6 + x - y$ must be multiples of 11 within the range $[-3, 15]$. The only multiple of 11 in this range is 0 or 11.

- Case 1: $6 + x - y = 0$
- Case 2: $6 + x - y = 11$

Let's solve for $x - y$ in each case:

- Case 1: $x - y = 0 - 6 = -6$
- Case 2: $x - y = 11 - 6 = 5$

So, for the number $1xy7$ to be divisible by 11, the value of $x - y$ can be either -6 or 5.

Evaluating the Given Options

The given options for the value of $x - y$ are:

1. -8
2. -2
3. -4
4. -6

Comparing these options with the possible values we found for $x - y$ (-6 or 5), we see that only -6 is present in the options.

- If $x - y = -8$, then $6 + x - y = 6 + (-8) = -2$, which is not a multiple of 11.
- If $x - y = -2$, then $6 + x - y = 6 + (-2) = 4$, which is not a multiple of 11.
- If $x - y = -4$, then $6 + x - y = 6 + (-4) = 2$, which is not a multiple of 11.
- If $x - y = -6$, then $6 + x - y = 6 + (-6) = 0$, which is a multiple of 11.

Thus, the value of $x - y$ must be -6 for the 4-digit number $1xy7$ to be divisible by 11.

Value of $x - y$	Value of $6 + x - y$	Divisible by 11?
-8	-2	No
-2	4	No
-4	2	No
-6	0	Yes
5 (not in options)	11	Yes

Therefore, the only value among the given options that makes the 4-digit number $1xy7$ divisible by 11 is -6 .

Revision Table: Divisibility Rules

Understanding divisibility rules is key to solving problems like this. Here's a quick look at some common divisibility rules:

Divisible by	Rule
2	The last digit is even (0, 2, 4, 6, or 8).
3	The sum of the digits is divisible by 3.
4	The number formed by the last two digits is divisible by 4.
5	The last digit is 0 or 5.
6	The number is divisible by both 2 and 3.
9	The sum of the digits is divisible by 9.
10	The last digit is 0.
11	The difference between the sum of the digits at odd places and the sum of the digits at even places is 0 or a multiple of 11.

Additional Information: Number Properties

A 4-digit number like $1xy7$ can be represented using place values. The number can be written as:

$$1 \times 1000 + x \times 100 + y \times 10 + 7 \times 1$$

This representation helps in understanding the contribution of each digit to the total value of the number. When dealing with divisibility rules, especially for numbers like 11, the position of the digits (odd place or even place from the right) is crucial, as seen in the solution above.

68. Answer: b

Explanation:

Calculating Power Supplied by Battery to Series Resistance

This problem involves calculating the power supplied by a battery to a circuit containing two resistances connected in series. When resistances are connected in series, the total resistance of the circuit is the sum of the individual resistances. The current flowing through each resistor in a series circuit is the same.

Understanding Series Resistance

In a series connection, components are connected end-to-end, forming a single path for the current. If we have resistances $R_1, R_2, \dots, R_n$ connected in series, the total equivalent resistance R_{eq} is given by the sum:

$$R_{eq} = R_1 + R_2 + \dots + R_n$$

Step-by-Step Solution to Find Power

Let's break down the problem:

1. Identify the given values:
 - Resistance 1, $R_1 = 2\Omega$
 - Resistance 2, $R_2 = 6\Omega$
 - Battery voltage, $V = 12\text{ V}$
2. Calculate the total equivalent resistance (R_{eq}) for the series combination. Since the two resistances are in series:

$$R_{eq} = R_1 + R_2$$

$$R_{eq} = 2\Omega + 6\Omega$$

$$R_{eq} = 8\Omega$$

3. Calculate the total current (I) drawn from the battery using Ohm's Law. Ohm's Law states that $V = IR$, where V is the voltage, I is the current, and R is the resistance. In this case, we use the total voltage and the total equivalent resistance:

$$V = I \times R_{eq}$$

$$12 \text{ V} = I \times 8\Omega$$

Solving for I :

$$I = \frac{12 \text{ V}}{8\Omega}$$

$$I = 1.5 \text{ A}$$

So, the total current flowing through the circuit and supplied by the battery is 1.5 A.

4. Calculate the power supplied by the battery (P). The power supplied by a source like a battery is given by the formula $P = VI$, where V is the voltage of the source and I is the total current drawn from the source.

$$P = V \times I$$

$$P = 12 \text{ V} \times 1.5 \text{ A}$$

$$P = 18 \text{ W}$$

Alternatively, we can use the formula $P = I^2 R$ or $P = \frac{V^2}{R}$ with the total current and total equivalent resistance, or the total voltage and total equivalent resistance, respectively:

$$P = I^2 \times R_{eq} = (1.5 \text{ A})^2 \times 8\Omega = 2.25 \text{ A}^2 \times 8\Omega = 18 \text{ W}$$

$$P = \frac{V^2}{R_{eq}} = \frac{(12 \text{ V})^2}{8\Omega} = \frac{144 \text{ V}^2}{8\Omega} = 18 \text{ W}$$

All methods yield the same result for the power supplied by the battery.

Summary of Calculation

Parameter	Value	Formula/Calculation
Resistance 1 (R_1)	2Ω	Given
Resistance 2 (R_2)	6Ω	Given
Battery Voltage (V)	12 V	Given
Equivalent Resistance (R_{eq})	8Ω	$R_1 + R_2 = 2 + 6$
Total Current (I)	1.5 A	$I = V/R_{eq} = 12/8$
Power Supplied (P)	18 W	$P = V \times I = 12 \times 1.5$

The power supplied by the battery to the series circuit is 18 W.

Revision Table: Key Electrical Concepts

Concept	Symbol	Unit	Description
Resistance	R	Ohm (Ω)	Opposition to electric current flow.
Voltage (Potential Difference)	V	Volt (V)	Energy per unit charge. Driving force for current.
Electric Current	I	Ampere (A)	Rate of flow of electric charge.
Power	P	Watt (W)	Rate at which energy is transferred or used.

Additional Information: Power in Electrical Circuits

Power is a fundamental concept in electrical circuits, representing the rate at which electrical energy is converted into other forms of energy (like heat or light) or transferred from a source (like a battery). The power (P) dissipated by a resistor or supplied by a source can be calculated using several equivalent formulas derived from Ohm's Law ($V = IR$):

- $P = VI$ (Power equals voltage times current)
- $P = I^2 R$ (Power equals current squared times resistance)
- $P = \frac{V^2}{R}$ (Power equals voltage squared divided by resistance)

These formulas are crucial for analyzing electrical circuits and understanding energy transfer and dissipation. In a series circuit, the total power supplied by the battery is equal to the sum of the power dissipated by each individual resistor in the circuit.

69. Answer: d

Explanation:

Understanding the Discovery of Bacteria

The question asks us to identify the scientist credited with the discovery of bacteria. This discovery is a significant milestone in the history of microbiology, opening up a new world of microscopic life.

Examining the Options

Let's look at the scientists listed in the options and their notable contributions:

- **Robert Koch:** Robert Koch was a German physician and microbiologist. He is known for establishing Koch's postulates, which are criteria used to determine whether a specific microorganism is the cause of a particular disease. He made significant discoveries regarding diseases like tuberculosis, cholera, and anthrax. While crucial to microbiology, he did not discover bacteria.
- **James Chadwick:** James Chadwick was an English physicist. He is famous for discovering the neutron in 1932, a fundamental particle within the atomic nucleus. His work is in the field of physics, not biology or microbiology.
- **Eugen Goldstein:** Eugen Goldstein was a German physicist. He is known for his work on cathode rays and the discovery of anode rays, also called canal rays. His contributions were in physics related to electricity and atomic structure, not the discovery of bacteria.
- **A V Leeuwenhoek:** Antoni van Leeuwenhoek was a Dutch businessman and scientist in the Golden Age of Dutch science and technology. Using his self-made microscopes, which were significantly more powerful than others at the time, he was

the first to observe and describe single-celled organisms, including bacteria. He referred to these microscopic creatures as "animalcules". His observations of pond water, saliva, and other substances led to the discovery of the microbial world.

The Pioneer of Microbiology: A V Leeuwenhoek

Based on the historical contributions of these scientists, Antoni van Leeuwenhoek is recognized as the first person to observe and describe bacteria. His meticulous observations using simple, yet powerful, microscopes provided the initial glimpse into the previously unseen world of microorganisms. He is often referred to as the "Father of Microbiology" for his pioneering work in microscopic observation.

Scientist	Known For	Discovery of Bacteria?
Robert Koch	Koch's postulates, disease identification (e.g., TB, cholera)	No
James Chadwick	Discovery of the neutron	No
Eugen Goldstein	Discovery of anode rays	No
A V Leeuwenhoek	First to observe and describe bacteria (animalcules)	Yes

Therefore, the scientist who discovered bacteria is A V Leeuwenhoek.

Revision Table: Key Discoveries

Scientist	Field	Major Contribution
Antoni van Leeuwenhoek	Microbiology	First observation of bacteria and other microorganisms
Robert Koch	Microbiology, Medicine	Koch's postulates, specific disease causation
James Chadwick	Physics	Discovery of the neutron
Eugen Goldstein	Physics	Discovery of anode rays

Additional Information on Bacteria Discovery

Antoni van Leeuwenhoek's discovery of bacteria in the 17th century was purely based on observation. He did not understand their function or their role in health and disease. It wasn't until much later, with the work of scientists like Louis Pasteur and Robert Koch in the 19th century, that the link between microorganisms (including bacteria) and diseases was firmly established. Leeuwenhoek's work laid the essential groundwork by simply revealing the existence of these microscopic life forms.

70. Answer: d

Explanation:

Understanding the Reserve Bank of India and its First Governor

The question asks about the individual who first held the office of the Governor of the Reserve Bank of India (RBI). The RBI is India's central bank, responsible for regulating the country's monetary policy.

Establishment of the RBI and its Early Leadership

The Reserve Bank of India was established on April 1, 1935, in accordance with the provisions of the Reserve Bank of India Act, 1934. As the central bank, it plays a crucial role

in maintaining financial stability and managing the currency. Like any major institution, it needed a head from its inception.

The first person appointed to lead the newly formed Reserve Bank of India as its Governor was an expatriate banker.

Analyzing the Options

Let's look at the provided options and their relation to the RBI governorship:

- **Sir James Braid Taylor:** Sir James Braid Taylor was actually the second Governor of the RBI, succeeding the first governor. He served from 1937 to 1943.
- **KR Puri:** K. R. Puri served as the tenth Governor of the RBI from 1975 to 1977. He was a much later governor compared to the RBI's founding.
- **HVR Iyengar:** H. V. R. Iyengar was the sixth Governor of the RBI, serving from 1957 to 1962. He also served long after the first governor.
- **Sir Osborne A Smith:** Sir Osborne Arkell Smith was an Australian banker. He was the first individual to be appointed as the Governor of the Reserve Bank of India when it was established in 1935. He served from April 1, 1935, to June 30, 1937.

Identifying the First RBI Governor

Based on historical records, the first Governor of the Reserve Bank of India was Sir Osborne A Smith. He was in office for a little over two years before being succeeded by Sir James Braid Taylor.

Governor	Term Start	Term End	Significance
Sir Osborne A Smith	April 1, 1935	June 30, 1937	First RBI Governor
Sir James Braid Taylor	July 1, 1937	February 17, 1943	Second RBI Governor
C. D. Deshmukh	August 11, 1943	June 30, 1949	First Indian RBI Governor

Therefore, the correct answer is Sir Osborne A Smith.

Revision Table: Key RBI Governors

Governor	Term	Notes
Sir Osborne A Smith	1935-1937	First Governor
Sir James Braid Taylor	1937-1943	Second Governor
C. D. Deshmukh	1943-1949	First Indian Governor
HVR Iyengar	1957-1962	Sixth Governor
KR Puri	1975-1977	Tenth Governor

Additional Information on RBI Governors

The Governor of the Reserve Bank of India is a very important position. The Governor is the chief executive of the central bank and is appointed by the Government of India. The Governor's role involves:

- Formulating, implementing, and monitoring the monetary policy.
- Being the banker to the government and banks.
- Managing foreign exchange.
- Issuing currency.
- Ensuring financial stability.

Understanding the history of the RBI's leadership, starting with the first governor Sir Osborne A Smith, provides insight into the evolution of India's financial system.

71. Answer: c

Explanation:

In all the figures except '3' the number of squares inside the figure is 9, while in figure 3, there are 10 squares inside the figure.

Hence, 'option 3' is the odd one out.

72. Answer: a

Explanation:

Understanding Electrical Resistance in Wires

The electrical resistance of a conductor, such as a cylindrical wire, depends on several factors: the material it is made of, its length, and its cross-sectional area. The relationship is described by a fundamental formula.

Formula for Wire Resistance

The resistance (R) of a wire is given by the formula:

$$R = \rho \frac{L}{A}$$

Where:

- ρ (rho) is the resistivity of the material (a property specific to the material).
- L is the length of the wire.
- A is the cross-sectional area of the wire.

For a cylindrical wire with radius r , the cross-sectional area A is given by:

$$A = \pi r^2$$

So, the resistance formula can also be written as:

$$R = \rho \frac{L}{\pi r^2}$$

Calculating Resistance of the First Wire

We are given a cylindrical wire with:

- Length = L
- Radius = r
- Resistance = R
- Material = Same for both wires, so resistivity ρ is constant.

Using the formula, the resistance of this first wire is:

$$R = \rho \frac{L}{\pi r^2} \text{ --- (Equation 1)}$$

Calculating Resistance of the Second Wire

The second wire is made of the same material (same ρ) but has different dimensions:

- Length $L' =$ thrice its original length $= 3L$
- Radius $r' =$ one-third its original radius $= \frac{r}{3}$

First, let's find the cross-sectional area A' of the second wire:

$$A' = \pi(r')^2 = \pi\left(\frac{r}{3}\right)^2 = \pi\frac{r^2}{9}$$

Now, let's find the resistance R' of the second wire using the resistance formula $R' = \rho\frac{L'}{A'}$:

$$R' = \rho\frac{3L}{\frac{\pi r^2}{9}}$$

To simplify, we can multiply the numerator by the reciprocal of the denominator:

$$R' = \rho \times 3L \times \frac{9}{\pi r^2}$$

$$R' = \rho\frac{27L}{\pi r^2}$$

Relating the New Resistance to the Original Resistance

We can rearrange the expression for R' to see its relationship with R . We can factor out the constant term:

$$R' = 27 \times \left(\rho\frac{L}{\pi r^2}\right)$$

From Equation 1, we know that $R = \rho\frac{L}{\pi r^2}$. Substituting this into the expression for R' :

$$R' = 27R$$

So, the resistance of the second wire is 27 times the resistance of the first wire.

Conclusion on Wire Resistance

When the length of a wire is tripled and its radius is reduced to one-third, its resistance becomes 27 times the original resistance, assuming the material remains the same.

Revision Table: Factors Affecting Wire Resistance

Factor	Effect on Resistance (R)	Relationship
Resistivity (ρ)	Higher resistivity means higher resistance	$R \propto \rho$
Length (L)	Longer wire means higher resistance	$R \propto L$
Cross-sectional Area (A)	Larger area means lower resistance	$R \propto \frac{1}{A}$
Radius (r)	Larger radius means lower resistance	$R \propto \frac{1}{r^2}$ (since $A = \pi r^2$)

Additional Information on Electrical Properties

Resistivity ρ is an intrinsic property of a material, indicating how strongly it resists electrical current flow. Good conductors like copper and aluminum have low resistivity, while insulators like rubber and glass have high resistivity. The unit of resistivity is ohm-meter ($\Omega \cdot m$).

The reciprocal of resistance is conductance (G), measured in Siemens (S). The reciprocal of resistivity is conductivity (σ), measured in Siemens per meter (S/m). The relationship between conductance and conductivity is $G = \sigma \frac{A}{L}$.

Temperature also affects the resistance of most materials. For metals, resistance generally increases with increasing temperature. For semiconductors and insulators, the effect can vary.

73. Answer: d

Explanation:

Understanding the Discount Percent Question

The question asks for the discount percent offered on the soap by the store. To find the discount percent, we generally need to know the amount of discount given and the original price of the item(s) purchased. The formula for discount percent is:

$$\left(\text{Discount Percent} = \left(\frac{\text{Discount Amount}}{\text{Original Price}} \right) \times 100 \right)$$

We need to evaluate whether each statement, independently or together, provides enough information to calculate this discount percent.

Analyzing Statement I: Free Soap Offer

Statement I says: "The store is giving 1 soap free on purchase of three."

Let's analyze this statement to see if it's sufficient to find the discount percent on the soap.

- This is a classic "Buy 3, Get 1 Free" offer.
- When you buy 3 soaps, you receive a total of $(3 + 1 = 4)$ soaps.
- You pay the price of 3 soaps but get 4 soaps.
- The value of the 1 free soap is the discount you receive.

Let the original price of one soap be (P) .

- The original price of the 4 soaps you receive is $(4 \times P = 4P)$.
- The amount you pay is for 3 soaps, which is $(3 \times P = 3P)$.
- The discount amount is the difference between the original price of what you received and what you paid: $(4P - 3P = P)$.
- The discount percent is calculated on the original price of the items received (4 soaps).

Using the formula for discount percent:

$$\left(\text{Discount Percent} = \left(\frac{\text{Discount Amount}}{\text{Original Price of Items Received}} \right) \times 100 \right)$$

$$\left(\text{Discount Percent} = \left(\frac{P}{4P} \right) \times 100 \right)$$

Since (P) is the price of one soap, it cannot be zero. We can cancel out (P) from the numerator and the denominator.

$$\left(\text{Discount Percent} = \left(\frac{1}{4} \right) \times 100 = 25\% \right)$$

Statement I alone allows us to calculate a specific discount percent (25%). Therefore, Statement I alone is sufficient to answer the question about the discount percent offered on the soap.

Analyzing Statement II: Fixed Discount Amount

Statement II says: "Rs. 10 discount is offered on purchase of soap worth 36."

Let's analyze this statement to see if it's sufficient to find the discount percent on the soap.

- This statement directly gives us the discount amount and the original price for a specific purchase of soap.
- Discount Amount = Rs. 10
- Original Price = Rs. 36

Using the formula for discount percent:

$$\left(\frac{\text{Discount Amount}}{\text{Original Price}} \right) \times 100$$

$$\left(\frac{10}{36} \right) \times 100$$

We can simplify the fraction $\frac{10}{36}$ by dividing both numerator and denominator by 2:

$$\frac{10}{36} = \frac{5}{18}$$

So, the discount percent is:

$$\left(\frac{5}{18} \right) \times 100 = \frac{500}{18} = \frac{250}{9}$$

The discount percent is approximately (27.78%) . Statement II alone allows us to calculate a specific discount percent. Therefore, Statement II alone is sufficient to answer the question about the discount percent offered on the soap.

Conclusion on Statement Sufficiency

Based on our analysis:

- Statement I alone is sufficient to determine the discount percent.
- Statement II alone is sufficient to determine the discount percent.

Since either statement by itself provides enough information to answer the question, the correct answer is that either statement I or statement II is sufficient.

Revision Table: Discount Calculations

Statement	Original Price	Amount Paid	Discount Amount	Discount Percent Formula	Calculated Discount Percent	Sufficiency
I (Buy 3 Get 1 Free)	Price of 4 soaps ($4P$)	Price of 3 soaps ($3P$)	Price of 1 soap (P)	$\left(\frac{P}{4P}\right) \times 100$	(25%)	Sufficient
II (Rs. 10 off on Rs. 36)	Rs. 36	Rs. $(36 - 10) = 26$	Rs. 10	$\left(\frac{10}{36}\right) \times 100$	$\left(\frac{250}{9}\right)\% \approx 27.78\%$	Sufficient

Additional Information on Data Sufficiency Questions

Data sufficiency questions test your ability to determine if the given statements provide enough information to answer a question, rather than actually solving the question yourself. Here's a general approach:

- Read the question carefully to understand what needs to be determined.
- Analyze Statement I alone: Assume Statement II is false or not available. Can you find a unique answer to the question using only Statement I? If yes, Statement I is sufficient.
- Analyze Statement II alone: Assume Statement I is false or not available. Can you find a unique answer to the question using only Statement II? If yes, Statement II is sufficient.
- Analyze Statements I and II together: If neither statement alone is sufficient, consider if both statements combined provide enough information. If yes, both together are sufficient.
- If even both statements together are not enough, then the data is insufficient.

Common outcomes in data sufficiency problems:

- Statement I alone is sufficient (and II is not).

- Statement II alone is sufficient (and I is not).
- Either Statement I or Statement II alone is sufficient.
- Both statements together are sufficient (but neither alone is).
- Neither statement I nor Statement II is sufficient (even together).

In this specific soap discount question, both statements individually provided enough information to calculate the discount percent, leading to the conclusion that either statement is sufficient.

74. Answer: d

Explanation:

Understanding the Problem: Calculating Airplane Distance

The question asks us to calculate the total distance covered by an airplane given its speed and the duration of its flight. The speed is provided in meters per second (m/s), and the time is given in hours. We need to find the distance in kilometers (km).

To solve this problem, we need to use the fundamental relationship between distance, speed, and time:

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

However, before we can directly apply this formula, we must ensure that the units are consistent. The speed is in m/s, and the time is in hours, but we need the final distance in km. Therefore, we need to convert the speed from m/s to km/h.

Converting Speed Units (m/s to km/h)

To convert speed from meters per second (m/s) to kilometers per hour (km/h), we need to consider the following conversion factors:

- 1 kilometer (km) = 1000 meters (m)
- 1 hour (h) = 60 minutes = 60 × 60 seconds = 3600 seconds (s)

So, 1 m/s can be converted as follows:

$$1 \text{ m/s} = \frac{1 \text{ m}}{1 \text{ s}}$$

To convert meters to kilometers, we divide by 1000:

$$1 \text{ m} = \frac{1}{1000} \text{ km}$$

To convert seconds to hours, we divide by 3600:

$$1 \text{ s} = \frac{1}{3600} \text{ h}$$

Substituting these into the m/s expression:

$$1 \text{ m/s} = \frac{\frac{1}{1000} \text{ km}}{\frac{1}{3600} \text{ h}} = \frac{1}{1000} \times \frac{3600}{1} \text{ km/h} = \frac{3600}{1000} \text{ km/h} = 3.6 \text{ km/h}$$

So, the conversion factor from m/s to km/h is 3.6. We can multiply the speed in m/s by 3.6 to get the speed in km/h.

Given speed = 50 m/s.

$$\text{Speed in km/h} = 50 \text{ m/s} \times 3.6 \text{ km/h per m/s}$$

$$\text{Speed in km/h} = 180 \text{ km/h}$$

Calculating the Distance Covered

Now that we have the speed in km/h and the time in hours, we can calculate the distance in km.

- Speed = 180 km/h
- Time = 5 hours

Using the formula: Distance = Speed $\times$ Time

$$\text{Distance} = 180 \text{ km/h} \times 5 \text{ hours}$$

$$\text{Distance} = 900 \text{ km}$$

Summary of Calculation Steps

Step	Description	Calculation	Result
1	Identify given values and required unit.	Speed = 50 m/s, Time = 5 hours, Need Distance in km.	-
2	Convert speed from m/s to km/h.	Multiply speed (m/s) by 3.6.	$50 \text{ m/s} \times 3.6 = 180 \text{ km/h}$
3	Calculate distance using the formula.	Distance = Speed $\times$ Time	$180 \text{ km/h} \times 5 \text{ hours}$
4	State the final answer.	Distance covered.	900 km

The airplane will cover a distance of 900 km in 5 hours at a speed of 50 m/s.

Revision Table: Airplane Speed and Distance Calculation

Concept	Formula/Relation	Key Units	Conversion Example
Speed, Distance, Time	$D = S \times T$	Distance (km, m), Speed (km/h, m/s), Time (h, s)	-
Unit Conversion (Speed)	m/s to km/h	$1 \text{ m/s} = 3.6 \text{ km/h}$	$50 \text{ m/s} \times 3.6 = 180 \text{ km/h}$
Distance Calculation	$D = S \times T$	Ensure S and T units are compatible (e.g., km/h and h)	$180 \text{ km/h} \times 5 \text{ h} = 900 \text{ km}$

Additional Information: Speed, Velocity, and Unit Conversion Tips

Speed vs. Velocity: In physics, speed is the magnitude of velocity. Speed is a scalar quantity (only magnitude), while velocity is a vector quantity (magnitude and direction). In this problem, we are only concerned with how fast the airplane is moving, so speed is the relevant concept.

Common Unit Conversions: It's helpful to remember common unit conversions for speed problems:

- m/s to km/h: Multiply by 3.6
- km/h to m/s: Divide by 3.6
- km to m: Multiply by 1000
- m to km: Divide by 1000
- hours to minutes: Multiply by 60
- minutes to seconds: Multiply by 60
- hours to seconds: Multiply by 3600

Mastering these conversions is key to solving many physics problems involving motion. Always check the units required for the final answer and ensure all given values are converted to compatible units before performing calculations.

75. Answer: d

Explanation:

Analyzing the Statement and Assumptions on Vitamin D

This question asks us to determine which of the given assumptions is implicit in the provided statement. An assumption is something taken for granted or accepted as true without proof. In statement and assumption questions, we need to find what lies beneath the surface of the statement – what the statement maker must believe for the statement to be true.

Understanding the Statement on Vitamin D Production

The statement says:

Statement: The human body produces Vitamin D when exposed to sunlight.

This is a factual statement describing one way the human body obtains Vitamin D: through production triggered by sunlight exposure.

Examining Assumption I: Vitamin D and Food Consumption

Let's look at the first assumption:

Assumption I: The human body will have Vitamin D even if it is not consumed via food.

The statement tells us that sunlight exposure **produces** Vitamin D in the body. If production happens via sunlight, it logically follows that the body can get Vitamin D from this source, independent of whether it is consumed through food. The statement provides a source of Vitamin D other than food. Therefore, assuming that the body can have Vitamin D without food consumption (because it can produce it with sunlight) is directly supported by and implicit in the statement.

Evaluating Assumption II: Global Vitamin D Deficiency

Now, let's consider the second assumption:

Assumption II: A large portion of the global population suffers from Vitamin D deficiency.

The statement only discusses the **production** of Vitamin D through sunlight. It does not provide any information about the prevalence of Vitamin D deficiency in any population, global or otherwise. The fact that the body **can** produce Vitamin D in sunlight doesn't tell us anything about whether people **actually** get enough sunlight, or if other factors lead to deficiency. This assumption introduces new information that is not implied or required for the original statement to be true.

Conclusion on Implicit Assumptions

Based on the analysis:

- Assumption I is implicit because the statement provides a mechanism (sunlight exposure leading to production) by which the body can obtain Vitamin D, implying that food consumption is not the only way to have it.
- Assumption II is not implicit because the statement focuses only on the production method and does not offer any information or implication regarding the health status or nutritional deficiencies of the population.

Therefore, only Assumption I is implicit in the statement.

Revision Table: Statement vs. Assumptions Analysis

Element	Content / Idea	Is it Implicit in the Statement?	Reasoning
Statement	Body produces Vitamin D with sunlight exposure.	N/A (This is the given fact)	The core premise provided.
Assumption I	Body has Vitamin D without food consumption.	Yes	Statement provides an alternative source (sunlight production) to food consumption for obtaining Vitamin D.
Assumption II	Large global population suffers from Vitamin D deficiency.	No	Statement does not contain information about population health, deficiency rates, or actual sunlight exposure levels of people.

Additional Information on Vitamin D

Vitamin D is a crucial nutrient for bone health, immune function, and overall well-being. While the statement correctly identifies sunlight exposure as a primary way our bodies produce Vitamin D, it's important to know other aspects:

- **Dietary Sources:** Vitamin D can also be obtained from food, although it's not naturally present in many foods. Fatty fish (like salmon, mackerel, and sardines), fish liver oils, and egg yolks contain Vitamin D. Many foods, such as milk, cereals, and orange juice, are fortified with Vitamin D.
- **Production Factors:** The amount of Vitamin D produced from sunlight depends on factors like the time of day, season, latitude, skin pigmentation, and sunscreen use.
- **Importance:** Vitamin D helps the body absorb calcium and phosphorus, which are vital for building and maintaining strong bones. Deficiency can lead to rickets in children and osteomalacia or osteoporosis in adults.
- **Deficiency Causes:** Deficiency can occur if people don't get enough sunlight exposure (due to lifestyle, location, or skin tone), consume insufficient dietary sources, or have medical conditions affecting absorption.

Understanding these points provides a broader context but does not change the analysis of what is strictly implicit in the original, specific statement provided in the question.

76. Answer: a

Explanation:

Understanding the Trigonometry Problem

The question asks us to find the value of $\cot \theta$ given that $\sin \theta = \frac{15}{17}$. This is a fundamental trigonometry problem involving trigonometric ratios in a right-angled triangle.

Relating $\sin \theta$ to a Right Triangle

In a right-angled triangle, the trigonometric ratio $\sin \theta$ is defined as the ratio of the length of the side opposite to the angle θ to the length of the hypotenuse. So, if $\sin \theta = \frac{15}{17}$, we can think of a right triangle where:

- Opposite side = 15 units
- Hypotenuse = 17 units

Finding the Missing Side using the Pythagorean Theorem

To find $\cot \theta$, which is defined as $\frac{\text{Adjacent}}{\text{Opposite}}$, we first need to find the length of the adjacent side. We can use the Pythagorean theorem, which states that in a right-angled triangle, the square of the hypotenuse is equal to the sum of the squares of the other two sides (opposite and adjacent).

The theorem is expressed as:

$$\text{Opposite}^2 + \text{Adjacent}^2 = \text{Hypotenuse}^2$$

Substituting the known values:

$$(15)^2 + \text{Adjacent}^2 = (17)^2$$

$$225 + \text{Adjacent}^2 = 289$$

Now, we solve for the Adjacent side:

$$\text{Adjacent}^2 = 289 - 225$$

$$\text{Adjacent}^2 = 64$$

Taking the square root of both sides (and assuming the side length is positive):

$$\text{Adjacent} = \sqrt{64}$$

$$\text{Adjacent} = 8 \text{ units}$$

Calculating $\cot \theta$

The trigonometric ratio $\cot \theta$ is defined as the ratio of the length of the side adjacent to the angle θ to the length of the side opposite to the angle θ .

$$\cot \theta = \frac{\text{Adjacent}}{\text{Opposite}}$$

Using the values we found:

$$\cot \theta = \frac{8}{15}$$

Summary of Sides and Ratios

For our right-angled triangle with angle θ :

Side	Length
Opposite	15
Hypotenuse	17
Adjacent	8

Based on these lengths, the trigonometric ratios are:

- $\sin \theta = \frac{\text{Opposite}}{\text{Hypotenuse}} = \frac{15}{17}$ (Given)
- $\cos \theta = \frac{\text{Adjacent}}{\text{Hypotenuse}} = \frac{8}{17}$
- $\tan \theta = \frac{\text{Opposite}}{\text{Adjacent}} = \frac{15}{8}$
- $\cot \theta = \frac{\text{Adjacent}}{\text{Opposite}} = \frac{8}{15}$
- $\sec \theta = \frac{\text{Hypotenuse}}{\text{Adjacent}} = \frac{17}{8}$
- $\csc \theta = \frac{\text{Hypotenuse}}{\text{Opposite}} = \frac{17}{15}$

Thus, the value of $\cot \theta$ is $\frac{8}{15}$.

Revision Table: Trigonometric Ratios

Ratio	Definition (Right Triangle)	Reciprocal Of
$\sin \theta$	$\frac{\text{Opposite}}{\text{Hypotenuse}}$	$\csc \theta$
$\cos \theta$	$\frac{\text{Adjacent}}{\text{Hypotenuse}}$	$\sec \theta$
$\tan \theta$	$\frac{\text{Opposite}}{\text{Adjacent}}$	$\cot \theta$
$\cot \theta$	$\frac{\text{Adjacent}}{\text{Opposite}}$	$\tan \theta$
$\sec \theta$	$\frac{\text{Hypotenuse}}{\text{Adjacent}}$	$\cos \theta$
$\csc \theta$	$\frac{\text{Hypotenuse}}{\text{Opposite}}$	$\sin \theta$

Additional Information on Trigonometry

Trigonometry is a branch of mathematics that studies relationships between side lengths and angles of triangles. The trigonometric ratios discussed here (sine, cosine, tangent, cosecant, secant, and cotangent) are fundamental. They are often remembered using mnemonics like SOH CAH TOA:

- SOH: $\sin \theta = \frac{\text{Opposite}}{\text{Hypotenuse}}$
- CAH: $\cos \theta = \frac{\text{Adjacent}}{\text{Hypotenuse}}$
- TOA: $\tan \theta = \frac{\text{Opposite}}{\text{Adjacent}}$

The reciprocal identities are also very useful:

- $\cot \theta = \frac{1}{\tan \theta}$
- $\sec \theta = \frac{1}{\cos \theta}$
- $\csc \theta = \frac{1}{\sin \theta}$

We could also solve this problem using trigonometric identities. For example, we know that $\sin^2 \theta + \cos^2 \theta = 1$. Given $\sin \theta = \frac{15}{17}$, we could find $\cos \theta$: $\cos^2 \theta = 1 - \sin^2 \theta = 1 - \left(\frac{15}{17}\right)^2 = 1 - \frac{225}{289} = \frac{289-225}{289} = \frac{64}{289}$. So, $\cos \theta = \sqrt{\frac{64}{289}} = \frac{8}{17}$ (assuming θ is in a quadrant where cosine is positive). Then, $\cot \theta = \frac{\cos \theta}{\sin \theta} = \frac{8/17}{15/17} = \frac{8}{15}$. This confirms our result obtained using the triangle method.

77. Answer: d

Explanation:

Solution: $2000 = (2/5) T \Rightarrow T = ₹5,000$.

$\therefore ₹5,000$

78. Answer: d

Explanation:

Understanding Harvest Festivals in India

Harvest festivals are celebrations that take place around the time of the main harvest in a particular region. They are often expressions of gratitude for a good harvest and prosperity. India, being an agricultural country, has many such festivals celebrated across different states.

Analyzing the Given Options

Let's examine each of the provided options to determine which one is a harvest festival:

1. **Teej**: Teej is primarily celebrated by women in North India, especially in Rajasthan, Uttar Pradesh, Madhya Pradesh, and Bihar. It is observed during the monsoon season and is associated with the reunion of Lord Shiva and Goddess Parvati. While celebrated during a season related to agriculture (monsoon leading to harvest), it is not considered a harvest festival itself but rather a festival related to the monsoon, marital bliss, and fasting.
2. **Janmashtami**: Janmashtami celebrates the birth of Lord Krishna. It is a major Hindu festival observed across India and by Hindu communities worldwide. It is a religious festival commemorating a divine event, not related to the agricultural harvest cycle.
3. **Deepawali (Diwali)**: Deepawali, the festival of lights, is one of the most widely celebrated festivals in India. It signifies the victory of light over darkness, good over evil, and knowledge over ignorance. While it marks the beginning of the financial year for many businesses and is celebrated after the harvest season in some regions, its primary significance is not tied directly to the harvest itself.
4. **Onam**: Onam is a major annual harvest festival celebrated in the state of Kerala, India. It falls during the Malayalam month of Chingam (August–September) and marks the homecoming of the mythical King Mahabali. Onam is intrinsically linked to the harvest season and is celebrated with great enthusiasm, featuring boat races,

flower carpets (Pookalam), traditional feasts (Onasadya), and cultural performances. It is clearly a harvest festival.

Identifying the Harvest Festival

Based on the analysis of each option, Onam is the festival directly associated with celebrating the harvest season, particularly in Kerala.

Festival	Primary Significance	Classification
Teej	Monsoon, Marital Bliss	Seasonal/Religious
Janmashtami	Birth of Lord Krishna	Religious
Deepawali	Festival of Lights, Victory of Good over Evil	Religious/Cultural
Onam	Harvest Celebration, Homecoming of King Mahabali	Harvest Festival

Therefore, among the given options, Onam is the harvest festival.

Revision Table: Key Indian Festivals

Your Personal Exams Guide

Festival	State(s) Primarily Celebrated	Timing (Approximate)	Key Aspect
Onam	Kerala	August/September (Chingam)	Harvest festival, King Mahabali's visit
Teej	Rajasthan, UP, MP, Bihar	July/August (Shravan)	Monsoon festival, Marital well-being
Janmashtami	Across India	August/September	Birth of Lord Krishna
Deepawali	Across India	October/November	Festival of Lights, Victory of light over darkness
Baisakhi	Punjab, Haryana	April	Harvest (Rabi crop), Sikh New Year
Pongal	Tamil Nadu	January	Harvest festival (Thai crop)
Makar Sankranti	Across India (different names)	January	Harvest festival, Sun's transit into Capricorn
Bihu	Assam	Various times (Rongali Bihu is harvest-related)	Assamese New Year, Harvest festivals

Additional Information on Indian Harvest Festivals

India's diverse geography and climate result in different harvesting periods, leading to various harvest festivals celebrated throughout the year. These festivals are deeply intertwined with local culture, traditions, and agricultural practices.

- **Regional Variations:** While Onam is specific to Kerala, other states have their own harvest festivals like Pongal in Tamil Nadu, Baisakhi in Punjab, Bihu in Assam, and Makar Sankranti celebrated under various names across many states.
- **Significance:** Harvest festivals are not just about celebrating the yield; they also involve rituals dedicated to deities of agriculture, cattle (important for farming), and

nature. They are a time for community gathering, feasting, dancing, and expressing gratitude for sustenance.

- **Cultural Richness:** These festivals showcase the rich cultural heritage of India through traditional music, dance forms, art (like Pookalam in Onam), specific foods prepared from the fresh harvest, and traditional attire.

Understanding these festivals provides insight into the agricultural backbone of India and the cultural significance of harvest in people's lives.

79. Answer: b

Explanation:

Solving the Equation and Evaluating the Expression

The problem asks us to first solve a given equation involving a variable 'x' and then use the value of 'x' to evaluate another expression.

The first part is the equation: $5050 \times 0.5x = 25250$

Our goal is to find the value of 'x'. To do this, we need to isolate 'x' on one side of the equation. Let's simplify the left side first or move terms around.

The equation can be written as:

$$5050 \times (0.5 \times x) = 25250$$

We can perform the multiplication on the left or divide both sides by 5050 first. Let's divide both sides by 5050:

$$\frac{5050 \times 0.5x}{5050} = \frac{25250}{5050}$$

This simplifies to:

$$0.5x = \frac{25250}{5050}$$

Now, let's calculate the value on the right side. Notice that 25250 is exactly 5 times 5050 (since $5050 \times 5 = 25250$).

$$0.5x = 5$$

To find 'x', we need to divide both sides by 0.5 (which is the same as multiplying by 2):

$$x = \frac{5}{0.5}$$

$$x = 10$$

So, the value of the variable 'x' is 10.

The second part of the problem asks us to evaluate the expression: $2505 \div x^2$

Now we substitute the value of $x = 10$ into this expression.

First, let's calculate x^2 :

$$x^2 = 10^2 = 10 \times 10 = 100$$

Now substitute this value back into the expression:

$$2505 \div x^2 = 2505 \div 100$$

Dividing a number by 100 involves moving the decimal point two places to the left.

$$2505 \div 100 = 25.05$$

So, the value of the expression $2505 \div x^2$ is 25.05.

Let's compare this result with the given options:

- Option 1: 0.2505
- Option 2: 25.05
- Option 3: 2.505
- Option 4: 250.5

Our calculated value, 25.05, matches Option 2.

The steps followed were:

1. Solve the initial equation $5050 \times 0.5x = 25250$ for the variable 'x'.
2. Substitute the found value of 'x' into the expression $2505 \div x^2$.
3. Evaluate the expression.

This systematic approach helps in tackling multi-step mathematical problems accurately.

Revision Table: Key Calculations

Step	Calculation	Result
Solve for $0.5x$	$0.5x = \frac{25250}{5050}$	$0.5x = 5$
Solve for x	$x = \frac{5}{0.5}$	$x = 10$
Calculate x^2	$x^2 = 10^2$	$x^2 = 100$
Evaluate $2505 \div x^2$	$2505 \div 100$	25.05

Additional Information: Solving Linear Equations

A linear equation in one variable, like the one we solved, can be written in the form $ax + b = c$, where 'a', 'b', and 'c' are numbers and 'x' is the variable.

In our equation, $5050 \times 0.5x = 25250$, we can simplify 5050×0.5 first:

$$5050 \times 0.5 = 2525$$

So the equation becomes $2525x = 25250$. This is now clearly in the form $ax = c$.

To solve for 'x', we divide both sides by 'a':

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

In our case:

$$x = \frac{25250}{2525}$$

$$x = 10$$

This confirms the value of x we found earlier. Solving linear equations involves using inverse operations (like division to undo multiplication, or subtraction to undo addition) to isolate the variable.

80. Answer: a

Explanation:

Understanding Coefficient of Volume Expansion

The coefficient of volume expansion describes how the volume of a substance changes in response to a change in temperature, while the pressure is held constant. It is typically denoted by the symbol β (beta). For an infinitesimal change in temperature dT and corresponding change in volume dV , the coefficient of volume expansion is defined as:

$$\beta = \frac{1}{V} \left(\frac{\partial V}{\partial T} \right)_p$$

For a finite temperature change ΔT , the change in volume ΔV is approximately given by:

$$\Delta V = V_0 \beta \Delta T$$

where V_0 is the initial volume. A lower coefficient of volume expansion means that the material's volume changes less for a given temperature change.

Analyzing the Given Materials

We are asked to identify the material from the given options that has a low coefficient of volume expansion. The options are Iron, Mercury, Brass, and Aluminium. These are common engineering materials and a liquid (Mercury).

Typical Coefficients of Volume Expansion

The coefficient of volume expansion varies significantly between different substances. For solids, the coefficient of volume expansion (β) is approximately three times the coefficient of linear expansion (α), i.e., $\beta \approx 3\alpha$. For liquids, the volume expansion is usually larger than for solids.

Let's look at approximate typical values for the coefficient of volume expansion (β):

Material	Approximate Coefficient of Linear Expansion (α) in ($10^{-6} \text{ }^{\circ}\text{C}^{-1}$)	Approximate Coefficient of Volume Expansion ($\beta \approx 3\alpha$ for solids) in ($10^{-6} \text{ }^{\circ}\text{C}^{-1}$)
Iron (Steel)	12	$3 \times 12 = 36$
Mercury (Liquid)	N/A	181
Brass	19	$3 \times 19 = 57$
Aluminium	23	$3 \times 23 = 69$

Comparing the Values

Comparing the approximate values for the coefficient of volume expansion:

- Iron: $36 \times 10^{-6} \text{ }^{\circ}\text{C}^{-1}$
- Mercury: $181 \times 10^{-6} \text{ }^{\circ}\text{C}^{-1}$
- Brass: $57 \times 10^{-6} \text{ }^{\circ}\text{C}^{-1}$
- Aluminium: $69 \times 10^{-6} \text{ }^{\circ}\text{C}^{-1}$

From this comparison, Iron has the lowest coefficient of volume expansion among the listed materials.

Conclusion

Based on the typical values, Iron exhibits the lowest coefficient of volume expansion compared to Mercury, Brass, and Aluminium. This means that Iron's volume changes least when its temperature changes, which is a property useful in applications requiring dimensional stability under varying temperatures.

Revision Table: Thermal Expansion Coefficients

Material	Type	β ($10^{-6} \text{ }^{\circ}\text{C}^{-1}$)	Relative Expansion
Iron	Solid	~36	Lowest
Mercury	Liquid	~181	Highest
Brass	Solid	~57	Medium
Aluminium	Solid	~69	High

Additional Information: Applications of Low Expansion Materials

Materials with low coefficients of thermal expansion are valuable in various applications where maintaining precise dimensions despite temperature fluctuations is critical. Examples include:

- **Construction:** Bridges and buildings use expansion joints to accommodate thermal expansion, but choosing materials with lower expansion can simplify design.
- **Scientific Instruments:** Precision measurement tools, telescopes, and laboratory equipment require materials that do not expand or contract significantly with temperature changes to maintain accuracy.
- **Cookware:** Some cookware materials are chosen for their relatively low thermal expansion and good thermal conductivity.
- **Engine Components:** In internal combustion engines, different parts expand at different rates. Materials with controlled expansion are crucial for maintaining proper clearances and preventing seizing.
- **Glassware:** Special types of glass, like borosilicate glass (e.g., Pyrex), have much lower thermal expansion than ordinary soda-lime glass, making them resistant to thermal shock (sudden temperature changes).

The coefficient of volume expansion is a key thermal property to consider when selecting materials for applications subjected to temperature variations.

81. Answer: a

Explanation:

Therefore, $\sec 45^\circ = (1)/(\cos 45^\circ) = \frac{1}{1/\sqrt{2}} = \sqrt{2}$. The value of $\tan 60^\circ$ is known to be $\sqrt{3}$. Therefore, the correct option is the one that matches $-\sqrt{3} + \sqrt{2}$. Therefore, the correct answer is $-\sqrt{3} + \sqrt{2}$.

82. Answer: d

Explanation:

Finding the Number Equally Distant from 50 and 84

The problem asks us to find a number that is the same distance away from 50 as it is from 84. Let's call the unknown number x .

According to the question:

- The number is as much greater than 50: This means the difference between the number and 50, which can be written as $x - 50$.
- The number is as much lesser than 84: This means the difference between 84 and the number, which can be written as $84 - x$.

The phrase "as much greater than 50 as it is lesser than 84" tells us that these two differences are equal. So, we can set up an equation:

$$x - 50 = 84 - x$$

Now, we need to solve this equation for x . To do this, we can gather the x terms on one side and the constant terms on the other side.

Add x to both sides of the equation:

$$x - 50 + x = 84 - x + x$$

$$2x - 50 = 84$$

Add 50 to both sides of the equation:

$$2x - 50 + 50 = 84 + 50$$

$$2x = 134$$

Finally, divide both sides by 2 to find the value of x :

$$\frac{2x}{2} = \frac{134}{2}$$

$$x = 67$$

So, the number is 67.

We can check this by verifying if 67 is the same distance from 50 and 84:

- Distance from 50: $67 - 50 = 17$
- Distance from 84: $84 - 67 = 17$

Since $17 = 17$, our answer is correct. The number is indeed 67.

Revision Table: Key Concepts

Concept	Description	Application in Problem
Algebraic Equation	A mathematical statement that two expressions are equal, often containing variables.	Setting up $x - 50 = 84 - x$ to represent the problem.
Solving Linear Equations	Finding the value(s) of the variable(s) that make the equation true.	Using addition and division to isolate x and find its value.
Representing Differences	Using subtraction to show how much greater or lesser one number is compared to another.	$x - 50$ for "greater than 50" and $84 - x$ for "lesser than 84".

Additional Information: Finding the Midpoint

Another way to think about a number that is "as much greater than A as it is lesser than B" is that the number is exactly in the middle of A and B. In other words, the number is the average (or midpoint) of A and B.

The formula for the average of two numbers, A and B, is:

$$\text{Average} = \frac{A + B}{2}$$

In this problem, $A = 50$ and $B = 84$. Let's use the average formula to find the number:

$$\text{Number} = \frac{50 + 84}{2}$$

$$\text{Number} = \frac{134}{2}$$

$$\text{Number} = 67$$

This method gives the same result and provides a quicker way to solve this specific type of problem.

83. Answer: b

Explanation:

Understanding Water's Peculiar Behavior Between 0°C and 4°C

Most substances expand when heated and contract when cooled. This is their typical thermal expansion behavior. However, water exhibits an unusual property, particularly in the temperature range between 0°C and 4°C. This behavior is known as the anomalous expansion of water.

Anomalous Expansion of Water Explained

When water is heated from its freezing point, 0°C, up to 4°C, it does not expand like most liquids. Instead, it contracts. Its density increases during this temperature rise, reaching its maximum density at 4°C. Beyond 4°C, water behaves normally; its density decreases, and its volume increases as the temperature rises further.

Why Volume Decreases Between 0°C to 4°C

For a given amount (mass) of water, density and volume are inversely related. Density (ρ) is defined as mass (m) per unit volume (V):

$$\rho = \frac{m}{V}$$

If the mass (m) is constant, then Volume (V) is inversely proportional to density (ρ):

$$V \propto \frac{1}{\rho}$$

Since the density of water increases as it is heated from 0°C to 4°C , its volume must decrease during this temperature range. At 4°C , the density is maximum, so the volume for a given mass is minimum.

Let's summarize the behavior:

- From 0°C to 4°C : Density increases, Volume decreases.
- At 4°C : Density is maximum, Volume is minimum.
- Above 4°C : Density decreases, Volume increases (normal expansion).

Therefore, the volume of a given amount of water decreases between 0°C to 4°C .

Revision Table: Water Properties

Temperature Range	Density Change	Volume Change (Fixed Mass)	Behavior Type
0°C to 4°C	Increases	Decreases	Anomalous
At 4°C	Maximum	Minimum	Turning Point
Above 4°C	Decreases	Increases	Normal

Additional Information: Significance of Anomalous Expansion

The anomalous expansion of water is crucial for aquatic life in cold climates. As winter approaches and the temperature drops, the surface water cools. When the surface water reaches 4°C , it is denser than the water above or below, so it sinks. Colder water (between 0°C and 4°C) or ice (at 0°C or below) is less dense than water at 4°C , so it floats. This means the warmest water in a lake during winter is at the bottom (4°C), while ice forms on the surface. This prevents lakes and rivers from freezing solid from bottom to top, allowing aquatic organisms to survive in the warmer water layer below the ice.

Understanding this property helps explain various natural phenomena involving water and temperature changes.

84. Answer: b

Explanation:

Understanding Abbreviations in Engineering Drawing

Engineering drawing is a universal language used by engineers, designers, and manufacturers to communicate precise information about objects. This information includes geometry, dimensions, tolerances, material, and surface finish. To make drawings concise and easy to read, especially for complex designs, engineers often use standard abbreviations and symbols.

What does LH mean in Engineering Drawing?

In the context of engineering drawing and design, specific abbreviations are used to denote characteristics or specifications of components. The abbreviation "LH" is a common one with a specific meaning.

Let's look at the options provided and analyze their relevance in engineering drawing:

- **Level Hide:** This term does not have a standard, widely recognized meaning as an abbreviation in engineering drawings for describing components or features.
- **Left Hand:** This is a standard abbreviation in engineering drawing. It is frequently used to specify the handedness of a component, most commonly threads. A "Left Hand" thread tightens when turned counter-clockwise, which is opposite to a standard "Right Hand" thread. It can also refer to components designed specifically for the left side of an assembly, such as doors, panels, or symmetrical parts with specific orientations.
- **Low Heat:** While "Heat" might be relevant in material properties or manufacturing processes, "Low Heat" is not a standard abbreviation used in engineering drawings to describe a feature or component's handedness or primary characteristic in this manner.
- **Limit of Height:** "Height" might be a dimension, but "Limit of Height" is not a standard abbreviation like "LH" in engineering drawing. Tolerances related to dimensions have specific symbols and notations (e.g., limit dimensions, tolerance zones).

Based on standard engineering practices and commonly accepted abbreviations, "LH" unambiguously signifies "Left Hand". This is particularly important when specifying parts that have handedness, such as screws, nuts, hinges, or assemblies where orientation matters.

Significance of Left Hand (LH) Designation

Identifying parts as "Left Hand" is crucial for manufacturing, assembly, and maintenance. Using the correct handedness ensures components fit together as intended and function correctly. For example, using a "Left Hand" thread where a "Right Hand" thread is required would prevent proper assembly.

Summary of LH Meaning

In the field of engineering drawing, the letters LH are an abbreviation that specifies the "Left Hand" nature of a component or feature. This is distinct from "Right Hand" (often abbreviated as RH or implied if not specified, as RH is standard for threads).

Abbreviation	Meaning in Engineering Drawing
LH	Left Hand
RH	Right Hand

Revision Table: Engineering Drawing Abbreviations

Understanding common abbreviations is key to reading and creating engineering drawings effectively. This table revisits the main point.

Term	Significance in Engineering Drawing
LH	Indicates a component or feature is designed for the left side or has a left-hand orientation (e.g., left-hand thread).

Additional Information: Handedness in Design

Handedness is a critical design consideration for many mechanical components and assemblies. Beyond threads, it applies to:

- Door hinges and handles: Specifying if a door opens to the left or right requires considering its hinges and hardware.

- Symmetrical parts with specific installation orientations: A part might look symmetrical but have features (like mounting holes or connection points) that dictate it fits only on the left or right side of a larger assembly.
- Mirror-image parts: Often, left and right versions of a part are manufactured separately, each requiring its own drawing and part number, distinguished by LH and RH designations.

Properly documenting handedness using abbreviations like LH in engineering drawings prevents errors in manufacturing and assembly processes.

85. Answer: d

Explanation:

Correct answer: 36

Percentages are widely used in many areas, including finance, statistics, and everyday life. Understanding how to calculate percentages, including sequential percentages like in this problem, is a fundamental skill.

Calculating a percentage of a number helps us find a part of a whole.

When dealing with "of" in percentage problems, it usually means multiplication.

Calculating a percentage increase or decrease involves finding the percentage of the original value and adding or subtracting it.

Sequential percentages (like 80% of 50%) multiply together; 80% of 50% is equivalent to $(0.80 \times 0.50) = 0.40$ or 40% of the original number.

This problem demonstrates a simple application of sequential percentage calculation, which is a common type of question in quantitative aptitude tests.

86. Answer: c

Explanation:

Understanding Wire Resistance Based on Dimensions

The resistance of a wire depends on its material, length, and cross-sectional area. We are given an initial wire and asked to find the resistance of another wire made of the same material but with different dimensions.

Formula for Electrical Resistance

The resistance (R) of a uniform conductor is directly proportional to its length (l) and inversely proportional to its cross-sectional area (A). The formula is given by:

$$R = \rho \frac{l}{A}$$

Where:

- ρ (rho) is the resistivity of the material (a property that depends only on the material and temperature).
- l is the length of the conductor.
- A is the cross-sectional area of the conductor.

For a wire with a circular cross-section of radius r , the area $A = \pi r^2$. So, the resistance formula becomes:

$$R = \rho \frac{l}{\pi r^2}$$

Analyzing the Given Wires

Let's denote the properties of the first wire with subscript 1 and the second wire with subscript 2.

Wire 1: Original Wire

- Length (l_1) = l
- Radius (r_1) = r
- Resistivity (ρ_1) = ρ (since the material is the same)
- Resistance (R_1) = R

Using the formula, the resistance of the first wire is:

$$R_1 = R = \rho \frac{l_1}{\pi r_1^2} = \rho \frac{l}{\pi r^2} \quad (\text{Equation 1})$$

Wire 2: New Wire

- Length (l_2) = half its length = $l/2$
- Radius (r_2) = half its radius = $r/2$
- Resistivity (ρ_2) = ρ (since the material is the same)
- Resistance (R_2) = ?

Using the same formula, the resistance of the second wire is:

$$R_2 = \rho \frac{l_2}{\pi r_2^2}$$

Substitute the values for l_2 and r_2 :

$$R_2 = \rho \frac{l/2}{\pi (r/2)^2}$$

Calculating the New Resistance

Now, let's simplify the expression for R_2 :

$$R_2 = \rho \frac{l/2}{\pi (r^2/4)}$$

Divide the term in the numerator by the term in the denominator:

$$R_2 = \rho \times \frac{l}{2} \times \frac{4}{\pi r^2}$$

Rearrange the terms:

$$R_2 = \rho \frac{4l}{2\pi r^2}$$

Simplify the constants ($4/2 = 2$):

$$R_2 = 2 \times \rho \frac{l}{\pi r^2}$$

Compare this with Equation 1 ($R = \rho \frac{l}{\pi r^2}$). We can see that the term $\rho \frac{l}{\pi r^2}$ is equal to R .

Therefore,

$$R_2 = 2R$$

The resistance of the new wire will be twice the resistance of the original wire.

Property	Wire 1 (Original)	Wire 2 (New)
Length	l	$l/2$
Radius	r	$r/2$
Area ($A = \pi r^2$)	πr^2	$\pi(r/2)^2 = \pi r^2/4$
Resistivity (ρ)	Same (ρ)	Same (ρ)
Resistance ($R = \rho l/A$)	$R = \rho \frac{l}{\pi r^2}$	$R_2 = \rho \frac{l/2}{\pi r^2/4} = 2\rho \frac{l}{\pi r^2} = 2R$

So, when the length is halved, the resistance is halved (proportionality with length). When the radius is halved, the area becomes one-fourth ($A \propto r^2$), and since resistance is inversely proportional to area, the resistance becomes four times ($1/(1/4) = 4$). Combining these effects: $R_2 = R \times (\text{length factor}) \times (\text{area factor}) = R \times (1/2) \times (4) = 2R$.

Revision Table: Wire Resistance Factors

Factor	Relationship with Resistance (R)	Explanation
Length (l)	$R \propto l$	Longer wire offers more opposition to current flow.
Cross-sectional Area (A)	$R \propto 1/A$	Wider wire offers less opposition to current flow.
Material (Resistivity ρ)	$R \propto \rho$	Different materials inherently resist current flow differently.
Temperature	Varies (usually R increases with T for metals)	Temperature affects the movement of charge carriers.

Additional Information: Resistivity and Conductivity

Resistivity (ρ) is an intrinsic property of a material at a given temperature, indicating how strongly it resists electric current. Materials with low resistivity are good conductors (like

copper, silver), while materials with high resistivity are poor conductors or insulators (like rubber, glass).

Conductivity (σ) is the reciprocal of resistivity ($\sigma = 1/\rho$). It measures how well a material conducts electric current. Materials with high conductivity are good conductors.

The resistance R is an extrinsic property, meaning it depends on the dimensions (length and area) as well as the material's resistivity. The resistivity ρ is an intrinsic property, depending only on the material and temperature, not its shape or size.

87. Answer: c

Explanation:

Understanding Kuchipudi Dance Origins

Kuchipudi is a prominent classical dance form of India. Like other classical dances, it has deep roots in tradition, mythology, and spirituality. Each classical dance form is typically associated with a specific region or state in India, where it originated and flourished over centuries.

Identifying the Home State of Kuchipudi

To determine the state where Kuchipudi originated, we need to look at its historical background and geographical association. The name "Kuchipudi" itself provides a clue. It is named after a village called Kuchipudi in the Krishna district of Andhra Pradesh.

Historically, this dance form evolved in the Kuchipudi village and was primarily performed by Brahmin males. Over time, it developed into a sophisticated art form incorporating dance, drama, and music.

Let's consider the given options:

- Arunachal Pradesh: This state is known for its vibrant folk dances, but not Kuchipudi as a classical form.
- Kerala: Kerala is the origin of classical dance forms like Kathakali and Mohiniyattam.
- Andhra Pradesh: This state is historically and geographically linked to the origin and development of Kuchipudi dance.

- Himachal Pradesh: This state is known for its rich folk dance traditions, distinct from classical Indian dance forms like Kuchipudi.

Based on historical evidence and the name itself, Kuchipudi dance is strongly associated with the state of Andhra Pradesh.

Conclusion on Kuchipudi's State of Origin

The classical Indian dance form Kuchipudi originated in the village of Kuchipudi in the state of Andhra Pradesh. Therefore, Andhra Pradesh is the correct state associated with the roots of Kuchipudi.

The final answer is **Andhra Pradesh**.

Revision Table: Indian Classical Dances and Their States

Here is a quick summary of some major classical Indian dance forms and their states of origin for revision.

Classical Dance Form	State of Origin
Bharatanatyam	Tamil Nadu
Kathak	Uttar Pradesh
Kathakali	Kerala
Kuchipudi	Andhra Pradesh
Odissi	Odisha
Sattriya	Assam
Manipuri	Manipur
Mohiniyattam	Kerala

Additional Information on Kuchipudi Dance

Kuchipudi is unique among classical dances due to its dramatic character and emphasis on dialogue and song. Key features include:

- **Bhamakalapam:** A famous dance-drama component of Kuchipudi, often featuring the story of Satyabhama.
- **Tarangam:** A unique item where the dancer performs on the edges of a brass plate, sometimes with a pot of water on their head or diyas (lamps) in their hands.
- **Vachikam:** The use of dialogue or recitation by the dancer.
- **Costumes:** Traditionally, men performed all roles, even female ones. Modern Kuchipudi includes both male and female dancers, with vibrant costumes and specific jewellery.
- **Music:** Performed to Carnatic music, typically accompanied by a mridangam, violin, flute, and sometimes a tambura.

Understanding the specific characteristics and origins helps appreciate the richness of India's classical dance heritage.

88. Answer: d

Explanation:

Understanding Safety Boots and Foot Protection

Safety boots and safety shoes are essential pieces of personal protective equipment (PPE) used in many workplaces. Their primary purpose is to protect workers from potential hazards encountered in the work environment. The question specifically asks what part of the body is protected from falling objects by safety boots or shoes in designated areas.

Purpose of Safety Footwear

Safety footwear, like safety boots or shoes, is designed with features such as reinforced toes (often steel or composite), puncture-resistant soles, and slip-resistant treads. These features are crucial for protecting against various workplace hazards, including:

- Falling or rolling objects
- Sharp objects that could puncture the sole
- Electrical hazards

- Slipping hazards
- Chemical spills

The question focuses on protection from falling objects, which is a common risk in construction sites, warehouses, factories, and other industrial settings.

Analyzing the Options for Foot Protection

Let's examine the given options in the context of protection offered by safety boots or shoes:

- **Option 1: eye**

Safety glasses, goggles, or face shields are used to protect the eyes from hazards like debris, chemicals, or intense light. Safety boots do not offer protection to the eyes.

- **Option 2: ear**

Earplugs or earmuffs are used to protect the ears from excessive noise. Safety boots do not offer protection to the ears.

- **Option 3: head**

Safety helmets or hard hats are used to protect the head from falling objects or impacts. Safety boots do not offer protection to the head.

- **Option 4: feet**

Safety boots and shoes are specifically designed to protect the feet from various hazards on the ground, including protection from heavy falling objects which could crush the toes or foot. The reinforced toe cap is a key feature for this type of protection.

Based on the design and purpose of safety boots, they are intended to safeguard the wearer's feet from hazards like falling objects.

Conclusion on Safety Boot Protection

When working in designated areas where there is a risk of objects falling, wearing safety boots or shoes is mandatory to protect your feet. The reinforced construction of the footwear is specifically engineered to withstand impacts and prevent injury to the feet from such incidents.

Revision Table: PPE and Body Part Protected

Personal Protective Equipment (PPE)	Body Part Primarily Protected	Example Hazard
Safety Boots/Shoes	Feet	Falling objects, sharp objects, slips
Safety Glasses/Goggles	Eyes	Flying debris, chemical splashes
Hard Hat	Head	Falling objects, impacts
Ear Plugs/Muffs	Ears	Excessive noise
Safety Gloves	Hands	Cuts, chemicals, burns

Additional Information: Importance of Foot Protection in Safety

Foot injuries are common in industrial and construction environments. They can range from minor cuts and bruises to severe fractures and crushing injuries caused by heavy falling objects. Proper foot protection is therefore a critical aspect of workplace safety programs. Employers often identify areas where foot hazards exist and designate them as mandatory safety footwear zones. Choosing the correct type of safety boot depends on the specific hazards present in the workplace. For protection against falling objects, footwear meeting standards for impact resistance is essential.

Understanding the specific protection each type of PPE offers helps ensure workers use the correct equipment for the task, thereby minimizing the risk of injury.

89. Answer: a

Explanation:

Solving Pipe and Cistern Problems: Filling a Tank

This problem involves calculating the time taken by pipes to fill a tank, considering their individual rates and changes in operation. We need to find the extra time Pipe B takes to complete the remaining part of the tank after Pipe A is closed.

Understanding Pipe Rates

The rate at which a pipe fills a tank is the fraction of the tank filled in one unit of time (here, one minute). If a pipe can fill a tank in 't' minutes, its rate is $\frac{1}{t}$ of the tank per minute.

- Pipe A fills the tank in 12 minutes.
- Rate of Pipe A = $\frac{1}{12}$ of the tank per minute.
- Pipe B fills the tank in 16 minutes.
- Rate of Pipe B = $\frac{1}{16}$ of the tank per minute.

Combined Work of Pipes A and B

Both pipes A and B are opened together for 4 minutes. To find the amount of work done (fraction of tank filled) in these 4 minutes, we first calculate their combined rate.

Combined rate of A and B = Rate of A + Rate of B

$$\text{Combined rate of A and B} = \frac{1}{12} + \frac{1}{16}$$

To add these fractions, we find a common denominator, which is the least common multiple (LCM) of 12 and 16. The LCM of 12 and 16 is 48.

$$\text{Combined rate of A and B} = \frac{1 \times 4}{12 \times 4} + \frac{1 \times 3}{16 \times 3} = \frac{4}{48} + \frac{3}{48} = \frac{4+3}{48} = \frac{7}{48} \text{ of the tank per minute.}$$

Work done by A and B together in 4 minutes = Combined rate $\times$ Time

$$\text{Work done in 4 minutes} = \frac{7}{48} \times 4 = \frac{28}{48}$$

$$\text{Simplifying the fraction: } \frac{28}{48} = \frac{7 \times 4}{12 \times 4} = \frac{7}{12}$$

So, in the first 4 minutes, $\frac{7}{12}$ of the tank is filled.

Calculating Remaining Work

After 4 minutes, Pipe A is closed. The remaining part of the tank needs to be filled by Pipe B alone. The total tank represents 1 whole unit of work.

$$\text{Remaining work} = \text{Total work} - \text{Work done in first 4 minutes}$$

$$\text{Remaining work} = 1 - \frac{7}{12}$$

$$\text{Remaining work} = \frac{12}{12} - \frac{7}{12} = \frac{12-7}{12} = \frac{5}{12}$$

So, $\frac{5}{12}$ of the tank is yet to be filled.

Time Taken by Pipe B to Fill Remaining Work

Pipe B will fill the remaining $\frac{5}{12}$ of the tank at its own rate. The rate of Pipe B is $\frac{1}{16}$ of the tank per minute.

$$\text{Time taken by B} = \text{Remaining work} \div \text{Rate of B}$$

$$\text{Time taken by B} = \frac{5}{12} \div \frac{1}{16}$$

Dividing by a fraction is the same as multiplying by its reciprocal:

$$\text{Time taken by B} = \frac{5}{12} \times \frac{16}{1} = \frac{5 \times 16}{12} = \frac{80}{12}$$

Simplifying the fraction $\frac{80}{12}$: Both 80 and 12 are divisible by 4.

$$\text{Time taken by B} = \frac{80 \div 4}{12 \div 4} = \frac{20}{3}$$

So, Pipe B will take $\frac{20}{3}$ minutes to fill the remaining part of the tank.

The question asks for the "extra time" B will take. Since B is the only pipe working after A is closed, the time B takes to fill the remaining part is the extra time required.

Therefore, the extra time B will take to fill the tank completely is $\frac{20}{3}$ minutes.

Step-by-Step Calculation Summary

1. Find the individual filling rate of each pipe.
2. Calculate the combined filling rate of both pipes when working together.
3. Determine the amount of tank filled when both pipes work for the specified time (4 minutes).
4. Calculate the remaining unfilled portion of the tank.
5. Determine the time taken by the remaining pipe (Pipe B) to fill this remaining portion at its individual rate.

Revision Table: Key Concepts

Concept	Formula/Definition	Application in this problem
Pipe Rate	$\frac{1}{\text{Time to fill tank alone}}$	Rate of A: $\frac{1}{12}$ Rate of B: $\frac{1}{16}$
Combined Rate (A+B)	Rate of A + Rate of B	$\frac{1}{12} + \frac{1}{16} = \frac{7}{48}$
Work Done	Rate $\times$ Time	Work in 4 mins: $\frac{7}{48} \times 4 = \frac{7}{12}$
Remaining Work	Total Work - Work Done	$1 - \frac{7}{12} = \frac{5}{12}$
Time for Remaining Work (by B)	Remaining Work $\div$ Rate of B	$\frac{5}{12} \div \frac{1}{16} = \frac{20}{3}$ mins

Additional Information: Pipe and Cistern Problems

Pipe and cistern problems are a type of time and work problem. The key idea is that filling a tank is considered 'work'. The rate of work is inversely proportional to the time taken to complete the work. If a pipe empties a tank, its rate is considered negative. If multiple pipes work together, their rates are added (for filling) or subtracted (if one is emptying).

Understanding LCM is crucial for efficiently adding or subtracting fractional rates.

Always pay attention to whether the question asks for the total time taken or the extra time taken after a certain period.

Your Personal Exams Guide

90. Answer: c

Explanation:

Understanding Ajanta Caves Paintings and Sculptures

The Ajanta Caves are a series of rock-cut Buddhist cave monuments located in the state of Maharashtra, India. These caves are famous for their stunning paintings and sculptures, which are considered masterpieces of ancient Indian art.

The question asks about the themes or tales depicted in the paintings and sculptures found in the Ajanta Caves. To answer this, we need to understand the historical and

religious context of these caves.

Exploring the Themes in Ajanta Art

The Ajanta Caves were built and decorated over a long period, primarily between the 2nd century BCE and 6th century CE. They served as monasteries and prayer halls for Buddhist monks. Naturally, the art within these caves reflects the beliefs, practices, and stories associated with Buddhism.

The paintings, often created using the fresco technique, cover the walls, ceilings, and pillars. They narrate various stories and events, including:

- **Jataka tales:** These are popular stories about the previous lives of Siddhartha Gautama, the Buddha. These tales illustrate moral lessons and the journey towards enlightenment.
- **Events from the life of the Buddha:** Depictions of important moments from Siddhartha Gautama's life, such as his birth, enlightenment, first sermon, and death (Mahaparinirvana).
- **Scenes of Buddhist deities and bodhisattvas.**
- **Everyday life and court scenes,** often integrated with the religious narratives.

The sculptures also depict various forms of the Buddha, bodhisattvas, and scenes from Buddhist mythology and tales.

Analyzing the Options

Let's consider the given options in the context of the Ajanta Caves:

- **Arabic:** Arabic culture and tales are associated with the Arabian peninsula and the spread of Islam, which occurred much later than the period when the Ajanta Caves were active. The themes would not be Arabic tales.
- **Islamic:** Islam emerged in the 7th century CE, after the main period of activity at Ajanta (2nd century BCE to 6th century CE). The art in Ajanta predates the arrival and influence of Islamic art in India.
- **Buddhist:** As discussed earlier, the Ajanta Caves were Buddhist monastic sites. The art directly relates to Buddhist teachings, the life of the Buddha, and Jataka tales. This aligns perfectly with the known history and purpose of the caves.
- **Maratha:** The Maratha Empire rose to prominence much later in Indian history (17th–18th centuries CE). The art in Ajanta belongs to a much earlier period and reflects the

prevailing religious traditions of that time, not Maratha history or culture.

Conclusion on Ajanta Caves Depictions

Based on the historical context and the nature of the art found within the Ajanta Caves, it is clear that the paintings and sculptures primarily depict Buddhist tales and themes. The Jataka tales, in particular, are a significant focus of the narrative paintings.

Therefore, the correct answer is that the Ajanta Caves feature paintings and sculptures that depict **Buddhist** tales.

Revision Table: Ajanta Caves Key Facts

Feature	Description
Location	Maharashtra, India
Period of Activity	Primarily 2nd Century BCE to 6th Century CE
Religion/Philosophy	Buddhism
Type of Monument	Rock-cut caves (monasteries, prayer halls)
Art Forms	Paintings (frescoes), Sculptures
Primary Themes	Jataka Tales, Life of the Buddha, Buddhist deities

Additional Information on Ajanta Caves Art

The Ajanta cave paintings are renowned for their vibrant colours and expressive figures. They provide valuable insights into the society, culture, and religious practices of ancient India. The techniques used in these paintings, particularly the fresco style where pigments are applied to wet plaster, demonstrate a high level of artistic skill. The detailed depictions of clothing, jewelry, hairstyles, and daily activities also make them an important historical source.

The caves are designated a UNESCO World Heritage Site, recognized for their significant contribution to the history of art and their representation of early Buddhist architecture and painting.

91. Answer: b

Explanation:

Understanding Coefficient of Volume Expansion in Materials

When materials are heated, they tend to expand in size. This phenomenon is known as thermal expansion. The extent to which a material expands upon heating is quantified by its coefficient of thermal expansion. For volume expansion, we use the coefficient of volume expansion, often denoted by the symbol β (beta).

The coefficient of volume expansion is defined as the fractional change in volume per degree Celsius (or Kelvin) change in temperature. A material with a high coefficient of volume expansion will experience a larger change in volume for a given change in temperature compared to a material with a low coefficient.

Comparing Volume Expansion Coefficients

Let's examine the approximate coefficient of volume expansion for the materials listed in the options: Brass, Alcohol, Glass, and Water. It's important to note that these values can vary slightly depending on the specific composition (for alloys like brass or glass) and the temperature range (especially for water).

Material	Approximate Coefficient of Volume Expansion (β) in ($10^{-6}/^{\circ}\text{C}$)
Brass	57
Alcohol (Ethanol)	1120
Glass (Soda-lime)	27
Water (around room temp)	210

By comparing the values in the table, we can clearly see which material has the highest coefficient of volume expansion:

- Glass has a relatively low coefficient ($\approx 27 \times 10^{-6}/^{\circ}\text{C}$), typical for solids.
- Brass, a metallic alloy, has a higher coefficient than glass but is still in the range for solids ($\approx 57 \times 10^{-6}/^{\circ}\text{C}$).
- Water has a significantly higher coefficient than solids, though its value varies with temperature ($\approx 210 \times 10^{-6}/^{\circ}\text{C}$ around room temperature).
- Alcohol (Ethanol in this case) has a very high coefficient of volume expansion compared to the other options ($\approx 1120 \times 10^{-6}/^{\circ}\text{C}$).

Liquids generally have higher coefficients of volume expansion than solids because the intermolecular forces are weaker, allowing molecules to move further apart more easily when heated.

Conclusion: Material with High Volume Expansion

Based on the comparison of their approximate coefficients of volume expansion, Alcohol exhibits the highest value among the given options. Therefore, Alcohol is the material having a high coefficient of volume expansion.

Revision Table: Key Concepts in Thermal Expansion

Concept	Description	Relevant Coefficient
Thermal Expansion	The tendency of matter to change in volume in response to changes in temperature.	Linear (α), Area (γ), Volume (β)
Coefficient of Volume Expansion (β)	Fractional change in volume per unit change in temperature. Units are typically $^{\circ}\text{C}^{-1}$ or K^{-1} .	$\beta = \frac{1}{V} \left(\frac{\partial V}{\partial T} \right)_P$
Relationship between α, γ, β (for isotropic solids)	For isotropic solids, $\gamma \approx 2\alpha$ and $\beta \approx 3\alpha$.	

Additional Information on Volume Expansion

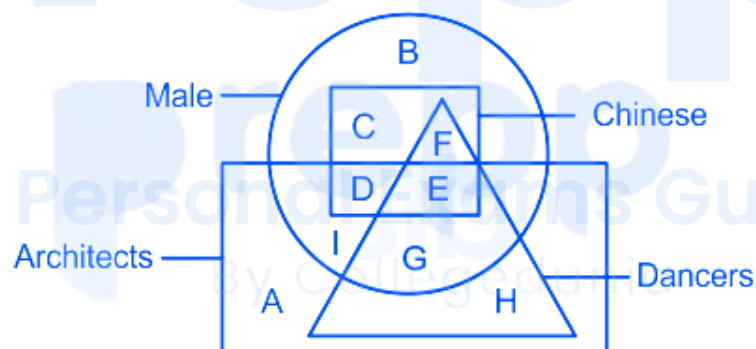
The coefficient of volume expansion for liquids is generally much larger than for solids. This is why liquids in thermometers are visibly responsive to temperature changes. The expansion of gases is even greater, described by different gas laws (like the ideal gas law), but thermal expansion also applies. For anisotropic solids (materials where properties vary with direction), the expansion might be different along different axes, requiring multiple coefficients of linear expansion.

It's also worth noting the anomalous expansion of water between 0°C and 4°C , where its volume decreases upon heating (or expands upon cooling), meaning its coefficient of volume expansion is negative in this specific temperature range. However, above 4°C , water expands normally upon heating, and its coefficient is positive, as shown in the table for room temperature.

92. Answer: d

Explanation:

The set of letters which represents dancers who are male, is shown by the shaded region:



Hence, 'GEF' is the correct answer.

93. Answer: a

Explanation:

Understanding Displacement with Constant Acceleration

The question asks about the relationship between displacement and time for a body that starts from rest and undergoes constant acceleration. This scenario is a classic example of motion under constant acceleration, which can be described using the equations of motion (also known as SUVAT equations).

Analyzing the Conditions

We are given two key pieces of information about the body's motion:

- The body starts from rest: This means the initial velocity (u) is zero, i.e., $u = 0$.
- The acceleration is constant: Let the constant acceleration be a .

Applying the Equations of Motion

The equations of motion relate initial velocity (u), final velocity (v), displacement (s), acceleration (a), and time (t). One of the fundamental equations is the relationship between displacement, initial velocity, acceleration, and time:

$$s = ut + \frac{1}{2}at^2$$

Deriving the Relationship

We can substitute the given condition that the body starts from rest ($u = 0$) into this equation:

$$s = (0)t + \frac{1}{2}at^2$$

This simplifies to:

$$s = \frac{1}{2}at^2$$

Since the acceleration a is constant, the term $\frac{1}{2}a$ is also a constant. Let's call this constant k , where $k = \frac{1}{2}a$. So, the equation becomes:

$$s = kt^2$$

This equation clearly shows that the displacement s is directly proportional to the square of the time t , because k is a constant value.

Mathematically, proportionality is represented as:

$$s \propto t^2$$

Analyzing the Options

Let's look at the given options in light of our derivation:

- 1. time squared:** Our derivation $s \propto t^2$ matches this option.
- 2. velocity:** Displacement is related to velocity, but specifically to the *initial* and *final* velocity or average velocity over time, not just "velocity" generally in a proportional way like this when starting from rest with constant acceleration. The final velocity $v = u + at = 0 + at = at$, so $v \propto t$. Displacement is proportional to velocity squared ($v^2 = u^2 + 2as$), but not velocity itself in this simple direct proportionality sense related to time.
- 3. work:** Work is related to force and displacement, but displacement itself is not directly proportional to work in this context.
- 4. time:** Displacement is proportional to time only when velocity is constant ($s = vt$) or when considering final velocity ($v \propto t$), but not when starting from rest with constant acceleration where $s \propto t^2$.

Based on the derivation using the equations of motion, the displacement is proportional to time squared when a body starts from rest with constant acceleration.

Revision Table: Key Equations of Motion (Constant Acceleration)

Equation	Relates
$v = u + at$	Final velocity (v), Initial velocity (u), Acceleration (a), Time (t)
$s = ut + \frac{1}{2}at^2$	Displacement (s), Initial velocity (u), Acceleration (a), Time (t)
$v^2 = u^2 + 2as$	Final velocity (v), Initial velocity (u), Acceleration (a), Displacement (s)
$s = \frac{(u+v)}{2}t$	Displacement (s), Initial velocity (u), Final velocity (v), Time (t)

Additional Information: Proportionality in Physics

Proportionality ($\propto$) is a way to describe how one quantity changes when another quantity changes. If quantity A is proportional to quantity B, it means $A = k * B$, where k is a constant. If quantity A is proportional to the square of quantity B, it means $A = k * B^2$, where k is a constant.

In this problem, we found that $s = \frac{1}{2}at^2$. Since $\frac{1}{2}a$ is constant, the displacement s is directly proportional to t^2 . This means if you double the time, the displacement increases by a factor of four (2^2); if you triple the time, the displacement increases by a factor of nine (3^2), and so on, assuming constant acceleration and starting from rest.

This principle is fundamental in understanding the motion of objects under gravity (like free fall, neglecting air resistance), where the acceleration due to gravity (g) is approximately constant. If an object is dropped from rest, its distance fallen is proportional to the square of the time it has been falling.

94. Answer: c

Explanation:

Understanding Software That Harms Computers

Let's break down the concept of software designed to damage or disrupt computer systems. The question asks for the specific term used for such programs.

A software program intentionally created to cause harm, disrupt operations, or gain unauthorized access to a computer system is known by a specific name. Let's look at the options provided and determine which one fits this description.

Analyzing the Options for Harmful Software

- **server:** A server is a computer or software that provides functionality for other programs or devices, called "clients." Servers are essential parts of networks and are not inherently designed to cause harm.
- **LAN:** LAN stands for Local Area Network. It is a type of computer network covering a limited area, like a home, school, or office building. It is a network type, not a software program designed to cause harm.
- **malware:** This term is short for **malicious software**. It is specifically designed to disrupt, damage, or gain unauthorized access to computer systems. Examples include viruses, worms, ransomware, spyware, and trojan horses. This definition perfectly matches the description in the question.
- **operating system:** An operating system (OS) is the core software that manages computer hardware and software resources and provides common services for

computer programs. Examples are Windows, macOS, Linux, and Android. While vulnerabilities in an OS can be exploited by malicious software, the OS itself is not designed to be harmful; it is essential for the computer's function.

Based on the definitions, the software program developed to harm other computers is called malware.

Term	Description	Harmful Software?
Server	Provides resources/services to clients	No
LAN	Type of computer network	No
Malware	Software designed to harm systems	Yes
Operating System	Manages hardware and software resources	No (essential software)

Defining Malware

Malware is a broad term covering various types of malicious software. Its primary purpose is often to:

- Steal sensitive data (like passwords or financial information).
- Damage or delete files.
- Disrupt the normal operation of a computer or network.
- Gain unauthorized control over a system.
- Display unwanted advertisements.

Understanding what malware is and how it operates is crucial for cybersecurity.

Conclusion on Harmful Programs

The term that accurately describes a software program developed with the intent to harm other computers is malware. The other options represent different concepts: servers and operating systems are fundamental computer components or software (not inherently harmful), and a LAN is a type of network.

Revision Table: Computer Terms

Term	Category	Purpose (Relevant to Question)
Server	Hardware/Software	Provide services
LAN	Network Type	Connect devices in a small area
Malware	Software	Cause harm/disruption
Operating System	Software	Manage hardware/software

Additional Information: Types of Malware

Malware comes in many forms, each with different characteristics and methods of attack. Some common types of malware include:

- **Viruses:** Self-replicating programs that attach themselves to other programs or files.
- **Worms:** Self-replicating malware that spreads through networks without needing to attach to files.
- **Trojan Horses:** Malware disguised as legitimate software; they don't replicate but cause harm once executed.
- **Ransomware:** Encrypts a user's files and demands payment for their release.
- **Spyware:** Collects information about the user's activity without their knowledge.
- **Adware:** Displays unwanted advertisements, sometimes collecting user data.

Protecting against malware requires vigilance, security software, and safe computing practices.

95. Answer: b

Explanation:

Understanding Temperature Conversion: Fahrenheit to Celsius

Temperature is a physical quantity that expresses the degree of hotness or coldness. It is measured using various scales, with Fahrenheit (°F) and Celsius (°C) being two of the most common ones. Converting a temperature from Fahrenheit to Celsius is a frequent requirement in many applications, from daily weather reports to scientific calculations.

The Formula for Fahrenheit to Celsius Conversion

To convert a temperature from the Fahrenheit scale to the Celsius scale, we use a specific formula. This formula accounts for the different reference points (freezing and boiling points of water) and the different size of the degrees between the two scales.

The formula for converting Fahrenheit ($^{\circ}\text{F}$) to Celsius ($^{\circ}\text{C}$) is:

$$^{\circ}\text{C} = \frac{5}{9} \times (^{\circ}\text{F} - 32)$$

Let's break down this formula:

- Subtract 32 from the Fahrenheit temperature. This is because the freezing point of water is 32°F but 0°C .
- Multiply the result by $\frac{5}{9}$. This ratio relates the size of a Celsius degree to a Fahrenheit degree. A 180-degree range separates the freezing and boiling points of water in Fahrenheit ($212 - 32 = 180$), while a 100-degree range separates them in Celsius ($100 - 0 = 100$). The ratio $\frac{100}{180}$ simplifies to $\frac{5}{9}$.

Step-by-Step Calculation: Converting 86°F to $^{\circ}\text{C}$

We are asked to convert 86°F to Celsius. We will use the formula mentioned above.

Given temperature in Fahrenheit, $^{\circ}\text{F} = 86$.

Using the formula:

$$^{\circ}\text{C} = \frac{5}{9} \times (^{\circ}\text{F} - 32)$$

Substitute the given Fahrenheit temperature into the formula:

$$^{\circ}\text{C} = \frac{5}{9} \times (86 - 32)$$

First, perform the subtraction inside the parentheses:

$$86 - 32 = 54$$

Now, substitute this result back into the formula:

$$^{\circ}\text{C} = \frac{5}{9} \times 54$$

Next, multiply $\frac{5}{9}$ by 54:

$$^{\circ}\text{C} = 5 \times \frac{54}{9}$$

Divide 54 by 9:

$$\frac{54}{9} = 6$$

Finally, multiply 5 by 6:

$$^{\circ}\text{C} = 5 \times 6$$

$$^{\circ}\text{C} = 30$$

So, 86°F is equal to 30°C.

Comparing with Options

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

- Option 1: 20°C
- Option 2: 30°C
- Option 3: 34°C
- Option 4: 10°C

Our calculated value of 30°C matches Option 2.

Summary of Calculation

Step	Calculation	Result
Formula	$^{\circ}\text{C} = \frac{5}{9} \times (^{\circ}\text{F} - 32)$	N/A
Subtract 32	$86 - 32$	54
Multiply by $\frac{5}{9}$	$\frac{5}{9} \times 54$	30
Final Answer	86°F converted to °C	30°C

Revision Table: Temperature Scales and Formulas

Common Temperature Scales and Conversion Formulas

Scale	Freezing Point of Water	Boiling Point of Water	Conversion to Celsius (°C)	Conversion from Celsius (°C)
Celsius (°C)	0°C	100°C	N/A	N/A
Fahrenheit (°F)	32°F	212°F	$^{\circ}\text{C} = \frac{5}{9} \times (^{\circ}\text{F} - 32)$	$^{\circ}\text{F} = \frac{9}{5} \times ^{\circ}\text{C} + 32$
Kelvin (K)	273.15 K	373.15 K	$^{\circ}\text{C} = K - 273.15$	$K = ^{\circ}\text{C} + 273.15$

Additional Information on Temperature Scales

Understanding different temperature scales is important in various fields. Here are a few key points:

- **Celsius Scale:** Widely used around the world, particularly in science and for everyday temperature measurement. It is based on the freezing (0°C) and boiling (100°C) points of pure water at standard atmospheric pressure.
- **Fahrenheit Scale:** Primarily used in the United States. Its original reference points were based on a brine solution and human body temperature, though it is now defined by the freezing (32°F) and boiling (212°F) points of water.
- **Kelvin Scale:** The standard unit of temperature in the International System of Units (SI). It is an absolute scale, meaning that 0 Kelvin (0 K) is absolute zero, the theoretical point where particles have minimum motion. There are no negative temperatures on the Kelvin scale. Kelvin degrees have the same magnitude as Celsius degrees.
- **Absolute Zero:** This is the lowest possible temperature, where the fundamental particles of nature have minimal motion and cannot get any colder. Absolute zero is 0 K, which is equivalent to -273.15°C or -459.67°F.

Being able to convert between these scales is crucial for comparing temperature measurements from different sources or systems.

96. Answer: c

Explanation:

Understanding the Location of UNESCO Headquarters

The question asks for the location of the headquarters of the United Nations Educational, Scientific and Cultural Organization, commonly known as UNESCO.

Identifying the Correct UNESCO Headquarters City

UNESCO is a specialized agency of the United Nations. Its primary goal is to contribute to the building of peace, the eradication of poverty, sustainable development, and intercultural dialogue through education, the sciences, culture, communication, and information.

The headquarters of this significant international organization are located in a major European city with a rich history of culture and intellectual life.

Option	City	Common Headquarters Located Here
1	New York City	United Nations (UN) Headquarters, UNICEF Headquarters
2	Geneva	Many UN agencies (e.g., WHO, ILO, UNHCR), World Trade Organization (WTO)
3	Paris	UNESCO Headquarters , OECD Headquarters
4	Washington DC	World Bank, International Monetary Fund (IMF), Organization of American States (OAS)

Based on the known locations of major international organization headquarters:

- New York City is famously home to the main United Nations headquarters.
- Geneva hosts numerous UN specialized agencies and international bodies, particularly those focused on health, labor, and trade.
- Washington DC is the base for major global financial institutions like the World Bank and IMF.
- Paris is the established location for the UNESCO headquarters.

Therefore, the headquarters of UNESCO are located in Paris.

Exploring UNESCO and its Mission

UNESCO was founded on 16 November 1945. Its constitution entered into force on 4 November 1946. The organization works globally on several key programmes:

- Education
- Natural Sciences
- Social and Human Sciences
- Culture
- Communication and Information

Locating its headquarters in Paris aligns with its mission focused on culture, education, and science, given the city's historical significance in these fields.

Revision Table: Key International Headquarters

Organization	Acronym	Headquarters City
United Nations Educational, Scientific and Cultural Organization	UNESCO	Paris, France
United Nations	UN	New York City, USA
World Health Organization	WHO	Geneva, Switzerland
International Labour Organization	ILO	Geneva, Switzerland
World Bank Group	WBG	Washington DC, USA
International Monetary Fund	IMF	Washington DC, USA

Additional Information on UNESCO

The headquarters building of UNESCO in Paris is a notable architectural structure. It was inaugurated in 1958. The organization plays a crucial role in promoting international

cooperation and safeguarding cultural and natural heritage around the world through initiatives like the World Heritage Sites program.

Understanding the headquarters location of key international bodies like UNESCO is important for general knowledge and for studying international relations and global affairs.

97. Answer: a

Explanation:

Understanding Profit Sharing Ratios

This problem involves dividing a total profit based on a given ratio between two individuals, A and B. A profit sharing ratio tells us how a total amount is to be split among partners. If the ratio is $x : 2$, it means that for every 'x' parts of the profit A receives, B receives 2 parts.

Analyzing the Given Profit Distribution Problem

We are given the following information:

- Total Profit = Rs. 240
- A's Share of Profit = Rs. 80
- Ratio of Profit Distribution between A and B = $x : 2$

Our goal is to find the value of x in the profit sharing ratio.

Setting up the Equation for Profit Ratio

The share of an individual in the profit is proportional to their part in the ratio, relative to the sum of the ratio parts. The total ratio parts in this case are $x + 2$.

A's share can be represented as the fraction of the total profit that corresponds to A's part in the ratio:

$$\text{A's Share} / \text{Total Profit} = \text{A's Ratio Part} / (\text{Sum of Ratio Parts})$$

Substituting the given values and the ratio:

$$\frac{80}{240} = \frac{x}{x+2}$$

Solving for the Unknown Value x

Now we need to solve the equation for x. First, simplify the fraction on the left side:

$$\frac{80}{240} = \frac{80 \div 80}{240 \div 80} = \frac{1}{3}$$

So the equation becomes:

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

To solve for x, we can cross-multiply:

$$1 \times (x + 2) = 3 \times x$$

$$x + 2 = 3x$$

Now, rearrange the equation to isolate x terms on one side:

$$2 = 3x - x$$

$$2 = 2x$$

Divide both sides by 2:

$$x = \frac{2}{2}$$

$$x = 1$$

The value of x in the profit sharing ratio is 1.

Verifying the Solution

If $x = 1$, the profit sharing ratio is 1 : 2. The total ratio parts are $1 + 2 = 3$.

A's share would be $\frac{1}{3}$ of the total profit.

$$\text{A's Share} = \frac{1}{3} \times 240 = \frac{240}{3} = 80$$

This matches the given information that A got Rs. 80 as his share. The solution is correct.

Revision Table: Key Concepts

Concept	Explanation
Profit Sharing Ratio	Describes how a total profit is divided among partners or individuals.
Total Ratio Parts	The sum of the individual parts in the ratio (e.g., for x:2, total parts are x+2).
Individual Share Calculation	Individual Share = $(\text{Individual Ratio Part} / \text{Total Ratio Parts}) \times \text{Total Profit}$

Additional Information on Profit Distribution

Profit distribution is a fundamental concept in partnerships and business. It ensures that profits are allocated fairly among contributors, often based on their investment, time contribution, or a pre-agreed ratio. Ratios provide a simple way to represent these proportions. When one value (like one partner's share or the total profit) is known along with the ratio, we can often find the other unknown values, including an unknown part within the ratio itself, as demonstrated in this problem.

98. Answer: a

Explanation:

Understanding Blood Relations Questions

Blood relations questions test your ability to understand and decipher family relationships based on given codes or symbols. In this type of question, specific symbols represent relationships like daughter, sister, husband, etc. You need to use these definitions to trace the relationship between two individuals mentioned in the question.

Decoding the Symbols

The question provides the meaning of three symbols:

- $G + H$ means G is daughter of H. (This tells us G is female.)
- $G - H$ means G is sister of H. (This tells us G is female.)
- $G * H$ means G is husband of H. (This tells us G is male and H is female.)

The Goal: Finding I as Daughter of H

We are looking for the option that correctly represents the relationship where I is the daughter of H. This means two things must be true:

1. I must be female.
2. H must be a parent of I (either mother or father).

Analyzing Each Option to Trace the Relationship

Let's break down each option symbol by symbol to find the relationship between I and H.

Option 1: $I - J + F * H$

We trace the relationships starting from the left:

- $I - J$: According to the rule ' $G - H$ means G is sister of H', this means I is the sister of J. (This tells us I is female).
- $J + F$: According to the rule ' $G + H$ means G is daughter of H', this means J is the daughter of F. (This tells us J is female and F is a parent of J).
- $F * H$: According to the rule ' $G * H$ means G is husband of H', this means F is the husband of H. (This tells us F is male and H is female. F and H are married).

Combining the relationships:

- I is the sister of J.
- J is the daughter of F.
Since I and J are siblings, and J is the daughter of F, I must also be the daughter of F.
- F is the husband of H.
Since F is the father of I (from the previous step) and F is married to H, H must be the mother of I.

Therefore, Option 1 shows that I is the daughter of H.

Option 2: $I + J - F * H$

Tracing the relationships:

- **I + J:** I is the daughter of J. (I is female. J is a parent of I).
- **J - F:** J is the sister of F. (J is female).
- **F * H:** F is the husband of H. (F is male, H is female).

Combining the relationships:

- I is the daughter of J.
- J is the sister of F.
So, I's mother (or father) is J, and J's sibling is F. This makes F the uncle or aunt of I.
- F is the husband of H.
F is I's uncle/aunt, and F is married to H. If F is I's uncle (J's brother), then H is I's aunt by marriage (uncle's wife). If F is I's aunt (J's sister), this contradicts $F * H$ where F is male. Since J is female and sister of F, F must be male or female. $F * H$ tells us F is male. So F is J's brother (I's uncle) and H is F's wife (I's aunt by marriage).

This option shows H is related to I as an aunt by marriage, not as a parent. Therefore, Option 2 does not show I is the daughter of H.

Option 3: $I - J * F + H$

Tracing the relationships:

- **I - J:** I is the sister of J. (I is female).
- **J * F:** J is the husband of F. (J is male, F is female).
- **F + H:** F is the daughter of H. (F is female. H is a parent of F).

Combining the relationships:

- I is the sister of J.
- J is the husband of F.
So I is the sister of F's husband.
- F is the daughter of H.
H is the parent of F. F is married to J, and I is J's sister. So H is the parent of F, and F is married to I's brother-in-law. H is the parent of F. F is married to J. I is sister of J. Thus, H is the parent of F, who is married to the brother of I. H is the mother or father of F, and thus H is mother-in-law or father-in-law of J. Since I is J's sister, H is the mother or father of I's sister-in-law.

This option shows H is a parent of F, who is married to I's brother. This does not show I is the daughter of H.

Option 4: $I * J - F + H$

Tracing the relationships:

- $I * J$: I is the husband of J. (I is male, J is female).
- $J - F$: J is the sister of F. (J is female).
- $F + H$: F is the daughter of H. (F is female. H is a parent of F).

Combining the relationships:

- I is the husband of J.
- J is the sister of F.
So, I's wife is J, and J's sister is F. F is I's sister-in-law.
- F is the daughter of H.
F is I's sister-in-law, and F is the daughter of H. Since F is the daughter of H, and J is the sister of F, J is also the daughter of H. Since J is the daughter of H, and I is the husband of J, H is the parent of I's wife. H is I's parent-in-law.

This option shows H is the parent of J, and I is married to J. This does not show I is the daughter of H. Also, the first relation $I * J$ establishes I as male, which contradicts the target of I being a daughter (female).

Summary of Relationships

Let's summarize the relationship between I and H derived from each option:

Option	Relationships	Relation between I and H	Is I Daughter of H?
1	$I - J, J + F, F * H$	I is daughter of H	Yes
2	$I + J, J - F, F * H$	H is I's aunt by marriage	No
3	$I - J, J * F, F + H$	H is parent of I's sister-in-law	No
4	$I * J, J - F, F + H$	H is I's parent-in-law	No

Based on the step-by-step analysis, only Option 1 results in I being the daughter of H.

Revision Table: Blood Relations Concepts

Concept	Explanation
Symbol Decoding	Assigning the given family relationship to each symbol in the expression.
Tracing Relationships	Working through the expression piece by piece to build the connection between the individuals.
Gender Identification	Some symbols (like +, -, *) explicitly define the gender of the first person in the relationship. This is crucial for checking if the final required relationship is possible.
Combining Relations	Linking the individual relationships (e.g., sister of daughter) to find the overall relation (e.g., daughter).

Additional Information: Solving Blood Relation Puzzles

Solving blood relation puzzles effectively involves a few key strategies:

- Read the symbol definitions carefully and maybe jot them down.
- Start tracing the relationship from one end of the expression towards the other.
- Pay close attention to the gender implied by each relationship symbol. This can help eliminate options quickly if the required gender doesn't match.
- Mentally or physically draw a small family tree if the relationships become complex.
- Break down complex expressions into smaller, manageable pairs (e.g., $A*B-C+D$ becomes $A*B$, then $(A*B)-C$, then $((A*B)-C)+D$).
- Always double-check the final derived relationship against the target relationship asked in the question.

Practice with different types of blood relation coding questions will improve your speed and accuracy.

99. Answer: c

Explanation:

Solving the Code Language Puzzle

This question involves decoding a simple code language based on the given phrases and their numerical codes. We need to find the specific code assigned to the word 'on'.

Let's analyze the given information step-by-step:

- Statement 1: 'water is liquid' is coded as **295**.
- Statement 2: 'oil is liquid' is coded as **549**.
- Statement 3: 'oil on water' is coded as **824**.

Step-by-Step Code Language Decoding

We will compare the statements to find common words and their corresponding codes.

Step 1: Compare Statement 1 and Statement 2

- Statement 1: water is liquid (295)
- Statement 2: oil is liquid (549)

The common words in both statements are 'is' and 'liquid'. The common codes in 295 and 549 are 5 and 9.

This means that the codes for 'is' and 'liquid' are 5 and 9, though we don't know which code corresponds to which word yet. However, the remaining words and codes must match.

- In statement 1, the remaining word is 'water' and the remaining code is 2. Thus, the code for '**water**' is **2**.
- In statement 2, the remaining word is 'oil' and the remaining code is 4. Thus, the code for '**oil**' is **4**.

Let's summarise the findings so far:

Word	Possible Code(s)
water	2
oil	4
is/liquid	5 or 9

Step 2: Use Findings with Statement 3

- Statement 3: oil on water (824)

From Step 1, we know the code for 'oil' is 4 and the code for 'water' is 2.

Statement 3 has the words 'oil', 'on', and 'water', with codes 8, 2, and 4.

We can substitute the known codes:

- 'oil' corresponds to code 4.
- 'water' corresponds to code 2.

The codes for 'oil on water' are 8, 2, 4. The codes 4 and 2 match the words 'oil' and 'water'. The only remaining code is 8, and the only remaining word is 'on'.

Therefore, the code for 'on' is 8.

Verification of Code Language

Let's check if the codes make sense with all statements:

- 'water is liquid' = 2 (water) + (5 or 9 for is/liquid) + (5 or 9 for liquid/is). Possible codes: 259 or 295. Given code is 295. Matches.
- 'oil is liquid' = 4 (oil) + (5 or 9 for is/liquid) + (5 or 9 for liquid/is). Possible codes: 459 or 495. Given code is 549, which is a rearrangement of 459. Matches.
- 'oil on water' = 4 (oil) + 8 (on) + 2 (water). Possible codes: Any arrangement of 4, 8, 2. Given code is 824, which is an arrangement of 4, 8, 2. Matches.

The codes we deduced are consistent with all given statements.

The code for 'on' is 8.

Revision Table: Code Language Decoding

Statement	Phrase	Code
1	water is liquid	295
2	oil is liquid	549
3	oil on water	824

Word	Decoded Code
water	2
oil	4
is/liquid	5/9 (in any order)
on	8

Additional Information: Code Language Concepts

Code language or coding-decoding questions are common in logical reasoning tests. They involve deciphering a pattern or rule that relates words/letters/numbers to a code. Here are some basic concepts:

- **Letter Coding:** Letters are substituted with other letters based on a pattern (e.g., shifting position in the alphabet).
- **Number Coding:** Words are assigned numerical codes. This can be based on letter positions, number of letters, or, as in this question, assigning unique numbers to individual words.
- **Symbol Coding:** Words are coded using symbols.
- **Substitution:** A specific word, letter, or number is replaced by another specific word, letter, or number.
- **Comparison Method:** This is the method used in the problem above. By comparing two or more coded messages, you find common elements (words and codes) to identify the code for individual words.

Practicing different types of coding-decoding problems helps in quickly identifying the underlying pattern.

100. Answer: d

Explanation:

Understanding the Famous Painting 'Mahishasura'

The question asks about the artist behind the renowned painting titled 'Mahishasura'. This painting is a significant work in modern Indian art, known for its powerful depiction and unique style.

Identifying the Painter: Tyeb Mehta and 'Mahishasura'

The painting 'Mahishasura', which depicts the slaying of the buffalo demon in Hindu mythology, is one of the most celebrated works by the Indian artist **Tyeb Mehta**. Mehta's style is characterized by his use of bold lines, fractured forms, and vibrant colours, often exploring themes of violence, anguish, and the human condition. 'Mahishasura' is a prime example of his distinctive approach, blending mythological narrative with a modernist sensibility.

About the Artist Tyeb Mehta

Tyeb Mehta (1925–2009) was a prominent Indian painter and a key figure in post-independence Indian art. He was associated with the Progressive Artists' Group but developed his own unique style. His works are highly valued and have achieved record-breaking prices at auctions, solidifying his place as one of India's most important modern artists.

Analysis of Other Options

Let's look at why the other artists listed are not the painter of 'Mahishasura':

- **MF Hussain:** Maqbool Fida Hussain (1915–2011) was another very famous Indian artist, often called the 'Picasso of India'. He is known for his prolific work across various media and his depiction of Indian themes, including mythology and everyday life, but 'Mahishasura' is not attributed to him.
- **Amrita Sher-Gil:** Amrita Sher-Gil (1913–1941) was a pioneering figure in modern Indian art. Her work is recognized for capturing the life of rural India and her unique blend of

European and Indian artistic styles. She died relatively young, and her famous works include 'Three Girls' and 'Self-Portrait', but not 'Mahishasura'.

- **Raja Ravi Varma:** Raja Ravi Varma (1848–1906) was a celebrated Indian painter known for merging European academic art with Indian traditions. He is famous for his realistic portraits and paintings of mythological subjects, which were widely reproduced and popularised through lithographs. While he painted many mythological themes, 'Mahishasura' in the context of the famous modern painting refers to Tyeb Mehta's work.

Based on the established attributions in art history, the famous painting 'Mahishasura' is the work of Tyeb Mehta.

Revision Table: Famous Indian Painters and Their Works

Painter	Era/Style	Notable Works (Examples)
Tyeb Mehta	Modern Indian Art, Fractured Forms	Mahishasura, Falling Figure, Celebration Triptych
MF Hussain	Modern Indian Art, Progressive Artists' Group	Mother India, The Horse Series, Gaja Gamini
Amrita Sher-Gil	Modern Indian Art, Blend of East-West	Three Girls, Bride's Toilet, South Indian Villagers Going to Market
Raja Ravi Varma	Traditional/Modern Transition, Academic Realism	Shakuntala, Damayanti and the Swan, Lady with a Fruit

Additional Information on Mahishasura by Tyeb Mehta

Tyeb Mehta painted 'Mahishasura' multiple times throughout his career, with variations. The most famous version is often cited as a masterpiece. His interpretation of the mythological theme is not a straightforward narrative but an exploration of power, conflict, and transformation through abstract and fragmented forms. The painting is highly regarded for its artistic merit and cultural significance, reflecting the artist's engagement with both Indian tradition and modernist aesthetics.

101. Answer: d

Explanation:

Understanding the Diesel Cycle Combustion Process

The diesel cycle is a thermodynamic cycle that approximates the process occurring in a diesel engine. It consists of a sequence of four distinct processes.

Ideal Diesel Cycle Processes

An ideal diesel cycle is typically represented by the following four processes:

1. **Process 1-2:** Reversible adiabatic compression (Isentropic compression). Air is compressed, increasing its temperature and pressure.
2. **Process 2-3:** Reversible isobaric heat addition (Constant pressure combustion). Fuel is injected into the hot compressed air and burns, adding heat to the system while the piston moves, maintaining constant pressure. This is the combustion phase.
3. **Process 3-4:** Reversible adiabatic expansion (Isentropic expansion). The hot gases expand, pushing the piston and doing work.
4. **Process 4-1:** Reversible isochoric heat rejection (Constant volume heat rejection). Heat is rejected from the system at constant volume.

Focus on Diesel Combustion

In the diesel cycle, combustion is modeled as occurring at constant pressure. Here's why:

- Diesel engines inject fuel towards the end of the compression stroke, just before the piston reaches Top Dead Centre (TDC).
- The fuel is injected gradually over a short period.
- As the fuel burns, the heat release causes the working fluid (air and combustion products) to expand.
- However, because the piston is starting to move downwards during the combustion process, it accommodates this expansion.
- The rate of fuel injection and combustion is controlled such that the pressure inside the cylinder remains approximately constant during this heat addition phase.

This is a key difference compared to the Otto cycle (used in petrol engines), where a spark ignites a pre-mixed fuel-air charge rapidly at constant volume (piston near TDC).

Comparison: Diesel vs. Otto Combustion

Feature	Diesel Cycle	Otto Cycle
Fuel Injection	Directly into cylinder near end of compression	Into intake manifold or cylinder during intake/compression
Ignition	Compression ignition (Autoignition)	Spark ignition
Combustion Process Model	Constant Pressure (Isobaric)	Constant Volume (Isochoric)
Heat Addition Process	Process 2-3 in the P-v diagram	Process 2-3 in the P-v diagram

Therefore, based on the ideal diesel cycle model, combustion takes place at constant pressure.

Pressure-Volume (P-v) Diagram Representation

On a P-v diagram:

- Process 1-2 (compression) is a curve upwards and left (increasing P , decreasing v).
- Process 2-3 (combustion/heat addition) is a horizontal line to the right (increasing v , constant P).
- Process 3-4 (expansion) is a curve downwards and right (decreasing P , increasing v).
- Process 4-1 (heat rejection) is a vertical line upwards (increasing P , constant v).

This visual representation clearly shows the isobaric (constant pressure) nature of the heat addition phase, which corresponds to combustion.

Conclusion on Diesel Combustion Pressure

In summary, the defining characteristic of the heat addition (combustion) process in the ideal diesel cycle is that it occurs while the pressure remains constant. This is achieved through controlled fuel injection as the piston begins to move downwards.

The combustion process in the diesel cycle occurs at constant pressure.

Revision Table: Key Concepts in Diesel Cycle

Process	Description	Thermodynamic Type
1-2	Compression	Isentropic (Constant Entropy)
2-3	Heat Addition (Combustion)	Isobaric (Constant Pressure)
3-4	Expansion	Isentropic (Constant Entropy)
4-1	Heat Rejection	Isochoric (Constant Volume)

Additional Information on Diesel Cycle Thermodynamics

The efficiency of an ideal diesel cycle is given by the formula:

$$\eta_{\text{diesel}} = 1 - \frac{1}{r^{\gamma-1}} \left(\frac{r_c^{\gamma} - 1}{\gamma(r_c - 1)} \right)$$

Where:

- η_{diesel} is the thermal efficiency.
- $r = \frac{V_1}{V_2}$ is the compression ratio.
- $r_c = \frac{V_3}{V_2}$ is the cut-off ratio (ratio of volumes after and before the heat addition process).
- $\gamma = \frac{C_p}{C_v}$ is the ratio of specific heats for the working fluid.

This formula highlights how both the compression ratio and the cut-off ratio influence the efficiency of the diesel cycle.

Real diesel engines deviate from the ideal cycle due to factors like friction, pressure drops during intake and exhaust, and non-instantaneous combustion and heat transfer.

102. Answer: c

Explanation:

Understanding Wire Gauge Measurement

The question asks about what characteristic of a wire is measured by its wire gauge. Wire gauge is a standard measurement system used to denote the diameter of electrical conductors or wires.

What is Wire Gauge?

Wire gauge is a way to specify the physical size, specifically the thickness or diameter, of a wire. Different gauge systems exist, but one common one is the American Wire Gauge (AWG).

In most wire gauge systems, the gauge number is inversely related to the wire's diameter. This means:

- A smaller gauge number indicates a larger wire diameter (and thus a larger cross-sectional area).
- A larger gauge number indicates a smaller wire diameter (and thus a smaller cross-sectional area).

The diameter or cross-sectional area of a wire is crucial because it affects its electrical resistance and its ability to carry current without overheating. A larger diameter (smaller gauge number) means lower resistance and a higher current carrying capacity.

Analyzing the Options

Let's look at the given options:

1. **Brittleness:** Brittleness refers to a material's tendency to break or fracture when subjected to stress. Wire gauge is not a measure of this property.
2. **Tensile strength:** Tensile strength is the maximum stress a material can withstand before breaking when stretched or pulled. Wire gauge does not directly measure tensile strength.
3. **Diameter:** As explained above, wire gauge is specifically a measure of the wire's diameter (or thickness). This is the correct characteristic.
4. **Hardness:** Hardness is a material's resistance to localized plastic deformation, usually by indentation. Wire gauge is not a measure of hardness.

Wire Gauge and Diameter Relationship

Here is a simplified example showing how AWG gauge numbers relate to diameter:

AWG Gauge vs. Diameter
(Approximate)

AWG Gauge	Diameter (mm)
10	2.588
12	2.053
14	1.628
16	1.291
18	1.024
20	0.812

Notice how a smaller gauge number (like 10 AWG) corresponds to a larger diameter, while a larger gauge number (like 20 AWG) corresponds to a smaller diameter.

Conclusion on Wire Gauge Measurement

Based on the standard definition and use of wire gauge systems, it is a direct measurement of the wire's diameter. This physical dimension is critical for determining the wire's electrical properties and suitability for various applications.

Revision Table: Wire Gauge Key Concepts

Key Aspects of Wire Gauge	
Concept	Explanation
What it measures	Wire Diameter (or thickness)
Common System	AWG (American Wire Gauge)
Gauge Number vs. Diameter	Inversely related: Smaller number = Larger diameter; Larger number = Smaller diameter
Importance	Affects electrical resistance and current carrying capacity

Additional Information: Wire Gauge and Applications

Understanding wire gauge is essential in many fields, particularly in electrical work. Choosing the correct wire gauge for an application is crucial for safety and performance. Using a wire that is too small (high gauge number) for the required current can lead to overheating, insulation damage, and potential fire hazards. Using a wire that is unnecessarily large (low gauge number) can be more costly and difficult to work with.

Other gauge systems exist besides AWG, such as the Standard Wire Gauge (SWG) used in the UK, which has a different numbering scheme. Regardless of the specific system, the purpose remains the same: to provide a standardized way to specify wire dimensions.

103. Answer: b

Explanation:

Understanding the Centre Punch and its Point Angle

A centre punch is a hand tool used in metalworking to make a large, deep indentation in a workpiece. This indentation serves as a starting point for drilling, helping to prevent the drill bit from wandering off-centre when starting a hole. The effectiveness of a centre punch is largely determined by the angle of its point.

Importance of Point Angle in Punching Tools

Different types of punches are used for various purposes, and their point angles are specifically designed for those tasks. The angle affects the size and depth of the mark left on the material. For layout work or making small marks, a sharp point is needed, while for creating a starting point for drilling, a wider angle is required to make a larger, more visible indentation.

Identifying the Centre Punch Point Angle

The standard point angle for a centre punch is 90 degrees. This relatively wide angle creates a large enough indentation to guide the tip of a drill bit accurately, ensuring the drill starts precisely where intended.

Let's look at the provided options and see how they relate to different punching tools:

- 30°: This angle is typically too acute for a standard centre punch. It is more commonly associated with tools like a prick punch, used for marking layout lines or transferring measurements.
- 90°: This is the standard point angle for a centre punch. It creates a broad, deep mark suitable for starting drill bits.
- 120°: This angle is even wider than a centre punch and is not standard for common punches used in layout or drilling starting points.
- 60°: This angle is commonly associated with a dot punch or prick punch used for marking layout lines. It creates a smaller, sharper mark than a centre punch.

Therefore, the correct point angle for a centre punch from the given options is 90°.

Comparing Different Punch Point Angles

Understanding the angles of different punches helps in selecting the right tool for the job. Here is a simple comparison:

Tool Type	Typical Point Angle	Primary Use
Prick Punch / Dot Punch	30° or 60°	Marking layout lines, transferring points
Centre Punch	90°	Creating starting points for drilling

Your Personal Exams Guide

The 90° angle of the centre punch ensures that the indentation is wide enough to accept the tip of most drill bits and provides a stable starting point.

Revision Table: Punch Angles and Uses

Punch Type	Common Point Angle	Application
Prick Punch	30°	Fine layout work, marking centres of small circles
Dot Punch	60°	General layout work, marking positions along lines
Centre Punch	90°	Preparing a starting point for drilling to prevent drill bit wander

Additional Information on Punching Tools

Punches are essential layout tools in metalworking and woodworking. Beyond the point angle, other characteristics like material, length, and diameter also vary depending on the intended use. Automatic centre punches are also available; they have a spring-loaded mechanism that strikes the punch automatically when pressure is applied, eliminating the need for a hammer. Regardless of whether it's manual or automatic, a centre punch typically has a 90° point for its primary function of preparing surfaces for drilling.

Always ensure the punch point is sharp and the hammer blow (if using a manual punch) is firm and perpendicular to the workpiece surface for an accurate mark.

104. Answer: b

Explanation:

Understanding Tools for Checking Flatness and Squareness

The question asks us to identify the instrument specifically used for checking the flatness and squareness of a surface. Various tools are used in workshops for checking dimensions, angles, and surface properties. Let's examine the options provided.

Analyzing the Options

- **Go Gauge:** A go gauge is a limit gauge used to check if a workpiece dimension is within the upper limit of the tolerance. It is a 'go' tool because if it fits, the dimension is acceptable. It does not check for flatness or squareness, only a specific linear dimension.
- **Try square:** A try square is a precision tool used primarily for checking if a surface or edge is at a 90-degree angle (square) to another surface or edge. The blade and stock of the try square form a perfect 90-degree angle. It can also be used to check for flatness by placing the edge of the blade on the surface and checking for any gap between the blade and the surface.

- **Vernier height gauge:** A vernier height gauge is a measuring instrument used to determine the height of objects or to scribe lines at a specific height from a reference surface plate. It uses a vernier scale for precise measurements. It is not used for checking angles like squareness or for assessing surface flatness directly.
- **Centre punch:** A centre punch is a hand tool used to make a small indentation in a workpiece. This indentation serves as a starting point for drilling, preventing the drill bit from wandering. It is a marking tool, not a measuring or checking instrument for flatness or squareness.

Identifying the Correct Instrument

Based on the analysis of the functions of each tool:

- The Go Gauge checks linear dimensions.
- The Vernier height gauge measures height and scribes lines.
- The Centre punch is used for marking.
- The Try square is designed specifically to check for 90-degree angles (squareness) and can also be used to check for flatness over a short span by observing light gaps under the blade.

Therefore, the instrument used for checking flatness and squareness of a surface among the given options is the try square.

Correct Answer Explanation: Try Square for Checking Flatness and Squareness

The **try square** is the appropriate tool for the stated purpose. It consists of a thick stock (handle) and a thin blade, rigidly fixed at a 90-degree angle. To check squareness, the stock is placed firmly against one surface (the reference edge), and the blade is brought into contact with the other surface. Any gap between the blade and the second surface indicates that the angle is not square (not 90 degrees). To check for flatness, the edge of the blade can be laid on the surface. If the surface is flat, the blade will rest evenly without any visible light passing underneath. While dedicated surface plates and straight edges are used for high-precision flatness checks, a try square is a common tool for checking flatness over smaller areas or when less precision is required.

Revision Table: Workshop Checking Tools

Instrument	Primary Use	Applicable for Flatness/Squareness?
Go Gauge	Checking dimension limits	No
Try square	Checking 90° angles (squareness), basic flatness checks	Yes
Vernier height gauge	Measuring height, scribing lines	No
Centre punch	Marking drilling points	No

Additional Information: Types of Squares and Related Tools

Beyond the basic try square, several other types of squares and related tools are used for checking angles and surfaces:

- **Engineer's Square:** Similar to a try square but typically more precise, used in engineering workshops.
- **Mitre Square:** Used for checking 45-degree and 135-degree angles.
- **Combination Square:** A versatile tool with multiple heads (square head, center head, protractor head) that slide along a ruler, allowing for checking squareness, mitre angles, center finding, and angle measurement.
- **Surface Plate:** A highly flat, rigid surface used as a reference plane for inspection and layout. Used in conjunction with other tools to check flatness and straightness with high accuracy.
- **Straight Edge:** A tool with a precisely straight edge used for checking the straightness or flatness of surfaces by comparing the surface against the straight edge.

Each tool serves a specific purpose in the process of measurement, marking, and checking in manufacturing and woodworking.

105. Answer: d

Explanation:

Understanding the Purpose of a Dynamo

A dynamo is a type of electrical generator that produces direct current (DC) using a commutator. Its primary function is to convert mechanical energy into electrical energy through the principle of electromagnetic induction. However, the question asks about its specific purpose in a practical context, and the options provide different applications or functions.

Analyzing the Options for Dynamo's Purpose

Let's examine each option to determine the most accurate purpose based on common applications of dynamos:

- **Option 1: Act as reservoir of electrical energy**

A reservoir stores energy. A dynamo generates energy, it does not store it. Devices like batteries or capacitors act as reservoirs of electrical energy. Therefore, this option is incorrect.

- **Option 2: Supply electric power**

This is a broad statement. A dynamo does supply electric power once it converts mechanical energy. However, it's often used for a more specific purpose, as suggested by other options. While true, it might not be the most precise answer depending on the intended context.

- **Option 3: Convert mechanical energy into electrical energy**

This describes the fundamental process by which a dynamo operates. It's the core function. However, the question asks for its "purpose," which might imply the *reason* it is used, rather than just *how* it works. In many systems, this energy conversion serves a further purpose.

- **Option 4: Continually recharge the battery**

This is a very common application of a dynamo (or more modern alternators, which serve a similar purpose and replaced dynamos in many vehicle applications). In vehicles or systems that use a battery, the dynamo generates electrical energy which is then used to power the electrical system and recharge the battery that gets discharged during use or starting. This purpose aligns well with the practical use of a dynamo in many applications where a continuous power source is needed alongside a battery storage system.

Why Recharging the Battery is a Key Purpose

In many real-world systems, such as older car electrical systems or bicycle lighting systems with a battery backup, the dynamo's output is directly used to power components and maintain the charge level of a connected battery. Without the dynamo, the battery would quickly discharge. Thus, a significant purpose of the dynamo in such setups is to ensure the battery remains charged, providing continuous power even when the dynamo is not operating at peak efficiency or when engine/wheel speed is low.

The operation relies on electromagnetic induction, where a changing magnetic field induces an electromotive force (EMF), as described by Faraday's Law:

$$\mathcal{E} = -\frac{d\Phi_B}{dt}$$

Where $\mathcal{E}$ is the induced EMF and $\frac{d\Phi_B}{dt}$ is the rate of change of magnetic flux.

The generated electrical energy powers the system and, importantly, replenishes the energy stored in the battery.

Comparing Option 3 and Option 4

While option 3 describes the essential **function** of energy conversion, option 4 describes a common and crucial **purpose** or **role** the dynamo plays within a larger system, particularly in systems that include a battery. In many applications, the energy conversion (Option 3) serves the explicit purpose of keeping the battery charged (Option 4) and powering the load.

Comparison of Dynamo Purposes		
Purpose	Description	Relevance
Convert mechanical energy into electrical energy	Fundamental operational principle.	How it works.
Continually recharge the battery	A key application role in battery-equipped systems.	Why it's used in certain contexts.

Based on the options provided and the common practical uses of a dynamo, "continually recharge the battery" represents a primary purpose in many systems where a dynamo is employed.

Revision Table: Dynamo Key Concepts

Concept	Description
Dynamo	A DC electrical generator.
Energy Conversion	Converts mechanical to electrical energy.
Principle	Electromagnetic Induction (Faraday's Law).
Output	Direct Current (DC) due to commutator.
Common Purpose	Supply power and recharge battery in systems like vehicles.

Additional Information on Dynamo and Battery Charging

Historically, dynamos were used in vehicles to generate electricity. The generated power was used to run headlights, ignition systems, etc., and any excess power was used to charge the vehicle's lead-acid battery. The battery was essential for starting the engine and powering accessories when the engine was off or running at low speed where the dynamo's output was insufficient.

The dynamo's voltage output varies with its rotation speed. A voltage regulator is typically used in the circuit to maintain a stable voltage for the electrical system and ensure proper battery charging without overcharging.

Modern vehicles largely use alternators instead of dynamos. An alternator produces alternating current (AC), which is then rectified into DC using diodes before being supplied to the vehicle's electrical system and battery. Alternators are generally more efficient and can produce higher output at lower engine speeds compared to dynamos.

Despite the shift to alternators in many applications, the fundamental purpose of generating power to run the system and charge the battery remains the same. Understanding the dynamo's role in maintaining battery charge is key to understanding its purpose in systems where it is employed.

106. Answer: a

Explanation:

Understanding Engine Brake Power Measurement

The question asks about the device used to measure the brake power of an engine. Brake power is the usable power delivered by an engine's crankshaft, measured at the output shaft. It is the power available after accounting for friction and other losses within the engine itself.

Identifying the Correct Device

Let's look at the options provided to determine which device is specifically designed for measuring engine brake power.

- **Dynamometer:** A dynamometer is a device used to measure the torque and rotational speed (RPM) of an engine or other rotating prime mover. Since power is the product of torque and angular velocity (related to RPM), a dynamometer can directly calculate or measure the brake power delivered by the engine. There are various types of dynamometers, such as eddy current, hydraulic, and electric dynamometers, but their core function is power measurement.
- **Energy meter:** An energy meter is typically used to measure electrical energy consumption over a period of time, usually in kilowatt-hours (kWh). It is not used for measuring the mechanical power output of an engine.
- **Torque wrench:** A torque wrench is a tool used to apply a specific, controlled amount of torque to a fastener, such as a nut or bolt. It measures the applied torque but is not designed for measuring the continuous torque or power output of a rotating engine shaft.
- **Indicator:** In the context of engines, an 'indicator' often refers to devices used to measure cylinder pressure versus crankshaft position or volume, like an engine indicator diagram. While this provides insight into the engine's combustion process and indicated power (power generated within the cylinders before losses), it does not directly measure the brake power delivered at the output shaft.

Based on the functions of these devices, the dynamometer is the correct instrument used to measure the brake power of an engine.

Device Functions Comparison

Device	Primary Function	Measures Brake Power?
Dynamometer	Measures torque and RPM to calculate power	Yes
Energy meter	Measures electrical energy consumption	No
Torque wrench	Measures applied torque on fasteners	No
Indicator (Engine)	Measures cylinder pressure/volume; related to indicated power	No (Directly brake power)

Calculating Brake Power

A dynamometer measures the torque (τ) exerted by the engine and its rotational speed (N) in RPM. The brake power (P_b) can then be calculated using the formula:

$$P_b = \frac{2\pi N\tau}{60}$$

(Power in Watts if τ is in Nm and N in RPM)

or often in horsepower:

$$P_b = \frac{2\pi N\tau}{33000}$$

(Power in HP if τ is in lb-ft and N in RPM)

The dynamometer applies a controllable load to the engine's output shaft, allowing measurements across various operating conditions.

Conclusion

The device specifically designed for measuring the brake power of an engine is a dynamometer. It provides the necessary data (torque and RPM) to determine the actual usable power output.

Revision Table: Key Engine Power Concepts

Engine Power Terms

Term	Definition	Measurement
Indicated Power (IP)	Power developed within the engine cylinders by combustion.	Calculated from cylinder pressure diagrams (e.g., using an indicator).
Frictional Power (FP)	Power lost due to friction between moving parts, pumping losses, etc.	Difference between IP and BP ($FP = IP - BP$).
Brake Power (BP)	Actual power available at the engine crankshaft (output shaft).	Measured using a dynamometer.

Additional Information on Dynamometers

Dynamometers are crucial tools in engine testing and development. They are used for:

- Determining maximum power and torque.
- Mapping engine performance curves.
- Testing fuel efficiency.
- Engine durability testing.
- Calibration of engine control systems.

Different types of dynamometers use various methods to absorb and measure the power, such as friction (Prony brake, historically), fluid shear (hydraulic dynamometer), eddy currents (eddy current dynamometer), or electrical generation/absorption (electric dynamometer).

107. Answer: c

Explanation:

Understanding the Diesel Fuel Feed Pump Drive

The diesel fuel feed pump is a vital component in a diesel engine's fuel system. Its main job is to draw fuel from the fuel tank and supply it at a low pressure to the high-pressure

fuel injection pump or common rail system. This ensures a constant and adequate supply of fuel is available for injection into the engine cylinders.

For the diesel fuel feed pump to operate, it requires a drive mechanism. This drive ensures that the pump runs whenever the engine is running and delivers fuel in proportion to the engine's needs, although it typically pumps more fuel than immediately required, with the excess being returned to the tank.

Analyzing the Drive Options for the Diesel Fuel Feed Pump

Let's consider the potential drive sources mentioned in the options:

- **Crankshaft:** The crankshaft rotates at engine speed. While it is the ultimate source of all engine drives, directly driving the low-pressure fuel feed pump from the crankshaft is uncommon because the required pump speed is often lower than full engine speed and needs to be synchronized with the overall fuel delivery timing, which is linked to the camshaft rotation.
- **Timing gears:** Timing gears connect the crankshaft to other rotating components like the engine camshaft and fuel injection pump camshaft. While the drive power is transmitted through these gears, the question asks for the specific component from which the drive is taken. The timing gears facilitate the drive but are not typically the direct component driving the feed pump.
- **Fuel injection camshaft:** In many diesel injection systems, especially those with mechanical fuel injection pumps (like inline or rotary pumps), a dedicated camshaft operates the pumping elements. The diesel fuel feed pump is often mechanically driven directly from this fuel injection camshaft. This arrangement ensures the fuel feed pump operates whenever the fuel injection system is active and provides a drive speed suitable for the feed pump's function, often proportional to the fuel injection pump speed.
- **Engine camshaft:** The engine camshaft operates the intake and exhaust valves and rotates at half the crankshaft speed. While some auxiliary pumps might be driven by the engine camshaft, the fuel feed pump, being closely tied to the fuel injection system's operation and timing, is frequently driven by the fuel injection camshaft if one exists, or sometimes by the engine camshaft in simpler systems. However, the option specifically mentioning the 'Fuel injection camshaft' is the most direct and common drive source for the fuel feed pump in systems utilizing such a camshaft.

Explanation of Fuel Injection Camshaft Drive

When the diesel fuel feed pump is driven by the fuel injection camshaft, it creates a mechanically synchronized system. As the fuel injection camshaft rotates to activate the high-pressure injection process, the feed pump also operates, ensuring that there is always fuel ready at the inlet of the high-pressure pump elements. This direct mechanical link simplifies the system and provides reliable fuel delivery.

Therefore, the drive for the diesel fuel feed pump is typically taken from the fuel injection camshaft in many diesel engine configurations.

Revision Table: Diesel Fuel Pump Drive Sources

Component	Typical Role	Likelihood of Driving Feed Pump
Crankshaft	Engine power source	Unlikely as direct drive
Timing gears	Connects crankshaft to camshafts	Facilitates drive, not direct source
Fuel injection camshaft	Drives injection pump elements	Likely direct drive source
Engine camshaft	Operates valves	Possible in some systems, but less direct relation to injection timing than the fuel injection camshaft.

Additional Information: Diesel Engine Fuel System Components

Understanding the different components of a diesel engine's fuel system helps clarify the role of the fuel feed pump and its drive.

- **Fuel Tank:** Stores the diesel fuel.
- **Fuel Filter:** Removes contaminants from the fuel before it reaches the pumps and injectors.
- **Fuel Feed Pump (Lift Pump):** Transfers fuel from the tank to the high-pressure pump. Operates at low pressure.

- **High-Pressure Fuel Pump:** Compresses the fuel to very high pressures required for injection. This pump's operation is often timed with the engine cycle.
- **Fuel Injectors:** Receive the high-pressure fuel and spray it into the combustion chamber at the precise moment determined by the engine control unit or the mechanical injection system.
- **Fuel Lines:** Carry fuel between the various components.
- **Return Line:** Carries excess fuel from the injection system back to the fuel tank.

The efficient operation of the diesel fuel feed pump is crucial for the proper functioning of the entire fuel injection system, ensuring the high-pressure pump is adequately supplied with fuel.

108. Answer: d

Explanation:

Understanding the Role of Lubricant Oil in Machine Operation

Lubricant oil plays a crucial role in the efficient and smooth operation of machines. Its primary function is to reduce friction between moving parts. When machine components rub against each other without proper lubrication, it causes significant friction, leading to heat generation, wear, and noise.

Let's examine the provided options in the context of lubricant oil's function:

- **Option 1: quickly**
While reduced friction can allow a machine to operate more efficiently and potentially at intended speeds, the primary purpose of lubricant oil isn't solely to make it run "quickly." Quickness is a result of many factors, including design and power input.
- **Option 2: With increased friction**
This is the opposite of what lubricant oil does. Lubricants are specifically designed to *decrease* friction, not increase it. Increased friction would lead to problems like overheating, wear, and seizing.
- **Option 3: With cleanliness**
Some lubricants can help carry away contaminants, contributing to cleanliness, but

this is often a secondary function compared to friction reduction, wear prevention, and heat dissipation. The main reason for needing the oil isn't *just* for cleanliness.

- **Option 4: Without noise**

When moving parts in a machine experience high friction, they can generate significant noise due to vibration, grinding, and scraping. By reducing friction, lubricant oil allows parts to slide or roll smoothly against each other, greatly minimizing or eliminating this noise. This is a direct and important benefit of proper lubrication.

Therefore, among the given options, the most direct and significant outcome achieved by using lubricant oil is enabling the machine to run with significantly reduced noise, which stems from its fundamental role in reducing friction.

Benefits of Using Lubricant Oil

Using appropriate lubricant oil provides multiple benefits for machine longevity and performance:

- Reduces friction between moving surfaces.
- Minimizes wear and tear on components, extending machine life.
- Dissipates heat generated by friction, preventing overheating.
- Helps seal against contaminants like dirt and moisture.
- Dampens shocks and vibrations.
- Reduces energy consumption by improving efficiency.
- Significantly reduces operational noise.

Analyzing the Options and Machine Performance

Considering the core functions of lubrication, option 4, "Without noise," is a direct consequence of reduced friction and smooth operation provided by the lubricant oil. While other benefits exist, the options focus on specific outcomes. Increased friction (option 2) is clearly incorrect. Running quickly (option 1) and cleanliness (option 3) are less direct or primary reasons compared to the fundamental benefits of reduced friction, wear, and noise.

The reduction in friction directly leads to quieter operation, making "Without noise" a valid and key benefit provided by lubricant oil.

Aspect of Lubrication	Impact on Machine Operation
Reduces Friction	Enables smooth movement, reduces wear, decreases heat and noise.
Prevents Wear	Extends the lifespan of machine components.
Manages Heat	Prevents thermal damage and maintains material properties.
Reduces Noise	Results in quieter and more pleasant operation.

Revision Table: Key Concepts of Lubrication

Concept	Explanation
Friction	Resistance to motion between two surfaces in contact. Lubrication minimizes this.
Wear	Surface damage or material removal from components due to friction. Lubrication prevents this.
Lubricant	A substance (like oil or grease) applied to reduce friction between surfaces.
Noise	Sound produced by vibrations, often caused by friction between moving parts. Lubrication reduces this.

Additional Information: Types and Applications of Lubricants

Lubricants come in various forms, including oils, greases, and dry lubricants, each suited for different applications based on factors like load, speed, temperature, and environmental conditions.

- **Oils:** Liquid lubricants, commonly used for high-speed applications or where cooling is required. They can be mineral-based, synthetic, or vegetable-based.
- **Greases:** Semi-solid lubricants, consisting of a base oil thickened with a soap or other agent. Ideal for applications where the lubricant needs to stay in place, like

bearings or gears.

- **Dry Lubricants:** Solid materials like graphite or PTFE, used when wet lubricants are unsuitable, such as in dusty environments or high-temperature applications.

Choosing the correct type of lubricant is essential for optimal machine performance, efficiency, and lifespan.

109. Answer: d

Explanation:

Understanding how to measure temperature accurately is crucial in many fields, especially when dealing with extreme heat, such as the temperatures found inside a furnace. Different instruments are designed for different temperature ranges and applications. This question asks about the specific instrument used for measuring high temperatures in a furnace.

Measuring High Temperatures in Furnaces

Furnaces operate at very high temperatures, often exceeding the range of standard thermometers. Instruments used in such environments must be able to withstand the heat and provide accurate readings without direct contact, or with specialized high-temperature resistant components.

Analyzing Temperature Measurement Instruments

Let's look at the options provided and understand what each instrument is typically used for:

- **Colorimeter:** This instrument is used to measure the absorbance of light by a liquid solution. It is used in chemistry and biology to determine the concentration of a colored substance. It is not used for measuring temperature.
- **Barometer:** A barometer is an instrument used to measure atmospheric pressure. It is primarily used in meteorology to forecast weather changes. It has no application in measuring temperature inside a furnace.
- **Thermometer:** A thermometer is a general instrument used for measuring temperature. While thermometers are widely used, standard types like liquid-in-

glass thermometers have a limited temperature range. For higher temperatures, specialized thermometers like thermocouples might be used, but often direct contact is still required, which can be challenging or impossible in a furnace environment.

- **Pyrometer:** A pyrometer is a type of thermometer used for measuring high temperatures from a distance, typically by detecting the thermal radiation emitted by the object. Because furnaces operate at extremely high temperatures and direct contact measurement can be difficult or damaging, pyrometers are ideal for this application. They can measure temperatures without touching the hot surface, making them safe and practical for furnace temperature monitoring.

Why Pyrometers are Used for Furnace Temperatures

Pyrometers work based on the principle that all objects above absolute zero emit thermal radiation. As the temperature of an object increases, the intensity and spectrum of the emitted radiation change. A pyrometer collects and measures this radiation, correlating it to the object's temperature. This non-contact method is essential for environments like furnaces where temperatures are very high and direct contact instruments might melt or be damaged. This makes the pyrometer the suitable instrument for measuring high temperatures in a furnace.

Instrument vs. Typical Measurement

Instrument	What it Measures	Suitable for High Furnace Temperatures?
Colorimeter	Light absorbance/concentration	No
Barometer	Atmospheric pressure	No
Thermometer (Standard)	Temperature (Limited range)	Generally No (for furnace heat)
Pyrometer	High temperature (via radiation)	Yes

Based on the function and typical application of each instrument, the pyrometer is specifically designed for measuring high temperatures remotely, making it the correct

instrument for use in a furnace.

Revision Table: Temperature and Pressure Instruments

Key Measurement Instruments		
Instrument Type	Primary Measurement	Common Application Example
Thermometer	Temperature	Measuring body temperature
Barometer	Atmospheric Pressure	Weather forecasting
Manometer	Pressure (relative)	Measuring gas pressure in a container
Pyrometer	High Temperature (remote)	Measuring furnace temperature

Additional Information: Types of Pyrometers for High Temperatures

Pyrometers come in different types, each utilizing slightly different principles or technologies to measure high temperatures:

- **Radiation Pyrometers:** These measure the total thermal radiation emitted by the object. They are sensitive to a broad range of wavelengths.
- **Optical Pyrometers:** These work by comparing the brightness of the hot object to a calibrated light source (often a filament). The observer adjusts the light source until it matches the brightness of the object at a specific wavelength (usually red).
- **Infrared Pyrometers:** These are a common type of radiation pyrometer that specifically measures infrared radiation emitted by the object. They are widely used in industrial applications for non-contact temperature measurement.

These variations allow pyrometers to be adapted for various industrial processes requiring high-temperature measurement, such as in steel production, glass manufacturing, and, as in the question, monitoring temperatures inside furnaces.

110. Answer: a

Explanation:

Understanding Belt Slackness in Power Transmission

Belt slackness, also known as excessive sag or looseness, is a common issue in belt drive systems. It occurs when the tension in the belt is lower than required for efficient power transmission. Proper belt tension is crucial; too little tension can lead to slipping and wear, while too much tension can overload bearings and reduce belt life.

Let's analyze the given options to understand what could potentially cause belt slackness.

Analyzing Potential Causes of Belt Slackness

We are looking for a component whose wear can lead to a loss of proper belt tension or alignment, resulting in slackness.

- **Worn shaft:** A shaft is a rotating component, often supporting pulleys in a belt drive. If the shaft itself is worn, especially at the points where bearings or pulleys are mounted, it can lead to improper pulley seating, misalignment, or excessive play. This misalignment or inability to maintain correct distance between pulleys can reduce the effective tension in the belt, causing slackness. A worn shaft could also contribute to vibration, further affecting belt tension.
- **Worn chain:** A chain is part of a chain drive system, which uses sprockets instead of pulleys and transmits power via interlocking links. Slackness in a chain drive is referred to as chain sag, and it's due to elongation of the chain links, not a worn shaft, gear, or sprocket in isolation, unless those components cause improper alignment. A worn chain is irrelevant to belt slackness.
- **Worn gear:** Gears are components of gear drives, which transmit power through the meshing of teeth. Wear in gears typically leads to issues like increased backlash, noise, and reduced efficiency, but it is unrelated to belt slackness in a belt drive system.
- **Worn sprocket:** A sprocket is used in chain drives or synchronous belt drives (timing belts). While a worn sprocket in a synchronous belt drive could affect meshing and cause issues, "belt slackness" usually refers to V-belts or flat belts where pulleys are used. Wear on a sprocket in a chain drive leads to chain sag and poor engagement, but again, this relates to chains, not belts.

Explaining the Impact of a Worn Shaft on Belt Tension

Considering the options, a worn shaft is the most plausible cause among the choices provided that could lead to belt slackness within a belt drive system. If the shaft supporting a pulley is worn where the pulley seats, the pulley might not be held firmly or squarely in place. This can cause:

- Pulley wobble or misalignment.
- Difficulty in maintaining the correct center distance between pulleys.
- Reduced effective tension in the belt, leading to slackness.

While belt stretch and pulley misalignment due to worn bearings are very common causes of belt slackness, a worn shaft is a fundamental structural issue that can precipitate these problems or directly affect pulley positioning, thereby causing slackness.

Comparative Analysis of Options

Component	System	Typical Result of Wear	Relevance to Belt Slackness
Worn shaft	Belt Drive (supports pulley)	Pulley misalignment, loose pulley fit, vibration	Can lead to reduced belt tension and slackness.
Worn chain	Chain Drive	Chain sag, skipping	Not applicable to belt drives.
Worn gear	Gear Drive	Backlash, noise	Not applicable to belt drives.
Worn sprocket	Chain Drive or Synchronous Belt Drive	Poor engagement, skipping	Not applicable to standard V/flat belt slackness in this context.

Based on this analysis, a worn shaft is the only option among the choices that is directly related to a potential cause of belt slackness in a typical belt drive system.

Revision Table: Key Concepts

Term	Definition/Explanation	Impact on Belt Drive
Belt Slackness	Insufficient tension in a belt drive belt.	Slipping, reduced power transmission, increased wear.
Belt Tension	The force applied to a belt to maintain proper grip and reduce sag.	Crucial for efficient and reliable operation.
Worn Shaft	Damage or material loss on a shaft surface.	Can affect seating and alignment of components like pulleys.

Additional Information: Causes of Belt Slackness

While a worn shaft can contribute, other common causes of belt slackness include:

- Belt stretch due to age or overloading.
- Improper initial tensioning during installation.
- Pulley misalignment (often due to worn bearings or mounting issues).
- Wear on pulley grooves (for V-belts) or surfaces (for flat belts).
- Loose motor or equipment mounts allowing pulleys to move closer together.

Regular inspection and maintenance are important for identifying and addressing these issues to prevent belt slackness and ensure efficient power transmission.

Your Personal Exams Guide

111. Answer: c

Explanation:

Understanding Ammeters and Electrical Measurement

The question asks what physical quantity an ammeter is used to measure. In the study of electricity, various instruments are used to quantify different aspects of an electrical circuit. An ammeter is a fundamental tool in this regard.

What is an Ammeter?

An ammeter is an electrical instrument specifically designed to measure the electric current flowing through a circuit. Electric current is the rate of flow of electric charge. It is typically measured in Amperes (A).

How an Ammeter Measures Current

To measure the current flowing through a particular component or section of a circuit, the ammeter must be connected in series with that component. This ensures that the entire current flowing through the component also passes through the ammeter. An ideal ammeter has zero resistance, so it does not affect the current being measured. Real ammeters have a very low internal resistance.

Distinguishing Ammeters from Other Instruments

It's important to distinguish an ammeter from other common electrical measuring instruments:

- **Voltmeter:** Measures the electric potential difference (voltage) across two points in a circuit. It is connected in parallel.
- **Ohmmeter:** Measures the electrical resistance of a component or circuit.
- **Wattmeter:** Measures the electric power consumed or produced by a component or circuit.

Let's look at a comparison in a table:

Instrument	Quantity Measured	Symbol for Quantity	SI Unit	How Connected
Ammeter	Electric Current	I	Ampere (A)	Series
Voltmeter	Electric Potential Difference (Voltage)	V	Volt (V)	Parallel
Ohmmeter	Electrical Resistance	R	Ohm (Ω)	Across component (when isolated)
Wattmeter	Electric Power	P	Watt (W)	Series for current, parallel for voltage

Based on this, the ammeter is specifically used for the measurement of electric current.

Analysing the Options

- **Resistance:** Measured by an ohmmeter. Incorrect.
- **Power:** Measured by a wattmeter. Incorrect.
- **Current:** Measured by an ammeter. Correct.
- **Voltage:** Measured by a voltmeter. Incorrect.

Therefore, an ammeter is used for the measurement of current.

Revision Table: Ammeter Measurement Summary

Key Aspect	Ammeter Details
Purpose	Measure electric current
Quantity Symbol	I
Unit	Ampere (A)
Circuit Connection	Series
Ideal Internal Resistance	Zero

Additional Information on Ammeters

Ammeter Range: Ammeters are available in different ranges (e.g., milliampere, ampere). Using an ammeter with an appropriate range is crucial for accurate measurement and to prevent damage to the instrument.

Types of Ammeters: Ammeters can be analog (using a needle on a scale) or digital (displaying the value numerically). They can also be classified based on the type of current they measure, such as DC ammeters and AC ammeters.

Connecting in Series: The rule for connecting an ammeter in series means that the current flows into one terminal and out the other, passing through the measuring mechanism. Breaking the circuit and inserting the ammeter is necessary for a series connection.

112. Answer: d

Explanation:

Understanding different cutting tools used in manufacturing and machining is essential. Each tool has a specific purpose, and choosing the right one ensures accuracy and efficiency. The question asks about a cutting tool used to finish and enlarge a hole.

Understanding Cutting Tools for Holes

Let's examine the functions of the tools provided in the options:

Drill Function

A drill is a rotary cutting tool used to create new, typically cylindrical, holes in solid materials. It removes material to form the hole. While drills can enlarge a pre-existing hole, their primary function is creating the initial hole, and they do not typically provide a high-precision finish.

Die Function

A die is a tool used to cut or form material, particularly for creating threads on the exterior of a cylindrical object (like a bolt or rod). It is not used for working on the inside of a hole.

Tap Function

A tap is a tool used to cut internal threads within a pre-drilled hole. This process is called tapping. Like a die, its purpose is thread cutting, not just general finishing or enlarging the hole itself.

Reamer Function

A reamer is a rotary cutting tool used to enlarge a hole to a precise size and provide a smooth finish. It removes only a small amount of material, typically after a hole has been drilled or punched. The goal of reaming is accuracy in diameter and a good surface finish inside the hole.

Why Reamer is the Correct Tool

Based on the functions described, the tool specifically designed to "finish and enlarge a hole" to a precise dimension with a good surface quality is the reamer. While a drill can enlarge a hole, it doesn't provide the finishing quality or dimensional accuracy that a reamer does. Dies and taps are used for cutting threads, not for finishing and simply enlarging a hole.

Revision Table: Tool Comparison

Tool	Primary Function	Used for Holes?	Finishing Capability
Drill	Create initial hole	Yes (creates hole)	Low precision finish
Die	Cut external threads	No (works on external surfaces)	Not applicable
Tap	Cut internal threads	Yes (threads inside hole)	Not applicable
Reamer	Finish and precisely enlarge hole	Yes (works inside hole)	High precision finish

Additional Information on Hole Finishing

Reaming is often the final step in creating an accurately sized and smooth hole. It is usually performed after drilling. The amount of material removed by a reamer is typically very small, ranging from a few thousandths to tens of thousandths of an inch, depending on the hole size and material. Reamers come in various types, including hand reamers, machine reamers, adjustable reamers, and shell reamers, each suited for different applications and levels of precision.

113. Answer: d

Explanation:

A vernier height gauge must be used on a **surface plate**, which provides the flat, stable reference datum from which heights are measured.

114. Answer: a

Explanation:

Understanding the Crankshaft and its Components

The question asks about the purpose of a hole drilled between the crankshaft's main journal and crank pin. To answer this, we need to understand what these parts are and how they function within an internal combustion engine.

Crankshaft Components Explained

- **Crankshaft:** This is a rotating shaft that translates the linear motion of the pistons into rotational motion. It is one of the most critical parts of an engine.
- **Main Journals:** These are the cylindrical surfaces on the crankshaft that rest in the main bearings of the engine block. The crankshaft rotates within these main bearings.
- **Crank Pins (or Connecting Rod Journals):** These are the cylindrical surfaces on the crankshaft to which the connecting rods are attached. They are offset from the main journals, creating the "throw" of the crankshaft, which determines the engine's stroke.
- **Connecting Rod:** This part connects the piston to the crank pin. It transmits the force from the piston to the crankshaft and translates the reciprocating motion of the piston into the rotary motion of the crankshaft.
- **Connecting Rod Bearings:** These are split bearings fitted around the crank pins, inside the big end of the connecting rod. They support the connecting rod on the crank pin and allow smooth relative motion.

The hole mentioned in the question connects the oil passage from the main journal to the crank pin.

Purpose of the Lubrication Hole

The primary function of this drilled passage is to facilitate the flow of lubricating oil. The engine's lubrication system supplies oil under pressure to the main bearings (where the main journals rotate). From the main journals, the oil enters the drilled passage within the crankshaft. This passage extends through the crankshaft's web or arm to reach the crank pin.

Once the oil reaches the crank pin, it lubricates the connecting rod bearings, which are under significant load and experience high relative motion. Proper lubrication is essential to prevent wear, reduce friction, dissipate heat, and dampen noise between the crank pin and the connecting rod bearings.

Analyzing the Options

Let's examine each option provided in the question based on the function of this specific hole in the crankshaft:

Option 1: Lubricating the connecting rod bearings

As explained above, the main purpose of the oil passage drilled between the main journal and the crank pin is specifically to deliver lubricating oil to the connecting rod bearings. This allows the connecting rod's big end to rotate smoothly around the crank pin.

Option 2: Reducing the crankshaft vibrations

Reducing crankshaft vibrations is primarily achieved through the design of the crankshaft, including counterweights that balance the rotating and reciprocating masses. A small drilled hole for lubrication does not significantly contribute to reducing overall crankshaft vibrations.

Option 3: Reducing crankshaft weight

While drilling a hole does remove some material and thus reduces weight, the primary engineering purpose of this specific hole is not weight reduction. The amount of weight removed by this narrow passage is negligible compared to the overall weight of the crankshaft and does not serve as a significant weight-saving measure in the way designed lightening holes in other components might.

Option 4: Balancing the crankshaft

Balancing the crankshaft is a critical process to ensure smooth engine operation. It is typically done by adding counterweights to the crankshaft webs or sometimes by removing material from specific areas during manufacturing. Drilling small lubrication holes is not a method used for balancing the crankshaft. Balancing requires precise calculation and material removal or addition in specific locations to counteract inertial forces.

Conclusion: The Primary Reason for the Hole

Based on the analysis of the crankshaft design and the function of its components, the hole drilled between the crankshaft main journal and crank pin serves the crucial role of transporting lubricating oil from the main bearing feed to the connecting rod bearings.

Revision Table: Crankshaft Hole Function

Crankshaft Hole Drilled Between Main Journal and Crank Pin	Purpose	Explanation
Connects oil passage	Lubricating Connecting Rod Bearings	Allows oil flow from main bearings to highly loaded connecting rod bearings, reducing friction and wear.
Does not significantly affect	Reducing Vibrations	Vibrations are handled by overall crankshaft design and balancing.
Does not significantly affect	Reducing Weight	Weight reduction is not the primary or significant purpose of this specific hole.
Does not affect	Balancing the Crankshaft	Balancing is done through counterweights or targeted material removal elsewhere.

Additional Information: Engine Lubrication System

The hole in the crankshaft is just one part of the complex engine lubrication system. Here's a brief overview:

- Oil is stored in the oil pan (or sump).
- An oil pump draws oil from the pan and sends it under pressure through oil passages (galleries) in the engine block and cylinder head.
- These passages lead to critical areas needing lubrication, including the main bearings.
- Oil enters the main bearings, lubricating the main journals of the crankshaft.

- From the main journal, oil enters the internal passages drilled within the crankshaft, flowing to the crank pins.
- The oil then lubricates the connecting rod bearings.
- Other passages send oil to lubricate camshaft bearings, valvetrain components, cylinder walls, piston rings, etc.
- The oil then drains back down into the oil pan, where the cycle repeats after being filtered.

Effective lubrication is vital for the longevity and performance of an engine, preventing seizure and excessive wear of moving parts like bearings.

115. Answer: b

Explanation:

Understanding Engine Stroke Definition

The question asks us to define the term "stroke" in the context of an engine. Engine stroke is a fundamental measurement related to the operation and design of an internal combustion engine.

Let's look at the options provided:

- Length of connecting rod
- The distance between TDC and BDC
- Internal diameter of cylinder
- Volume of cylinder

To determine the correct definition of the engine stroke, we need to understand the basic movement of the piston inside the cylinder.

Defining Engine Stroke, TDC, and BDC

In a reciprocating engine, a piston moves up and down within a cylinder. This movement is crucial for the engine's operation. The stroke is directly related to the distance covered by the piston during one complete movement in a single direction.

There are two critical points in the piston's travel:

- **Top Dead Centre (TDC):** This is the position of the piston when it is at its uppermost point within the cylinder. At TDC, the piston is closest to the cylinder head.
- **Bottom Dead Centre (BDC):** This is the position of the piston when it is at its lowermost point within the cylinder. At BDC, the piston is furthest from the cylinder head.

The piston travels from TDC to BDC, and then from BDC back to TDC. Each of these single movements is called a stroke. The distance the piston travels during one stroke is the same whether it's moving down (TDC to BDC) or moving up (BDC to TDC).

Therefore, the engine stroke is defined as the linear distance covered by the piston as it moves from Top Dead Centre (TDC) to Bottom Dead Centre (BDC), or vice versa. This distance is a fixed dimension for a given engine design.

Analyzing the Given Options

- **Option 1: Length of connecting rod**

The connecting rod connects the piston to the crankshaft. While its length is important for engine geometry, the stroke is not equal to the length of the connecting rod. The stroke is twice the radius of the crankshaft throw (half the diameter of the circle traced by the crankpin).

- **Option 2: The distance between TDC and BDC**

As explained above, the distance between the piston's uppermost point (TDC) and lowermost point (BDC) is exactly the definition of one engine stroke. This is the correct interpretation.

- **Option 3: Internal diameter of cylinder**

The internal diameter of the cylinder is known as the "bore" of the engine. Bore and stroke are two distinct measurements used to define the size and characteristics of an engine cylinder. The stroke is the distance the piston travels along the cylinder's axis, while the bore is the diameter perpendicular to the axis.

- **Option 4: Volume of cylinder**

The volume of the cylinder can refer to different things, such as the clearance volume (volume above the piston at TDC), the swept volume (volume displaced by the piston

from TDC to BDC), or the total cylinder volume (swept volume + clearance volume).
None of these volumes directly represent the linear distance defined as the stroke.

Based on the definitions and analysis, the correct answer is the distance between TDC and BDC.

In summary, the engine stroke is a crucial dimension in engine design, representing the path length of the piston during its vertical travel within the cylinder.

Term	Definition
Engine Stroke	The distance between Top Dead Centre (TDC) and Bottom Dead Centre (BDC).
Bore	The internal diameter of the cylinder.
TDC (Top Dead Centre)	Piston position at its highest point.
BDC (Bottom Dead Centre)	Piston position at its lowest point.
Swept Volume	Volume displaced by the piston moving from TDC to BDC (calculated using bore and stroke).

Revision Table: Key Engine Dimensions

Dimension	Description	Relevance to Stroke
Stroke (L)	Distance piston travels between TDC and BDC.	Is the definition itself.
Bore (D)	Internal cylinder diameter.	Used with stroke to calculate swept volume.
TDC	Highest piston position.	One endpoint of the stroke.
BDC	Lowest piston position.	The other endpoint of the stroke.
Swept Volume (V_s)	Volume displaced by piston in one stroke ($V_s = \frac{\pi}{4} D^2 L$).	Calculated using stroke (L) and bore (D).

Additional Information on Engine Stroke

The stroke length, along with the bore diameter, significantly influences the characteristics of an engine. Engines are often described by their bore and stroke dimensions. For example, a "square" engine has a bore approximately equal to its stroke.

- **Stroke length affects engine speed:** A longer stroke means the piston travels a greater distance, which generally limits the maximum engine speed (RPM) compared to a shorter stroke for the same piston speed.
- **Stroke length affects torque:** A longer stroke can contribute to producing more torque at lower engine speeds.
- **Stroke length affects engine height:** Longer stroke engines tend to be taller.
- **Stroke length relates to crankshaft throw:** The stroke length is equal to two times the crankshaft throw (the distance from the center of the main journal to the center of the crankpin journal).

Understanding the definition of stroke is fundamental to understanding engine dimensions, displacement, and performance characteristics.

116. Answer: c

Explanation:

Understanding Scavenging in Engines

Scavenging is a process used in internal combustion engines, particularly common in two-stroke engines. Its main purpose is to clear the cylinder of exhaust gases from the previous combustion cycle and replace them with a fresh charge of air or fuel-air mixture before the next cycle begins.

Think of it like clearing the air in a room before bringing in fresh air. If the smoky air from the previous cycle isn't properly removed, it mixes with the new fuel-air mixture. This dilution affects the quality of the mixture and prevents complete and efficient burning during combustion.

Why Scavenging is Performed: Maintaining Engine Performance

The effectiveness of scavenging directly impacts how much fresh charge enters the cylinder and how much residual exhaust gas remains. A cylinder filled with a high proportion of fresh, undiluted charge will combust more completely and produce more power compared to one where the fresh charge is mixed with exhaust gases.

Therefore, efficient scavenging is crucial for maximizing the amount of usable fuel-air mixture that participates in combustion in each cycle. This leads to stronger power strokes and higher power output from the engine.

Poor scavenging results in:

- Less fresh charge in the cylinder.
- Dilution of the fresh charge by residual exhaust gases.
- Incomplete or less efficient combustion.
- Reduced power output.

By ensuring effective removal of exhaust gases and introduction of a fresh charge, scavenging helps to maintain and optimize the engine's power output.

Analyzing the Options

Let's look at the given options in the context of scavenging:

- **Oil pressure:** This relates to the lubrication system of the engine, ensuring moving parts are adequately lubricated. Scavenging has no direct role in maintaining engine oil pressure.
- **Fuel consumption:** While efficient scavenging contributes to better combustion and thus potentially better fuel efficiency, scavenging is primarily aimed at ensuring the cylinder is ready for the *next* combustion cycle to produce power. Fuel consumption is more of a consequence of overall engine efficiency, which scavenging helps improve, but it's not the direct parameter being maintained by the scavenging process itself.
- **Power output:** This is directly affected by scavenging. Effective scavenging ensures a high volume of fresh, undiluted mixture is available for combustion, leading to maximum possible power output from the engine cycle. Scavenging is performed precisely to achieve this optimal cylinder condition for generating power.
- **Speed:** Engine speed is a result of the power produced and the load on the engine. Scavenging helps in producing the power, but it doesn't directly maintain a specific speed. The power output facilitated by scavenging allows the engine to achieve and maintain speed under load.

Based on the direct impact of scavenging on the combustion quality and the volume of combustible mixture, its primary purpose is to maintain or enhance the power output of the engine.

Term	Relation to Scavenging
Power Output	Directly maintained/enhanced by effective scavenging.
Fuel Consumption	Indirectly improved by efficient combustion resulting from scavenging.
Oil Pressure	No direct relation.
Speed	A result of the power output enabled by scavenging.

Revision Table: Scavenging Fundamentals

Concept	Explanation	Importance
What is Scavenging?	Process of removing exhaust gases from cylinder and replacing with fresh charge.	Essential for preparing cylinder for next combustion cycle.
Why is it done?	To prevent dilution of fresh charge by residual exhaust.	Ensures effective combustion and high power output.
Primary Goal	To maximize the amount of fresh, combustible mixture in the cylinder.	Directly impacts engine power and efficiency.

Additional Information on Engine Scavenging

Scavenging is particularly critical in two-stroke engines because the intake and exhaust processes happen almost simultaneously and very quickly. Unlike four-stroke engines where intake and exhaust have dedicated strokes, in two-stroke engines, these events occur during the piston's movement near the bottom dead center (BDC).

Different methods of scavenging exist, depending on the engine design:

- **Cross-flow scavenging:** Intake and exhaust ports are on opposite sides of the cylinder. Baffles on the piston or in the ports direct the flow.
- **Loop scavenging:** Intake and exhaust ports are on the same side of the cylinder. The fresh charge is directed upwards, loops around the cylinder head, and pushes exhaust gases out of the exhaust ports below.
- **Uniflow scavenging:** Intake ports are located at one end of the cylinder (usually the bottom), and exhaust valves are at the other end (usually the cylinder head). The fresh charge flows in one direction, pushing exhaust gases out ahead of it. This is considered the most efficient method.

Efficient scavenging ensures that the cylinder is effectively 'cleaned' of burned gases, allowing the fresh charge to occupy the maximum possible volume, leading to better combustion efficiency and, consequently, higher engine power output and potentially improved fuel economy, although the primary objective is power.

117. Answer: a

Explanation:

Understanding Exhaust Heat Recovery

Exhaust gases produced by engines, whether in vehicles or power plants, carry a significant amount of thermal energy. This energy is typically wasted as the gases are released into the atmosphere. Recovering this heat can improve overall system efficiency and reduce energy consumption.

Why Recover Heat from Exhaust Gases?

Recovering heat from exhaust gases can lead to various benefits:

- **Increased energy efficiency:** The recovered heat can be used for other purposes instead of requiring additional energy input.
- **Reduced fuel consumption:** By using waste heat, less primary fuel might be needed.
- **Lower emissions:** Improved efficiency can sometimes lead to reduced overall emissions.
- **Potential for power generation:** In some large-scale applications, recovered heat can be used to generate electricity.

Identifying the Tool for Heat Recovery

The question asks specifically about the device used to recover heat from exhaust gases. Let's examine the provided options:

- **Heat-exchanger:** A heat-exchanger is a device that efficiently transfers heat from one medium (like hot exhaust gas) to another medium (like air, water, or another gas) without mixing the two. This is precisely the function required for recovering heat.
- **Manifold:** An exhaust manifold collects the exhaust gases from the multiple cylinders of an engine into a single pipe. Its primary function is collection, not heat recovery.
- **Muffler:** A muffler is designed to reduce the noise level of the exhaust gases as they exit the engine. It achieves this through a series of chambers, baffles, or perforated tubes, but it does not recover heat for beneficial use.

- **Tail pipe:** The tail pipe is the final section of the exhaust system through which the gases are expelled into the atmosphere. It is simply a conduit and does not perform heat recovery.

Based on the function of each component, the device specifically designed for transferring and recovering heat from one fluid stream to another is a heat-exchanger.

How Heat-Exchangers Recover Exhaust Heat

In the context of exhaust heat recovery, a heat-exchanger would be positioned in the exhaust path. The hot exhaust gases flow on one side, and a working fluid flows on the other side. Heat is transferred from the hot exhaust gases through a separating wall to the working fluid. The heated working fluid can then be used for various applications, such as heating cabin air, preheating combustion air, generating steam, or even converting into mechanical or electrical energy in advanced systems like Organic Rankine Cycles (ORCs).

Therefore, much of the heat in the exhaust gases can indeed be recovered by the use of a heat-exchanger.

Revision Table: Exhaust System Components

Component	Primary Function	Role in Heat Recovery
Exhaust Manifold	Collects exhaust gases from cylinders	None
Heat-exchanger	Transfers heat between fluids	Primary device for heat recovery
Muffler	Reduces exhaust noise	None
Tail Pipe	Expels exhaust gases to atmosphere	None

Additional Information on Heat-Exchanger Types for Exhaust Heat Recovery

Various types of heat-exchangers can be used for exhaust heat recovery, depending on the application and required efficiency. Some common types include:

- **Shell and Tube Heat-Exchangers:** Suitable for high pressures and temperatures, often used in industrial applications.
- **Plate Heat-Exchangers:** Compact and efficient for liquid-to-liquid or liquid-to-gas applications, though less common directly in high-temperature exhaust streams without pre-cooling.
- **Fin-and-Tube Heat-Exchangers:** Often used for transferring heat between gas (exhaust) and air or liquid, providing a large surface area for heat transfer.
- **Regenerative Heat-Exchangers:** The hot and cold fluids flow alternately through the same space, which stores and releases heat.

The specific design chosen depends on factors like the temperature of the exhaust gases, the desired amount of heat recovery, the working fluid being heated, pressure constraints, and cost.

118. Answer: d

Explanation:

Depth Micrometer Use for Bore Measurement

A depth micrometer is a precision measuring instrument specifically designed to accurately measure the depth of holes, slots, recesses, and similar features. It works by having a base that rests on the surface surrounding the feature being measured, and a spindle that extends perpendicularly from this base. As the thimble is turned, the spindle moves in or out, allowing for precise depth readings referenced against the sleeve's scale.

Understanding Depth Measurement Capabilities

The question asks what a depth micrometer can be used to measure, specifically concerning a bore. Let's examine the options provided:

- **Diameter of bore:** Measuring the diameter (the distance across the center of the bore) requires tools like internal micrometers or bore gauges. A depth micrometer's design does not allow for measuring this dimension.
- **Thickness of the bore:** This term is somewhat ambiguous, but typically refers to the wall thickness of the material around the bore. Measuring wall thickness involves determining both the outer and inner diameters and calculating the difference,

usually with outside and internal micrometers respectively. A depth micrometer is not used for this.

- **Roundness of the bore:** Roundness relates to how perfectly circular the bore is. Measuring deviations from a true circle requires specialized equipment like roundness testers or comparators. A depth micrometer measures linear depth, not geometric form.
- **Length of the bore:** In the context of a bore, its "length" is synonymous with its depth – the measurement from the surface down to the bottom of the hole. This is precisely what a depth micrometer is designed to measure accurately.

Comparing Tools for Bore Dimensions

Different measuring instruments serve distinct purposes in engineering and machining. The following table highlights common measurements related to bores and the appropriate tools:

Measurement Needed	Typical Tool Used	Depth Micrometer Suitability
Depth of a Bore (Length)	Depth Micrometer, Depth Gauge	Suitable
Internal Diameter of Bore	Internal Micrometer, Bore Gauge, Vernier Caliper	Not Suitable
External Diameter	Outside Micrometer, Vernier Caliper	Not Suitable
Wall Thickness	Requires Internal and External Diameter Measurements	Not Suitable
Roundness / Circularity	Roundness Tester, Gauge Pins	Not Suitable

Final Confirmation: Bore Length Measurement

Based on its mechanical design and intended application, a depth micrometer excels at measuring the distance from a surface to the bottom of a feature. Therefore, it is the appropriate tool for measuring the length, or depth, of a bore.

119. Answer: a

Explanation:

Understanding Clearance Measurement Between Mating Parts

When parts are designed to fit together, there is often a small gap or space left between them. This gap is called clearance. Measuring this clearance accurately is important for the proper function, assembly, and maintenance of machines and components. Different tools are used depending on the type and size of the clearance or dimension being measured.

Analyzing Tools for Measuring Clearance

Let's look at the options provided and understand what each tool is typically used for:

1. **Feeler gauge:** A feeler gauge is a tool used to measure narrow gaps or clearances between two parts. It consists of a set of thin blades of precisely measured thickness, often hinged together in a case like a pocket knife. To measure a clearance, blades of different thicknesses are inserted into the gap until a blade that fits snugly is found. The thickness marked on that blade represents the clearance. This makes the feeler gauge ideal for measuring the clearance between mating parts, such as valve clearances in an engine or bearing clearances.
2. **"Go" gauge:** A "Go" gauge is a type of limit gauge used for checking whether a part's dimension is within the upper limit of tolerance. If the "Go" gauge fits into or over the part, it means the dimension is not too small (for internal features) or not too large (for external features). It doesn't measure the actual dimension or the clearance, but rather checks if a part is acceptable based on a size limit.
3. **Dial gauge:** A dial gauge (also known as a dial indicator) is a precision measuring instrument used to measure small linear distances or deflections. It typically has a pointer that moves around a dial, amplifying the small movement of a spindle. Dial gauges are often used to check for runout, flatness, parallelism, or relative movement between surfaces. While it can measure small distances, it's not the standard tool for directly measuring a static clearance or gap between two fixed mating parts in the way a feeler gauge is used.

4. **Caliper gauge:** The term "caliper gauge" can be broad, but calipers in general (like vernier calipers or digital calipers) are used to measure the external or internal dimensions of a part, or the distance between two surfaces. While calipers measure dimensions, they are not specifically designed or effective for measuring the narrow gap or clearance *between* two already assembled mating parts. They measure the size of the parts themselves.

Conclusion on Measuring Clearance

Based on the function of each tool, the feeler gauge is specifically designed for and commonly used to measure the clearance or gap between mating parts. Its set of blades allows for the precise determination of narrow space thickness.

Therefore, the clearance between mating parts is measured by a feeler gauge.

Tool Usage Summary for Clearance Measurement

Tool	Primary Use	Measures Clearance Directly?
Feeler gauge	Measuring narrow gaps/clearances	Yes
"Go" gauge	Checking upper size limit (tolerance)	No
Dial gauge	Measuring small displacements, runout, flatness, etc.	Indirectly/Not standard for static gap
Caliper gauge	Measuring external/internal dimensions	No (not for assembled clearance)

Revision Table: Understanding Measuring Instruments

Reviewing different measuring instruments helps solidify understanding.

- **Feeler Gauge:** Measures small clearances/gaps using blades of known thickness.
- **Limit Gauges ("Go"/"No-Go"):** Check if a dimension is within tolerance limits, not the actual size.
- **Dial Indicator:** Measures small linear displacements, variations in surface profile.
- **Calipers (Vernier, Digital):** Measure external, internal dimensions, depth, step.

- **Micrometer:** Measures external, internal dimensions with higher precision than calipers.

Additional Information on Mating Parts and Tolerances

The concept of clearance is closely related to fits and tolerances in engineering. When designing mating parts, engineers specify tolerances, which are the permissible variations in dimensions. The relationship between the dimensions and tolerances of mating parts determines the type of fit:

- **Clearance Fit:** Always results in a gap (clearance) between the mating parts.
- **Interference Fit:** Always results in the parts interfering, requiring force to assemble.
- **Transition Fit:** May result in either a small clearance or a small interference.

Measuring the actual clearance after manufacturing is crucial to ensure the intended fit is achieved and that the assembly will function correctly, preventing issues like excessive wear or seizing.

120. Answer: b

Explanation:

Understanding the Purpose of a Diesel Engine Injector

A diesel engine operates differently from a petrol (gasoline) engine, particularly in how fuel is introduced into the combustion chamber. In a diesel engine, air is first compressed to a very high temperature and pressure. Fuel is then injected into this hot, compressed air, which causes it to auto-ignite and burn, producing power.

The component responsible for injecting the fuel into the cylinder at the precise moment and in the correct amount and pattern is called the ****fuel injector****.

Analyzing the Role of the Injector

Let's look at the options provided and determine the primary purpose of the injector in a diesel engine:

- Option 1: Lubricant in the bearings

Injecting lubricant into bearings is the function of the engine's lubrication system, which includes components like the oil pump, oil lines, and various lubrication points. The fuel injector plays no role in lubricating bearings.

- **Option 2: Fuel in the cylinders**

This option correctly describes the main function of a fuel injector in a diesel engine. The injector atomizes the diesel fuel and sprays it directly into the combustion chamber (cylinder) where the air is already compressed and hot. This injection must happen at a very high pressure to ensure proper mixing and combustion.

- **Option 3: Lubricant in the cylinder**

Lubricating the cylinder walls is also part of the engine's lubrication system, typically achieved through oil splash from the crankshaft or oil delivered to the cylinder walls. Injectors are not used for this purpose; they are specifically designed for fuel delivery.

- **Option 4: Coolant in the cylinder jacket**

Introducing coolant into the cylinder jacket is the function of the engine's cooling system, which uses a water pump, radiator, thermostat, and coolant passages (jackets) around the cylinders to regulate engine temperature. The fuel injector is not involved in the cooling process.

Based on the fundamental operation of a diesel engine, the sole purpose of the fuel injector is to precisely deliver fuel into the cylinder for combustion.

Conclusion on Diesel Injector Function

The core function of the injector in a diesel engine is the high-pressure injection of fuel directly into the cylinder's combustion chamber at the appropriate time during the engine cycle. This injection process is critical for efficient combustion and engine performance.

Engine System	Component	Primary Function
Fuel System	Injector	Injects fuel into cylinder
Lubrication System	Oil Pump, Passages	Lubricates bearings, cylinder walls, etc.
Cooling System	Water Pump, Radiator, Jacket	Circulates coolant around cylinders

Revision Table: Diesel Engine Components

Component	Purpose
Piston	Compresses air/fuel mixture, transmits power
Cylinder	Confines combustion, guides piston
Crankshaft	Converts linear piston motion to rotary motion
Valves (Intake/Exhaust)	Control flow of air into and exhaust out of cylinder
Fuel Pump (High Pressure)	Supplies high-pressure fuel to injector
Injector	Injects fuel into cylinder

Your Personal Exams Guide

Additional Information on Diesel Fuel Injection

Modern diesel engines use sophisticated fuel injection systems, often electronically controlled, such as Common Rail Direct Injection (CRDI) or Unit Injector systems. These systems allow for very precise control over the timing, pressure, and pattern of fuel injection, which significantly impacts engine efficiency, power output, and emissions.

Key aspects of diesel injection:

- High injection pressure (hundreds or even thousands of bar) to atomize fuel finely.
- Injection timing is crucial and synchronized with piston position.
- The injection spray pattern is designed to mix effectively with the hot, compressed air.

121. Answer: a

Explanation:

Understanding Indicated Work and Mean Effective Pressure

The question asks about the formula for calculating the indicated work per cycle in an engine. Indicated work represents the total work done by the working fluid (like combustion gases) inside the cylinder during one complete cycle. It's calculated based on the pressure and volume changes within the cylinder.

Defining Indicated Work Per Cycle

In thermodynamics, the work done by a system during a process involving pressure and volume change is generally given by the integral of pressure with respect to volume ($\int P dV$). For a complete cycle in an engine cylinder, the net work done is the area enclosed by the pressure-volume (P-V) diagram.

The indicated work per cycle is a key performance parameter for internal combustion engines. It tells us how much energy is converted from the combustion process into mechanical work delivered to the piston.

What is Mean Effective Pressure (MEP)?

Mean effective pressure (MEP) is a concept used to represent the average pressure acting on the piston during the power stroke over the entire displacement volume. It's a hypothetical constant pressure that would produce the same net work per cycle if it acted on the piston during the expansion stroke and there was zero pressure during the compression stroke. It is calculated as:

$$\text{MEP} = \frac{\text{Indicated Work Per Cycle}}{\text{Stroke Volume}}$$

MEP is useful because it allows comparison of engines of different sizes. A higher MEP generally indicates a more efficient use of the cylinder volume to produce work.

Calculating Indicated Work Using MEP and Stroke Volume

From the definition of Mean Effective Pressure, we can rearrange the formula to find the Indicated Work Per Cycle:

$$\text{Indicated Work Per Cycle} = \text{Mean Effective Pressure} \times \text{Stroke Volume}$$

The stroke volume (also known as displacement volume) is the volume swept by the piston as it moves from Top Dead Centre (TDC) to Bottom Dead Centre (BDC). It is calculated as:

$$\text{Stroke Volume} = \text{Area of Piston} \times \text{Stroke Length}$$

$$\text{Stroke Volume} = \left(\frac{\pi}{4} \times \text{Bore Diameter}^2 \right) \times \text{Stroke Length}$$

Therefore, the indicated work per cycle is directly proportional to both the Mean Effective Pressure and the Stroke Volume.

Analyzing the Options

Let's look at the given options:

- Option 1: Mean effective pressure x stroke volume
- Option 2: Mean effective pressure x bore diameter
- Option 3: Mean effective pressure x bore arc
- Option 4: Mean effective pressure x stroke length

Based on our derivation from the definition of Mean Effective Pressure, the indicated work per cycle is calculated by multiplying the Mean Effective Pressure by the Stroke Volume. Options 2, 3, and 4 use geometric parameters (bore diameter, bore arc, stroke length) individually instead of the stroke volume, which represents the actual volume displaced by the piston over the power stroke. Therefore, these options are incorrect.

The correct relationship for indicated work per cycle is indeed Mean Effective Pressure multiplied by Stroke Volume.

Conclusion on Indicated Work Calculation

The formula for indicated work per cycle is a fundamental concept in engine performance analysis. It links the average effective pressure acting on the piston (Mean Effective Pressure) to the volume over which this pressure acts (Stroke Volume).

Term	Definition	Formula Relevance
Indicated Work Per Cycle	Net work done by the gas on the piston in one cycle.	The value being calculated.
Mean Effective Pressure (MEP)	Hypothetical constant pressure producing same net work over stroke volume.	A key factor in the formula.
Stroke Volume	Volume displaced by the piston from TDC to BDC.	The volume over which MEP acts to produce work.

Revision Table: Engine Performance Metrics

Metric	Description	Related Concepts
Indicated Work	Work done by gas on piston. Calculated from cylinder pressure.	P-V Diagram, Net Work Per Cycle
Brake Work	Work delivered by the engine crankshaft. Always less than indicated work due to friction.	Mechanical Efficiency, Friction Work
Mean Effective Pressure (MEP)	Average effective pressure over stroke volume. Used for comparing engine performance.	Indicated Work, Stroke Volume, Engine Capacity
Thermal Efficiency	Ratio of work output to heat input.	Combustion, Fuel Energy
Mechanical Efficiency	Ratio of brake work to indicated work. Accounts for friction.	Friction Losses

Additional Information: Factors Affecting Indicated Work

The indicated work per cycle is influenced by several factors:

- **Combustion Process:** The amount of heat released during combustion and the timing of combustion significantly impact the cylinder pressure and hence the Mean Effective Pressure.

- **Engine Design:** Factors like compression ratio, valve timing, intake and exhaust systems affect how effectively the engine can fill and empty the cylinder and utilize the fuel, influencing MEP and thus indicated work.
- **Operating Conditions:** Engine speed, load, fuel quality, and ambient conditions (temperature, pressure) all affect the combustion efficiency and the mass of air-fuel mixture entering the cylinder, impacting indicated work.
- **Losses:** Blowdown losses (pressure escaping during exhaust valve opening) and pumping losses (work required to push exhaust gases out and pull fresh mixture in) reduce the net indicated work, although MEP calculation usually accounts for these by using the net area of the P-V diagram.

Understanding indicated work is crucial for optimizing engine performance and efficiency.

122. Answer: c

Explanation:

Air Filtration for High-Speed Diesel Engines

High-speed diesel engines, commonly found in vehicles and various machinery, require clean air for efficient combustion and longevity. Dust and other airborne particles entering the engine can cause significant wear on internal components, particularly the cylinders, pistons, and fuel injection system. Therefore, effective air filtration is crucial for these engines.

Understanding Engine Air Filter Types

Several types of air filters are used in internal combustion engines, differing in their filtration mechanism and suitability for various operating conditions. Let's look at the types mentioned in the options:

- **Dry filters:** These filters use pleated paper or synthetic media to trap particles as air passes through. They are very common in modern vehicles due to high filtration efficiency, relatively low maintenance (easy replacement), and compact size.
- **Heaters:** This option refers to devices used to heat air or fuel, often to aid cold starting or improve combustion efficiency. Heaters are not air filters; their function is related to temperature control, not particle removal from the air intake.

- **Oil bath filters:** In this type, the incoming air is directed downwards over a pool of oil. Larger particles are trapped in the oil, and the air then passes through a filter mesh or element wetted by the oil mist, which captures finer particles. These filters were historically popular, especially in heavy-duty or dusty environments.
- **Impingement filters:** These filters work on the principle of inertia, where particles are collected as they impact (impinge upon) an oiled surface or a series of baffles. While used in some industrial applications, they are less common for primary engine air intake filtration compared to dry or oil bath types.

Why Oil Bath Filters for High-Speed Diesel?

Historically, oil bath filters were widely used in diesel engines, particularly in applications involving dusty conditions like agriculture, construction, and some older vehicles. Their design allows them to handle large amounts of dust before maintenance is required, as dirt settles in the oil reservoir. This made them robust and suitable for the demanding environments often encountered by high-speed diesel engines operating off-road or in heavy-duty scenarios. While dry filters are now very prevalent and offer high efficiency, oil bath filters were traditionally considered a standard for certain high-speed diesel applications due to their durability and dust-holding capacity.

Considering the general types historically or typically used in challenging conditions associated with many diesel engine applications, oil bath filters fit the description.

Filter Type	Mechanism	Common Use in Engines
Dry Filter	Mechanical trapping by media (paper/synthetic)	Very common in modern engines (gas and diesel)
Oil Bath Filter	Air passes over/through oil pool and oiled element	Historically common, especially in heavy-duty/dusty diesel applications
Impingement Filter	Particles hit oiled surfaces/baffles	Less common for primary engine air intake
Heater	Heats air or fuel	Not an air filter; aids starting/combustion

Revision Table: Diesel Engine Air Filters

To summarize the types of air filters and their relevance to high-speed diesel engines:

- Dry filters are widely used and highly effective today.
- Oil bath filters were a common choice, particularly where dealing with significant dust was a primary concern.
- Impingement filters are not typically the main air filter type for engines.
- Heaters are not air filters.

Based on the common historical and ongoing use in certain heavy-duty applications where diesel engines operate, oil bath filters have been a general type used.

Additional Information on Engine Air Filters

The choice of air filter type depends on several factors, including the operating environment (dust level), required filtration efficiency, space constraints, cost, and maintenance requirements. Filter efficiency is often rated based on the size and percentage of particles they can remove. Proper maintenance, such as regular inspection and cleaning or replacement, is essential for any air filter type to ensure it continues to protect the engine effectively and maintain optimal performance and fuel efficiency.

123. Answer: d

Explanation:

Understanding Cyaniding and Nitriding Heat Treatments

Heat treatment processes are crucial in materials science and engineering to alter the physical, mechanical, and sometimes chemical properties of a material, usually a metal. These processes involve heating the material to a specific temperature, holding it there, and then cooling it at a controlled rate. Cyaniding and nitriding are two specific types of heat treatments.

What are Cyaniding and Nitriding?

Both cyaniding and nitriding are surface hardening techniques. They involve introducing elements into the surface layer of a material to increase its hardness, wear resistance, and fatigue life without affecting the core properties significantly.

- **Cyaniding:** This process involves heating the steel in a cyanide bath. It adds both carbon and nitrogen to the surface layer, forming a hard case. The temperature used is typically between 850°C and 950°C. It's relatively fast compared to other methods but involves toxic cyanide salts.
- **Nitriding:** This process involves heating the steel in a nitrogen-rich atmosphere, usually ammonia gas (NH_3). Nitrogen atoms diffuse into the surface, forming hard nitride compounds. Nitriding is typically carried out at lower temperatures, between 490°C and 560°C, and takes longer than cyaniding. It results in a very hard case with excellent wear and fatigue resistance and less distortion compared to higher-temperature processes.

What is Case Hardening?

Case hardening, also known as surface hardening, is a heat treatment process that results in a hard outer layer (the "case") on a metal component while leaving the core relatively soft and tough. This combination of a hard surface and a tough core is desirable for many applications where wear resistance and impact strength are required simultaneously, such as gears, shafts, and bearings.

The goal of case hardening is to enrich the surface layer with elements that can form hard phases upon subsequent heat treatment (like quenching). Common elements used are carbon (in carburizing) and nitrogen (in nitriding), or a combination of both (in cyaniding or carbonitriding).

Since both cyaniding and nitriding involve adding elements (carbon and nitrogen, or nitrogen alone) to the surface of a metal to harden it, they are fundamentally methods of achieving case hardening.

Comparison with Other Heat Treatments

Let's briefly look at the other options provided to understand why they are not the correct answer:

- **Normalising:** This heat treatment is used to refine the grain structure, improve uniformity, and relieve internal stresses in a material. It involves heating the metal to a suitable temperature above its critical range and cooling it in air. It is a bulk heat treatment, not specifically for surface hardening.
- **Hardening (Bulk Hardening):** This process involves heating the steel to a suitable temperature above its critical range and then rapidly cooling it (quenching) in a

medium like water, oil, or polymer solution. This transforms the structure into martensite, making the entire part hard and brittle. This hardens the bulk of the material, not just the surface (though the cooling rate might lead to variations).

- **Tempering:** Tempering is a heat treatment applied after hardening (bulk hardening). It involves reheating the hardened steel to a temperature below its critical range and then cooling it slowly. This reduces the brittleness induced by hardening and improves toughness and ductility, often at the cost of some hardness.

Based on the definitions, cyaniding and nitriding are processes specifically designed to harden the surface or "case" of a material, which is the definition of case hardening.

Process	Primary Goal	Area Affected	Elements Added to Surface
Normalising	Refine grain, relieve stress	Bulk	None
Hardening (Bulk)	Increase overall hardness	Bulk	None (structure change)
Tempering	Improve toughness/ductility (after hardening)	Bulk	None
Cyaniding	Increase surface hardness	Surface (Case)	Carbon & Nitrogen
Nitriding	Increase surface hardness	Surface (Case)	Nitrogen
Case Hardening	Increase surface hardness while keeping core tough	Surface (Case)	Carbon, Nitrogen, etc. (Method dependent)

Conclusion on Cyaniding and Nitriding

Cyaniding and nitriding are processes where the surface of a metal is chemically altered by adding carbon and nitrogen (cyaniding) or just nitrogen (nitriding) at elevated temperatures. This leads to the formation of a very hard layer on the surface. This is precisely what is meant by case hardening.

Revision Table: Heat Treatment Methods Summary

Heat Treatment Type	Purpose	Example Processes
Bulk Treatments	Alter properties throughout the material volume	Normalising, Annealing, Hardening, Tempering
Surface Treatments (Case Hardening)	Alter properties only on the surface layer	Carburizing, Nitriding, Carbonitriding, Cyaniding, Induction/Flame Hardening

Additional Information on Case Hardening Methods

Beyond cyaniding and nitriding, other common case hardening methods include:

- **Carburizing:** Adds carbon to the surface of low-carbon steel. Can be gas carburizing, pack carburizing, or liquid carburizing. Requires higher temperatures than nitriding.
- **Carbonitriding:** Similar to carburizing but adds both carbon and nitrogen. Often done at lower temperatures than carburizing and results in a shallower case.
- **Induction Hardening and Flame Hardening:** These are surface hardening methods that use rapid heating and quenching of the surface layer, rather than adding elements chemically.

Each case hardening method has its advantages and disadvantages regarding case depth, hardness profile, distortion, processing time, and cost, making the choice dependent on the specific application requirements.

124. Answer: d

Explanation:

Understanding Otto Cycle Efficiency

The Otto cycle is an idealized thermodynamic cycle that describes the functioning of a typical spark ignition piston engine. It is the most common cycle found in automobile engines. The efficiency of this cycle determines how much of the heat energy from burning fuel is converted into useful work.

Factors Affecting Otto Cycle Efficiency

The thermal efficiency (η_{th}) of an ideal Otto cycle is given by the formula:

$$\eta_{th} = 1 - \frac{1}{r^{\gamma-1}}$$

Where:

- r is the compression ratio
- γ is the heat capacity ratio (c_p/c_v)

Let's break down the terms:

- **Compression Ratio (r):** This is the ratio of the volume of the cylinder and combustion chamber when the piston is at its lowest point (bottom dead center, BDC) to the volume when the piston is at its highest point (top dead center, TDC). A higher compression ratio means the air-fuel mixture is compressed into a smaller volume before ignition.
- **Heat Capacity Ratio (γ):** This is a property of the working fluid (air in an ideal cycle). For air, $\gamma \approx 1.4$.

Analyzing the Impact of Compression Ratio on Efficiency

Looking at the formula, $\eta_{th} = 1 - \frac{1}{r^{\gamma-1}}$, we can see how the compression ratio (r) directly affects the thermal efficiency (η_{th}).

Since $\gamma > 1$, the term $r^{\gamma-1}$ increases as the compression ratio (r) increases. Consequently, the term $\frac{1}{r^{\gamma-1}}$ decreases as r increases. As this subtractive term gets smaller, the overall efficiency η_{th} increases.

Therefore, increasing the compression ratio is a primary way to increase the thermal efficiency of an ideal Otto cycle.

Examining Other Options

- **Fuel intake:** While the amount of fuel intake affects the power output of the engine (more fuel means more energy released), it does not directly change the fundamental thermodynamic efficiency of the cycle itself.
- **Stroke length:** Stroke length, along with bore size, determines the engine's displacement (cylinder volume). While it is related to the volumes used in the

compression ratio calculation, simply increasing stroke length without changing other parameters might not necessarily increase the *ratio* in a way that optimizes efficiency. The compression ratio is the critical factor, not just the absolute volumes.

- **Cylinder size:** Similar to stroke length, cylinder size (bore and stroke) determines the engine's displacement. Increasing cylinder size generally increases the engine's power output by allowing more fuel and air to be burned per cycle, but it doesn't inherently change the *efficiency* of the thermodynamic cycle unless it allows for a change in the achievable compression ratio or reduces heat losses relative to work done (which is a real-world effect not captured by the ideal cycle).

Based on the ideal Otto cycle theory and its efficiency formula, increasing the compression ratio is the most significant factor among the given options for increasing thermal efficiency.

Conclusion on Otto Cycle Improvement

The thermal efficiency of an Otto cycle is fundamentally linked to its compression ratio. A higher compression ratio leads to a higher theoretical maximum efficiency.

Impact of Otto Cycle Parameters		
Parameter	Primary Effect on Engine	Effect on Ideal Otto Cycle Efficiency
Compression Ratio	Increases pressure & temperature at the end of compression	Increases significantly ($\eta_{th} = 1 - 1/r^{\gamma-1}$)
Fuel Intake Amount	Increases power output (energy per cycle)	Negligible on ideal η_{th}
Stroke Length / Cylinder Size	Increases engine displacement, potentially power output	Indirect effect on ideal η_{th} ; relevant as parts of compression ratio calculation

Therefore, higher efficiency can be achieved of an OTTO cycle by increasing the compression ratio.

Revision Table: Otto Cycle Basics

Key Stages of the Ideal Otto Cycle

Process	Description	Thermodynamic Change
1-2: Isentropic Compression	Piston moves from BDC to TDC, compressing air-fuel mixture.	Volume decreases, Pressure increases, Temperature increases. No heat transfer.
2-3: Isochoric Heat Addition	Spark plug ignites mixture at TDC. Volume constant.	Volume constant, Pressure increases sharply, Temperature increases sharply. Heat is added.
3-4: Isentropic Expansion	Hot gases push piston from TDC to BDC (Power Stroke).	Volume increases, Pressure decreases, Temperature decreases. No heat transfer.
4-1: Isochoric Heat Rejection	Exhaust valve opens at BDC. Volume constant.	Volume constant, Pressure decreases sharply, Temperature decreases sharply. Heat is rejected.

Additional Information: Practical Limitations on Compression Ratio

While higher compression ratio theoretically increases Otto cycle efficiency, practical engines face limitations:

- **Knocking/Detonation:** High compression ratios can cause the air-fuel mixture to auto-ignite before the spark occurs, leading to engine knock and potential damage. This requires higher octane fuel.
- **Increased Stresses:** Higher pressures and temperatures put more stress on engine components.
- **Heat Losses:** As temperature increases with higher compression, heat losses to the cylinder walls also increase, reducing actual efficiency compared to the ideal cycle.

Engine designers balance the benefits of higher compression ratio with these practical constraints.

125. Answer: a

Explanation:

Understanding Piston Rings and Materials

Piston rings are critical components in internal combustion engines. They are fitted into grooves on the outer surface of a piston and perform several vital functions:

- Sealing the combustion chamber to prevent gases from escaping into the crankcase, maintaining compression.
- Transferring heat from the piston to the cylinder wall.
- Regulating engine oil consumption by scraping excess oil from the cylinder walls.
- Supporting the piston in the cylinder bore.

Due to these demanding roles, the material used for piston rings must possess specific properties, including wear resistance, elasticity, strength at high temperatures, and the ability to conform to the cylinder bore.

Why Cast Iron is Commonly Used for Piston Rings

Cast iron, specifically gray cast iron or nodular cast iron, is the most common material for piston rings because it exhibits an excellent balance of the required properties. Here's a look at its key advantages:

- **Wear Resistance:** Cast iron contains free graphite. This graphite acts as a solid lubricant, significantly reducing friction and wear between the ring and the cylinder liner. This is crucial for long engine life.
- **Elasticity and Conformity:** The material must be elastic enough to spring outwards and maintain contact pressure against the cylinder wall. Cast iron provides the necessary elasticity to conform to the shape of the bore, even if it's slightly worn or distorted.
- **Strength and Durability:** It has sufficient strength to withstand the high pressures and temperatures within the combustion chamber and the mechanical stresses from piston movement.
- **Thermal Stability:** Cast iron maintains its properties well under the varying high temperatures encountered in an engine cylinder.
- **Cost-Effectiveness:** Compared to many other engineering alloys, cast iron is relatively inexpensive to produce and process.

Exploring Other Potential Materials

Let's consider why the other options provided are generally less suitable for standard piston ring applications compared to cast iron:

- **Brass:** Brass is a copper alloy. While it has good corrosion resistance and machinability, it typically lacks the high-temperature strength, hardness, and wear resistance required for piston rings operating under extreme conditions of heat, pressure, and friction against a cylinder liner.
- **Aluminium:** Aluminium alloys are lightweight and have good thermal conductivity. However, they have a much lower melting point and lower strength at high temperatures compared to cast iron. Aluminium piston rings would suffer from excessive wear, deformation, and potential seizure under engine operating conditions. Its thermal expansion is also much higher, making clearance control difficult.
- **Copper:** Copper is a relatively soft metal with good thermal conductivity but poor wear resistance and low strength at the high temperatures found in the combustion zone. Like brass, it is not durable enough to withstand the continuous sliding contact and forces experienced by piston rings.

Based on the functional requirements and material properties, cast iron stands out as the preferred material for manufacturing piston rings in most applications.

Property	Cast Iron (Typical)	Brass (Typical)	Aluminium (Typical)	Copper (Typical)
Wear Resistance	Excellent (due to graphite)	Fair	Poor	Poor
High-Temperature Strength	Good	Fair/Poor	Poor	Poor
Elasticity for Ring Tension	Good	Fair	Poor	Fair
Thermal Expansion	Moderate	Moderate	High	High
Cost	Moderate	High	Moderate	Very High

Revision Table: Key Facts About Piston Rings Material

Here's a quick summary of why cast iron is the standard choice for piston rings:

- Primary material: Cast iron (Gray or Nodular).
- Key property 1: High wear resistance (due to free graphite).
- Key property 2: Good elasticity and conformability.
- Key property 3: Strength and stability at high temperatures.
- Alternative materials like Brass, Aluminium, Copper are generally unsuitable due to lower strength, wear resistance, or thermal limitations.

Additional Information: Piston Ring Types and Coatings

While cast iron is common, piston rings can be made from or coated with other materials for specific applications, especially in high-performance engines. Some examples include:

- **Steel:** Used for some modern piston rings, often with coatings. Offers higher strength.
- **Coatings:** Materials like chromium plating, molybdenum (moly) spray, or Ceramic-Metallic (Cermet) coatings are applied to the surface of cast iron or steel rings to further enhance wear resistance, reduce friction, and improve durability, particularly on the top compression ring which faces the harshest conditions.
- **Types of Rings:** Engines use different types of rings – compression rings (usually top two) for sealing the combustion chamber and oil control rings (usually bottom) for managing lubrication. The design and material requirements can vary slightly between ring types.

Your Personal Exams Guide

126. Answer: d

Explanation:

Understanding Cutting Tools for Threads

Cutting threads is a fundamental process in manufacturing and metalworking to join parts together using fasteners like bolts and nuts. Threads can be cut either on the outside surface of a rod or bolt (external threads) or on the inside surface of a hole (internal threads). Different specialized cutting tools are used for these specific tasks.

Identifying the Tool for Outside Threads

The question asks for the cutting tool used specifically for cutting **outside threads**. Let's look at the options provided and understand what each tool is used for:

- **Reamer:** A reamer is a rotary cutting tool used in metalworking to enlarge and finish a hole to a precise size and shape. It removes only a small amount of material and is not used for cutting threads.
- **Tap:** A tap is a cutting tool used to create **internal threads** within a hole. This process is called tapping. Taps are typically used to prepare holes for bolts or screws.
- **Drill:** A drill is a cutting tool used to create cylindrical holes in solid materials. It removes material from a solid block to make a hole and is the first step before tapping or reaming, but it does not cut threads.
- **Die:** A die is a cutting tool used to create **external threads** on a cylindrical rod or bar. This process is called threading. Dies are typically used to make bolts or threaded rods.

Based on the function of each tool, the tool specifically designed for cutting **outside threads** is a **die**.

Comparison of Threading Tools

Here's a simple comparison of the primary tools used for cutting threads:

Tool	Purpose	Type of Thread Cut	Workpiece
Tap	Cut threads	Internal (in a hole)	Hole
Die	Cut threads	External (on a rod)	Rod or Bolt

Therefore, a cutting tool used to cut outside threads is called a die.

Revision Table: Threading Tools Explained

Let's review the main tools used in threading and hole finishing:

- **Die:** Cuts external threads on rods/bolts.
- **Tap:** Cuts internal threads in holes.
- **Drill:** Makes holes (prepares for tapping or reaming).
- **Reamer:** Finishes and enlarges existing holes to precise dimensions.

Understanding the difference between these tools is crucial for various machining operations.

Additional Information: Types of Dies and Taps

Both dies and taps come in various forms depending on the material, thread size, and application. For example:

- **Dies:** Include split dies (adjustable), solid dies, and die nuts. They can be used with a die handle or in a machine.
- **Taps:** Include hand taps (taper, plug, bottoming), machine taps, and forming taps. They are used with tap wrenches or in tapping machines.

The choice of tool depends on the specific threading requirement and the type of material being worked on. The basic principle, however, remains the same: dies cut outside threads, and taps cut inside threads.

127. Answer: a

Explanation:

Understanding Water Circulation in Thermo-siphon Cooling Systems

Thermo-siphon cooling systems are a simple and reliable method used in various applications, particularly in older engines or certain solar water heating systems, to circulate cooling fluid without the need for a mechanical pump. The fundamental principle behind the water circulation in such systems is a natural phenomenon driven by changes in the properties of the cooling fluid, specifically water.

How Thermo-siphon Cooling Works

The operation of a thermo-siphon system relies on the natural convection currents that occur when a fluid is heated. Here's a breakdown:

- The cooling water surrounds the hot parts of the engine (like cylinders).
- As the water absorbs heat from the engine, its temperature increases.

- When water heats up, its density decreases; it becomes lighter per unit volume compared to cooler water.
- The hotter, less dense water rises to the top of the cooling passages and flows into the radiator.
- In the radiator, the water is cooled by the surrounding air.
- As the water cools down, its density increases; it becomes heavier.
- The cooler, denser water sinks to the bottom of the radiator and flows back into the engine block to absorb more heat.
- This continuous cycle of heating, rising, cooling, and sinking creates a natural water circulation path, driven solely by the difference in water density caused by temperature variations.

Analyzing the Options for Water Circulation

Let's look at the provided options in the context of a thermo-siphon cooling system:

Option 1: The change in density of the water

This option directly describes the driving force in a thermo-siphon system. As explained above, the temperature difference between the water absorbing heat from the engine and the water cooling in the radiator leads to a difference in density. This density difference creates a pressure gradient that causes the water to circulate naturally.

Option 2: A belt driven water impeller

A belt-driven water impeller is a component of a mechanical water pump, commonly found in pump-driven cooling systems, not thermo-siphon systems. An impeller actively pushes the coolant through the system, which is different from the natural circulation in a thermo-siphon.

Option 3: Conduction currents

Conduction is a mode of heat transfer where heat energy is transmitted through direct contact of particles. While conduction is involved in the initial transfer of heat from the engine to the water, it does not cause the bulk movement or circulation of the water itself. Water circulation is caused by convection, which is the movement of fluid due to density differences resulting from temperature variations.

Option 4: A gear driven water pump

Similar to a belt-driven impeller, a gear-driven water pump is a mechanical device that forces water circulation. Thermo-siphon systems are specifically designed to avoid the need for such mechanical pumps, relying instead on natural convection.

Based on the analysis, the water circulation in a thermo-siphon cooling system is caused by the natural convection process driven by the change in density of the water as it heats and cools.

Revision Table: Thermo-siphon vs. Pumped Cooling

Feature	Thermo-siphon System	Pumped System
Circulation Mechanism	Natural convection (density change)	Mechanical pump (impeller)
Driving Force	Density difference due to temperature	Mechanical force from pump
Complexity	Simpler (fewer parts)	More complex (includes pump, belts/gears)
Efficiency at High Load	Less efficient (circulation rate is self-regulated)	More efficient (circulation rate is controlled)
Maintenance	Generally lower maintenance (no pump issues)	Requires pump maintenance

Additional Information on Natural Convection and Cooling Systems

Natural convection is a fundamental principle in heat transfer and fluid dynamics. It's the basis for many natural phenomena, like weather patterns, and engineered systems, such as the thermo-siphon cooling system. Understanding how density changes with temperature is key to grasping how these systems operate passively.

In contrast to natural convection, forced convection uses external means, like a pump or fan, to induce fluid movement. Pumped cooling systems use forced convection to ensure adequate coolant flow even at low engine speeds or high heat loads, where natural convection alone might not be sufficient. However, the simplicity and reliability of thermo-

siphon systems made them suitable for earlier, lower-performance engines or specific applications where simplicity was prioritized over maximum cooling efficiency.

128. Answer: a

Explanation:

Ignition Condenser Function Explained

The **ignition condenser** is a vital component in traditional ignition systems, often found in older vehicles. Its main purpose relates to managing the electrical circuit involving the contact breaker points and the ignition coil.

Understanding the Role of the Ignition Condenser

The core function of the **ignition condenser** is to manage the electrical energy generated when the contact breaker points open. Let's break down how it works:

- **Point Operation:** The contact breaker points are mechanical switches that control the flow of current to the ignition coil's primary winding. They open and close in synchronization with the engine's rotation.
- **Coil Energizing:** When the points are closed, current flows through the primary winding, building up a magnetic field within the ignition coil.
- **Field Collapse:** When the points open, the primary circuit is broken. This rapid interruption causes the stored magnetic field in the coil to collapse quickly.
- **Voltage Generation:** The collapsing magnetic field induces a very high voltage in the coil's secondary winding. This high voltage is what eventually creates the spark across the spark plug gap.
- **The Problem of Arcing:** As the points separate, the collapsing magnetic field can create a voltage surge that tries to jump the gap between the points. This jump creates an electrical arc, or spark, directly across the points.
- **Condenser's Solution:** The **ignition condenser** is connected in parallel with the points. When the points open, the condenser provides a path for this voltage surge to dissipate safely. It absorbs the electrical energy that would otherwise cause the arc.

By absorbing this energy, the condenser effectively suppresses the spark that would form between the points. This action is crucial for several reasons:

- It significantly **reduces arcing at contacts** (the breaker points).
- Minimizing arcing prevents rapid burning, pitting, and erosion of the point surfaces, extending their service life.
- It aids in a faster collapse of the magnetic field, which helps produce a stronger, more intense spark at the spark plug for efficient combustion.

Evaluating the Options

Based on this explanation, let's look at why the correct option is the best fit and why others are not:

- **1. Reduces arcing at contacts:** This accurately describes the primary function of the ignition condenser. It absorbs the voltage spike that occurs when the points open, preventing a damaging arc across the point surfaces.
- **2. Protects plugs from load:** While a functioning ignition system is necessary for the spark plugs, the condenser's role is specifically related to the points and the coil's primary circuit, not directly protecting the plugs from electrical load.
- **3. Reduces secondary spark:** This is incorrect. By ensuring a rapid collapse of the coil's magnetic field, the condenser actually helps to **strengthen** the secondary spark delivered to the spark plug.
- **4. Increases contact arcing:** This is the opposite of the condenser's actual function. It is designed specifically to *decrease* or prevent arcing at the points.

Conclusion

The essential role of the **ignition condenser** is to protect the breaker points from the damaging effects of electrical arcing by absorbing the inductive voltage kick when the points open, thereby ensuring a cleaner break and a stronger subsequent spark.

129. Answer: b

Explanation:

Understanding Twist Drill Grooves: Identifying the Flutes

Let's break down the parts of a twist drill and understand what the grooves running along its body are called.

A twist drill is a cutting tool commonly used to create cylindrical holes. It has a specific geometry designed for cutting material and removing chips.

The question asks about the grooves provided on the entire length of the body of a twist drill. These grooves are a very important feature of the drill bit.

Let's look at the functions of these grooves:

- They form the cutting edges (lips) of the drill bit.
- They provide channels for chips (cut material) to escape from the hole during drilling.
- They allow cutting fluid or coolant to reach the cutting edges, reducing heat and friction.

These specific grooves that spiral around the body of the twist drill are known by a particular term.

Analysing the Options

Let's examine the provided options:

1. **Webs:** The web is the thin central section of the drill bit that separates the flutes. It runs the entire length but is not the groove itself.
2. **Flutes:** Flutes are the helical or spiral grooves cut into the body of the drill. They create the cutting edges and allow chip evacuation.
3. **Lips:** The lips are the cutting edges formed at the point of the drill where the flutes meet the web. They are part of the drill point, not the grooves running along the body's length.
4. **Margins:** The margin is the narrow strip along the edge of the flute, slightly larger in diameter than the rest of the flute. It helps guide the drill in the hole and maintains the hole diameter. It's part of the flute structure but not the main groove itself.

Based on the function and definition, the grooves provided on the entire length of the body of a twist drill are called flutes.

Twist Drill Parts

Part	Description	Relation to Grooves
Flutes	Spiral grooves along the body	Are the grooves in question; provide chip evacuation and form cutting edges.
Web	Central section separating flutes	Located between the flutes; not a groove.
Lips	Cutting edges at the point	Formed by the intersection of flutes and web; part of the drill point.
Margins	Narrow strip along flute edge	Part of the flute structure; guide the drill.

Therefore, the correct term for the grooves running along the entire length of the twist drill body is Flutes.

Revision Table: Twist Drill Terminology

Key Twist Drill Terms for Exam Prep

Term	Definition
Body	The main part of the drill bit, from the point to the start of the shank.
Point	The cutting end of the drill bit.
Shank	The part of the drill bit that is held by the drill chuck or holder.
Flutes	Helical grooves along the body that remove chips and form cutting edges.
Web	The central core separating the flutes.
Lips	The main cutting edges at the point.
Margin	Narrow raised section on the flute edge that guides the drill.
Land	The surface of the margin.

Additional Information on Twist Drill Flutes

The shape and number of flutes on a twist drill can vary depending on the material being drilled and the desired outcome. For example:

- Drills for general purpose often have two flutes.
- Drills for softer materials or high feed rates might have wider flutes for better chip evacuation.
- Drills for harder materials might have flutes designed for increased rigidity.

The helix angle of the flutes also influences the drill's performance. A larger helix angle is good for soft materials and efficient chip lifting, while a smaller helix angle provides stronger cutting edges suitable for harder materials.

130. Answer: a

Explanation:

Understanding the Combination Set Tool

A combination set is a versatile measuring and marking tool commonly used in metalworking, woodworking, and mechanical tasks. It typically consists of a steel rule, a square head, a protractor head, and sometimes a center head. These components can be attached to the steel rule, allowing the tool to perform various functions.

Components of a Standard Combination Set

Let's look at the typical parts that make up a standard combination set:

- **Steel Rule:** This is the base of the set, acting as a standard measuring rule. It has graduations along its length for taking linear measurements. The other heads slide along this rule.
- **Square Head:** This head is used for marking or checking 90-degree angles (squares) and 45-degree angles (miters) relative to an edge or surface. It also helps in setting the rule perpendicular to an edge.
- **Protractor Head:** This head allows the user to measure or lay out angles between 0 and 180 degrees. It has a rotating turret with a scale.
- **Center Head:** (Often included, but not always standard in every set) This head is used to locate the center of round or square stock.

Analyzing the Given Options

Now let's examine the options provided and determine which one is not a part of a typical combination set:

1. **Socket spanners:** Socket spanners, also known as socket wrenches, are tools used for tightening or loosening nuts and bolts. They are part of a mechanic's toolkit, not a measuring or marking tool like a combination set.
2. **Steel rule:** As discussed, the steel rule is a fundamental component of the combination set, serving as the base for the other heads.
3. **Square head:** The square head is a key attachment used for 90 and 45-degree measurements and marking within the combination set.
4. **Protractor head:** The protractor head is another essential attachment used for measuring and marking angles with the combination set.

Conclusion

Based on the standard components of a combination set, the steel rule, square head, and protractor head are all integral parts. Socket spanners, however, are completely different tools used for fastening applications and are not included in a combination set.

Revision Table: Combination Set Components

Component	Part of Combination Set?	Typical Function
Socket spanners	No	Tightening/loosening fasteners
Steel rule	Yes	Linear measurement base
Square head	Yes	90° and 45° layout/checking
Protractor head	Yes	Angle measurement/layout

Additional Information: Related Measuring Tools

Understanding different measuring and layout tools is important for various trades. Here are a few other related tools:

- **Calipers:** Used for precise measurement of dimensions, such as external or internal diameters and depths.
- **Micrometers:** Provide even more precise measurements than calipers, often used for checking tolerances.
- **Vernier Height Gauge:** Used for measuring vertical heights accurately and marking layout lines on vertical surfaces or workpieces.
- **Surface Plate:** A flat, rigid surface used as a base for inspection, marking out, and calibration. Tools like height gauges are used on surface plates.

131. Answer: d

Explanation:

Understanding Vernier Caliper Measurement

A vernier caliper is a precision measuring instrument used in various fields, including manufacturing, engineering, and science. It is primarily designed to measure dimensions with high accuracy.

What a Vernier Caliper Measures

A standard vernier caliper can measure several types of dimensions. These include:

- **External dimensions:** Such as the diameter of a rod or the length/width/thickness of an object.
- **Internal dimensions:** Such as the internal diameter of a pipe or hole.
- **Depth:** The depth of holes or recesses.
- **Step:** The distance between two planes at different levels.

The question asks what the vernier caliper is used for measuring. Let's examine the given options in the context of these capabilities:

Analyzing the Options

- **Roundness:** Roundness refers to how closely a shape approximates a perfect circle. While a vernier caliper can measure the diameter at different points, it is not the

primary or most accurate tool for assessing the overall roundness of an object. Specialized instruments like roundness testers are used for this purpose.

- **Shape:** 'Shape' is a very general term. A vernier caliper measures specific dimensions (like length, width, diameter, thickness), but it does not describe or measure the overall complex shape of an object.
- **Flatness:** Flatness refers to how flat a surface is. This is typically measured using tools like surface plates, dial indicators, or optical flats and interferometers. A vernier caliper is not designed to measure the flatness of a surface.
- **Thickness:** Thickness is a specific external dimension of an object, representing the distance between two parallel surfaces. Measuring thickness is one of the fundamental uses of a vernier caliper. You place the object's thickness between the external jaws of the caliper and read the measurement.

Conclusion on Vernier Caliper Use

Based on the capabilities of a vernier caliper, measuring thickness is a direct and common application of this instrument.

Therefore, the vernier caliper is commonly used for measuring thickness.

Revision Table: Vernier Caliper Uses

Measurement Type	Vernier Caliper Used?	Explanation
Thickness	Yes	Measures the external dimension between two surfaces.
Roundness	No (Primary)	Not designed for assessing overall circular perfection.
Shape	No (General)	Measures dimensions, not overall complex form.
Flatness	No	Not designed for surface flatness assessment.

Additional Information on Measurement Tools

While the vernier caliper is versatile, other instruments are used for different types of measurements:

- **Micrometer:** Offers higher precision than a vernier caliper, often used for smaller external, internal, or depth measurements.
- **Dial Indicator:** Used for measuring small distances and variations, often to check for flatness, runout, or comparison measurements.
- **Coordinate Measuring Machine (CMM):** A highly accurate machine used to measure complex shapes and dimensions in three dimensions.

The vernier caliper remains a fundamental tool for precise linear measurements, including thickness.

132. Answer: c

Explanation:

Understanding how to accurately measure the internal diameter of a cylinder bore is crucial in many mechanical and engineering fields, particularly in engine rebuilding and maintenance. The question asks for the specific instrument used for this measurement.

Analyzing Instruments for Cylinder Bore Measurement

Let's look at the options provided and consider their typical uses:

1. **Vernier caliper:** This is a versatile instrument used for measuring external dimensions, internal dimensions, and depth. While a vernier caliper can measure the *outside* diameter of a cylinder or the *depth* of a bore, its jaws might not fit correctly or provide sufficient accuracy for precisely measuring the internal diameter (bore) deep within a cylinder, especially larger ones or those requiring high precision.
2. **Outside micrometer:** An outside micrometer is designed to measure the external dimensions of an object by placing it between an anvil and a spindle. It is not suitable for measuring internal diameters like a cylinder bore.
3. **Bore dial gauge:** A bore dial gauge, also known as a cylinder bore gauge or dial bore gauge, is specifically designed for measuring the internal diameter of bores, cylinders, and holes. It works by transferring the measurement from a set of contact points through a mechanical linkage to a dial indicator, showing the deviation from

a set master size. This allows for very precise measurements of internal diameters, including variations in roundness and taper along the bore.

4. **Telescopic gauge:** A telescopic gauge (also called a telescoping gauge or tele gauge) is a transfer-type measuring tool used to obtain the size of a bore or groove. It has plungers that are extended into the bore and locked. The gauge is then removed, and the distance across the locked plungers is measured using an outside micrometer. While useful for bore measurement, it is not a direct reading instrument; it requires a secondary measurement with a micrometer and is often considered less precise and efficient for cylinder bores compared to a bore dial gauge.

Why Bore Dial Gauge is Ideal for Cylinder Bores

Based on the function and design of these instruments, the bore dial gauge is the most appropriate tool specifically used for accurately measuring the internal diameter of a cylinder bore. It allows for direct measurement and can detect variations like taper and out-of-roundness, which are critical parameters in engine cylinder inspection.

Here's a simple comparison:

Instrument	Primary Use	Suitable for Cylinder Bore?	Type of Measurement
Vernier Caliper	External, Internal, Depth	Limited precision/reach for deep bores	Direct Reading
Outside Micrometer	External Dimensions	No	Direct Reading
Bore Dial Gauge	Internal Bores/Cylinders	Yes (Specifically designed)	Direct Reading (deviation from master)
Telescopic Gauge	Internal Bores/Grooves	Yes (Requires secondary tool)	Transfer (Requires micrometer)

Therefore, the instrument specifically designed and commonly used for measuring cylinder bore diameter with precision is the bore dial gauge.

Revision Table: Measuring Instruments

Instrument	Key Function	Example Use Case
Vernier Caliper	Measures dimensions directly (external, internal, depth).	Measuring the length of a bolt or the outside diameter of a pipe.
Micrometer (Outside)	Measures external dimensions with high precision.	Measuring the thickness of a sheet metal or the diameter of a shaft.
Bore Dial Gauge	Measures internal bore diameters and detects variations.	Checking wear in engine cylinders or hydraulic bores.
Telescopic Gauge	Transfers internal dimension for measurement with micrometer.	Measuring the diameter of a small hole or recess.

Additional Information on Bore Measurement

When measuring a cylinder bore, it's important to take measurements at different depths and orientations (e.g., parallel and perpendicular to the crankshaft axis) to check for taper and out-of-roundness. These measurements are critical for determining if a cylinder requires honing or boring to a larger size during an engine rebuild.

A bore dial gauge is typically used in conjunction with a setting ring or a known master bore to establish a zero point or a reference size. The dial indicator then shows the deviation from this reference, allowing for precise measurement of the actual bore size and its consistency.

133. Answer: d

Explanation:

Understanding the Twist Drill Point Angle

A **twist drill** is a common cutting tool used for drilling holes. One of its most important features is the shape of its tip, which includes the **point angle**. The point angle is the angle

formed at the tip of the drill, between the cutting edges.

This angle is crucial because it affects how the drill enters the material, how effectively it cuts, and the shape and amount of chip produced. Different materials require different point angles for optimal performance and tool life.

The General Purpose Twist Drill Point Angle

The question asks about the standard or **general purpose point angle** for twist drills. While drills are available with various point angles, there is a specific angle widely accepted and used for drilling a broad range of common materials, particularly mild steel and cast iron.

The widely accepted **general purpose point angle** for twist drills is **118 degrees** (118°).

This angle provides a good balance of features:

- It is sharp enough to cut effectively into many common materials.
- It is strong enough to resist breaking during normal drilling operations.
- It helps center the drill reasonably well, reducing the need for extensive center punching in some cases.

Drills with a 118° point angle are the most common type found in general workshops and toolkits because of their versatility across various materials.

Comparison with Other Point Angles

While 118° is the general purpose angle, other angles are used for specific applications:

- **Lower Angles (e.g., 90°):** Used for softer materials like plastics, nylon, or very soft metals. A sharper point penetrates more easily but is weaker.
- **Higher Angles (e.g., 135° to 150°):** Used for harder materials like stainless steel or high-strength alloys. A blunter point is stronger and less likely to chip when dealing with tough materials, though it requires more force to penetrate and may need a pilot hole.

Therefore, for general purpose drilling in typical workshop materials, the 118° point angle is the standard choice.

Common Twist Drill Point Angles and Uses

Point Angle	Typical Material Use
90°	Plastics, soft woods, nylon
118°	General purpose (mild steel, cast iron, aluminum, brass)
135° – 150°	Hard steels, stainless steel, titanium

Revision Table: Twist Drill Point Angle

Key Facts on Twist Drill Point Angles

Concept	Description
Point Angle Definition	Angle at the tip between the cutting edges.
General Purpose Angle	118°
Use of 118° Angle	Common metals like mild steel, cast iron, aluminum. Balances strength and sharpness.
Use of Smaller Angles (< 118°)	Softer materials (plastics, wood). Sharper point, easier penetration, but weaker.
Use of Larger Angles (> 118°)	Harder materials (stainless steel). Stronger point, requires more force.

Additional Information on Twist Drill Angles

Beyond the point angle, other angles are important for a twist drill's performance:

- **Helix Angle (or Rake Angle):** This is the angle of the flutes relative to the drill's axis. It affects how chips are cleared and the cutting edge's sharpness. Standard helix angles vary, but common ones include normal (around 25 – 30° for general use), fast (higher angle for softer materials), and slow (lower angle for harder materials or brass).
- **Lip Relief Angle:** This is the angle ground behind the cutting edge (lip). It provides clearance so that the cutting edge is the only part contacting the bottom of the hole. Too little relief causes rubbing and heat; too much makes the cutting edge weak.

- **Chisel Edge Angle:** The angle of the web at the drill point. A larger chisel edge angle requires more thrust force. Splitting the point can reduce the chisel edge and the required thrust.

Optimizing these angles, along with the point angle, is crucial for efficient and effective drilling in different materials.

134. Answer: c

Explanation:

Understanding the materials used in engine components like pistons is crucial in the study of mechanical engineering and automotive technology. Pistons are key parts of reciprocating engines, moving up and down within the cylinder bore and transmitting force from expanding combustion gases to the crankshaft.

Exploring Common Piston Materials

Pistons operate under extremely harsh conditions, facing high temperatures, pressures, and mechanical stress. Therefore, the material chosen must possess specific properties to withstand these demands. Several materials can potentially be used, but some are more suitable for common applications than others.

- **Brass:** While brass has good wear resistance, it is generally not used for main engine pistons due to its density and lower strength compared to other options under high-temperature, high-pressure engine conditions.
- **Galvanised Copper:** Copper alloys can have good thermal conductivity, but galvanization involves coating with zinc, which has a relatively low melting point and might not be suitable for the high temperatures inside a combustion chamber. Copper itself is also heavy and less strong than other common piston materials.
- **Steel:** Steel pistons are very strong and durable, especially in high-performance or diesel engines where mechanical loads and temperatures are very high. However, steel is significantly heavier than aluminum alloys, which can impact engine efficiency and performance due to increased reciprocating mass. Steel also has lower thermal conductivity compared to aluminum.
- **Aluminum Alloy:** Aluminum alloys are the most common material for pistons in gasoline engines and many diesel engines. They offer an excellent balance of

properties essential for piston operation.

Why Aluminum Alloy is Preferred for Pistons

Aluminum alloys are widely used for pistons due to several advantageous properties:

- **Lightweight:** Aluminum is much lighter than steel or brass. This reduces the reciprocating mass of the piston, which allows for higher engine speeds, less vibration, and improved fuel efficiency.
- **Excellent Thermal Conductivity:** Aluminum conducts heat away from the combustion chamber very effectively. This is vital for keeping the piston cool and preventing overheating, which can lead to engine damage or knocking.
- **Good Strength-to-Weight Ratio:** Although not as strong as steel, aluminum alloys (often containing silicon, copper, magnesium, etc.) can be engineered to have sufficient strength to withstand combustion pressures while remaining light.
- **Machinability:** Aluminum alloys are relatively easy to machine into the complex shapes required for pistons.

While aluminum alloys expand more with heat than steel (requiring careful design of piston clearance), their overall benefits make them the material of choice for the vast majority of modern vehicle pistons.

Summary of Piston Material Properties

Material	Density	Strength at High Temp	Thermal Conductivity	Typical Use
Aluminum Alloy	Low	Moderate to Good	Excellent	Most Gasoline Engines, Many Diesel Engines
Steel	High	Excellent	Good	High-Performance, Heavy-Duty Diesel Engines
Brass	High	Low to Moderate	Good	Less common for main Pistons
Galvanised Copper	High	Low (due to Zinc)	Excellent (Copper component)	Not suitable for main Pistons

Considering the required properties for a piston operating under high temperature and pressure, a material that is lightweight, strong, and has high thermal conductivity is ideal. Aluminum alloys fit these criteria effectively for most applications.

Revision Table: Key Piston Material Facts

Concept	Detail
Piston Function	Transmit force from combustion to crankshaft
Ideal Piston Material Properties	Lightweight, High Strength, High Thermal Conductivity
Most Common Piston Material	Aluminum Alloy
Reason for Aluminum Alloy Use	Balance of low weight, good strength, excellent heat transfer

Additional Information: Piston Design Considerations

Beyond the material choice, the design of the piston is critical for its performance and longevity. Factors like piston rings, piston skirt shape, and cooling channels are all important. The thermal expansion difference between the aluminum piston and the cast iron or aluminum cylinder bore is also a key design consideration, impacting the required clearance between the piston and cylinder wall. This clearance is usually smaller for aluminum cylinder blocks compared to cast iron blocks.

135. Answer: c

Explanation:

Understanding Pneumatics and Driven Systems

The question asks about the medium that drives the systems studied in pneumatics. Pneumatics is a fascinating field of engineering.

What is Pneumatics?

Pneumatics is a branch of engineering science that deals with systems using gas or pressurized air. It is a type of fluid power, similar to hydraulics, but it uses a compressible gas (like air) instead of an incompressible liquid (like oil or water).

- Pneumatic systems often use compressed air to perform work.
- This work can involve moving cylinders, operating motors, or controlling processes.
- Common applications include industrial automation, manufacturing machinery, vehicle braking systems, and even dental drills.

Analyzing the Options for Pneumatics

Let's look at the provided options to determine which one correctly identifies the driving medium in pneumatics:

1. Oil: Systems driven by oil are typically studied under [hydraulics](#), not pneumatics. Oil is a liquid and relatively incompressible.
2. Water: Similar to oil, systems driven by water fall under [hydraulics](#) or hydro-power, using a liquid medium.
3. Air: As discussed, pneumatics specifically deals with systems powered by gas, most commonly compressed air. This aligns directly with the definition of pneumatics.
4. Solid matter: Solid matter is not a fluid medium used to drive systems in the way air, oil, or water are in fluid power systems like pneumatics or hydraulics.

Based on the definition and comparison of the options, the study of pneumatics deals with systems driven by **Air**.

Why Air is Used in Pneumatics

Compressed air is a popular choice for driving systems in many applications due to several advantages:

- **Availability:** Air is readily available everywhere and is free.
- **Storage:** Compressed air can be easily stored in tanks.
- **Safety:** Unlike hydraulic oil, air is not flammable, reducing fire hazards. Leaks are also less messy.

- **Cleanliness:** Exhausted air can be vented directly into the atmosphere, making pneumatic systems relatively clean.
- **Speed:** Pneumatic systems can achieve high speeds for operations.

However, air is compressible, which can sometimes lead to less precise positioning compared to hydraulic systems.

Conclusion on Pneumatic Systems

In summary, pneumatics is the field focused on using compressed gas, primarily air, to create motion and perform work. Therefore, systems driven by air are the subject of study in pneumatics.

Revision Table: Fluid Power Comparison

System Type	Driving Medium	Compressibility	Primary Focus
Pneumatics	Gas (typically Air)	Compressible	Using compressed gas for power & control
Hydraulics	Liquid (Oil or Water)	Relatively Incompressible	Using pressurized liquid for power & control

Additional Information: Components of Pneumatic Systems

A typical pneumatic system involves several key components working together:

- **Compressor:** Draws in ambient air and compresses it to a higher pressure.
- **Air Receiver (Tank):** Stores the compressed air.
- **Air Preparation Unit:** Filters, regulates the pressure, and sometimes lubricates the air.
- **Control Valves:** Direct the flow of compressed air to actuators.
- **Actuators:** Convert the energy of compressed air into mechanical motion (e.g., cylinders for linear motion, motors for rotary motion).
- **Piping/Hoses:** Carry the compressed air between components.

Understanding these components helps in grasping how pneumatic systems operate using compressed air as the driving force.

136. Answer: a

Explanation:

Understanding Clutch Free Play in Vehicles

The question asks about the purpose of having "free play" in a vehicle's clutch system. Free play refers to the small amount of distance the clutch pedal can be pushed before the clutch mechanism actually begins to disengage the engine from the transmission.

What is Clutch Free Play?

Clutch free play is the clearance or slack in the clutch pedal linkage. When you first press the clutch pedal, there is a short distance you can move it without any resistance before you feel the pressure required to disengage the clutch.

This free play is a critical adjustment in manual transmission vehicles.

Why is Clutch Free Play Needed?

The primary reason for having clutch free play is to ensure that the clutch disc is fully engaged with the flywheel and pressure plate when the pedal is in its resting position (up). If there is no free play, or insufficient free play, the pressure plate might not exert full clamping force on the clutch disc even when the pedal is released.

Let's look at the consequences of insufficient free play:

- If there is no free play, the release bearing (or throw-out bearing) remains in constant contact with the pressure plate's diaphragm springs.
- This constant contact applies a slight force that can prevent the pressure plate from fully clamping the clutch disc.
- When the clutch disc is not fully clamped, it will slip against the flywheel and pressure plate, even when the pedal is not pressed.
- Clutch slip generates significant friction and heat. This excessive heat can quickly damage the clutch disc, pressure plate, flywheel, and even the release bearing.

Therefore, maintaining the correct amount of clutch free play prevents the clutch from slipping unintentionally and heating up excessively, ensuring the clutch components last longer and function correctly.

Analyzing the Options

Let's consider the given options in the context of why clutch free play is needed:

- **to prevent the clutch from heating up:** As explained above, sufficient free play ensures full clutch engagement, preventing slippage and the resulting heat build-up. This is a direct and major consequence of incorrect free play.
- **to increase speed of engine rpm:** Clutch operation is about connecting or disconnecting the engine from the wheels. Free play adjustment does not directly influence the engine's ability to increase RPM; that depends on engine tuning and throttle input.
- **to improve efficiency of braking:** The clutch system is part of the powertrain, managing power transmission from the engine to the wheels. It is separate from the braking system, which is responsible for slowing or stopping the vehicle. Clutch free play has no bearing on braking efficiency.
- **to facilitate lubrication by oil:** The clutch components operate dry (except for the release bearing which might be sealed and pre-greased). They are not lubricated by engine oil or transmission oil circulating around the friction surfaces. Free play is a mechanical adjustment, not related to lubrication.

Based on this analysis, the need for clutch free play is directly related to preventing clutch slip and subsequent overheating.

Revision Table: Clutch Free Play

Your Personal Exams Guide

Term	Description	Importance of Correct Adjustment
Clutch Free Play	Slack in the clutch pedal linkage before disengagement begins.	Ensures full clutch engagement when pedal is released.
Insufficient Free Play	Too little or no slack.	Causes clutch slip, overheating, and premature wear of clutch components (disc, pressure plate, flywheel, release bearing).
Excessive Free Play	Too much slack.	Requires the pedal to be pushed further down to disengage the clutch fully, potentially making gear shifting difficult or causing grinding.

Additional Information: Clutch Operation

The clutch is a mechanical device that allows a driver to connect or disconnect the engine's rotating crankshaft from the transmission's input shaft. This is essential for changing gears smoothly and for stopping the vehicle without stopping the engine.

Key components involved in clutch operation include:

- **Flywheel:** Bolted to the engine crankshaft.
- **Clutch Disc (Friction Plate):** Sandwiched between the flywheel and pressure plate. It has friction material on both sides.
- **Pressure Plate:** Bolts to the flywheel assembly and uses strong springs (often a diaphragm spring) to clamp the clutch disc against the flywheel.
- **Release Bearing (Throw-out Bearing):** Activated by the clutch fork (or hydraulic slave cylinder) when the pedal is pressed. It pushes against the pressure plate's diaphragm springs to release the clamping force.
- **Clutch Fork / Slave Cylinder:** Linkage that translates pedal movement to the release bearing.

When the clutch pedal is released, the pressure plate firmly clamps the clutch disc to the flywheel. The engine's rotation is then transmitted through the clutch disc to the transmission input shaft. When the pedal is pressed, the release bearing moves, causing

the pressure plate to lift away from the clutch disc, breaking the connection and allowing the transmission to be shifted into a different gear.

Maintaining the correct clutch free play is a simple but crucial maintenance task to ensure the longevity and proper function of the entire clutch system.

137. Answer: d

Explanation:

Let's understand the primary function of the alternator in a vehicle's electrical system. The alternator is a crucial component that plays a vital role in keeping the vehicle running smoothly, especially when it comes to electrical power.

Understanding the Vehicle Alternator Function

The vehicle's electrical system relies on a continuous supply of power. While the battery provides power to start the engine and run accessories when the engine is off, it is not designed to provide all the power needed by the vehicle continuously. This is where the alternator comes in.

The alternator is essentially an electrical generator. It converts mechanical energy supplied by the engine (via a belt connected to the crankshaft) into electrical energy. This electrical energy is then used to power various electrical components in the vehicle, such as the lights, radio, power windows, ignition system, and onboard computers.

Charging the Vehicle Battery

One of the most critical functions of the alternator is to charge the vehicle's battery. As the engine runs, the alternator generates voltage. This generated voltage is typically higher than the battery's voltage (around 12 volts), which allows current to flow into the battery, recharging it. This ensures that the battery is topped up and ready to provide power again when needed, such as for the next engine start.

If the alternator were not working correctly, the vehicle would eventually run solely on battery power, which would quickly deplete the battery, leading to the vehicle stalling or being unable to start.

Analyzing the Options

Let's look at the given options in the context of the alternator's function:

- Measure the current: Devices like ammeters or current sensors measure current. This is not the primary function of the alternator itself.
- Discharge the battery: The alternator is designed to supply power and charge the battery, not discharge it. Components that draw power (like lights or the starter) discharge the battery.
- Measure the voltage: Devices like voltmeters or voltage regulators measure or control voltage. While the alternator's output voltage is regulated, measuring voltage is not the alternator's primary function.
- Charge the battery: As discussed, generating electrical power to recharge the battery is a core function of the alternator.

Based on the function of generating electrical power to meet the vehicle's electrical demands and recharge the battery, the correct answer is that the alternator functions to charge the battery.

Revision Table: Vehicle Electrical Components

Component	Primary Function
Battery	Stores electrical energy; Provides power for starting and accessories when the engine is off.
Alternator	Generates electrical power while the engine is running; Charges the battery and powers electrical systems.
Starter Motor	Uses battery power to crank the engine for starting.
Voltage Regulator	Controls the voltage output of the alternator to protect electrical components and the battery.

Additional Information on Vehicle Charging System

The vehicle charging system consists of more than just the alternator and battery. It also includes the voltage regulator and the wiring connecting these components. The voltage regulator is crucial because it ensures the alternator's output voltage stays within a safe

range (typically 13.5 to 14.5 volts). If the voltage is too low, the battery won't charge properly. If it's too high, it can damage the battery and other electrical components.

In modern vehicles, the voltage regulator is often integrated into the alternator or controlled by the engine's computer (ECU). The alternator output varies with engine speed; the voltage regulator adjusts the magnetic field inside the alternator to maintain a consistent output voltage across different RPMs.

138. Answer: b

Explanation:

Understanding Air Cleaners and Their Function

An air cleaner, often called an air filter, is a crucial component in engines and many types of machinery. Its primary purpose is to filter the air entering the system, ensuring that it is clean and free from contaminants. Dirty air containing dust, dirt, sand, and other particles can cause significant wear and damage to internal parts, such as engine cylinders, pistons, and bearings.

Purpose of Air Cleaners

The main function of an air cleaner is to protect the sensitive internal components of an engine or machine from abrasive particles present in the atmosphere. By removing these contaminants, the air cleaner helps to:

- Prevent premature wear of engine parts.
- Maintain engine efficiency and performance.
- Extend the lifespan of the engine or machine.
- Ensure proper combustion by supplying clean air.

Role of Oil in Oil-Filled Air Cleaners

Some types of air cleaners, particularly older designs or those used in very dusty or harsh environments, use oil as part of the filtration process. These are often called oil-bath or oil-filled air cleaners. The air passes through an oil bath or comes into contact with an oil-saturated element. The oil plays a vital role in trapping dust and dirt particles.

Here's how the oil works to prevent dust and dirt:

- Air is drawn into the cleaner, often directed downwards towards the oil reservoir.
- Larger, heavier particles tend to drop directly into the oil bath due to inertia.
- The air then passes through a filter element (like wire mesh or foam) which is often wet with oil or through which the air bubbles through the oil.
- As the air passes through this oil-laden medium or bubbles through the oil, dust and dirt particles stick to the oil.
- The oil traps these contaminants, preventing them from continuing into the engine intake.
- The cleaned air then proceeds into the engine.

The oil effectively acts as an adhesive agent that captures and holds onto the fine particles that would otherwise pass through a dry filter.

Analyzing the Given Options

Let's examine the provided options in the context of why an air cleaner is filled with oil:

1. **To lubricate:** The primary purpose of the oil in an air cleaner is not lubrication of engine components. Lubrication is handled by the engine's oil circulation system. While oil is a lubricant, its function within the air cleaner is different.
2. **To prevent dust and dirt from atmosphere:** This aligns directly with the function of oil in trapping contaminants as described above. The oil enhances the air cleaner's ability to capture particles.
3. **To pressurize:** An air cleaner does not pressurize the incoming air. Air intake systems might be pressurized in turbocharged or supercharged engines, but this function is not performed by the air cleaner itself or the oil within it.
4. **To accelerate:** The air cleaner is designed to filter air, not accelerate its flow. In fact, the presence of the filter element and oil can slightly impede airflow, although air cleaner design aims to minimize this restriction while maximizing filtration efficiency.

Based on the mechanism of oil-filled air cleaners, the oil is used specifically for its ability to capture and retain dust and dirt particles from the incoming air.

Conclusion

The main reason an air cleaner is filled with oil in certain designs is to effectively capture and prevent dust and dirt from the atmosphere from entering the engine or machinery. The oil serves as a sticky medium that traps particulate matter, thereby protecting the system.

Revision Table: Air Cleaner Functions

Component/Substance	Primary Role in Air Cleaner
Air Filter Element (e.g., paper, foam, mesh)	Provides a large surface area to trap particles.
Housing	Directs airflow and holds the filter element and oil.
Oil (in oil-filled types)	Acts as an adhesive medium to trap dust and dirt particles.

Additional Information: Types of Air Cleaners

Air cleaners come in various types, each with a different filtration method:

- **Dry Type Air Cleaners:** These are the most common type today. They use a pleated paper or synthetic filter element to trap particles. They require periodic replacement.
- **Oil-Bath Air Cleaners:** As discussed, these use an oil reservoir and an oil-wetted element to trap contaminants. They require the oil to be changed and the element cleaned periodically.
- **Foam Air Filters:** These use a foam element, often saturated with a special oil, to trap particles. Common in motorcycles and some older vehicles.
- **Water Bath Air Cleaners:** Less common, used in some industrial applications where moist filtration is suitable.

Each type has advantages and disadvantages depending on the application and environment.

139. Answer: a

Explanation:

Connecting Rod Materials: Understanding Engine Components

A connecting rod is a crucial part of a piston engine. It connects the piston to the crankshaft. Its main job is to convert the linear up-and-down motion of the piston into the rotational motion of the crankshaft. Because it transfers the power from the piston to the crankshaft, the connecting rod experiences very high stresses, including tension, compression, and bending. It also undergoes cyclic loading, which makes fatigue strength a critical consideration.

Choosing the right material for the connecting rod is essential for the engine's reliability, durability, and performance. The material must be strong enough to withstand the forces involved and resistant to fatigue failure over millions of cycles.

Analyzing Connecting Rod Material Options

Let's look at the properties of the common engineering materials listed in the options and see how suitable they are for making a connecting rod:

- **Alloy Steel:** Steels alloyed with elements like chromium, nickel, molybdenum, or vanadium have significantly improved mechanical properties compared to plain carbon steel. These properties include higher strength, toughness, fatigue resistance, and the ability to be heat-treated to achieve optimal hardness and strength profiles. Forged alloy steel is a very common material for connecting rods in passenger cars, trucks, and performance engines due to its excellent balance of strength, durability, and cost-effectiveness.
- **Aluminium:** Aluminium alloys are much lighter than steel. This weight reduction can be beneficial in reducing reciprocating mass, which can allow for higher engine speeds and potentially improve engine response. However, aluminium alloys generally have lower strength and fatigue resistance compared to alloy steels at engine operating temperatures. While some connecting rods, particularly for racing applications where weight saving is paramount and service life is shorter, are made from high-strength aluminium alloys, they are less common in standard production engines due to durability concerns under prolonged stress cycles.
- **Cast Iron:** Cast iron is relatively inexpensive and easy to cast into complex shapes. However, it is generally brittle and has lower tensile strength and fatigue resistance

compared to forged steel or aluminium alloys. Grey cast iron and ductile cast iron are used in some less stressed components, but they are typically not suitable for the high dynamic loads experienced by connecting rods in modern engines.

- **Mild Steel:** Mild steel (low carbon steel) is ductile and easy to machine or forge. However, it has lower strength and fatigue limit compared to alloy steels. It would not be able to withstand the high forces and repeated stress cycles that a connecting rod experiences in most internal combustion engines without failing prematurely.

Why Alloy Steel is Preferred for Connecting Rods

Based on the requirements of high strength, excellent fatigue resistance, and reasonable cost for mass production, alloy steel stands out as a highly suitable material for connecting rods in a wide range of engines. Forging is a common manufacturing process used with alloy steel to create connecting rods, which refines the grain structure and further improves their strength and fatigue life.

Material Suitability for Connecting Rods

Material	Strength	Fatigue Resistance	Weight	Cost	Suitability for Connecting Rods
Alloy Steel	High	High	Medium	Medium	High (Common)
Aluminium	Medium (lower at temp)	Medium (lower than steel)	Low	High	Medium (Specialized applications)
Cast Iron	Low	Low	High	Low	Low (Generally unsuitable)
Mild Steel	Medium	Low to Medium	Medium	Low	Low (Generally unsuitable)

Considering the functional requirements and material properties, alloy steel is the most common and appropriate material for manufacturing connecting rods in typical automotive engines.

Revision Table: Connecting Rod Essentials

Component	Function	Typical Material	Key Property Needed
Connecting Rod	Connects piston to crankshaft; converts linear to rotary motion	Forged Alloy Steel	High Strength & Fatigue Resistance

Additional Information on Connecting Rods and Materials

While alloy steel is prevalent, variations exist:

- **Forged Steel Connecting Rods:** Most common type, offering excellent strength and durability.
- **Billet Steel Connecting Rods:** Machined from solid billets of high-strength steel, often used in high-performance or racing engines for maximum strength, though more expensive.
- **Powdered Metal Connecting Rods:** An alternative manufacturing process gaining traction for cost-effectiveness in some production engines, though material properties need careful control.
- **Aluminium Connecting Rods:** Used in some racing engines or specific applications where weight reduction is critical, despite having a shorter fatigue life compared to steel.

The specific composition of the alloy steel used (e.g., 4340, 5140, 4140 steel) and the subsequent heat treatment process are optimized to achieve the desired balance of strength, toughness, and fatigue life for the specific engine application.

140. Answer: b

Explanation:

Identifying the Correct Tool for Laying Out Diameter

When working in workshops or technical settings, precisely marking out shapes and sizes on materials is a crucial step before cutting or machining. For tasks involving circles or

diameters, a specific tool is typically used for marking or "laying out." Let's examine the given options to understand which tool is best suited for laying out diameter.

Understanding the Function of Each Tool Option

Let's look at what each tool is primarily used for:

- **Scriber:** This tool has a sharp point and is used to draw lines on surfaces, especially metal, acting much like a pencil does on paper. It's for general marking of lines.
- **Divider:** A divider is a tool with two legs, usually hinged, that can be adjusted to set a specific distance between their points. It is commonly used for transferring measurements, comparing distances, or scribing circles and arcs.
- **Outside micrometer:** This is a precision measuring instrument used to measure the external dimensions of objects, such as the diameter of a rod or the thickness of a plate. It provides accurate numerical readings.
- **Tri square (or Try square):** This tool consists of two parts, a stock and a blade, joined at a right angle (90 degrees). It is used for checking the squareness of an edge or surface and for marking lines perpendicular to an edge.

Why the Divider is Used for Laying Out Diameter

Laying out a diameter essentially involves marking a circle or an arc of a specific size. The divider is the ideal tool for this purpose among the given options. To lay out a circle with a specific diameter, you would set the distance between the legs of the divider to half the desired diameter (which is the radius). Then, with one leg placed at the center point, the other leg is used to scribe the circle or arc onto the material surface.

While a scriber makes the mark, a divider controls the radius/diameter of the mark based on a central point. A micrometer measures an existing diameter, and a tri square checks or marks 90-degree angles.

Therefore, the tool specifically designed and commonly used for the task of laying out diameter (by marking the corresponding circle or radius) is the Divider.

Comparison of Tools for Laying Out Diameter

Tool	Primary Function	Use for Laying Out Diameter?
Scriber	Marking general lines	Indirectly, as part of other tools (like dividers or compasses) or freehand marking guide. Not the primary tool for precise diameter layout.
Divider	Transferring distances, scribing circles/arcs	Yes, commonly used to lay out circles/arcs of a specific radius/diameter from a center point.
Outside micrometer	Measuring external dimensions	No, used for measurement, not marking/layout.
Tri square	Checking/marking 90° angles	No, unrelated to circular layout.

Based on their functions, the Divider is the tool used for laying out diameter.

Revision Table: Key Tool Functions

Tool Name	Main Purpose
Scriber	Marking lines on material surfaces
Divider	Transferring dimensions, drawing circles and arcs, dividing lines/arcs
Outside Micrometer	Precise measurement of external dimensions (e.g., diameter, thickness)
Tri Square	Checking and marking 90-degree angles (squareness)

Additional Information on Layout Tools

Layout tools are essential in manufacturing and fabrication for transferring design specifications onto raw materials. Accurate layout ensures that subsequent operations like cutting, drilling, and machining are performed correctly according to the desired dimensions and shapes. Besides the tools mentioned, other layout tools include center punches (for marking hole centers), surface plates (as a reference datum), height gauges (for vertical measurements and scribing), and compasses (similar to dividers but

often used for larger circles and may accommodate a pencil or pen). Each tool serves a specific purpose in the overall process of preparing a workpiece for manufacturing.

141. Answer: b

Explanation:

Understanding Gunmetal Composition

Gunmetal is a type of bronze, an alloy known for its strength and resistance to corrosion. Historically, its name comes from its use in making cannons (guns).

The primary components of Gunmetal are Copper, Tin, and Zinc. While it is fundamentally a bronze (which is an alloy of Copper and Tin), the addition of Zinc distinguishes it and affects its properties, often improving castability and reducing porosity.

Typical Composition of Gunmetal Alloy

The exact percentages of the elements in Gunmetal can vary depending on the specific application or standard (like red brass, which is sometimes classified with gunmetals), but a common composition is approximately:

Element	Approximate Percentage (%)
Copper (Cu)	88
Tin (Sn)	8 - 10
Zinc (Zn)	2 - 4

Minor amounts of other elements like lead might also be present in some specific types, but the core composition revolves around Copper, Tin, and Zinc.

Analyzing the Composition Options

Let's evaluate the given options based on the known composition of Gunmetal:

1. **Lead and nickel:** This combination does not represent Gunmetal. Alloys primarily of lead and nickel (like Monel metal, which also contains copper and iron) have different properties and uses.
2. **Copper, Tin and zinc:** This combination accurately describes the main components of Gunmetal. As discussed, it's an alloy of copper, tin, and zinc, falling under the broader category of bronze (copper and tin alloy) with the addition of zinc.
3. **Copper, zinc and nickel:** An alloy primarily of copper, zinc, and nickel is known as Nickel Silver or German Silver. Despite the name, it contains no silver but is valued for its hardness and resistance to corrosion.
4. **Lead and zinc:** An alloy of lead and zinc is not Gunmetal. Such alloys are typically used in solders or die-casting, depending on the proportions, and have very different properties from Gunmetal.

Based on this analysis, the composition of Gunmetal is Copper, Tin, and Zinc.

Revision Table: Understanding Alloys

Alloy Name	Primary Components	Notes
Gunmetal	Copper, Tin, Zinc	A type of bronze with added zinc; good strength & corrosion resistance.
Bronze	Copper, Tin	Broad category, historically used for tools, weapons, statues.
Brass	Copper, Zinc	Varies widely in composition; more ductile than bronze, used for plumbing, musical instruments.
Nickel Silver / German Silver	Copper, Zinc, Nickel	No silver; known for hardness, corrosion resistance, silvery appearance.

Additional Information on Gunmetal Properties and Uses

Gunmetal possesses several useful properties that make it suitable for various applications:

- **Good Strength and Wear Resistance:** It is tougher and stronger than many other bronzes.
- **Corrosion Resistance:** It holds up well in environments, especially those involving water or steam.
- **Good Castability:** The addition of zinc helps in producing sound castings with fewer defects.

Due to these properties, Gunmetal is commonly used in:

- Valve bodies and fittings, especially in steam and water systems.
- Pumps and pump components.
- Bearings and bushings.
- High-quality castings requiring durability and corrosion resistance.

Understanding the specific composition of an alloy like Gunmetal is crucial for selecting the correct material for engineering applications, ensuring it meets the required mechanical, thermal, and chemical properties.

142. Answer: a

Explanation:

Understanding Drill Bit Point Angles

A drill bit's point angle is a crucial feature that affects its performance when drilling different materials. It is the angle formed at the very tip of the drill bit between the cutting edges.

Standard Drill Bit Angles

Different materials and drilling applications require specific point angles for optimal performance, lifespan, and chip evacuation. Some common point angles include:

- **118 degrees:** This is the most common general-purpose point angle. It is suitable for drilling in a wide range of materials including mild steel, wood, and plastics. This angle provides a good balance between centering ability and cutting edge strength.
- **135 degrees:** This angle is used for harder materials like stainless steel, cast iron, and titanium. The blunter angle provides a stronger tip, reducing the risk of chipping, and

the shorter cutting edges require less thrust force. These bits often have a split-point tip to improve self-centering.

- **90 degrees:** Used primarily for drilling soft materials like plastics and some very soft metals.
- **Other angles:** Specialized drill bits exist with angles tailored for specific tasks, such as spade bits for large holes in wood or brad-point bits for precise woodworking.

Analyzing the Given Options

The question asks for the angle of the drill bit. We are given several options representing common drill bit angles:

- 118°
- 115°
- 119°
- 59°

Considering the standard angles for common drill bits, 118 degrees is the most widely recognized general-purpose angle.

Let's look at a comparison of common drill bit angles and their typical applications:

Point Angle	Typical Application
90°	Soft materials (plastics, some woods)
118°	General purpose (mild steel, wood, plastics)
135°	Harder materials (stainless steel, cast iron)

Based on the standard industry practices and common types of drill bits, 118° is a very common and standard angle.

Revision Table: Key Drill Bit Concepts

Term	Explanation
Point Angle	Angle formed by the cutting edges at the tip.
Flute	Spiral grooves that remove chips from the hole.
Shank	Part that fits into the drill chuck.
Lip/Cutting Edge	The sharp part that cuts the material.
Body	The part between the point and the shank.

Additional Information on Drill Bit Performance

The performance of a drill bit is influenced by several factors beyond just the point angle:

- **Material:** The material the drill bit is made from (e.g., High-Speed Steel (HSS), Cobalt HSS, Carbide) affects its hardness and heat resistance, crucial for drilling different workpieces.
- **Coatings:** Coatings like Titanium Nitride (TiN), Titanium Carbonitride (TiCN), or Aluminum Titanium Nitride (AlTiN) can increase hardness, reduce friction, and improve wear resistance, extending the bit's life and performance.
- **Web Thickness:** The thickness of the web at the center of the drill bit affects its strength and rigidity.
- **Lip Relief Angle:** The angle behind the cutting edge allows it to penetrate the material without rubbing excessively.
- **Feed Rate and Speed:** The speed at which the drill rotates (RPM) and the rate at which it is advanced into the material (feed rate) must be appropriate for the material being drilled and the drill bit size to prevent overheating or breakage.
- **Sharpening:** Properly sharpening a drill bit restores its cutting edges and point angle, significantly extending its usability. Sharpening jigs can help maintain the correct point angle and lip relief.

Choosing the right drill bit for the material and application ensures efficient drilling, reduces wear on the bit and machine, and produces a cleaner, more accurate hole.

143. Answer: d

Explanation:

Understanding Engine Speed Control Mechanisms

Controlling the speed of an engine is crucial for its safe and efficient operation. Different mechanisms are employed for this purpose, each with a specific function.

Identifying the Engine Speed Control Mechanism

The question asks for the name of a mechanism specifically used to control the speed of an engine. Let's look at the options provided:

- **Speedometer:** This device measures the instantaneous speed of a vehicle or machine. It does not control the speed; it only indicates it.
- **Regulator:** This is a broad term for a device that maintains a particular parameter (like voltage, pressure, or temperature) within a specified range. While a governor is a type of regulator, the term 'regulator' itself is not specific enough in the context of engine speed control compared to 'governor'.
- **Odometer:** This device measures the total distance traveled by a vehicle. It has no function in controlling engine speed.
- **Governor:** A governor is a mechanical or electronic device used to measure and regulate the speed of an engine. It works by controlling the flow of fuel or air-fuel mixture to the engine, thereby maintaining the speed within a desired range or preventing it from exceeding a maximum limit.

Based on the functions described, the mechanism specifically designed to control engine speed is the governor.

How a Governor Controls Engine Speed

Governors typically work by sensing the engine speed. If the speed increases beyond a set point (due to reduced load, for example), the governor reduces the fuel supply. If the speed decreases below a set point (due to increased load), the governor increases the fuel supply. This feedback mechanism helps in maintaining a relatively constant speed despite variations in load.

Comparison of Options

Let's compare the functions of the given options:

Mechanism	Primary Function	Controls Engine Speed?
Speedometer	Measures speed	No
Regulator	Maintains a parameter (general term)	Indirectly (Governor is a type of regulator)
Odometer	Measures distance	No
Governor	Controls/Regulates engine speed	Yes

From the comparison, it is clear that the governor is the mechanism used specifically for controlling the speed of an engine.

Revision Table: Engine Instruments

Term	Function
Speedometer	Measures instantaneous speed
Odometer	Measures total distance traveled
Governor	Controls/Regulates engine speed
Regulator	General term for a device that maintains a parameter within a range

Additional Information: Types of Engine Governors

Governors can be classified based on their operating principle. Some common types include:

- **Mechanical Governors:** These use weights (like centrifugal weights) that move outwards as speed increases. This movement is linked to the fuel control mechanism.
- **Hydraulic Governors:** These use hydraulic pressure to operate the fuel control based on speed sensing.
- **Electronic Governors:** These use electronic sensors to measure speed and electronic control units (ECUs) to control fuel injection or ignition timing to regulate speed.

Governors are used in various applications, including power generation (maintaining constant frequency), vehicles (limiting maximum speed or idle speed control), and industrial machinery.

144. Answer: b

Explanation:

Understanding How to Measure a Cylinder Bore

Measuring a cylinder bore accurately is crucial in many fields, especially in automotive repair and machining. A cylinder bore is the diameter of the cylindrical chamber in an engine block or similar component. Precise measurement helps determine wear, check for roundness and taper, and ensure proper fit for pistons or other parts.

Analyzing Instruments for Cylinder Bore Measurement

Let's look at the instruments provided in the options and evaluate their suitability for measuring a cylinder bore:

- **Screw gauge:** A screw gauge is primarily used for measuring very small external or internal diameters, or the thickness of thin objects. While some types might measure small internal diameters, they are generally not suitable or precise enough for the larger diameters and required accuracy when measuring typical engine cylinder bores.
- **Dial indicator:** A dial indicator is a measuring tool that displays small linear distances with precision. It doesn't directly measure length on its own, but it is a key component of many measuring instruments, including bore gauges. A dial bore gauge uses a dial indicator to show the deviation from a set master size. This instrument is specifically designed for accurately measuring internal diameters, such as cylinder bores, and checking for variations like ovality and taper.
- **Caliper:** Calipers (like Vernier or digital calipers) can measure internal dimensions using their internal jaws. They are useful for general measurements and quick checks. However, for high-precision measurements of cylinder bores, especially for checking consistency along the bore length (taper) and roundness (ovality), a

caliper is typically less accurate and less effective than a dedicated bore gauge that incorporates a dial indicator.

- **Micrometer:** Micrometers are known for their high precision. Inside micrometers or bore micrometers are instruments used to measure internal diameters. While a bore micrometer can measure the size, a dial bore gauge provides a more convenient way to check for size variations throughout the bore.

Identifying the Correct Instrument for Cylinder Bore Measurement

Considering the options and common practices for precise cylinder bore measurement, an instrument incorporating a dial indicator, specifically a dial bore gauge, is widely used for this purpose. The dial indicator on the bore gauge allows for easy reading of variations and deviations from the desired size, making it ideal for checking the condition of the bore.

Instrument	Typical Use	Suitability for Cylinder Bore
Screw Gauge	Small dimensions (thickness, small diameters)	Generally not suitable for typical cylinder bores
Dial Indicator	Used as a component in instruments (like bore gauges) to show deviation/variation; indirect measurement of length variation.	Key component of a Dial Bore Gauge, specifically designed for cylinder bores.
Caliper	General external/internal measurements	Less precise for critical cylinder bore checks (taper, ovality)
Micrometer	Precise external/internal measurements (e.g., bore micrometer)	Precise for size, but dial bore gauge better for checking variations

Therefore, among the given options, the instrument most directly associated with the precise measurement and inspection of a cylinder bore, often used to detect variations in size, is one that utilizes a dial indicator (as part of a bore gauge).

Revision Table: Instruments for Measurement

Instrument Type	Measures	Principle
Screw Gauge	Small lengths, diameters, thickness	Screw principle (converting rotation to linear movement)
Dial Indicator	Small linear displacement/deviation	Gear train amplifies plunger movement to pointer rotation
Caliper	External and internal dimensions, depth, step	Sliding jaws with scale (Vernier, digital, dial)
Micrometer	Precise lengths, external and internal diameters	Screw principle (precise thread)

Additional Information: Dial Bore Gauges

A dial bore gauge is specifically designed for measuring the diameter of bores (cylindrical holes) and checking for their geometry, such as taper and ovality. It consists of a handle, a set of interchangeable extension rods and anvils for different size ranges, and a measuring head connected to a dial indicator. The measuring head has a fixed anvil and a movable anvil. When inserted into the bore, the movable anvil retracts. The dial indicator shows the difference between the bore size and a master setting (often a setting ring or a micrometer). By taking readings at different depths and angles within the bore, one can accurately determine the bore's size and identify any imperfections like taper (diameter changes along the length) or ovality (out-of-roundness).

145. Answer: b

Explanation:

Understanding Electronic Ignition System Triggering

An electronic ignition system is a modern type of ignition system used in internal combustion engines. Unlike older systems that rely on mechanical contact breakers, electronic systems use electronic components to control the timing of the ignition spark. A critical part of this system is determining exactly when to fire the spark plug. This timing is controlled by a signal, and the component that generates this signal is called the trigger.

How the Trigger Works in Electronic Ignition

The primary function of the **trigger** in an electronic ignition system is to sense the position of the engine's crankshaft or distributor shaft. Based on this position, the trigger sends a signal to the electronic control unit (ECU) or ignition module. This signal tells the ECU when to interrupt the current flowing through the ignition coil. When the current is interrupted, a high voltage is induced in the coil's secondary winding, which then creates the spark at the spark plug.

Common types of triggers used in electronic ignition systems include:

- **Magnetic Pick-up:** Uses a sensor coil and a rotating reluctor wheel or permanent magnet to generate a voltage pulse as teeth pass the sensor.
- **Hall Effect Sensor:** Uses a semiconductor material that produces a voltage when a magnetic field is applied. A rotating blade or shutter interrupts the magnetic field, generating a square wave signal.
- **Optical Sensor:** Uses an LED and a phototransistor. A rotating disc with slots interrupts the light beam, generating pulses.

Regardless of the specific technology, the component providing this timed signal is broadly referred to as the **trigger**.

Analyzing the Options

Let's look at the provided options in the context of triggering an electronic ignition system:

- **Contact breaker:** This is a mechanical switch used in older, conventional ignition systems. It physically opens and closes contacts to interrupt the coil current. Electronic ignition systems specifically replace the mechanical contact breaker with electronic switching controlled by a trigger signal. Therefore, the electronic system is not triggered on and off through a contact breaker itself, but rather it replaces the function of the contact breaker.

- **Trigger:** As explained above, the trigger is the component designed specifically to generate the timing signal for the electronic ignition system, initiating the switching process.
- **Diode:** A diode is an electronic component that allows current to flow in only one direction. While diodes are used within the circuitry of an electronic ignition system, they are not the component that initiates or triggers the system's on/off switching action based on engine position.
- **Permanent magnet:** A permanent magnet might be part of a magnetic pick-up trigger mechanism (e.g., attached to the reluctor wheel or the sensor itself), but the permanent magnet alone does not trigger the system. It's the interaction of the magnet and the sensor/coil that generates the trigger signal. The trigger is the more accurate term for the component or mechanism that provides the timing signal.

Based on the function, the electronic ignition system relies on the **trigger** mechanism to determine the precise moment to switch the ignition coil's current, thereby initiating the spark.

Component	Role in Ignition System	Used for Triggering Electronic Ignition?
Contact breaker	Mechanical switch to interrupt coil current (conventional system)	No (replaced by electronic triggering)
Trigger (Magnetic Pick-up, Hall Effect, Optical)	Generates timing signal based on engine position	Yes
Diode	Allows current flow in one direction (part of circuitry)	No (not the primary timing signal source)
Permanent magnet	Component used in some trigger types (e.g., magnetic pick-up)	No (part of the trigger mechanism, not the trigger itself)

Conclusion on Electronic Ignition Triggering

The electronic ignition system's timing, and thus its on/off switching cycle for generating the spark, is precisely controlled by the **trigger** component or mechanism. This trigger

provides the essential signal that synchronizes the ignition spark with the engine's operation.

Revision Table: Electronic Ignition Key Components

Component	Function
Trigger	Provides engine position timing signal
Ignition Module/ECU	Receives trigger signal, controls coil switching
Ignition Coil	Transforms low voltage into high voltage for spark
Spark Plug	Creates spark in combustion chamber

Additional Information on Ignition Systems

Ignition systems have evolved significantly. Early systems used mechanical contact breakers and capacitors. Electronic ignition systems improved reliability and timing accuracy by replacing mechanical points with electronic switches (like transistors or SCRs) controlled by precise timing signals from a trigger. Further advancements led to fully digital systems controlled by the engine control unit (ECU), integrating ignition timing with other engine parameters for optimal performance and emissions control. The trigger remains a fundamental part of sensing engine timing in these modern systems.

146. Answer: b

Explanation:

Understanding Battery Cell Connections for Increased Voltage

The question asks how to connect battery cells to increase the output voltage. When working with battery cells or any voltage sources like power supplies, the way they are connected significantly impacts the total voltage and current capabilities of the combined system.

Series Connection Explained

Connecting battery cells in **series** means connecting the positive terminal of one cell to the negative terminal of the next cell, and so on. This type of connection is used when you need a higher total voltage than what a single cell can provide.

- In a series connection, the voltages of the individual cells add up to give the total output voltage.
- The current capacity (measured in ampere-hours, Ah) remains the same as that of a single cell (assuming all cells have the same capacity).
- For example, if you have three 1.5V cells, connecting them in series will give you a total voltage of $1.5V + 1.5V + 1.5V = 4.5V$.

The formula for total voltage (V_{total}) in a series connection of n cells, each with voltage V_{cell} , is:

$$V_{total} = V_{cell1} + V_{cell2} + \dots + V_{celln}$$

If all cells are identical (V_{cell}), the formula simplifies to:

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

Therefore, to increase the output voltage, battery cells are connected in series.

Comparing Series and Parallel Connections

Let's look at how parallel connection differs and why it doesn't increase voltage.

- Connecting battery cells in **parallel** means connecting all the positive terminals together and all the negative terminals together.
- This type of connection is used when you need to increase the total current capacity or runtime, not the voltage.
- In a parallel connection, the total voltage remains the same as the voltage of a single cell (assuming all cells have the same voltage).
- The current capacities of the individual cells add up to give the total output current capacity.
- For example, if you have three 1.5V cells, each with 1000mAh capacity, connecting them in parallel will still give you 1.5V, but the total capacity will be $1000mAh + 1000mAh + 1000mAh = 3000mAh$.

Analyzing the Options

- **Inclined:** This term is not relevant to electrical connections of battery cells.
- **Series:** As discussed, connecting cells in series adds their voltages, thus increasing the total output voltage. This aligns with the goal stated in the question.
- **Parallel:** Connecting cells in parallel keeps the voltage the same but increases the current capacity. This does not meet the requirement of increasing voltage.
- **Series-parallel:** This connection method involves groups of series-connected cells which are then connected in parallel, or vice versa. It is used to achieve a desired combination of voltage and current capacity. While it involves series connections, the primary method to **only** increase voltage is the series connection itself. If the goal is purely voltage increase, series is the direct method.

Based on the analysis, connecting battery cells in series is the correct method to increase the output voltage.

Revision Table: Battery Connection Types

Connection Type	Effect on Voltage	Effect on Current Capacity	Primary Use
Series	Increases (adds up)	Remains the same (as single cell)	Increase total voltage
Parallel	Remains the same (as single cell)	Increases (adds up)	Increase total runtime/current capacity

Additional Information on Battery Connections

When connecting battery cells, especially in series or parallel, it's important to consider factors beyond just voltage and current:

- **Matching Cells:** For optimal performance and longevity, especially in series, it's best to use cells of the same type, capacity, age, and state of charge. Mismatched cells can lead to reduced performance or damage.
- **Internal Resistance:** Each battery cell has internal resistance. When connected in series, these resistances add up, which can affect the maximum current delivery and efficiency. When connected in parallel, the total internal resistance decreases.
- **Safety:** Incorrect connections can be dangerous, leading to short circuits, overheating, or even fire. Always follow recommended practices and use

appropriate protection circuits where necessary, especially for complex series or parallel arrangements like in battery packs.

- **Balancing:** For battery packs made of multiple cells (common in laptops, electric vehicles, etc.), especially lithium-ion cells in series, balancing circuits are often used. These circuits ensure that all cells in the series string have similar voltages and states of charge, which improves the pack's overall health and lifespan.

147. Answer: c

Explanation:

Understanding the Engine Thermostat's Purpose

The engine thermostat is a critical component in a vehicle's cooling system. Its main job is to regulate the flow of coolant through the engine and radiator. By doing this, the thermostat helps the engine reach and maintain its optimal operating temperature.

How the Engine Thermostat Works

The thermostat is typically located in the coolant passage between the engine and the radiator. It contains a temperature-sensitive element (often wax or bimetallic). When the engine is cold, the thermostat remains closed, preventing coolant from flowing to the radiator. This allows the engine to warm up quickly. As the engine temperature rises and reaches a specific point (its opening temperature, typically around 180–200°F or 82–93°C), the thermostat opens. This allows hot coolant to flow to the radiator, where it is cooled by air passing over the fins, and then recirculated back to the engine.

The thermostat constantly adjusts its opening based on the coolant temperature, maintaining a steady temperature range for the engine.

Why Maintaining Optimal Engine Temperature is Key

Operating an engine at the correct temperature is vital for several reasons:

- **Efficiency:** Engines are designed to run most efficiently within a specific temperature range. Running too cold or too hot reduces fuel efficiency.

- **Performance:** Optimal temperature ensures proper combustion and lubrication, leading to better engine performance.
- **Emissions:** Efficient combustion at the correct temperature reduces harmful emissions.
- **Wear and Tear:** Maintaining a stable temperature reduces thermal stress on engine components and ensures proper lubrication, extending engine life.

Analyzing the Options

Let's look at the provided options in the context of the engine thermostat's function:

- **Warm:** While the thermostat helps the engine get warm, its purpose is not just to keep it "warm". It's about reaching a specific, higher temperature range.
- **Cool:** This is incorrect. The thermostat prevents the engine from getting **too** hot by allowing cooling, but its purpose is not to keep the engine "cool". Engines need to be hot to operate efficiently.
- **At desired temperature:** This accurately describes the thermostat's role. It actively works to keep the engine's coolant within a specific, optimal temperature range specified by the manufacturer.
- **Hot:** While the engine operates at high temperatures, the thermostat's purpose is not simply to keep it "hot". It's to keep it hot enough for efficiency but **prevent** it from getting excessively hot, which would cause damage.

Therefore, the most accurate description of the thermostat's purpose is to keep the engine at its desired, optimal operating temperature.

Conclusion on Engine Thermostat Purpose

The primary function of an engine thermostat is to ensure the engine operates within its most efficient and safe temperature range. It manages coolant flow to achieve this balance, preventing the engine from running too cold or overheating.

Revision Table: Engine Thermostat

Component	Primary Function	Why it's Important
Engine Thermostat	Regulates coolant flow based on temperature	Keeps engine at desired/optimal operating temperature for efficiency, performance, and longevity.

Additional Information: Engine Cooling System

The thermostat is part of a larger engine cooling system, which typically includes:

- **Radiator:** Dissipates heat from the coolant to the surrounding air.
- **Water Pump:** Circulates coolant throughout the engine and cooling system.
- **Coolant (Antifreeze):** A mixture of water and antifreeze that absorbs heat from the engine and prevents freezing/boiling.
- **Radiator Fan:** Provides airflow through the radiator when the vehicle is stationary or moving slowly.
- **Hoses:** Connect the components of the cooling system.
- **Coolant Reservoir:** Holds excess coolant and allows for expansion/contraction.

All these components work together with the thermostat to maintain the engine's thermal health.

148. Answer: d

Explanation:

Understanding the symptoms of a failing or poorly operating clutch system is crucial for vehicle maintenance and safe driving. A clutch is a mechanical device that connects and disconnects the engine from the transmission, allowing you to change gears smoothly or stop the vehicle without stalling the engine.

Let's analyze the given options to determine which one is typically **not** a symptom of poor clutch operation.

Analyzing Clutch Operation Symptoms

We will examine each symptom listed in the options and discuss its relation to a poorly operating clutch.

1. Vibration of engine:

Engine vibration during clutch engagement or disengagement can be a symptom of a faulty clutch. This vibration might be caused by uneven wear on the clutch plate, a

damaged pressure plate, or problems with the flywheel. If the clutch is not engaging smoothly, it can cause the engine to shake as it tries to transfer power to the transmission.

2. Difficulty in gear shifting when pedal press:

Difficulty shifting gears, especially when the clutch pedal is fully pressed, is a classic symptom of clutch problems. This usually indicates that the clutch is not fully disengaging the engine from the transmission. Potential causes include improper clutch adjustment, air in the hydraulic system, a faulty master or slave cylinder, or a worn-out clutch plate or pressure plate that isn't releasing completely.

3. pedal gets stuck when pressed:

A clutch pedal that sticks to the floor or doesn't return properly after being pressed is a clear sign of a problem within the clutch system. This could be due to issues with the clutch cable (if applicable), the hydraulic system (leaks, faulty cylinders, air), or mechanical problems within the pedal assembly or linkage. A stuck pedal means you cannot properly disengage the clutch, leading to difficulties in shifting or stopping.

4. increase in engine noise when pedal press:

An increase in engine noise is typically related to the engine itself (combustion, exhaust, internal mechanics) or the engine's RPM. While some noise (like a chirping or grinding) might be heard when pressing the clutch pedal due to issues with the throw-out bearing or pilot bearing, a general "increase in engine noise" (implying louder engine operation itself) when the pedal is simply pressed is not a primary or direct symptom of *poor clutch operation* compared to the other options. Symptoms like slipping clutch cause the engine RPM (and thus noise) to increase *without* a corresponding increase in vehicle speed when trying to accelerate, but the symptom here is specifically tied to the act of pressing the pedal, not necessarily related to acceleration or slipping. Therefore, compared to vibration, difficulty shifting, or a stuck pedal, increased engine noise when just pressing the pedal is the least indicative symptom of a general clutch operating problem.

Based on the analysis, options 1, 2, and 3 are well-known symptoms associated with poor clutch operation. Option 4, "increase in engine noise when pedal press," is not a typical or direct symptom of poor clutch operation itself in the same way the others are described.

Symptom	Related to Poor Clutch Operation?	Explanation
Vibration of engine	Yes	Caused by uneven clutch wear, flywheel issues, etc., affecting smooth engagement.
Difficulty in gear shifting when pedal press	Yes	Clutch not fully disengaging, preventing smooth gear changes.
Pedal gets stuck when pressed	Yes	Problem in clutch hydraulic system, linkage, or pedal mechanism.
Increase in engine noise when pedal press	No (Generally)	More often related to engine issues or specific bearing noise, not general poor clutch operation itself.

Conclusion on Clutch Symptoms

The question asks which of the listed items is **not** a symptom of poor clutch operation. Based on our analysis, vibration, difficulty shifting, and a stuck pedal are common signs of clutch problems. An increase in engine noise when the pedal is pressed is not a standard or direct symptom of poor clutch operation compared to the others.

Revision Table: Clutch System Review

Clutch Component	Function	Symptom if Faulty
Clutch Plate (Friction Disc)	Connects engine to transmission via friction	Slipping, Shuddering, Vibration
Pressure Plate	Clamps clutch plate against flywheel	Difficulty shifting, Shuddering, Vibration
Flywheel	Provides surface for clutch plate, stores rotational energy	Vibration, Shuddering
Release Bearing (Throw-out Bearing)	Engages/disengages clutch when pedal is pressed	Noise (chirping, grinding) when pedal pressed/released
Clutch Fork/Lever	Moves release bearing	Sticking pedal, Difficulty disengaging
Hydraulic System (Master/Slave Cylinder) or Cable	Transmits pedal force to the clutch fork	Sticking pedal, Pedal feels soft/hard, Difficulty disengaging

Additional Information on Clutch Issues

Recognizing the signs of clutch trouble early can prevent more significant damage to the transmission and engine. If you notice any of the common symptoms like difficulty shifting, a strange smell (burning clutch), slipping (engine RPM increases but speed doesn't), or unusual noises when using the clutch, it's advisable to have the clutch system inspected by a qualified mechanic.

Driving with a severely worn or malfunctioning clutch can lead to:

- Damage to the transmission gears.
- Overheating of the engine.
- Potential stranding of the vehicle.

Proper clutch operation involves pressing the pedal fully when shifting gears and releasing it smoothly. Avoid resting your foot on the clutch pedal while driving, as this can cause premature wear (riding the clutch).

149. Answer: d

Explanation:

Understanding Why a Fuse Blows in an Electrical Circuit

A fuse is a critical safety device used in electrical circuits. Its primary purpose is to protect the circuit and connected appliances from damage caused by excessive current flow. When the current exceeds a safe limit, the fuse interrupts the circuit, preventing potential hazards like overheating, fires, or equipment failure.

How a Fuse Protects an Electrical Circuit

A fuse contains a thin wire designed to melt and break the circuit when heated sufficiently by the passage of current. The amount of heat generated in the wire is proportional to the square of the current (I) and the resistance (R) of the wire, according to Joule's law of heating: $\text{Heat} \propto I^2 R t$, where t is time. If the current (I) becomes too high, the heat generated quickly raises the temperature of the fuse wire to its melting point, causing it to 'blow' or rupture.

What Causes High Current in a Circuit?

In a standard electrical circuit operating normally, the total resistance is determined by the components connected (like appliances, lights, etc.) and the wiring. According to Ohm's Law ($V = IR$), the current (I) flowing through the circuit is directly proportional to the voltage (V) across it and inversely proportional to the total resistance (R).

A common condition that leads to a dangerously high current is a **short circuit**. A short circuit occurs when a low-resistance path is created, allowing current to bypass the normal load. This often happens when two wires that should be separated touch each other directly, or when faulty wiring creates a direct connection to the ground. With a very low resistance path, and the voltage remaining relatively constant, the current flow increases dramatically ($I = V/R$, if R is very small, I is very large).

Analyzing the Options for Why a Fuse Blows

Let's examine each option in the context of how fuses work and what causes high current:

- Option 1: **Voltage caused by a short circuit.** While a short circuit involves voltage (it's powered by a voltage source), it's the resulting *high current*, not a high voltage itself caused by the short circuit, that blows the fuse. In fact, the voltage across the shorted part might drop. The fuse blows due to the thermal effect of excessive current. This option is not the direct cause described.
- Option 2: **Current flow caused by an open circuit.** An open circuit occurs when the circuit is broken, such as flipping a switch off or a wire breaking. In an open circuit, the resistance is effectively infinite, and therefore, no current flows ($I = V/\infty = 0$). A fuse requires current flow to heat up and blow, so an open circuit condition would prevent the fuse from blowing. This option is incorrect.
- Option 3: **Voltage caused by an open circuit.** Similar to option 2, an open circuit means no current flow, so the fuse will not blow. Also, voltage itself doesn't blow the fuse; it's the current-induced heat. This option is incorrect.
- Option 4: **Current flow caused by a short circuit.** As explained earlier, a short circuit creates a very low-resistance path, leading to a massive surge in current flow. This high current causes the fuse wire to heat up rapidly and melt, thereby breaking the circuit and preventing damage. This option accurately describes the common scenario leading to a fuse blowing.

Based on the function of a fuse and the nature of electrical faults, a fuse blows specifically because of the high current that flows through it, which is typically a result of a short circuit condition in the electrical circuit.

Condition	Resistance	Current Flow	Effect on Fuse
Normal Operation	Standard (Load resistance)	Normal	No effect (if current is within rating)
Short Circuit	Very Low	Very High	Blows (melts)
Open Circuit	Infinite	Zero	No effect (no current)

Revision Table: Fuse Blowing Explained

Here is a summary of key concepts related to why a fuse blows:

- A fuse is a safety device in an electrical circuit.
- It protects against excessive current.
- Excessive current heats the fuse wire.
- High heat melts the fuse wire, breaking the circuit ('blowing' the fuse).
- A **short circuit** is a common cause of excessive current.
- A short circuit provides a low-resistance path for current.
- High current due to a short circuit leads to the fuse blowing.

Additional Information: Circuit Protection

Circuit protection is vital for electrical safety. Fuses are one type of overcurrent protection device. Another common device is a circuit breaker, which serves the same purpose but can be reset instead of needing replacement. Understanding the difference between normal operation, a short circuit, and an open circuit is fundamental to understanding electrical safety and fault conditions.

- **Overcurrent:** Any current exceeding the normal operating current of a circuit. This can be due to overload (connecting too many devices) or a fault like a short circuit.
- **Fuse Rating:** Fuses have a specific current rating. They are designed to blow when the current consistently exceeds this rating, usually by a certain margin.
- **Circuit Breakers:** These use electromagnetic or thermal mechanisms to trip and open the circuit during an overcurrent condition. They can be manually reset once the fault is cleared.

Protecting electrical circuits with appropriate devices like fuses or circuit breakers is essential to prevent damage to wiring and appliances and, most importantly, to prevent electrical fires.

150. Answer: a

Explanation:

Understanding Rivet Specification

A rivet is a permanent mechanical fastener. Before it is installed, it is typically a smooth cylindrical shaft with a head on one end. The other end is called the tail. When installing a

rivet, the tail end is expanded or deformed (often by hammering or using a rivet gun) to create a new head, permanently fastening two or more pieces of material together.

Key Parts of a Rivet

Every standard rivet has three main parts:

- **Head:** This is the pre-formed shape at one end of the rivet. Heads come in various shapes like round, flat, countersunk, etc.
- **Shank:** This is the cylindrical body of the rivet, located between the head and the tail. This part passes through the holes in the materials being joined.
- **Tail:** This is the end opposite the head. During installation, this end is modified to create the second head.

How Rivet Diameter is Specified

When you need to select a rivet for a particular application, you need to know its size. The size of a rivet is commonly defined by two main dimensions: its diameter and its length.

The diameter of a rivet is always specified by the diameter of its **shank**. This is the critical dimension because it determines the size of the hole that needs to be drilled in the materials being joined. The shank must fit snugly into this hole to ensure a strong joint.

Here's a simple way to think about it:

- The **shank diameter** tells you how wide the cylindrical part is, which directly corresponds to the required hole size.
- The **length** tells you how long the shank is, measured from the underside of the head to the tip of the tail. This length is important to ensure there is enough material to form the second head after passing through the joined materials.
- The **head size** (diameter and height) is also standardized, but it's the shank diameter that is the primary specification for the rivet's thickness or size grade.
- The **tail** is the part deformed during installation and doesn't specify the rivet's initial dimensions.

Why Shank Diameter Matters for Specification

Using the shank diameter as the primary specification for a rivet's thickness is standard practice in engineering and manufacturing. This is because the fit of the shank in the

drilled or punched hole is crucial for the strength and integrity of the riveted joint. A correct fit prevents looseness and distributes the load effectively.

For example, if you need to rivet two plates together using a 5 mm hole, you would choose a rivet with a 5 mm shank diameter.

Rivet Part	Role in Specification
Head	Specifies shape (e.g., round, flat) and contributes to length measurement point.
Shank	Specifies the diameter of the rivet. Determines hole size.
Tail	The end deformed during installation; not used for initial diameter specification.
Length	Specifies the length of the shank.

Therefore, a rivet is specified by its diameter of the shank.

Revision Table: Rivet Specification

Concept	Description
Rivet Definition	Permanent mechanical fastener used to join materials.
Rivet Parts	Head, Shank, Tail.
Diameter Specification	Always specified by the shank diameter.
Length Specification	Measured from the underside of the head to the tail end.
Importance of Shank Diameter	Ensures proper fit in the joint's hole for strength.

Additional Information on Rivet Applications

Rivets are used in many industries where strong, permanent joints are needed, and welding or screwing might not be suitable or possible. Some common applications include:

- Aerospace (aircraft structures)
- Construction (bridges, buildings)
- Automotive (chassis, body panels)
- Manufacturing (various products and assemblies)
- Fabrication work

Different types of rivets exist based on their material (steel, aluminum, copper), head shape (dome, countersunk, universal), and installation method (solid, blind, tubular). Regardless of the type, the fundamental specification of its diameter is tied to the shank.

151. Answer: b

Explanation:

Understanding the Principle of Brakes

Brakes are essential components in vehicles and machinery that allow us to slow down or stop movement. They work by opposing motion, and the primary physical principle behind most braking systems is friction.

What is Friction?

Friction is a force that resists the relative motion or tendency for motion of two surfaces in contact. When two surfaces rub against each other, friction acts to slow down or stop their movement.

How Brakes Utilize Friction

In typical braking systems, friction is deliberately created between two surfaces. For example, in disc brakes, brake pads are pressed against a rotating disc (rotor) attached to the wheel. In drum brakes, brake shoes are pressed against the inside of a rotating drum.

- When the brake pedal is applied, hydraulic pressure (or mechanical linkage) forces the friction material (pads or shoes) into contact with the rotating component (disc or drum).
- The resulting friction between these surfaces opposes the rotation of the wheel.
- This friction converts the kinetic energy (energy of motion) of the vehicle or machine into heat energy, which is then dissipated.
- As the kinetic energy is reduced, the vehicle or machine slows down or stops.

Let's look at the given options in the context of how brakes work:

- **Vibration:** While braking might cause some vibration under certain conditions (like warped rotors), vibration itself is not the principle by which brakes primarily function to stop motion.
- **Friction:** This is the fundamental principle. The force of friction between brake components is what opposes the motion and slows down the object.
- **Suction:** Suction involves creating a pressure difference (a vacuum). This principle is not used in standard braking systems to stop motion, although vacuum boosters are used in some vehicles to *assist* the driver in applying brake pressure, they don't provide the stopping force themselves.
- **Dragging:** Dragging is a descriptive term for a type of resistance, and it is a consequence of friction. When brakes are applied, the friction between surfaces causes a "dragging" effect that opposes motion. However, friction is the underlying physical principle creating this drag.

Therefore, the most accurate answer describing the principle on which brakes work is friction.

Revision Table: Key Concepts of Brakes

Concept	Explanation	Role in Braking
Friction	A force that opposes relative motion between surfaces in contact.	The primary force used to slow down/stop motion.
Brake Pads/Shoes	Components made of friction material.	Press against rotating parts to create friction.
Brake Disc/Drum	Rotating components attached to the wheel.	Surface against which pads/shoes press.
Kinetic Energy	Energy of motion.	Converted into heat energy by friction during braking.
Heat Dissipation	Process of removing heat generated by friction.	Prevents brake fade and damage.

Additional Information on Braking Principles

While friction is the main principle, braking systems involve several other concepts and components to function effectively and safely:

- **Hydraulics:** Most modern vehicle brakes use hydraulic fluid to transmit the force from the brake pedal to the brake pads or shoes. This allows a small force applied by the driver to be multiplied into a large force at the wheels.
- **Brake Fade:** This occurs when brakes become less effective due to overheating caused by excessive friction. The heat reduces the friction coefficient of the materials.
- **Braking Systems:** There are different types of braking systems, including disc brakes, drum brakes, and electromagnetic brakes (used in some trains and electric vehicles). Each uses friction or a related principle (like electromagnetic resistance) in a specific way.
- **Anti-lock Braking System (ABS):** This system prevents the wheels from locking up during braking by modulating the brake pressure. While it works in conjunction with the friction braking system, its principle is based on preventing skidding by maintaining traction.

Understanding friction is fundamental to understanding how most brakes achieve their function of slowing down or stopping motion.

152. Answer: a

Explanation:

Understanding Air Chisel Operation

An air chisel is a powerful tool primarily used for breaking, cutting, or shaping metal, stone, or concrete. Unlike manual chisels that rely on physical force applied by hand or hammer, an air chisel utilizes a different mechanism to deliver rapid, forceful blows to the chisel bit. Understanding how it operates is key to identifying the correct power source.

How Air Chisels Function

Air chisels belong to a category of tools known as pneumatic tools. The term "pneumatic" relates to or uses compressed air or other gas. Air chisels work by converting the energy from compressed air into rapid mechanical movement. Inside the tool, a piston moves back and forth rapidly, striking the end of the chisel bit. This high-speed hammering action is what allows the air chisel to break or cut through materials effectively.

Analyzing the Operation Methods

Let's examine the given options to determine which best describes how an air chisel is operated.

- **Pneumatically:** This term means operated by gas pressure, typically compressed air. As discussed, air chisels use compressed air to power a piston that strikes the chisel bit. This aligns perfectly with the working principle of an air chisel.
- **Magnetically:** This refers to operation using magnetic fields. While magnetism is used in some tools or processes, it is not the primary operating principle for delivering the impact needed by an air chisel.
- **Frictionally:** This relates to operation based on friction. Friction is a force that opposes motion between surfaces. While friction might be present as a side effect during the cutting process, it is not the mechanism that powers the rapid hammering action of the air chisel itself.
- **Hydraulically:** This means operated by liquid pressure, typically oil or water. Hydraulic tools use fluid pressure to move pistons or motors. While powerful, hydraulic systems

are distinct from pneumatic systems, which use gas pressure. Air chisels specifically use compressed air.

Based on the analysis, the operation of an air chisel relies on compressed air to generate the necessary force and rapid movement. This is the definition of pneumatic operation.

Conclusion on Air Chisel Operation

An air chisel breaks or cuts metal objects apart by using compressed air to power its striking mechanism. Therefore, it is operated pneumatically.

Revision Table: Tool Operation Methods

Method	Description	Examples of Tools
Pneumatic	Uses compressed gas (like air) to produce mechanical energy.	Air chisel, nail gun, impact wrench, air drill.
Hydraulic	Uses pressurized liquid (like oil) to transmit power.	Hydraulic jack, hydraulic press, excavators (many functions).
Electric	Uses electrical energy to power a motor.	Electric drill, circular saw, angle grinder.
Manual	Operated solely by human physical strength.	Hammer, hand saw, screwdriver.

Additional Information on Pneumatic Tools

Pneumatic tools are widely used in construction, automotive repair, manufacturing, and other industries. They offer several advantages:

- **High Power-to-Weight Ratio:** They can deliver significant power while being relatively lightweight compared to electric tools of similar output.
- **Durability:** With fewer moving parts than electric motors, pneumatic tools are often more durable and less prone to overheating.
- **Safety:** They do not contain electric motors, reducing the risk of electric shock, especially in wet environments. They also don't generate sparks, making them safer in environments with flammable materials.

- **Simplicity:** Their design is often simpler, leading to easier maintenance.

However, they require an air compressor and air hoses to operate, which adds to the initial setup cost and limits mobility to the length of the hose or the portability of the compressor.

153. Answer: a

Explanation:

Understanding Diesel Engine Smoke Emissions

Diesel engines are known for their efficiency and torque, but they can sometimes produce visible exhaust smoke, especially under certain operating conditions. Understanding when and why this smoke occurs is important for diagnosing engine health and understanding emissions.

Causes of Smoke in Diesel Engines

Smoke from a diesel engine is primarily caused by incomplete combustion of fuel. This happens when there isn't enough oxygen available in the combustion chamber to burn all the injected fuel effectively. The result is the formation of soot particles, which are seen as black smoke.

Other types of smoke (white or blue) can indicate different issues, but excessive black smoke is the most common visible emission during operation.

Analyzing Operating Conditions and Smoke Emissions

Let's consider the given operating conditions:

- **Accelerating:** When a diesel engine accelerates, the driver demands more power quickly. The engine control unit (ECU) responds by injecting a larger amount of fuel into the cylinders. While the turbocharger spools up to supply more air, there is a brief period where the sudden increase in fuel injection outpaces the available air supply. This creates a temporary rich fuel-air mixture, leading to incomplete combustion and the emission of excessive black smoke (soot).

- **Decelerating:** During deceleration, the fuel supply is typically reduced or even cut off (fuel cut-off) while the engine is still rotating and drawing in air. This results in a very lean mixture or just air passing through the engine. Therefore, incomplete combustion is unlikely, and excessive smoke is not usually emitted.
- **Starting:** During engine starting, especially in cold conditions, the combustion process might be inefficient. White or light blue smoke can be emitted due to unburnt fuel (white smoke) or burning of lubricating oil (blue smoke), but this is different from the black smoke caused by a rich mixture during acceleration. Once the engine warms up and runs steadily, this type of smoke usually disappears.
- **Idling:** At idle speeds, the engine is running under very low load and requires only a small amount of fuel to maintain its speed. The amount of air available is relatively high compared to the fuel injected. The fuel-air mixture is lean, and combustion is generally more complete than during high load or transient conditions. Excessive black smoke is not typical during stable idling in a well-maintained engine.

Why Acceleration Causes Excessive Smoke

The key reason excessive smoke is often observed during acceleration in diesel engines is the rapid transition from a lower load (and potentially leaner mixture) to a higher load where a large amount of fuel is suddenly injected. The air supply, regulated by the turbocharger, takes a moment longer to catch up to the increased fuel demand. This transient rich condition directly causes the production of soot.

Comparing the conditions, acceleration presents the scenario where the fuel-air ratio is most likely to become temporarily rich, leading to excessive incomplete combustion and visible black smoke.

Revision Table: Diesel Engine Smoke by Condition

Operating Condition	Fuel Injection	Air Supply	Fuel-Air Mixture Tendency	Typical Smoke Emission
Accelerating	Rapidly Increases	Increases (with delay)	Temporarily Rich	Excessive Black Smoke (Soot)
Decelerating	Reduced/Cut Off	High relative to fuel	Very Lean	Minimal/No Smoke
Starting (Cold)	Injected	Available	Inefficient Combustion	White/Blue Smoke possible
Idling	Low	Sufficient	Lean	Minimal/No Smoke (Black)

Additional Information: Types of Diesel Smoke and Causes

- **Black Smoke:** Indicates incomplete combustion, typically due to a rich fuel-air mixture (too much fuel for the available air). Common causes include clogged air filters, faulty injectors, turbocharger issues, or heavy acceleration.
- **White Smoke:** Often seen during cold starts. It is usually unburnt fuel that has vaporized but not ignited. Can also be caused by coolant entering the combustion chamber or poor compression.
- **Blue Smoke:** Indicates that engine oil is being burnt. This can happen if oil is leaking past piston rings, valve guides, or seals and entering the combustion chamber.

While modern diesel engines with advanced fuel injection systems and emission controls (like Diesel Particulate Filters – DPFs) significantly reduce visible smoke, the fundamental physics of combustion under transient conditions like acceleration still make this the most likely time for excessive black smoke to occur, especially in older or poorly maintained engines, or when the emission control system is compromised.

154. Answer: d

Explanation:

Understanding the Diesel Engine Fuel System

In a diesel engine, the fuel system is responsible for delivering fuel from the fuel tank to the combustion chamber under high pressure at the correct time. This system involves several components working together, including the fuel tank, fuel lines, filters, feed pump, and the fuel injection pump (FIP).

Role of the Feed Pump in a Diesel Engine

The feed pump, also known as the fuel transfer pump, is a vital part of the diesel fuel system. Its main function is to draw fuel from the fuel tank and supply it under a relatively low pressure to the fuel injection pump (FIP). The FIP then increases the pressure significantly before injecting the fuel into the engine cylinders. The feed pump ensures that the FIP has a constant supply of fuel to operate correctly.

Analyzing Feed Pump Location Options

Let's examine the potential locations mentioned in the options:

- **Fuel tank:** The fuel tank stores the fuel. While the feed pump draws fuel from the tank, it is not usually mounted directly on the tank itself, especially in engine applications where it's part of the engine-mounted fuel system.
- **Fuel filter:** Fuel filters are used to remove contaminants from the fuel. The feed pump typically pumps fuel through one or more filters, but the pump is a separate component and not mounted on the filter housing.
- **Cylinder head:** The cylinder head is located on top of the engine block and houses components like valves and fuel injectors (in direct injection diesel engines). The feed pump is not located on the cylinder head.
- **F.I.P (Fuel Injection Pump):** The Fuel Injection Pump is a complex pump responsible for pressurizing and distributing fuel to the injectors. The feed pump is usually mounted directly on the FIP housing or is integrated within the FIP unit. It is often driven by the same camshaft or mechanism that drives the FIP itself. This close proximity ensures a direct and reliable supply of fuel to the high-pressure pump.

Why the Feed Pump is Located on the F.I.P

Mounting the feed pump directly onto the Fuel Injection Pump (FIP) is a common design choice in many diesel engines. This integration simplifies the fuel lines between the two pumps and allows them to be driven by a common power source, often a gear or camshaft connected to the engine. The feed pump ensures the FIP reservoir is kept full

and pressurized, preventing cavitation and ensuring consistent performance of the high-pressure injection system.

Component	Typical Function	Likely Location Relative to Feed Pump
Fuel Tank	Stores diesel fuel	Feed pump draws from it (distant)
Fuel Filter	Removes contaminants	Feed pump pumps through it (between tank and FIP)
Cylinder Head	Houses valves, injectors	Distant from feed pump
F.I.P (Fuel Injection Pump)	Pressurizes & distributes fuel	Feed pump is typically mounted on it or integrated

Based on the function of the feed pump to supply fuel to the FIP, and common diesel engine design practices, the feed pump is most commonly located on the F.I.P.

Diesel Fuel System Revision Table

Component	Function
Fuel Tank	Stores diesel fuel.
Feed Pump (Low Pressure)	Draws fuel from tank, supplies to FIP at low pressure.
Fuel Filters	Clean fuel by removing particles.
Fuel Injection Pump (FIP) (High Pressure)	Pressurizes fuel for injection & controls timing/quantity.
Fuel Injectors	Spray atomized fuel into cylinder/combustion chamber.
Fuel Lines	Carry fuel between components.

Additional Information on Diesel Engine Fuel Supply

The pressure created by the feed pump is relatively low, typically just enough to overcome the resistance in the fuel lines, filters, and lift the fuel from the tank. The real high pressure needed for injection is generated by the Fuel Injection Pump (FIP) or by individual injectors in modern common rail systems. In systems with a separate feed pump, its reliable operation is critical for the FIP to function correctly and ensure efficient engine performance.

155. Answer: b

Explanation:

Understanding Tolerance in Engineering

In engineering and manufacturing, achieving the exact theoretical dimension of a part is practically impossible due to variations in processes, tools, and materials. To ensure that parts can still function correctly, engineers define acceptable limits for these variations. This allowed variation is known as **tolerance**.

Basic Dimension and Tolerance

The **basic dimension** is the theoretical exact size of a feature. The actual size of the manufactured part will vary slightly from this basic dimension. Tolerance specifies how much this variation is allowed to be.

Exploring Different Types of Tolerance

Tolerance can be applied to the basic dimension in different ways. The question asks about the specific case where tolerance is given on only one side of the basic dimension.

What is Unilateral Tolerance?

Unilateral tolerance is a system where the total tolerance is applied in only one direction from the basic dimension. This means the acceptable range of sizes is either entirely above the basic dimension or entirely below it, but not crossing the basic dimension unless one limit is exactly the basic dimension itself.

For example, if the basic dimension is 20 mm, a unilateral tolerance might be specified as:

- $20.00^{+0.10}$ mm (Tolerance is +0.10 mm, applied only above 20.00 mm. Acceptable range: 20.00 mm to 20.10 mm)
- $20.00_{-0.05}$ mm (Tolerance is -0.05 mm, applied only below 20.00 mm. Acceptable range: 19.95 mm to 20.00 mm)

In unilateral tolerance, one of the tolerance limits (either upper or lower) often coincides with the basic dimension.

What is Bilateral Tolerance?

In contrast, **bilateral tolerance** is a system where the total tolerance is divided and applied in both directions from the basic dimension. The acceptable range of sizes extends both above and below the basic dimension.

For example, if the basic dimension is 20 mm, a bilateral tolerance might be specified as:

- 20.00 ± 0.05 mm (Tolerance is ± 0.05 mm. Acceptable range: 19.95 mm to 20.05 mm)
- $20.00^{+0.07}_{-0.03}$ mm (Tolerance is +0.07 mm above and -0.03 mm below. Acceptable range: 19.97 mm to 20.07 mm)

In bilateral tolerance, the basic dimension is typically somewhere within the acceptable range, often (but not always) at the center.

Other Concepts: Allowance and Tolerance System

- **Allowance System:** Allowance refers to the intentional difference between the dimensions of mating parts, such as a shaft and a hole, to achieve a desired fit (e.g., clearance fit, interference fit). It is related to tolerance but focuses on the relationship between parts rather than the variation of a single feature.
- **Tolerance System:** This is a broader term that refers to the overall framework, standards, and practices used for specifying and managing tolerances in engineering designs and manufacturing. It encompasses various aspects, including unilateral and bilateral tolerance methods.

Conclusion

Based on the definitions, when tolerance is specified only on one side of the basic dimension, it is clearly defined as unilateral tolerance.

Type of Tolerance	Description	Example (Basic Dim = 20)
Unilateral Tolerance	Variation allowed on only one side of the basic dimension.	$20.00^{+0.10}$, $20.00_{-0.05}$
Bilateral Tolerance	Variation allowed on both sides of the basic dimension.	20.00 ± 0.05 , $20.00^{+0.07}_{-0.03}$

Revision Table: Key Tolerance Terms

Term	Meaning
Tolerance	Permissible variation in a dimension.
Basic Dimension	Theoretical exact size.
Unilateral Tolerance	Tolerance on one side of the basic dimension.
Bilateral Tolerance	Tolerance on both sides of the basic dimension.
Allowance	Intentional difference between mating parts.

Additional Information on Tolerance Application

The choice between unilateral and bilateral tolerance often depends on the design requirements and manufacturing processes. Unilateral tolerance is frequently used when a critical size limit is essential, such as maintaining a minimum wall thickness or ensuring a specific type of fit by fixing one boundary of the tolerance zone at the basic size.

For instance, in a clearance fit, if the hole size is given a unilateral tolerance above its basic dimension and the shaft size is given a unilateral tolerance below its basic dimension, it guarantees that there will always be some clearance between the parts.

156. Answer: c

Explanation:

Understanding Engine Mechanical Efficiency

The question asks about the ratio between Brake Horsepower (BHP) and Indicated Horsepower (IHP) in an engine. Let's break down what these terms mean and how their ratio is defined.

What are BHP and IHP?

- **Indicated Horsepower (IHP):** This is the theoretical power developed inside the combustion chambers of an engine. It is calculated based on the pressure-volume diagram of the engine cycle. Think of it as the total power generated by the expansion of gases during combustion.
- **Brake Horsepower (BHP):** This is the actual usable power delivered by the engine at the crankshaft (or flywheel). It is the power measured by a dynamometer and is less than the IHP due to losses within the engine.

The Ratio: BHP to IHP

The difference between the Indicated Horsepower (IHP) and the Brake Horsepower (BHP) is the power lost due to friction within the engine. This friction comes from various moving parts like pistons against cylinder walls, bearings, gears, etc., as well as the power needed to drive engine accessories like the oil pump, water pump, etc. This lost power is called Frictional Horsepower (FHP).

$$\text{Frictional Horsepower (FHP)} = \text{IHP} - \text{BHP}$$

The ratio of the power delivered (BHP) to the power developed inside the cylinders (IHP) tells us how effectively the engine converts the internal power into useful output power. This ratio is a measure of the engine's **mechanical efficiency**.

The formula for mechanical efficiency (η_{mech}) is:

$$\eta_{mech} = \frac{\text{BHP}}{\text{IHP}}$$

Analyzing the Options

Let's look at the given options in the context of engine efficiencies:

- **Engine efficiency:** This is a very broad term and doesn't specifically define the ratio of BHP to IHP. Engine efficiency can refer to overall efficiency, thermal efficiency, or mechanical efficiency depending on the context.

- **Power efficiency:** Similar to engine efficiency, this is a general term and not a specific technical definition for the ratio of BHP to IHP.
- **Mechanical efficiency:** As explained above, this is the precise term used to define the ratio of Brake Horsepower (usable power) to Indicated Horsepower (power developed inside cylinders). It quantifies the power losses due to friction and other mechanical resistances.
- **Thermal efficiency:** This measures how effectively the engine converts the heat energy from burning fuel into mechanical work (power). It is typically the ratio of IHP or BHP to the heat supplied by the fuel. This is different from the ratio of BHP to IHP.

Based on the definitions, the ratio between BHP and IHP is specifically called Mechanical efficiency.

Summary of Engine Efficiencies

Efficiency Type	Definition	Formula (Common)
Mechanical Efficiency	Ratio of Brake Power to Indicated Power. Accounts for frictional losses.	$\eta_{mech} = \frac{BHP}{IHP}$
Thermal Efficiency (Indicated)	Ratio of Indicated Power to Heat Input from fuel.	$\eta_{th,i} = \frac{IHP}{\text{Heat Input}}$
Thermal Efficiency (Brake)	Ratio of Brake Power to Heat Input from fuel.	$\eta_{th,b} = \frac{BHP}{\text{Heat Input}}$

Conclusion

The ratio between BHP and IHP is a direct measure of the mechanical efficiency of an engine, indicating how much of the power generated internally is delivered as usable power at the output shaft after overcoming frictional losses.

Revision Table: Key Engine Performance Terms

Key Terms in Engine Performance

Term	Description	Formula/Relation
Indicated Horsepower (IHP)	Theoretical power developed inside the engine cylinders.	Based on cylinder pressure & volume
Brake Horsepower (BHP)	Actual usable power delivered at the engine crankshaft.	Measured by dynamometer
Frictional Horsepower (FHP)	Power lost due to friction and accessories.	$FHP = IHP - BHP$
Mechanical Efficiency (η_{mech})	Ratio of useful power output to total power developed.	$\eta_{mech} = \frac{BHP}{IHP}$

Additional Information: Engine Power and Efficiency

Understanding different types of engine power and efficiency is crucial for analyzing engine performance. While mechanical efficiency focuses on internal power losses, thermal efficiency looks at how efficiently the fuel's energy is converted into heat and then into work.

Brake thermal efficiency ($\eta_{th,b}$) is often considered the overall engine efficiency because it relates the final useful output (BHP) to the energy input (heat from fuel). It can also be expressed using indicated thermal efficiency and mechanical efficiency:

$$\eta_{th,b} = \eta_{th,i} \times \eta_{mech}$$

This relationship highlights that the final efficiency is a product of how well heat is converted to power inside the cylinder (indicated thermal efficiency) and how much of that power makes it out of the engine after overcoming mechanical losses (mechanical efficiency).

157. Answer: c

Explanation:

Understanding Governors: Speed Sensitive Devices Explained

A governor is a crucial component in many engines and machinery. Its primary function is to regulate or control the speed of an engine or machine. It does this automatically by adjusting the fuel supply or load to maintain a desired speed, despite variations in load or other operating conditions.

Let's analyze the options provided:

- **Pressure sensitive device:** A pressure sensitive device reacts to changes in pressure. While pressure can be a factor in engine operation, it's not the primary characteristic that defines a governor's function.
- **Temperature sensitive device:** A temperature sensitive device reacts to changes in temperature. While temperature affects engine performance, a governor's main role is not to respond to temperature fluctuations.
- **Speed sensitive device:** A speed sensitive device reacts to changes in speed. This perfectly describes a governor, as it is designed to detect and respond to variations in the rotational speed of an engine or mechanism.
- **Vacuum sensitive device:** A vacuum sensitive device reacts to changes in vacuum pressure. Similar to general pressure, vacuum is not the primary parameter a governor monitors and controls.

Based on its function, a governor is fundamentally a device that senses and reacts to changes in speed to maintain stability. For example, in an engine, if the load decreases, the engine speed would naturally increase. A governor senses this increase in speed and reduces the fuel supply, bringing the speed back down to the desired level. Conversely, if the load increases, speed tends to drop, and the governor increases the fuel supply.

Therefore, the most accurate description of a governor from the given options is that it is a speed sensitive device.

Governor Sensitivity Comparison

Device Type	Primary Sensitivity	Relevant to Governor?
Pressure Sensitive	Pressure	No (Not Primary)
Temperature Sensitive	Temperature	No (Not Primary)
Speed Sensitive	Speed	Yes (Primary Function)
Vacuum Sensitive	Vacuum Pressure	No (Not Primary)

Revision Table: Key Concepts about Governors

Concept	Description
What is a Governor?	A device that automatically regulates the speed of an engine or machine.
Primary Function	To maintain a constant or desired operating speed despite load variations.
How it Works (General)	Senses speed, compares it to a set point, and adjusts power input (like fuel) or load.
Key Characteristic	Speed sensitive.

Additional Information on Governor Types and Applications

Governors come in various types, including:

- **Mechanical Governors:** Often use centrifugal force generated by rotating weights to sense speed.
- **Hydraulic Governors:** Use hydraulic pressure to control the speed mechanism, often actuated by a speed-sensing element.
- **Electronic Governors:** Use electronic sensors to measure speed and microprocessors to control actuators (like fuel injectors or throttle bodies).

Governors are widely used in:

- Engines (vehicles, generators, industrial machinery)
- Turbines (steam, gas, hydro)
- Pumps and compressors
- Any system where maintaining a stable operating speed is critical for efficiency, safety, or performance.

Understanding that a governor is primarily a speed sensitive device is key to understanding its fundamental role in control systems.

158. Answer: a

Explanation:

Understanding Diesel Engine Electric Starting Systems

Starting a diesel engine, especially a large one, requires a significant amount of initial torque to overcome the compression resistance and inertia. This is where the electric starting system plays a crucial role. It uses an electric motor, known as the starter motor, to crank the engine until it starts running on its own power.

Powering the Diesel Engine Starter Motor

The starter motor in a diesel engine needs a powerful burst of electrical energy to turn the crankshaft quickly enough to initiate combustion. This energy must be readily available even when the engine is not running.

Let's look at the options provided:

- **Battery:** A battery is an electrochemical device that stores electrical energy in chemical form and releases it as electrical energy when needed. Vehicle batteries, particularly those for diesel engines, are designed to provide a high surge of current for a short period, which is exactly what a starter motor requires.
- **Rectifier:** A rectifier is a device that converts alternating current (AC) into direct current (DC). It is typically part of the charging system (working with the alternator/generator) but does not store energy or provide power for starting.

- **Dynamo:** A dynamo (or DC generator) is a device that generates electrical power by converting mechanical energy into electrical energy. In a vehicle, it is driven by the engine after it has started. It cannot provide power to start the engine because it requires the engine to be running to generate electricity.
- **Generator:** Similar to a dynamo (though often referring to AC generators or alternators in modern vehicles), a generator produces power while the engine is running. It is part of the charging system but cannot power the starter motor to initiate the engine's movement.

Based on the function of each component, the device that provides the necessary electrical power to the starter motor to crank the diesel engine during the starting process is the battery.

Component	Function in Starting	Provides Power for Starter?
Battery	Stores and provides electrical energy on demand.	Yes, provides high current surge.
Rectifier	Converts AC to DC.	No, part of charging system, not power source for starting.
Dynamo/Generator	Generates electricity when driven by engine.	No, requires engine to be running to produce power.

Therefore, the battery is the essential component that supplies power to the starter motor for starting a diesel engine.

Revision Table: Diesel Engine Starting Components

Component	Role in Starting System
Battery	Provides initial high current to the starter motor.
Starter Motor	Electric motor that physically cranks the engine crankshaft.
Starter Solenoid	Acts as a heavy-duty relay and engages the starter pinion with the flywheel.
Ignition Switch	Activates the starter circuit when turned to the "start" position.

Additional Information on Diesel Engine Starting

Starting diesel engines can sometimes be more challenging than petrol engines, especially in cold weather. This is because diesel engines rely on compression ignition, meaning the air in the cylinder is compressed to a very high temperature, and fuel is injected into this hot air to ignite. The starter motor must turn the engine fast enough to achieve this high compression temperature.

- **Glow Plugs:** Many diesel engines use glow plugs in each cylinder. These are electric heaters that warm the combustion chamber before starting, especially in cold conditions, making ignition easier. While important for starting, glow plugs are powered by the battery but do not power the starter motor itself.
- **Starting Aid Systems:** Some heavy-duty diesel engines use starting aids like ether (starting fluid) sprays or block heaters (to warm the engine coolant and block) in extremely cold environments to assist the starting process, but the primary power for cranking still comes from the battery to the starter motor.

Understanding the function of each part of the electric starting system, particularly the high-power output requirement of the starter motor, clarifies why the battery is the necessary power source.

159. Answer: c

Explanation:

Coolants for Reaming Aluminum Workpieces Explained

Reaming is a finishing operation used to enlarge or clean up a hole to a precise size and smooth finish. When reaming aluminum, selecting the correct coolant is crucial for achieving the desired accuracy, surface finish, and tool life. Aluminum is a relatively soft material, and its tendency to gall (weld itself to the cutting tool) requires a coolant that provides excellent lubrication in addition to cooling.

Why Coolant is Important for Reaming Aluminum

Using a suitable coolant during the reaming process for aluminum offers several benefits:

- **Cooling:** It dissipates the heat generated by friction and cutting, preventing the workpiece and the reamer from overheating.
- **Lubrication:** It reduces friction between the reamer and the workpiece material, which is especially important for preventing galling when machining aluminum. Good lubrication helps the reamer cut cleanly and prevents material build-up on the cutting edges.
- **Chip Removal:** It helps to flush away the chips from the cutting zone, preventing them from interfering with the cutting action or scratching the finished surface.
- **Improved Surface Finish:** Proper cooling and lubrication lead to a smoother and more accurate finished hole.
- **Extended Tool Life:** Reduced heat and friction prolong the life of the reamer.

Analyzing Coolant Options for Aluminum Reaming

Let's consider the effectiveness of common coolants when reaming aluminum:

- **Air pressure:** While useful for clearing chips in some operations, air pressure provides minimal cooling and no lubrication. It is not suitable for reaming aluminum where lubrication is key to prevent galling.
- **Water:** Water provides good cooling but lacks significant lubricating properties. Pure water can even cause corrosion. Soluble oils mixed with water are better, but plain water is generally not ideal for reaming aluminum due to the galling issue.
- **Kerosene:** Kerosene is known for its excellent lubricating properties, particularly when machining non-ferrous metals like aluminum. Its low viscosity allows it to penetrate the cutting zone effectively, providing good lubrication and helping to prevent material from sticking to the reamer (galling). It also provides some cooling.
- **Lard oil:** Lard oil is a traditional cutting oil known for its good lubricating qualities. However, it is less common now due to its odor, stability issues, and cost compared to synthetic or semi-synthetic oils. While it can provide lubrication, kerosene is a widely accepted and effective choice specifically for aluminum reaming.

Based on the need for excellent lubrication to prevent galling when reaming aluminum, kerosene stands out as a highly effective and commonly recommended coolant.

Coolant Option	Cooling	Lubrication	Chip Removal	Suitability for Aluminum Reaming
Air pressure	Poor	None	Good (for clearing)	Poor (No lubrication)
Water	Good	Poor	Good (flushing)	Poor (Risk of galling)
Kerosene	Moderate	Excellent	Good (flushing)	Excellent (Prevents galling)
Lard oil	Moderate	Good	Moderate	Good (Traditional)

Therefore, kerosene is the most appropriate coolant among the given options for reaming an aluminum workpiece.

Revision Table: Aluminum Reaming Coolants

Key points regarding coolants for aluminum reaming:

- Aluminum requires a coolant with strong lubricating properties.
- Kerosene effectively prevents galling and provides good lubrication.
- Water provides cooling but lacks lubrication for aluminum.
- Air pressure is not suitable as it provides no lubrication.

Additional Information: Machining Aluminum

Aluminum alloys are widely used due to their light weight, strength, and corrosion resistance. However, they present specific challenges during machining, particularly the tendency to build up material on the cutting tool edges (galling or built-up edge). This makes lubrication critical for processes like drilling, tapping, and reaming. Other common coolants or cutting fluids used for aluminum machining include:

- Soluble oils (emulsions of oil and water)
- Synthetic cutting fluids
- Straight oils (like mineral oil, sometimes with additives)

The choice often depends on the specific aluminum alloy, the machining operation, desired surface finish, and environmental considerations. However, for operations like reaming where a smooth, accurate finish is paramount and galling is a significant risk,

high-lubricity coolants are preferred. Kerosene is a classic example of such a coolant used for aluminum.

160. Answer: a

Explanation:

Understanding the Diesel Engine Suction Stroke

Internal combustion engines, including diesel engines, typically operate through a cycle of four strokes. These strokes are: suction (or intake), compression, power (or expansion), and exhaust. Each stroke plays a crucial role in converting fuel energy into mechanical work. Let's focus on the initial stroke in a diesel engine – the **suction stroke**.

What Enters the Cylinder During Suction?

During the **suction stroke** of a diesel engine, the piston moves downwards inside the **cylinder**. This downward movement increases the volume within the cylinder, creating a pressure lower than the atmospheric pressure. The intake valve opens, and the higher atmospheric pressure forces a working fluid into the **cylinder**.

A key difference between a diesel engine and a petrol (or spark-ignition) engine lies in what this working fluid is during the **suction stroke**.

- In a petrol engine, a pre-mixed mixture of air and fuel is drawn into the **cylinder** during suction.
- In a diesel engine, however, only **pure air alone** is drawn into the **cylinder** during the **suction stroke**.

The diesel fuel is injected directly into the combustion chamber later, during the compression stroke, not mixed with the air before it enters the cylinder.

Why Only Pure Air in a Diesel Engine Suction Stroke?

The reason only **pure air alone** is drawn into the **cylinder** of a diesel engine during the **suction stroke** is related to its combustion process. Diesel engines rely on compression ignition. During the compression stroke, the piston moves upwards, compressing the air

drawn in during suction to a very high pressure and temperature. Once the air reaches a sufficiently high temperature, diesel fuel is injected into this hot, compressed air. The fuel ignites spontaneously upon contact with the hot air without needing a spark plug (as in petrol engines).

If fuel were present during the entire compression stroke from the beginning (like in a petrol engine), it could potentially ignite prematurely under the very high compression ratios used in diesel engines, leading to engine knocking or damage. By drawing in only **pure air alone** during suction and injecting fuel later, the diesel engine controls the ignition process precisely.

Comparing Suction Stroke: Diesel vs. Petrol Engines

Here is a simple comparison of what is drawn into the cylinder during the suction stroke for the two common types of internal combustion engines:

Engine Type	Substance Drawn In During Suction Stroke
Diesel Engine	Pure air alone
Petrol Engine	Mixture of air and fuel

The Diesel Engine Suction Stroke Process

Let's break down the steps during the suction stroke in a **diesel engine**:

1. The piston is at the top of its travel (Top Dead Center or TDC).
2. The intake valve opens.
3. The piston moves downwards towards the bottom of its travel (Bottom Dead Center or BDC).
4. The volume inside the **cylinder** increases, creating vacuum pressure.
5. Atmospheric pressure pushes **pure air alone** into the **cylinder** through the open intake valve.
6. The piston reaches BDC, and the cylinder is filled with **pure air**.
7. The intake valve closes, sealing the cylinder for the next stroke (compression).

Therefore, the substance drawn into the cylinder of a diesel engine during the suction stroke is indeed **pure air alone**.

Revision Table: Diesel Engine Strokes

Stroke	Piston Movement	Valves (Intake/Exhaust)	Process	Cylinder Content (Start of Stroke)	Cylinder Content (End of Stroke)
Suction (Intake)	TDC to BDC (Down)	Intake Open, Exhaust Closed	Draws in working fluid	Residual Exhaust Gas	Pure Air Alone
Compression	BDC to TDC (Up)	Intake Closed, Exhaust Closed	Compresses working fluid	Pure Air	Compressed, Hot Air
Power (Expansion)	TDC to BDC (Down)	Intake Closed, Exhaust Closed	Fuel injected and ignites, expanding gases push piston	Compressed Air + Injected Fuel	Hot Exhaust Gases
Exhaust	BDC to TDC (Up)	Intake Closed, Exhaust Open	Expels exhaust gases	Hot Exhaust Gases	Residual Exhaust Gas

Your Personal Exams Guide

Additional Information on Diesel Engine Operation

Understanding the **suction stroke** of a diesel engine is fundamental to grasping its overall operation. Here are some related points:

- **Fuel Injection:** As mentioned, diesel fuel is injected later in the cycle (during compression or slightly before power stroke). This is done using a high-pressure fuel injector.
- **Compression Ratio:** Diesel engines typically have much higher compression ratios than petrol engines (e.g., 14:1 to 25:1 compared to 8:1 to 12:1 for petrol). This high compression is necessary to raise the air temperature sufficiently for auto-ignition of the diesel fuel.

- **Air-Fuel Ratio Control:** Unlike petrol engines which often run close to stoichiometric ratios for efficient combustion and emissions control using a catalytic converter, diesel engines often run lean (excess air). The power output in a diesel engine is primarily controlled by the amount of fuel injected, not by throttling the air intake (like in many petrol engines).
- **Turbocharging:** Diesel engines often use turbochargers to force more air into the **cylinder** during the **suction stroke** (this is called forced induction), increasing power output and efficiency.

The correct answer is that **pure air alone** is drawn into the cylinder of a diesel engine during the suction stroke.

161. Answer: d

Explanation:

Understanding Micrometer Measurement Precision

A micrometer, often called a micrometer screw gauge, is a precision measuring instrument used to measure small distances or thicknesses with high accuracy. It works based on the principle of a screw, converting rotational motion into linear motion. The accuracy of a micrometer is significantly higher than that of instruments like rulers or Vernier calipers.

What a Micrometer Measures Upto

The question asks what a micrometer can measure "upto". In the context of precision instruments, "measure upto" usually refers to the smallest measurement it can reliably distinguish, which is defined by its precision or resolution. This is known as the instrument's least count.

Let's analyze the given options in relation to a micrometer's capability:

- **0.005 mm:** This is a typical value for the least count of many standard micrometers. While a micrometer can measure down to this dimension, it's a specific value, not the general principle of its measurement limit.

- **Microns size:** A micron (μm) is 10^{-6} meters or 0.001 millimeters. Micrometers are indeed capable of measuring dimensions in the micron range. However, "microns size" is a range of units, not the fundamental limit defined by the instrument's mechanism.
- **0.1 mm:** This value is typically the least count of a standard Vernier caliper, which is less precise than a micrometer. A micrometer's precision is much finer than 0.1 mm.
- **Least count:** The least count of a measuring instrument is the smallest value that can be measured by it. A micrometer is designed to measure with a precision determined by its screw pitch and the divisions on its thimble. Its ability to measure small dimensions is limited by this least count. Therefore, it can measure up to the precision of its least count, meaning it can distinguish measurements down to this smallest increment.

Considering the options, the most accurate description of the limit of what a micrometer can measure with precision is its least count. The least count represents the smallest difference in measurement that the instrument can detect.

A standard micrometer often has a least count of 0.01 mm or 0.001 mm (1 micron), depending on the specific design and number of divisions on the thimble and sleeve. Instruments designed for higher precision might have even smaller least counts, such as 0.005 mm.

Therefore, a micrometer's measuring capability goes down to the level of its least count, allowing it to measure dimensions with that specific resolution.

Conclusion on Micrometer Measurement

Based on the analysis of the options and the definition of least count, the micrometer can measure up to its least count, which defines its minimum distinguishable measurement.

Revision Table: Micrometer Capabilities

Here is a summary of key points regarding micrometer measurement:

- Instrument Type: Precision measuring tool
- Principle: Screw mechanism
- Typical Use: Measuring small dimensions (thickness, diameter)
- Precision: High
- Measurement Limit (Precision): Defined by Least Count

Additional Information: Least Count Explained

The least count is a crucial concept for understanding the precision of any measuring instrument, including a micrometer. It determines the smallest change in the quantity being measured that the instrument can detect and display. For a micrometer, the least count is typically calculated based on the pitch of the screw and the number of divisions on the circular scale (thimble).

Calculation example:

- $\text{Least Count} = (\text{Pitch of the screw}) / (\text{Number of divisions on the thimble})$

If the screw pitch is 0.5 mm and there are 50 divisions on the thimble, the least count would be $0.5 \text{ mm} / 50 = 0.01 \text{ mm}$. If the screw pitch is 1 mm and there are 100 divisions, the least count is $1 \text{ mm} / 100 = 0.01 \text{ mm}$. High-precision micrometers achieve smaller least counts through finer screw pitches or more thimble divisions.

162. Answer: b

Explanation:

Material Selection for Twist Drills

Twist drills are essential cutting tools used in machining to create holes in various materials. The choice of material for a twist drill is crucial because it must withstand high temperatures, abrasive wear, and significant mechanical stress during the drilling process. The material needs to maintain its hardness even when heated due to friction (this property is known as 'red hardness'), resist chipping, and have good toughness to prevent breakage.

Evaluating High-Speed Steel (H.S.S.) for Twist Drills

High-Speed Steel, commonly abbreviated as H.S.S., is a type of tool steel known for its excellent hardness, wear resistance, and ability to retain hardness at elevated temperatures. These characteristics make H.S.S. an ideal material for twist drills, allowing them to operate at higher speeds and feeds compared to drills made from simpler steels without becoming dull or damaged quickly.

H.S.S. drills are widely used across various industries due to their balance of performance, durability, and cost-effectiveness. They can effectively drill through metals like steel, aluminum, and brass, as well as plastics and wood.

Comparison of Materials for Twist Drill Applications

Let's consider the suitability of the other options:

- **Alloy Steel:** While alloy steels can be heat-treated to achieve higher hardness than plain carbon steels, they generally do not possess the same level of red hardness and wear resistance as High-Speed Steel. They might be used for less demanding drilling tasks but are not the preferred choice for high-performance drills.
- **Carbide Steel (Tungsten Carbide):** Carbide materials, specifically Tungsten Carbide cemented with cobalt, are significantly harder and more wear-resistant than H.S.S. However, they are also much more brittle. While tungsten carbide is used for drill *bits* or *inserts* in very specific, high-production applications where rigidity is high and impact forces are low, it's less common for standard twist drills due to its brittleness, which makes it prone to chipping or breaking under the torsional and bending stresses encountered during typical drilling operations.
- **Mild Steel:** Mild steel, also known as low-carbon steel, is relatively soft and ductile. It lacks the necessary hardness and wear resistance to function effectively as a twist drill material. Drills made from mild steel would quickly lose their cutting edge and deform, making them unsuitable for any practical drilling task.

Conclusion on Twist Drill Material

Considering the demanding requirements of drilling, including the need for hardness, wear resistance, and red hardness, High-Speed Steel (H.S.S.) emerges as the most suitable and commonly used material for manufacturing twist drills. Its properties provide the best combination for efficiency and durability in most drilling applications.

163. Answer: c

Explanation:

Understanding Air Locks in Fuel Lines

An air lock in a fuel line occurs when air becomes trapped within the fuel delivery system. This trapped air obstructs the flow of fuel, preventing it from reaching the engine's combustion chamber efficiently or at all. Fuel systems are designed to pump liquid fuel, and the presence of air interferes with this process.

Effects of Air Lock on Engine Performance

When an air lock is present, the engine experiences a disruption in its fuel supply. This leads to several problems, most notably affecting how the engine starts and runs.

- **Fuel Flow Disruption:** Air is much less dense than fuel and behaves differently under pressure. A pocket of air can break the continuous column of liquid fuel, making it difficult for the fuel pump to deliver a steady supply.
- **Fuel Starvation:** The engine may not receive enough fuel, especially during demanding conditions like starting or under acceleration.

Analyzing the Given Options

Let's look at how an air lock relates to the symptoms described in the options:

- **Engine idling high:** An air lock typically reduces fuel flow, leading to a lean mixture or insufficient fuel, which would more likely cause rough idling or stalling rather than high idling. High idling is often related to issues like vacuum leaks or incorrect idle speed settings.
- **Injector pressure is too high:** An air lock reduces the volume of fluid being pumped. While air is compressible, its presence in the line usually disrupts the pressure delivery, often leading to *low* or erratic pressure at the injectors, not excessively high pressure.
- **Late starting and erratic running:** This is a direct consequence of insufficient or inconsistent fuel supply. During starting, the engine needs a rich fuel-air mixture; if fuel flow is poor due to an air lock, starting will be difficult or delayed (late starting). Once running, the inconsistent fuel delivery caused by the trapped air will lead to the engine running unevenly, misfiring, or surging (erratic running).
- **Excess delivery from FI pump:** An air lock restricts flow, preventing the fuel injection (FI) pump from delivering its intended volume of fuel. It leads to *reduced* delivery, not excess.

Based on this analysis, the most accurate description of the effect of an air lock in a fuel line is late starting and erratic running.

Revision Table: Air Lock Effects

Issue	Effect of Air Lock
Fuel Flow	Restricted, inconsistent
Starting	Difficult, late, may require prolonged cranking
Running	Erratic, rough, misfires, surging, may stall
Fuel Pressure	Low or erratic
Engine Idling	Rough or may stall (not high idling)

Therefore, an air lock in the fuel line directly leads to difficulties in starting and uneven engine operation.

Additional Information: Fuel System Components

The fuel delivery system in a vehicle is complex and includes several components that work together to supply fuel to the engine. Understanding these components helps understand issues like air locks.

- **Fuel Tank:** Stores the fuel.
- **Fuel Pump:** Moves fuel from the tank to the engine. Can be mechanical or electric.
- **Fuel Filter:** Removes contaminants from the fuel.
- **Fuel Lines:** The pipes or hoses that carry fuel. This is where an air lock occurs.
- **Fuel Injectors or Carburetor:** Deliver fuel into the engine's intake system or combustion chamber.
- **Fuel Pressure Regulator:** Maintains correct fuel pressure in the system.

Air can enter the fuel system due to various reasons, such as running out of fuel, improper maintenance (like replacing a fuel filter without bleeding the system), or leaks in the fuel lines.

164. Answer: d

Explanation:

Understanding the Universal Surface Gauge and its Parts

The question asks about a specific part of a universal surface gauge used for drawing parallel lines along an edge, particularly a dotted or irregular edge. A universal surface gauge is a versatile tool used in machining and layout work primarily for marking lines parallel to an edge or surface and for checking the height of objects.

Let's look at the key components and their functions:

- **Base:** Provides stability for the gauge. Universal surface gauges have a base that can be used on flat or cylindrical surfaces.
- **Spindle:** A vertical column extending from the base. The scriber and other attachments are mounted on the spindle.
- **Sleeve or Clamp:** Holds the scriber onto the spindle and allows its vertical adjustment.
- **Scriber:** A sharp, pointed tool used for scratching or marking lines on a workpiece surface.
- **Rocker Arm:** Allows the spindle to be tilted or adjusted for fine settings.
- **Fine Adjusting Screw:** Provides fine, precise movement of the spindle or scriber for accurate settings.
- **Guide Pins (or Guide Pegs):** These are pins located at the base of the surface gauge. They can be adjusted or retracted. Their primary function is to register against the edge of a reference surface or workpiece, allowing the gauge to slide along that edge while keeping the scriber at a constant distance, thus drawing a line parallel to the edge. They are especially useful when the edge is not perfectly straight or is interrupted (like a dotted edge in the context of following a template edge).

Now let's analyze the given options:

1. **Fine adjusting screw:** While crucial for setting the precise height of the scriber, the fine adjusting screw itself does not guide the gauge along an edge to draw a parallel line. Its function is height adjustment.
2. **Center punch:** A center punch is a separate hand tool used for making indentations on a workpiece surface, typically to locate the center of a hole before drilling. It is not a part of the standard universal surface gauge assembly used for drawing lines.
3. **Scriber:** The scriber is the part that does the marking (draws the line), but it is the guidance mechanism that ensures the line is parallel to the edge. The scriber itself does not provide the guidance function along the edge.

4. Guide pins: As discussed, the guide pins are specifically designed to bear against a reference edge, enabling the surface gauge to be moved smoothly along that edge. This action maintains a constant distance between the scribe (which is set at a specific height and horizontal distance from the edge) and the edge, resulting in a line drawn parallel to the reference edge. This is exactly the function described in the question, particularly useful for following an existing edge, even one that might be slightly irregular or dotted.

Therefore, the part of the universal surface gauge that helps draw a parallel line along a dotted edge by registering against the edge is known as the guide pins.

Here is a summary table of the parts mentioned:

Part	Primary Function
Scriber	Marks lines on the workpiece surface
Fine Adjusting Screw	Allows for precise height adjustment of the scribe
Guide Pins	Guides the gauge along a reference edge to draw parallel lines
Center Punch	Marks point locations (not part of surface gauge for line marking)

Revision Table: Key Universal Surface Gauge Components

Revisiting the main components and their roles can help solidify understanding for exam preparation.

- **Base:** Stability, foundation for the tool.
- **Spindle:** Vertical support for holding and positioning the scribe.
- **Scriber:** Marks the workpiece surface.
- **Guide Pins:** Enables drawing lines parallel to an edge by serving as a guide/register.
- **Fine Adjusting Screw:** Allows for very small, precise adjustments of the scribe's position.

Additional Information: Using Guide Pins for Parallel Lines

When using the guide pins on a universal surface gauge to draw a parallel line, the workpiece is typically placed on a surface plate or other flat reference surface. The guide pins are adjusted to contact the edge of the workpiece or a reference block. The scribe is

then set to the desired height and distance from the guide pins. By sliding the base of the surface gauge along the reference edge, with the guide pins maintaining contact, the scribe draws a line at a consistent distance from the edge, ensuring parallelism ($\parallel$). This technique is fundamental in layout work for transferring dimensions or marking features relative to an existing edge.

165. Answer: a

Explanation:

Rate of Work Definition

The question asks to identify the term that describes the rate at which work is done. Understanding this concept is fundamental in physics, especially when studying mechanics and energy.

Defining the Rate of Doing Work

In physics, **work** is done when a force causes a displacement. The **rate** at which this work is performed signifies how quickly the work is being completed. This specific rate is given a distinct name.

The rate of doing work is formally defined as the amount of work done divided by the time taken to do that work. Mathematically, this relationship is expressed using the formula:

$$P = \frac{W}{t}$$

Where:

- P represents Power
- W represents Work done
- t represents the time taken

The standard unit for the rate of doing work (Power) in the International System of Units (SI) is the Watt (W), where 1 Watt is equal to 1 Joule per second (1 J/s).

Analyzing the Options

Let's examine the given options to understand why Power is the correct term:

- **Power:** This term directly corresponds to the rate of doing work or the rate at which energy is transferred.
- **Energy:** Energy is the capacity or ability to do work. While related to work, it is not the rate at which work is done. Energy is measured in Joules (J).
- **Force:** Force is an influence that can change an object's motion; it's a push or a pull. It is measured in Newtons (N). Force is required to do work, but it doesn't represent the rate of doing work.
- **Torque:** Torque is a measure of the twisting force that causes rotation. It is a rotational equivalent of linear force, measured in Newton-meters (N·m). It's involved in rotational work but isn't the rate itself.

Conclusion: Rate of Work is Power

Based on the definitions, the term that specifically means the **rate of doing work** is **Power**. It quantizes how fast work is accomplished.

166. Answer: b

Explanation:

Understanding Maintenance Free Batteries

A maintenance free battery is a type of rechargeable battery, commonly used in vehicles, designed to require minimal or no attention throughout its lifespan. Unlike traditional flooded lead-acid batteries that need periodic topping up of distilled water, these batteries are sealed and engineered to significantly reduce water loss.

Analyzing the Options for Maintenance Free Batteries

Let's look at each option provided regarding the characteristics of a maintenance free battery:

- **Option 1: Does not contain water**
This statement is incorrect. Maintenance free batteries are a type of lead-acid

battery, and their electrolyte is a mixture of sulfuric acid and water. While they are designed to minimize water loss, they still contain water as part of the electrolyte.

- **Option 2: Has lead-calcium plate grid**

This statement is correct. A key feature of maintenance free batteries is the use of lead-calcium alloy for the plate grids instead of the traditional lead-antimony alloy. The lead-calcium alloy significantly reduces the electrolysis of water during charging, which is the primary cause of water loss (gassing) in traditional batteries. This reduced gassing means less water evaporation, hence the "maintenance free" aspect.

- **Option 3: Does not contain acid**

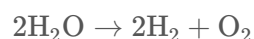
This statement is incorrect. Maintenance free batteries, like all lead-acid batteries, use sulfuric acid (H_2SO_4) as a crucial component of their electrolyte. The acid is essential for the chemical reactions that store and release electrical energy.

- **Option 4: Has lead-antimony plate grid**

This statement is incorrect. Traditional flooded lead-acid batteries use lead-antimony alloys for plate grids. Antimony enhances the mechanical strength of the plates but also increases the rate of water electrolysis during charging, leading to more gassing and the need for regular water replenishment. Maintenance free batteries specifically avoid or minimize antimony to reduce water loss.

Why Lead-Calcium Grid is Key to Maintenance Free Design

The primary reason traditional batteries require maintenance (adding water) is the process of electrolysis that occurs during charging. When the battery is overcharged or reaches a certain voltage, water (H_2O) in the electrolyte breaks down into hydrogen (H_2) and oxygen (O_2) gases. This process is called gassing.



Lead-antimony grids catalyze this reaction, increasing gassing. Lead-calcium grids, on the other hand, significantly reduce this catalytic effect, minimizing water loss through gassing. This is the fundamental difference that makes a battery "maintenance free".

Therefore, the characteristic that defines a maintenance free battery among the given options is the presence of a lead-calcium plate grid.

Comparison: Maintenance Free vs. Traditional Battery Grids

Feature	Maintenance Free Battery	Traditional Flooded Battery
Grid Material	Lead-Calcium alloy	Lead-Antimony alloy
Gassing during Charge	Very low	Significant
Water Loss	Minimal	Significant, requires top-up
Maintenance Required	None (sealed)	Regular water topping

Revision Table: Maintenance Free Battery Features

Key Feature	Description
Plate Grid Alloy	Typically Lead-Calcium
Water Consumption	Very low
Gassing	Significantly reduced
Design	Usually sealed, no access to cells
Electrolyte	Sulfuric acid and water

Additional Information on Battery Technology

Beyond the grid material, maintenance free batteries often incorporate other design features to minimize water loss and enhance performance, such as optimized charging voltage regulation in the vehicle's electrical system. Some advanced maintenance free batteries use Absorbent Glass Mat (AGM) or Gel technologies to suspend the electrolyte, further preventing spillage and improving performance, especially in demanding applications like start-stop vehicles. While AGM and Gel batteries are also maintenance-free, the use of a lead-calcium grid is a fundamental characteristic distinguishing a basic maintenance-free battery from a traditional one.

Understanding the role of the lead-calcium grid is crucial to understanding why these batteries don't need water top-ups like their traditional counterparts.

167. Answer: c

Explanation:

Understanding the Vernier Depth Gauge for Measuring

A Vernier depth gauge is a precision measuring instrument used specifically for measuring the depth of holes, slots, recesses, and steps. It consists of a base, which rests on the reference surface, and a graduated beam or scale that slides perpendicular to the base. A Vernier scale on the beam allows for precise readings.

Let's analyze the given options in the context of what a Vernier depth gauge is designed to measure:

- **Pitch circle diameter:** The pitch circle diameter (PCD) is a concept related to gears and is measured using specialized tools like a gear tooth Vernier caliper or by calculations based on other measurements. A standard Vernier depth gauge is not used for this purpose.
- **Round dimension:** Measuring the diameter of a round object is typically done with instruments like outside Vernier calipers or micrometers. A Vernier depth gauge measures linear distance from a reference plane downwards or upwards, not the diameter of a round feature on a single plane.
- **Step, depth of blind hole:** This option aligns perfectly with the primary function of a Vernier depth gauge. A 'step' refers to a change in the surface level, and the depth gauge can measure the vertical distance between the two levels. A 'blind hole' is a hole that does not go all the way through a workpiece. The Vernier depth gauge's base sits on the top surface, and the beam extends down to the bottom of the blind hole to measure its depth accurately.
- **Internal dimension:** Measuring the internal diameter or width of a hole or slot is typically done using internal calipers, inside micrometers, or the internal jaws of a standard Vernier caliper. A Vernier depth gauge measures depth from a surface, not the internal size across a feature.

Therefore, the most appropriate use of a Vernier depth gauge among the given options is for measuring the step and the depth of a blind hole.

How Vernier Depth Gauge Measures Depth and Step

When measuring the depth of a blind hole or a recess, the base of the Vernier depth gauge is placed firmly on the surface from which the measurement is to be taken. The beam is then extended into the hole or recess until it reaches the bottom. The reading is

taken from the main scale and the Vernier scale, providing a highly accurate measurement of the depth.

For measuring a step, the base is placed on the higher surface, and the beam is extended down to the lower surface of the step. The reading indicates the height of the step.

Comparison with Other Measuring Tools

Understanding the specific function of a Vernier depth gauge is clearer when compared to other common precision measuring instruments:

Instrument	Primary Use	Examples of Measurement
Vernier Caliper	External, internal, depth, and step measurements	Outer diameter of a shaft, inner diameter of a hole, total depth of a through hole, height of a step.
Micrometer	Highly precise external, internal, or depth measurements	External diameter of a wire (outside micrometer), internal bore size (inside micrometer), depth of a slot (depth micrometer).
Vernier Depth Gauge	Depth of holes, slots, recesses, and step heights	Depth of a blind hole, height of a step on a workpiece, depth of a keyway slot.
Gear Tooth Vernier Caliper	Measuring gear tooth thickness (chordal thickness and height)	Specific measurement related to gear geometry.

This table highlights that while a standard Vernier caliper can measure depth (usually of through holes using the tail end), the Vernier depth gauge is specifically designed with a stable base for accurate depth measurements, especially for blind features and steps.

Revision Table: Vernier Depth Gauge Usage

Concept	Description
Vernier Depth Gauge	Precision tool for measuring linear depth from a reference surface.
Measured Features	Depth of blind holes, recesses, slots, step heights.
Key Components	Base, Beam (scale), Vernier scale, Fine adjustment screw.
Principle	Measures distance perpendicular to the base surface.

Additional Information: Precision Measuring Tools

Precision measuring tools like the Vernier depth gauge, Vernier caliper, and micrometers are essential in manufacturing, inspection, and quality control. They allow for accurate measurement of dimensions that are critical for the proper function and fit of components.

- The precision of these tools depends on the least count, which is the smallest measurement increment they can accurately resolve. For a Vernier scale, the least count is calculated based on the main scale division and the number of divisions on the Vernier scale.
- Proper handling and calibration are crucial to ensure the accuracy of readings from Vernier depth gauges and other precision instruments.
- Modern digital depth gauges are also available, offering direct digital readouts, which can reduce reading errors but operate on the same fundamental principle of measuring depth from a reference surface.

Understanding the specific application of each measuring tool, such as the Vernier depth gauge for measuring depth and steps, ensures that the correct instrument is selected for the required measurement task, leading to accurate results.

168. Answer: b

Explanation:

Understanding Fuel Injection Systems in Large Engines

Fuel injection systems are crucial components in internal combustion engines, responsible for delivering fuel into the combustion chamber or intake manifold. Different types of systems have been developed and used over the years, each with its own advantages and disadvantages, often tailored to specific engine sizes and applications.

Analyzing Fuel Injection Methods

Let's look at the fuel injection methods mentioned in the options to understand which one fits the description of being originally used in large stationary and marine engines and is now seldom used.

- **Mechanical Injection:** This is a broad category that refers to systems that use purely mechanical means (like cams and plungers) to pressurize and inject fuel. While older mechanical systems existed, mechanical principles are still used in various modern injection systems, including many solid injection systems. So, this category is too general and not "seldom used."
- **Air Injection:** This method uses a blast of highly compressed air to force the fuel into the combustion chamber. The fuel is typically pumped into a small chamber near the injector, and then a valve opens to allow high-pressure air (often around 60–90 bar) to atomize and carry the fuel into the cylinder.
- **Solid Injection (or Airless Injection):** This is the predominant type of fuel injection used today in diesel engines. In this system, the fuel itself is pressurized to a very high level (hundreds or thousands of bar) by a pump, and this high pressure is used to atomize and inject the fuel directly into the combustion chamber without the aid of compressed air.
- **Airless Injection:** This is another name for Solid Injection, as explained above. Since Solid Injection is widely used today, Airless Injection is also a commonly used system, not one that is seldom used.

The History of Air Injection in Large Engines

The air injection system was notably used in the early days of diesel engines, particularly in large stationary and marine engines developed around the turn of the 20th century. Rudolf Diesel's original engines, for example, utilized air injection. This system had advantages at the time, such as relatively good fuel atomization even with low injection

pressures of the fuel pump itself, and it allowed for reliable operation with the technology available then.

However, air injection systems were complex, requiring a multi-stage air compressor running off the engine, bulky high-pressure air storage tanks, and elaborate valve mechanisms. They were also less efficient because power was needed to drive the compressor. These drawbacks led to the development and eventual widespread adoption of solid injection (airless) systems, which were simpler, more compact, and more efficient as engine and fuel pump technology advanced.

Today, air injection systems are indeed seldom used, having been largely replaced by various forms of solid or airless injection systems like common rail, unit injectors, and distributor pumps, even in large marine and power generation applications.

Comparing Injection Systems

Here's a brief comparison:

Feature	Air Injection	Solid/Airless Injection
Injection Medium	High-pressure air	High-pressure fuel
Complexity	High (requires compressor, air tanks)	Lower (requires high-pressure fuel pump)
Efficiency	Lower (compressor uses engine power)	Higher
Atomization	Relied on air blast	Relies on very high fuel pressure and nozzle design
Historical Use	Early large engines, now seldom used	Dominant in modern diesel engines

Based on this analysis, the fuel injection system that was originally used in large stationary and marine engines and is now seldom used is Air injection.

Revision Table: Fuel Injection Concepts

Term	Description
Fuel Injection	Process of atomizing and delivering fuel into an engine.
Air Injection	Injection method using high-pressure air blast. Historically significant for diesel engines.
Solid Injection	Injection method using high fuel pressure directly (airless). Modern standard.
Stationary Engine	Engine used in a fixed location (e.g., power generation).
Marine Engine	Engine used to power ships or boats.

Additional Information: Evolution of Diesel Injection

The evolution from air injection to solid injection was a major step in diesel engine technology. Early air injection engines were robust but bulky and less fuel-efficient due to the power consumed by the air compressor. The development of high-pressure fuel pumps and advanced nozzle designs allowed for effective fuel atomization using fuel pressure alone. This led to simpler, more compact, and more efficient engines, paving the way for the widespread use of diesel engines in vehicles and other applications where the bulk and complexity of air injection were prohibitive. Modern solid injection systems, like common rail and unit injector systems, achieve extremely high pressures and precise control over injection timing and quantity, leading to improved performance, fuel economy, and reduced emissions.

169. Answer: b

Explanation:

Understanding Petrol Engine Combustion

In a petrol engine, also known as a gasoline engine, energy is produced by burning a mixture of air and fuel. This burning process, called combustion, happens in a very specific part of the engine.

Components Involved in Petrol Engine Combustion

Several key components work together to make combustion happen:

- **Cylinder:** This is a hollow tube or chamber where the piston moves back and forth. It's the main space where the combustion reaction occurs.
- **Piston:** A movable component that slides inside the cylinder. Its movement compresses the air-fuel mixture and is pushed down by the force of combustion, turning the crankshaft.
- **Spark Plug:** Located at the top of the cylinder, it generates an electric spark to ignite the compressed air-fuel mixture.
- **Valves:** Control the flow of air and fuel into the cylinder and exhaust gases out of the cylinder.

The Role of the Cylinder in Combustion

The combustion of the air-fuel mixture is an explosive process that generates high temperature and pressure. The cylinder provides a sealed environment where this reaction can safely and effectively take place. The air and fuel are drawn into the cylinder, mixed, and then compressed by the piston. Once compressed, the spark plug ignites the mixture, causing it to burn rapidly and expand. This rapid expansion of gases pushes the piston downwards, generating power.

Therefore, the cylinder acts as the containment chamber for the combustion event. The combustion reaction itself occurs within the volume defined by the cylinder walls, the piston top, and the cylinder head (where the spark plug and valves are located).

Analyzing the Options for Combustion Location

Let's look at the provided options:

- **Intestinal cells:** These are biological components found in living organisms for digestion. They are completely unrelated to engine operation.
- **Cylinder:** As explained, this is the chamber where the air-fuel mixture is contained and burned in a petrol engine. This aligns with the process of combustion.
- **Piston:** The piston is a component that moves inside the cylinder and is driven by the combustion pressure. However, the combustion doesn't take place *within* the piston itself, but in the space above it, *inside* the cylinder.

- **The elimination cells:** This term is not relevant to the components or processes of a petrol engine. It seems to refer to biological or other systems.

Based on the analysis, the cylinder is where the combustion of the air-fuel mixture takes place.

Conclusion on Petrol Engine Combustion Location

The combustion process in a petrol engine occurs within the cylinder. The cylinder is designed to withstand the heat and pressure generated by the burning air-fuel mixture and channel the resulting energy to move the piston and produce power.

Revision Table: Petrol Engine Combustion Key Points

Component	Role in Combustion	Location of Combustion
Cylinder	Contains the air-fuel mixture and the combustion reaction.	Inside the Cylinder
Piston	Compresses the mixture and is moved by combustion.	Space above the piston (within cylinder)
Spark Plug	Ignites the mixture.	Located in the Cylinder Head, sparks inside cylinder.

Additional Information: Internal Combustion Engines

Petrol engines are a type of internal combustion engine (ICE). This means that the combustion of the fuel happens **inside** the engine itself, within the cylinders, unlike external combustion engines (like steam engines) where fuel is burned outside to heat a working fluid.

Most petrol engines in cars operate on a four-stroke cycle (Intake, Compression, Power, Exhaust). The combustion happens during the Power stroke, after the air-fuel mixture has been compressed during the Compression stroke and ignited by the spark plug.

The efficiency and performance of a petrol engine depend heavily on how effectively the combustion process occurs within the cylinder.

170. Answer: b

Explanation:

Understanding the Engine Cooling Radiator Principle

The radiator is a key component in the engine cooling system of a vehicle. Its main job is to remove excess heat from the engine coolant (usually a mixture of water and antifreeze) that has circulated through the engine block and cylinder head, absorbing heat generated by combustion.

The fundamental principle behind how a radiator works is heat transfer. It facilitates the movement of heat from the hot coolant to the cooler air flowing around it. To make this heat transfer as efficient as possible, the radiator is designed to maximize the contact area between the hot coolant and the air.

How the Radiator Achieves Heat Dissipation

A typical radiator consists of:

- **Core:** Made up of many small tubes or flattened passages through which the hot coolant flows.
- **Fins:** Thin metal strips attached to the tubes. These fins greatly increase the surface area exposed to the air.
- **Tanks:** Located at the top and bottom or sides of the core, which store and distribute the coolant into the tubes.

As hot coolant flows through the tubes, heat is transferred by conduction through the tube walls and into the fins. The fins, having a very large surface area, then transfer this heat to the air flowing over them, primarily through convection and some radiation. The cooler air absorbs the heat, and the cooled coolant then returns to the engine to absorb more heat.

Analyzing the Options for Radiator Principle

Let's look at the given options in the context of the radiator's design and function:

- Option 1: Increase the air speed as it flows over the hot surface

Increasing air speed (often done by a cooling fan) definitely improves heat transfer, but this describes a method to enhance the process, not the core principle of the radiator's *design*. The radiator is designed to *have* a large surface area for air to flow *over*, regardless of the speed.

- **Option 2: Spread out the hot water over a large area**

This accurately describes how the radiator's internal structure (tubes and fins) works. By making the coolant flow through numerous narrow passages and using fins, the hot coolant's heat is effectively spread out over a vast surface area (the combined area of the tubes and fins), allowing for efficient heat exchange with the air. This is the primary principle of the radiator's construction for heat dissipation.

- **Option 3: Acts as a reservoir for the water**

While the radiator holds coolant as part of the system's capacity, its main purpose is heat exchange, not storage. The expansion tank or overflow reservoir is primarily designed to accommodate coolant expansion and act as a reservoir.

- **Option 4: Cause a heat flow by convection currents**

Heat transfer from the fins to the air happens through convection (and radiation). However, the radiator doesn't *cause* the convection currents; it provides the hot surface from which air absorbs heat, leading to convection. The fundamental principle is creating the large surface area for this heat transfer to occur efficiently, rather than initiating the convection itself.

Based on the analysis, the core principle of the radiator's design for effective heat removal is to maximize the surface area over which heat can be transferred from the hot coolant to the surrounding air.

Option	Analysis	Relevance to Radiator Principle
1. Increase air speed	Enhances heat transfer rate	Method to improve function, not core design principle
2. Spread out hot water over a large area	Maximizes surface area for heat exchange	Core design principle for efficient heat dissipation
3. Acts as a reservoir	Secondary function (holds coolant)	Not the primary principle of its heat exchange role
4. Cause convection currents	Heat transfer mechanism involved	Result of heat transfer from large surface, not the design principle itself

Conclusion on Radiator Principle

The most accurate description of the principle behind an engine cooling system radiator is its design that facilitates spreading out the hot coolant's heat over a vast surface area using tubes and fins. This large surface area is crucial for efficiently transferring heat to the air, thereby cooling the engine coolant.

Revision Table: Engine Radiator Principle

Key Concept	Explanation
Radiator Function	Removes heat from engine coolant.
Primary Principle	Maximizing surface area for heat transfer.
How Area is Increased	Using tubes and fins in the radiator core.
Heat Transfer Methods	Conduction (coolant to fins), Convection (fins to air).
Goal	Efficiently cool the engine coolant.

Additional Information: Engine Cooling System Components

The radiator is part of a larger engine cooling system. Other key components include:

- **Water Pump:** Circulates the coolant through the engine and radiator.
- **Thermostat:** Regulates the engine's temperature by controlling coolant flow to the radiator.
- **Cooling Fan:** Provides airflow through the radiator fins, especially at low vehicle speeds or when idling.
- **Hoses:** Connect the various components of the cooling system.
- **Coolant (Antifreeze):** The fluid that absorbs heat from the engine and carries it to the radiator.
- **Expansion Tank / Overflow Reservoir:** Accommodates changes in coolant volume due to temperature variations and acts as a reserve.

All these components work together to maintain the engine within its optimal operating temperature range, preventing overheating which can cause significant engine damage.

171. Answer: c

Explanation:

Understanding Governors and Airtight Systems

A governor is a device used to regulate the speed of an engine or other prime mover. It automatically maintains a constant speed under varying load conditions by controlling the supply of fuel or working fluid.

Governors can be broadly classified based on their operating principle. Common types include mechanical, hydraulic, and pneumatic governors. Each type uses a different medium or mechanism to sense the speed and actuate the control element.

Why Airtightness is Crucial for Certain Governor Types

The question highlights a specific requirement for a governor's effective operation: the system must be airtight. Let's consider why this might be the case for different governor types.

- **Mechanical Governors:** These rely on centrifugal force, springs, and linkages. While mechanical integrity and lubrication are important, airtightness of the main speed-

sensing mechanism is generally not a primary requirement for effective operation, unless considering sealed components for specific purposes like oil retention.

- **Hydraulic Governors:** These use an incompressible liquid, typically oil, as the working fluid to transmit forces and actuate control. While leaks in a hydraulic system lead to pressure loss and inefficiency, the term "airtight" doesn't specifically apply to a liquid system in the same critical way it would to a system using a compressible fluid like air or gas. Fluid leaks are the concern here, not air leaks affecting the fluid itself.
- **Pneumatic Governors:** These governors use compressed air or gas as the working medium. They often sense engine speed by responding to air pressure signals (like intake manifold pressure or a signal from a pneumatic speed sensor) and use air pressure to operate an actuator that controls fuel supply. Air is a compressible fluid. Leaks in a pneumatic system cause significant pressure drops. If the system is not airtight, the control signals (pressure) will be inaccurate, and the force produced by the actuator will be diminished. This makes effective and precise speed control impossible. Therefore, airtightness is absolutely essential for the proper and effective operation of a pneumatic governor.
- **Servo Governors:** A servo governor is a type of governor that uses a servo mechanism, typically hydraulic or pneumatic, to amplify the small force generated by the speed-sensing element (like a mechanical flyweight) to operate a larger control valve. The term "servo" describes the amplifying mechanism, not the primary working fluid type. However, if the servo mechanism is pneumatic, then airtightness would indeed be critical, linking back to the pneumatic principle. The question asks for a *type* of governor where operation is only effective if airtight, and Pneumatic is a fundamental type whose operational principle relies on maintained air pressure.

Based on this analysis, the governor type whose operation is critically dependent on the system being airtight is the pneumatic governor, because leaks directly compromise the air pressure required for accurate sensing and actuation.

Conclusion

For a governor to operate effectively when its mechanism relies on maintaining precise air pressure or flow, the system must be sealed against air leaks. This fundamental requirement is characteristic of pneumatic systems.

Governor Type	Working Medium	Airtightness Requirement
Mechanical	Weights, Springs, Linkages	Generally not critical for primary function
Hydraulic	Liquid (e.g., Oil)	Sealing against fluid leaks is critical
Pneumatic	Air or Gas	Critically essential for effective operation
Servo	Uses amplifier (often Hydraulic or Pneumatic)	Depends on the servo fluid type; critical if pneumatic

Revision Table: Governor Types and Characteristics

Governor Type	Principle of Operation	Key Feature
Mechanical	Balances centrifugal force against spring force	Simple, direct mechanical linkage
Hydraulic	Uses oil pressure to move control elements	Provides significant power amplification
Pneumatic	Uses air pressure signals for sensing and control	Relies on intake manifold pressure or pneumatic sensor

Additional Information on Governor Systems

Governors are vital components in engines and power generation systems, ensuring stable speed regardless of the load. Stability, sensitivity, and isochronism (ability to maintain constant speed at all loads) are key performance characteristics of governors.

Pneumatic governors were commonly used in older diesel engines and some gasoline engines, particularly in vehicles, due to their simplicity and low cost compared to hydraulic systems. They often work in conjunction with the fuel injection pump or carburetor to regulate fuel supply.

172. Answer: c

Explanation:

Diesel Engine Ignition: Understanding the Process

In a diesel engine, the process of igniting the fuel is fundamentally different from that in a petrol (gasoline) engine. Petrol engines use a spark plug to ignite the fuel-air mixture. Diesel engines, however, rely on a different principle related to the properties of compressed air.

How Diesel Engine Ignition Works

Let's break down the ignition process in a diesel engine:

1. **Air Intake:** During the intake stroke, the piston moves down, drawing only fresh air into the cylinder.
2. **Compression Stroke:** The piston then moves up, compressing this air into a very small volume. As the air is compressed rapidly, its temperature rises significantly. This is due to the work done on the air and its thermal properties. The compression ratio in a diesel engine is much higher than in a petrol engine, leading to a much higher final temperature.
3. **Fuel Injection:** Near the end of the compression stroke, when the air is at a very high temperature (typically around 700–900 °C) and high pressure, diesel fuel is injected into the cylinder as a fine spray by the injector.
4. **Ignition:** When the finely atomized diesel fuel comes into contact with the extremely hot, compressed air, it spontaneously ignites. This self-ignition process is why diesel engines do not need a spark plug. The burning of the fuel causes a rapid expansion of gases, pushing the piston down and generating power.

Therefore, the ignition in a diesel engine is directly caused by the high temperature of the air resulting from compression.

Analyzing the Options

Let's look at why the other options are not the primary cause of ignition in a diesel engine:

- **Spark plug:** Spark plugs are used in petrol engines to create a spark that ignites the fuel-air mixture. Diesel engines do not have spark plugs.
- **Injector:** The injector delivers the fuel into the combustion chamber as a fine spray. While essential for introducing the fuel, the injector itself does not cause the ignition.
- **Glow Plug:** Glow plugs are used in some diesel engines, especially in colder climates, to heat the combustion chamber or the air within it *before* starting the engine. This helps the initial ignition process when the engine is cold and compression alone might not raise the air temperature high enough. However, once the engine is running normally, the high temperature from compression is sufficient for ignition, and glow plugs are not actively involved in every combustion cycle. They assist cold starting, they don't cause regular ignition.
- **Compressed air temperature properties:** As explained, the rapid compression of air significantly increases its temperature due to its thermodynamic properties. This high temperature is sufficient to ignite the injected diesel fuel. This is the fundamental mechanism of ignition in a diesel engine.

Based on the working principle of diesel engines, fuel ignition is a result of the fuel being injected into the cylinder filled with highly compressed and thus very hot air.

Ignition Comparison: Diesel vs. Petrol Engines

Feature	Diesel Engine	Petrol Engine
Ignition Source	High temperature of compressed air	Spark plug
Fuel Introduction	Fuel injected into hot compressed air	Fuel mixed with air before intake (usually)
Compression Ratio	High (typically 14:1 to 25:1)	Lower (typically 8:1 to 12:1)
Fuel Type	Diesel fuel	Petrol (gasoline)

Revision Table: Key Diesel Engine Concepts

Concept	Description
Compression Ignition	Ignition caused by the heat generated from compressing air.
Compression Stroke	Piston moves up, reducing cylinder volume and increasing air pressure and temperature.
Fuel Injection	Process of spraying diesel fuel into the combustion chamber at high pressure.
High Compression Ratio	Ratio of cylinder volume when piston is at bottom to volume when piston is at top; higher in diesel engines.

Additional Information: Diesel Engine Operation

Diesel engines are known for their efficiency and torque. Their robust construction is necessary to withstand the higher pressures generated during the compression and combustion strokes compared to petrol engines. The process of fuel ignition by compressed air temperature is a key characteristic that defines the diesel cycle.

While the primary ignition source is the heat of compression, factors like fuel quality, injection timing, and the design of the combustion chamber are also crucial for efficient and clean combustion in a diesel engine.

Your Personal Exams Guide

173. Answer: c

Explanation:

Understanding the SI Unit of Pressure

Pressure is a fundamental physical quantity that describes the force applied perpendicular to the surface of an object per unit area over which that force is distributed. It is a scalar quantity.

The formula for pressure (P) is given by:

$$P = \frac{F}{A}$$

where F is the magnitude of the normal force and A is the area of the surface.

The SI Unit of Pressure: Pascal

The International System of Units (SI) defines standard units for various physical quantities. For pressure, the standard SI unit is the **pascal** (Pa).

One pascal is defined as one newton per square meter. Mathematically, this relationship is expressed as:

$$1 \text{ Pa} = 1 \frac{\text{N}}{\text{m}^2}$$

This definition directly follows from the formula for pressure, where force is measured in newtons (N) and area is measured in square meters (m^2).

Analyzing the Given Options for Pressure Units

Let's examine the other options provided to understand why Pascal is the correct SI unit of pressure:

- **Newton (N):** Newton is the SI unit of **force**. Force is related to pressure, but it is not pressure itself.
- **Joule (J):** Joule is the SI unit of **energy or work**. Energy is a completely different physical quantity from pressure.
- **Dyne:** Dyne is a unit of **force** in the CGS (centimeter-gram-second) system of units. It is not an SI unit, and it represents force, not pressure.

Based on the definitions in the International System of Units, the pascal is the designated unit for measuring pressure.

Common Units and Quantities

Quantity	SI Unit	Other Units
Pressure	Pascal (Pa)	bar, atm, psi, mmHg, torr
Force	Newton (N)	Dyne, pound-force
Energy / Work	Joule (J)	erg, calorie, kilowatt-hour
Area	Square Meter (m^2)	Square centimeter, square foot, acre

Therefore, the correct SI unit for pressure is the pascal.

Revision Table: Key SI Units

Important SI Base and Derived Units

Quantity	SI Unit	Symbol
Length	Meter	m
Mass	Kilogram	kg
Time	Second	s
Electric Current	Ampere	A
Thermodynamic Temperature	Kelvin	K
Amount of Substance	Mole	mol
Luminous Intensity	Candela	cd
Force	Newton	$N (kg \cdot m/s^2)$
Energy / Work	Joule	$J (N \cdot m)$
Pressure	Pascal	$Pa (N/m^2)$

Additional Information on Pressure and Units

Pressure is used in many fields, including meteorology (atmospheric pressure), fluid mechanics (fluid pressure), and engineering. While Pascal is the SI unit, other units are frequently used depending on the application and region.

Some common non-SI units of pressure include:

- **Bar:** Often used in meteorology; 1 bar = 100,000 Pa.
- **Atmosphere (atm):** Represents the average atmospheric pressure at sea level; 1 atm $\approx$ 101,325 Pa.
- **Pounds per square inch (psi):** Common in the United States, especially for tire pressure; 1 psi $\approx$ 6,895 Pa.
- **Millimeters of mercury (mmHg) or Torr:** Used in medicine and vacuum measurements; 760 mmHg = 1 atm.

Understanding the relationships between Pascal and these other units is crucial for conversions and practical applications.

174. Answer: a

Explanation:

Understanding Diesel Engine Noise: Fuel Knock

The question asks about a specific type of noise produced inside the cylinder of a diesel engine during combustion, particularly when the fuel quality isn't optimal, specifically mentioning low ignition quality fuel. This noise is a critical indicator of abnormal combustion.

What is Diesel Knock?

In a diesel engine, fuel is injected into compressed hot air, and it ignites spontaneously. For smooth operation, the fuel should ignite after a short delay and burn progressively. However, if the fuel has low ignition quality, the ignition delay period is longer. During this extended delay, more fuel accumulates in the combustion chamber. When ignition finally occurs, this larger amount of fuel burns very rapidly, almost explosively. This rapid and uncontrolled pressure rise creates a shock wave that travels through the combustion chamber, hitting the cylinder walls and piston, which is heard as a loud, sharp metallic sound. This sound is commonly known as **diesel knock** or **fuel knock**.

Diesel knock is analogous to engine knocking or pinging in spark-ignition (petrol) engines, although the underlying causes are slightly different. In diesel engines, the knock is due to a rapid pressure rise caused by the accumulation and sudden combustion of fuel following a long ignition delay.

Causes of Diesel Knock

Several factors can contribute to diesel knock, but a primary one highlighted in the question is the fuel's ignition quality.

- **Low Ignition Quality Fuel:** Fuel with low cetane number has a longer ignition delay. This leads to more fuel accumulating before ignition, causing rapid, knocking

combustion.

- **Low Compression Temperature:** If the air temperature in the cylinder is too low at the end of compression (e.g., due to low engine load, cold weather, or poor injection timing), the ignition delay can increase, leading to knock.
- **Improper Injection Timing:** If fuel is injected too early, it has a longer time to mix and potentially accumulate before sufficient temperature for ignition is reached, leading to knocking combustion when ignition finally occurs.

Analyzing the Options

Let's look at the provided options in the context of the noise produced in a diesel engine cylinder due to low ignition quality fuel:

Option	Explanation	Relevance to the Noise
Fuel knock	Describes the rapid, uncontrolled combustion noise caused by delayed ignition and sudden pressure rise in diesel engines, often due to low ignition quality fuel.	Highly relevant. This is the standard term for the phenomenon described.
Fuel explosion	While the combustion during knock is rapid and could be described as explosive, "Fuel explosion" is not the specific engineering term used for this particular noise and combustion anomaly in engines.	Too general and not the correct technical term.
Fuel thrust	Refers to the force generated by expelling mass, typically in jet or rocket engines for propulsion. It has no relation to combustion noise inside a cylinder.	Completely irrelevant.
Fuel burst	This term is not a standard engineering term related to combustion noise or anomalies in internal combustion engines. It doesn't describe the specific knocking phenomenon.	Not a standard technical term.

Based on the technical definition and common terminology in internal combustion engines, the noise produced in a diesel cylinder during combustion, especially with low

ignition quality fuel leading to delayed and rapid ignition, is correctly termed **fuel knock** or diesel knock.

Revision Table: Diesel Engine Combustion Noise

Term	Description	Associated Cause
Fuel Knock (Diesel Knock)	Sharp, metallic noise during combustion due to rapid pressure rise from delayed ignition and sudden burning of accumulated fuel.	Low fuel ignition quality (low cetane number), low compression temperature, early injection timing.
Normal Combustion	Smooth, progressive burning of fuel starting after a short ignition delay.	High fuel ignition quality (high cetane number), adequate compression temperature, correct injection timing.

Additional Information: Impact of Fuel Ignition Quality

The ignition quality of diesel fuel is measured by its cetane number. A higher cetane number indicates a shorter ignition delay and better ignition quality. Fuels with a high cetane number ignite more readily under compression and burn smoothly, reducing the tendency for diesel knock. Using fuel with an inappropriately low cetane number for a given engine or operating condition is a common cause of significant knocking. Persistent heavy knocking can lead to engine damage over time due to the excessive stresses caused by the rapid pressure fluctuations.

175. Answer: b

Explanation:

Understanding Jenny Spanners and Their Names

Jenny spanners, also known as odd-leg calipers or sometimes hermaphrodite calipers, are essential layout and measurement tools used in metalworking and engineering

workshops. They are primarily used for marking lines parallel to an edge or finding the center of round stock.

Analyzing the Options

Let's look at the given options to find the correct alternative name for Jenny spanners:

- Outside caliper
- Joint spanner
- Inside caliper
- Odd-leg caliper

Explanation of the Correct Answer

A Jenny spanner has two legs, but they are different. One leg is a caliper leg with a pointed tip, and the other leg is a divider leg that is bent outwards near the point. This unique configuration is why it is often referred to by names that describe its appearance or function. The term "odd-leg" directly describes the fact that its two legs are different in shape and purpose.

Another less common name sometimes used for Jenny spanners is "hermaphrodite caliper" because it combines features of both inside and outside calipers or dividers. However, among the given options, "Odd-leg caliper" is the most widely recognized alternative name.

Let's examine why the other options are incorrect:

- **Outside caliper:** Used to measure the external size of an object. It has two curved legs that point inwards.
- **Joint spanner:** This term does not refer to a type of caliper or measuring tool. It might refer to a type of wrench used for tightening joints.
- **Inside caliper:** Used to measure the internal size of an opening or groove. It has two legs that curve outwards.

Based on the structure and function of a Jenny spanner, "Odd-leg caliper" accurately describes it and is a common alternative name.

Revision Table: Types of Calipers

Tool Name	Primary Use	Leg Shape
Outside Caliper	Measuring external dimensions	Legs curved inwards
Inside Caliper	Measuring internal dimensions	Legs curved outwards
Divider Caliper	Marking distances, scribing circles/arcs	Two pointed legs
Jenny Caliper (Odd-leg caliper)	Marking lines parallel to edge, finding center of rounds	One caliper leg, one divider leg (bent)

Additional Information on Layout Tools

Layout tools like Jenny spanners are crucial before machining or shaping material. They help in accurately marking where cuts or holes should be made. Other common layout tools include:

- **Surface Plate:** A flat, rigid surface used as a reference plane.
- **Surface Gauge:** Used on a surface plate to scribe lines parallel to the surface plate or to check heights.
- **Scribers:** Hardened steel points used for scratching lines on metal surfaces.
- **Punches (Center Punch, Prick Punch):** Used to make indentations for drilling or marking.
- **Rules and Squares:** For measuring lengths and checking squareness.

Understanding the specific function and common names of each tool, like the Jenny spanner or Odd-leg caliper, is important for accurate work in manufacturing and fabrication.