

Answers

1. Answer: b

Explanation:

Understanding Average Speed in River Flow Problems

This problem asks us to find the average speed of a boat traveling the same distance upstream and downstream in a steady-flowing river. The key here is that the distance covered in both directions is the same, but the speeds are different.

Given Information:

- Upstream speed ($v_{upstream}$): 12 km/h
- Downstream speed ($v_{downstream}$): 24 km/h
- Distance traveled upstream is equal to the distance traveled downstream.

Calculating Average Speed for Equal Distances

When an object travels the same distance at two different speeds, the average speed is not simply the arithmetic mean of the two speeds. Instead, we use a specific formula derived from the definition of average speed (Total Distance / Total Time).

Let the distance traveled in one direction be d . The total distance for the round trip is $d + d = 2d$.

The time taken for the upstream journey ($t_{upstream}$) is:

$$t_{upstream} = \frac{\text{Distance}}{\text{Speed}} = \frac{d}{v_{upstream}} = \frac{d}{12}$$

The time taken for the downstream journey ($t_{downstream}$) is:

$$t_{downstream} = \frac{\text{Distance}}{\text{Speed}} = \frac{d}{v_{downstream}} = \frac{d}{24}$$

The total time (T) for the whole journey is:

$$T = t_{upstream} + t_{downstream} = \frac{d}{12} + \frac{d}{24}$$

To add these fractions, find a common denominator (24):

$$T = \frac{2d}{24} + \frac{d}{24} = \frac{2d+d}{24} = \frac{3d}{24} = \frac{d}{8}$$

The average speed ($v_{average}$) is Total Distance divided by Total Time:

$$v_{average} = \frac{\text{Total Distance}}{\text{Total Time}} = \frac{2d}{\frac{d}{8}}$$

Simplify the expression:

$$v_{average} = 2d \times \frac{8}{d}$$

$$v_{average} = 2 \times 8$$

$$v_{average} = 16 \text{ km/h}$$

Alternatively, when the distance is the same for two parts of a journey with speeds v_1 and v_2 , the average speed can be directly calculated using the formula:

$$v_{average} = \frac{2v_1v_2}{v_1+v_2}$$

Substitute the given speeds $v_1 = 12 \text{ km/h}$ and $v_2 = 24 \text{ km/h}$:

$$v_{average} = \frac{2 \times 12 \times 24}{12+24}$$

$$v_{average} = \frac{2 \times 288}{36}$$

$$v_{average} = \frac{576}{36}$$

$$v_{average} = 16 \text{ km/h}$$

Both methods yield the same result.

Summary of Average Speed Calculation

Parameter	Value
Upstream Speed ($v_{upstream}$)	12 km/h
Downstream Speed ($v_{downstream}$)	24 km/h
Total Distance	2d (where d is distance one way)
Total Time	$\frac{d}{12} + \frac{d}{24} = \frac{d}{8}$ hours
Average Speed	$\frac{2d}{d/8} = 16 \text{ km/h}$

Revision Table: Average Speed Concepts

Scenario	Average Speed Formula	Explanation
Same Distance ($d_1 = d_2 = d$) at Speeds v_1, v_2	$\frac{2v_1v_2}{v_1+v_2}$	This is the harmonic mean of the speeds. Used when distance is constant but time varies.
Same Time ($t_1 = t_2 = t$) at Speeds v_1, v_2	$\frac{v_1+v_2}{2}$	This is the arithmetic mean of the speeds. Used when time is constant but distance varies.
Different Distances ($d_1, d_2, \dots$) and Times ($t_1, t_2, \dots$)	$\frac{d_1+d_2+\dots}{t_1+t_2+\dots}$	General definition: Total Distance / Total Time.

Additional Information: Boat Speed and River Flow

In problems involving boats and rivers, we often consider the boat's speed in still water and the speed of the river current.

- Let v_b be the speed of the boat in still water.
- Let v_s be the speed of the stream (river current).

When the boat travels downstream, the current helps it, so the effective speed is:

$$v_{\text{downstream}} = v_b + v_s$$

When the boat travels upstream, the current opposes it, so the effective speed is:

$$v_{\text{upstream}} = v_b - v_s$$

In this problem, we were directly given the effective upstream and downstream speeds (12 km/h and 24 km/h). Had we been given v_b and v_s , we would first calculate v_{upstream} and $v_{\text{downstream}}$ and then use these values in the average speed formula for equal distances.

For example, from the given speeds:

- $v_b - v_s = 12$

- $v_b + v_s = 24$

Adding these two equations gives:

$$(v_b - v_s) + (v_b + v_s) = 12 + 24$$

$$2v_b = 36$$

$$v_b = 18 \text{ km/h}$$

Substituting $v_b = 18$ into $v_b + v_s = 24$ gives:

$$18 + v_s = 24$$

$$v_s = 6 \text{ km/h}$$

So, the boat's speed in still water is 18 km/h, and the speed of the stream is 6 km/h. However, finding these speeds was not necessary to calculate the average speed for the total journey using the upstream and downstream speeds given.

2. Answer: d

Explanation:

Finding the Missing Number in a Series

The question asks us to identify the missing term in the given number series:

-1.3, -0.8, -0.3, ?, 0.7, 1.2

Analyzing the Number Series Pattern

Let's examine the relationship between consecutive terms in the series. We can calculate the difference between each term and the one preceding it.

- Difference between the second and first term:

$$-0.8 - (-1.3) = -0.8 + 1.3 = 0.5$$

- Difference between the third and second term:

$$-0.3 - (-0.8) = -0.3 + 0.8 = 0.5$$

We observe a consistent difference of 0.5 between the terms we have checked so far. This suggests that the series is an arithmetic progression with a common difference of 0.5.

Calculating the Missing Term

Assuming the pattern of adding 0.5 continues, the missing term should be obtained by adding 0.5 to the third term (-0.3).

Missing term = Third term + Common difference

Missing term =

$$-0.3 + 0.5 = 0.2$$

Verifying the Pattern

Let's check if adding the common difference (0.5) to our calculated missing term (0.2) gives the next term in the series (0.7).

Calculated missing term + Common difference =

$$0.2 + 0.5 = 0.7$$

This matches the fifth term in the given series (0.7), confirming our calculated missing number and the pattern.

Conclusion on the Missing Number

Based on the analysis, the missing number in the series -1.3, -0.8, -0.3, ?, 0.7, 1.2 is 0.2.

Term Position	Term Value	Difference from Previous Term
1st	-1.3	-
2nd	-0.8	$-0.8 - (-1.3) = 0.5$
3rd	-0.3	$-0.3 - (-0.8) = 0.5$
4th (Missing)	0.2	$0.2 - (-0.3) = 0.5$
5th	0.7	$0.7 - 0.2 = 0.5$
6th	1.2	$1.2 - 0.7 = 0.5$

Revision Table: Number Series Concepts

Concept	Description
Number Series	A sequence of numbers that follow a specific pattern or rule.
Arithmetic Progression (AP)	A sequence where the difference between consecutive terms is constant (called the common difference).
Common Difference	The constant value added to each term to get the next term in an arithmetic progression.

Additional Information: Solving Number Series Questions

Solving number series questions often involves identifying the underlying pattern. Common patterns include:

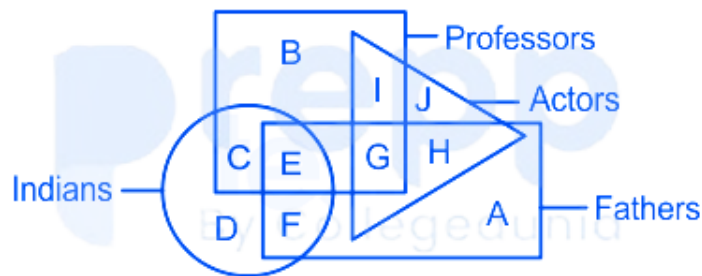
- Arithmetic progressions (constant difference).
- Geometric progressions (constant ratio).
- Differences between terms forming another pattern (e.g., squares, cubes, prime numbers).
- Alternating patterns.
- Patterns involving multiplication, division, squares, cubes, etc.

To solve these questions, start by finding the differences or ratios between consecutive terms. If a simple pattern isn't obvious, look at the differences of the differences, or consider other operations like multiplication, division, or sequences of squares/cubes.

3. Answer: b

Explanation:

The set of letters which represents Indians who are either professors or fathers, is shown below:



Hence, 'CEF' is the correct answer.

4. Answer: b

Explanation:

Identifying the Author of 'The Monk Who Sold His Ferrari'

The question asks us to identify the author of the well-known book, 'The Monk Who Sold His Ferrari'.

Let's examine the given options:

- Arvind Adiga
- Robin Sharma
- Manohar Malgonkar
- Arundhati Roy

Determining the Correct Author

To find the author, we need to recall or look up who wrote 'The Monk Who Sold His Ferrari'. This book is a popular inspirational fable that has been translated into many languages. It is widely attributed to a specific author known for his works on self-help and leadership.

Upon researching the authorship of this particular book, it is confirmed that **Robin Sharma** is the author of 'The Monk Who Sold His Ferrari'.

Analysis of Options

Let's briefly consider the other options to confirm why they are not the author of this book:

- **Arvind Adiga:** An Indian writer whose notable work includes 'The White Tiger'. He is not the author of 'The Monk Who Sold His Ferrari'.
- **Robin Sharma:** A Canadian writer, best known for his 'The Monk Who Sold His Ferrari' series. He is indeed the author.
- **Manohar Malgonkar:** An Indian author of novels and historical works. His writings are distinct from Sharma's self-help genre.
- **Arundhati Roy:** An Indian author known for 'The God of Small Things'. She writes fiction and essays, different from the genre of 'The Monk Who Sold His Ferrari'.

Based on this analysis, **Robin Sharma** is definitively the author of 'The Monk Who Sold His Ferrari'.

Conclusion

The author of the book 'The Monk Who Sold His Ferrari' is **Robin Sharma**.

Summary of Authors and a Famous Work

Author	Known For (Example Work)	Author of 'The Monk Who Sold His Ferrari'?
Arvind Adiga	The White Tiger	No
Robin Sharma	The Monk Who Sold His Ferrari	Yes
Manohar Malgonkar	The Princes	No
Arundhati Roy	The God of Small Things	No

Revision Table: Key Book Authorship

Recalling Authors for Revision

Book Title	Author
The Monk Who Sold His Ferrari	Robin Sharma
The White Tiger	Arvind Adiga
The God of Small Things	Arundhati Roy

Additional Information: About 'The Monk Who Sold His Ferrari'

'The Monk Who Sold His Ferrari' is a widely acclaimed book by Robin Sharma. It is a fable that tells the story of Julian Mantle, a successful but stressed lawyer who sells his possessions and goes on a spiritual journey to the Himalayas. The book shares timeless lessons on living a fulfilling life, personal development, and achieving inner peace. It is considered a classic in the self-help and motivational genre.

5. Answer: b

Explanation:

Analyzing Statement and Assumptions in Logical Reasoning

This question requires us to analyze a given statement and determine which of the subsequent assumptions can be logically drawn from it. We must consider the statement true, regardless of common knowledge.

Understanding the Statement on Trainee Accountant Appointment

The statement is an extract from an appointment letter for a trainee accountant. It explicitly states:

- The appointment is initially on a **temporary basis**.
- The employment status will **automatically convert to permanent**.
- The condition for conversion to permanent status is the **submission of the pass result of the M.Com final exam**.

This statement sets a clear condition for the trainee accountant to become permanent.

Evaluating Assumption I: M Com Final Exam & Accountant Capability

Assumption I states: "M Com final exam is a reasonable proof of an accountant's capability."

- The statement links passing the M.Com final exam directly to gaining permanent employment as an accountant.
- The employer is making the M.Com pass result the sole criterion for upgrading the trainee accountant's status to permanent.
- This action by the employer implies that they consider passing this specific exam as sufficient evidence of the skills and knowledge required for a permanent accountant role.
- Therefore, the employer views the M.Com final exam result as a valid indicator or proof of an accountant's capability, at least for this position.

Based on the employer's condition for permanent status, this assumption is implicit in the statement.

Evaluating Assumption II: All Trainees on Temporary Basis

Assumption II states: "All trainees are appointed on temporary basis."

- The statement specifically mentions the appointment of "You" as a trainee accountant on a temporary basis.
- It provides information only about **this particular trainee's** appointment status.
- The statement gives no information about the appointment status of other trainees, if any, within the same organization or elsewhere.
- We cannot generalize from one specific instance mentioned in the statement to make a conclusion about all trainees.
- It is possible that other trainees are appointed on a permanent basis from the start, or under different terms and conditions not mentioned here.

Since the statement provides no information about the appointment status of trainees other than the one addressed, this assumption cannot be definitely drawn from the statement. It is not implicit.

Conclusion on Implicit Assumptions

Based on the analysis:

- Assumption I is implicit because the statement uses the M.Com pass result as the condition for permanent employment, implying it is considered proof of capability.
- Assumption II is not implicit because the statement only details one specific trainee's appointment and does not provide information about all trainees.

Therefore, only assumption I is implicit in the given statement.

Assumption	Analysis	Implicit?
I. M Com final exam is a reasonable proof of an accountant's capability.	Statement links M.Com pass to permanent accountant role, implying employer values it as capability proof.	Yes
II. All trainees are appointed on temporary basis.	Statement only discusses one specific trainee's status, not all trainees. Cannot generalize.	No

Revision Table: Statement and Assumption Analysis

Key Concept	Explanation
Statement	A piece of information given as true for the purpose of the question.
Assumption	Something supposed or taken for granted; a hidden premise in the statement.
Implicit Assumption	An assumption that is not directly stated but must necessarily be true if the statement is true.
Drawing Conclusions vs. Making Assumptions	Conclusions are derived directly from the statement; assumptions are underlying beliefs that make the statement valid or meaningful.

Additional Information on Logical Reasoning Questions

Statement and Assumption questions test your ability to identify underlying beliefs or premises that are taken for granted when making a statement. Key points to remember:

- Focus **only** on the information given in the statement. Do not use outside knowledge or common sense unless specifically instructed.
- An assumption is something the speaker or writer **believes to be true** when making the statement.
- Test each assumption by asking: "Does the statement make sense if this assumption is NOT true?" If the statement falls apart without the assumption, then the assumption is implicit.
- Look for keywords in the statement (like "automatically converted on submission") that indicate the basis for the action described.
- Be careful not to confuse assumptions with conclusions, inferences, or implications that go beyond what is strictly necessary for the statement to be made.

Practicing various types of statement and assumption questions helps in understanding the subtle differences and improving analytical skills.

6. Answer: a

Explanation:

Understanding the Grain Bag Profit Problem

The question asks us to determine the selling price per bag for the remaining grain bags so that the trader achieves a specific overall profit percentage on the entire lot of 60 bags. We are given the initial cost, the number of bags sold at a certain profit, and the target overall profit percentage.

Calculating the Total Cost Price

The trader buys 60 bags of grain at Rs. 400 each. The total cost price is calculated by multiplying the number of bags by the cost per bag.

Total Cost Price = Number of Bags $\times$ Cost Price per Bag

Total Cost Price = $60 \times 400 = \text{Rs. } 24000$

Calculating the Target Overall Selling Price

The trader wants to make a 16.4% overall profit on selling all 60 bags. This profit is calculated on the total cost price.

Overall Profit Amount = 16.4% of Total Cost Price

$$\text{Overall Profit Amount} = \frac{16.4}{100} \times 24000$$

$$\text{Overall Profit Amount} = 16.4 \times 240 = \text{Rs. } 3936$$

The target total selling price for all 60 bags is the sum of the total cost price and the overall profit amount.

Target Total Selling Price = Total Cost Price + Overall Profit Amount

$$\text{Target Total Selling Price} = 24000 + 3936 = \text{Rs. } 27936$$

Calculating Revenue from the First 18 Bags

The trader sells 18 bags at an 8% profit. First, let's find the selling price per bag for these 18 bags.

Profit per Bag (on first 18) = 8% of Cost Price per Bag

$$\text{Profit per Bag} = \frac{8}{100} \times 400 = 8 \times 4 = \text{Rs. } 32$$

Selling Price per Bag (for first 18) = Cost Price per Bag + Profit per Bag

$$\text{Selling Price per Bag} = 400 + 32 = \text{Rs. } 432$$

The total revenue generated from selling these 18 bags is the number of bags multiplied by their selling price per bag.

Total Revenue from 18 Bags = Number of Bags $\times$ Selling Price per Bag

$$\text{Total Revenue from 18 Bags} = 18 \times 432 = \text{Rs. } 7776$$

Calculating Required Revenue from Remaining Bags

The total number of bags is 60. After selling 18 bags, the number of remaining bags is:

$$\text{Remaining Bags} = \text{Total Bags} - \text{Bags Sold}$$

$$\text{Remaining Bags} = 60 - 18 = 42$$

The total revenue needed from all 60 bags is Rs. 27936. The revenue already earned from the first 18 bags is Rs. 7776. The revenue that must be generated from the remaining 42 bags is the difference.

Required Revenue from Remaining Bags = Target Total Selling Price - Total Revenue from 18 Bags

$$\text{Required Revenue from Remaining Bags} = 27936 - 7776 = \text{Rs. } 20160$$

Calculating Selling Price per Bag for Remaining Bags

To find the selling price per bag for the remaining 42 bags, we divide the required revenue from these bags by the number of remaining bags.

Selling Price per Bag (for remaining bags) = Required Revenue from Remaining Bags / Number of Remaining Bags

$$\text{Selling Price per Bag} = \frac{20160}{42}$$

Let's perform the division:

$$\frac{20160}{42} = \frac{10080}{21} = \frac{3360}{7} = 480$$

So, the selling price per bag for the remaining 42 bags should be Rs. 480.

Summary of Calculations

Item	Calculation	Value
Total Cost Price (60 bags)	60×400	Rs. 24000
Overall Profit (16.4%)	16.4% of 24000	Rs. 3936
Target Total Selling Price (60 bags)	$24000 + 3936$	Rs. 27936
Profit per bag (first 18)	8% of 400	Rs. 32
Selling Price per bag (first 18)	$400 + 32$	Rs. 432
Total Revenue (first 18 bags)	18×432	Rs. 7776
Remaining Bags	$60 - 18$	42
Required Revenue (remaining 42 bags)	$27936 - 7776$	Rs. 20160
Selling Price per bag (remaining 42 bags)	$20160/42$	Rs. 480

The trader should sell the remaining 42 bags at Rs. 480 each to achieve an overall profit of 16.4% on the 60 bags.

Revision Table: Key Profit Concepts

Concept	Definition	Formula
Cost Price (CP)	The price at which an item is bought.	Given or Calculated
Selling Price (SP)	The price at which an item is sold.	Given or Calculated
Profit	When Selling Price > Cost Price.	$\text{Profit} = \text{SP} - \text{CP}$
Profit Percentage	Profit expressed as a percentage of the Cost Price.	$\text{Profit \%} = \left(\frac{\text{Profit}}{\text{CP}} \right) \times 100$
Overall Profit	The total profit made on a complete transaction or lot of items.	Sum of Profit from individual parts
Overall Profit Percentage	Overall Profit expressed as a percentage of the Total Cost Price.	$\text{Overall Profit \%} = \left(\frac{\text{Overall Profit}}{\text{Total CP}} \right) \times 100$

Additional Information: Calculating Profit & Selling Price

Understanding how profit percentages relate to cost price and selling price is crucial in business math problems. Here's a bit more detail:

- When an item is sold at a profit, the selling price is the cost price plus the profit amount. If the profit percentage is $P\%$, the selling price is $CP + P\% \text{ of } CP = CP \times (1 + P/100)$.
- If you know the selling price and the profit percentage, you can find the cost price using the formula: $CP = \frac{SP}{1 + P/100}$.
- In this problem, we calculated the selling price of the first 18 bags using the formula: $SP = 400 \times (1 + 8/100) = 400 \times (1.08) = 432$.
- For the remaining bags, we worked backward. We found the required total selling price for the remaining bags and then divided by the number of bags to get the price per bag. This price is designed to ensure the overall profit target is met.
- It's important to distinguish between profit on a portion of goods and the overall profit on the entire transaction. The overall profit percentage is always calculated on the total initial cost of all items considered in the transaction.

7. Answer: d

Explanation:

Understanding Inertia and Mass

Inertia is a fundamental property of matter that describes an object's resistance to changes in its state of motion. This means an object at rest tends to stay at rest, and an object in motion tends to stay in motion with the same speed and in the same direction, unless acted upon by an external force. This principle is famously described by Newton's first law of motion, also known as the law of inertia.

The amount of inertia an object possesses is directly related to its mass. Mass is a measure of the amount of matter in an object. The greater the mass of an object, the more difficult it is to change its velocity (either in speed or direction). Therefore, an object with greater mass has greater inertia.

Let's consider the relationship between inertia and the given options:

- **Velocity:** Velocity describes the speed and direction of an object's motion. Inertia is about *resisting* changes in velocity, not about the velocity itself. An object with zero velocity (at rest) still has inertia.
- **Acceleration:** Acceleration is the rate of change of velocity. While inertia is overcome by applying a force which causes acceleration (according to Newton's second law, $F = ma$), acceleration itself is not a measure of inertia. Inertia is the resistance to being accelerated.
- **Volume:** Volume is the amount of space an object occupies. While mass and volume are often related for a given substance (through density), volume alone does not determine inertia. A large balloon (high volume) has much less mass and inertia than a small rock (low volume) because the rock contains much more matter.
- **Mass:** Mass is the measure of the amount of matter in an object. The more mass an object has, the more force is required to produce a given acceleration, meaning it has greater resistance to changes in motion. This resistance is inertia. The relationship is directly proportional: $\text{Inertia} \propto \text{Mass}$.

Therefore, an object with greater mass has greater inertia.

Comparing Physical Properties

Here is a quick comparison of the concepts discussed:

Property	Description	Relation to Inertia
Inertia	Resistance to changes in motion	The property itself
Mass	Amount of matter	Directly proportional; measure of inertia
Velocity	Speed and direction of motion	Inertia resists changes in velocity
Acceleration	Rate of change of velocity	Force overcomes inertia to cause acceleration
Volume	Amount of space occupied	Generally not a direct measure of inertia (unless density is constant)

Conclusion on Inertia

Based on the principles of physics, particularly Newton's laws of motion, the inertia of an object is a direct consequence of its mass. A larger mass means a larger inertia, making it harder to start moving if it's at rest, or harder to stop or change direction if it's in motion.

Revision Table: Key Concepts in Motion

Concept	Brief Definition	Related Law/Principle
Inertia	Resistance to changes in state of motion	Newton's First Law
Mass	Amount of matter; measure of inertia	Fundamental property
Force	A push or pull; causes changes in motion	Newton's First & Second Laws
Acceleration	Rate of change of velocity	Newton's Second Law $F = ma$

Additional Information: Inertial vs. Gravitational Mass

It's interesting to note that physics distinguishes between inertial mass and gravitational mass, although they are experimentally found to be equivalent. Inertial mass is defined by the object's resistance to acceleration (m in $F = ma$). Gravitational mass is defined by the strength of the gravitational force it experiences or exerts (the m in $F = G \frac{Mm}{r^2}$). The equivalence principle states that these two masses are the same, which is a cornerstone of general relativity. For the purpose of understanding inertia as resistance to change in motion, we are concerned with inertial mass.

8. Answer: b

Explanation:

Sphere Diameter Calculation from Surface Area

Let's find the diameter of a sphere when we know its surface area. The question gives us the surface area of the sphere as 1386 cm^2 and asks for the diameter in cm.

Understanding Sphere Formulas

To solve this problem, we need to use the formula for the surface area of a sphere and the relationship between the radius and diameter.

- The surface area (A) of a sphere with radius (r) is given by the formula:

$$A = 4\pi r^2$$

- The diameter (d) of a sphere is twice its radius (r):

$$d = 2r$$

Step-by-Step Calculation

We are given the surface area $A = 1386 \text{ cm}^2$. We can use the surface area formula to first find the radius (r), and then use the radius to find the diameter (d).

Step 1: Find the Radius (r)

Substitute the given surface area into the formula $A = 4\pi r^2$. We will use the value of $\pi \approx \frac{22}{7}$ for this calculation.

$$\begin{aligned} 1386 &= 4 \times \frac{22}{7} \times r^2 \\ 1386 &= \frac{88}{7} \times r^2 \end{aligned}$$

Now, we need to isolate r^2 . Multiply both sides by $\frac{7}{88}$.

$$r^2 = 1386 \times \frac{7}{88}$$

Let's simplify the fraction $\frac{1386}{88}$ first. Both numbers are divisible by 2:

$$r^2 = \frac{1386 \div 2}{88 \div 2} \times 7 = \frac{693}{44} \times 7$$

Now, let's see if 693 is divisible by 11 (since 44 is 4×11). The sum of alternating digits in 693 is $(6 + 3) - 9 = 9 - 9 = 0$. Since the result is 0, 693 is divisible by 11.

$$693 \div 11 = 63$$

$$44 \div 11 = 4$$

So the expression becomes:

$$r^2 = \frac{63}{4} \times 7$$

$$r^2 = \frac{63 \times 7}{4}$$

$$r^2 = \frac{441}{4}$$

To find r, take the square root of both sides:

$$r = \sqrt{\frac{441}{4}}$$

We know that $\sqrt{441} = 21$ and $\sqrt{4} = 2$.

$$r = \frac{21}{2} \text{ cm}$$

So, the radius of the sphere is 10.5 cm.

Step 2: Find the Diameter (d)

The diameter is twice the radius.

$$d = 2r$$

Substitute the value of r we found:

$$d = 2 \times \frac{21}{2}$$

$$d = 21 \text{ cm}$$

Conclusion

The diameter of the sphere with a surface area of 1386 cm^2 is 21 cm.

Revision Table: Sphere Properties

Property	Formula	Variables
Radius	r	Distance from center to surface
Diameter	$d = 2r$	Distance across sphere through center
Surface Area	$A = 4\pi r^2$ or $A = \pi d^2$	Total area of the sphere's surface
Volume	$V = \frac{4}{3}\pi r^3$ or $V = \frac{1}{6}\pi d^3$	Space occupied by the sphere

Additional Information: Sphere Calculations

Spheres are fundamental 3D shapes, and understanding their properties like surface area and volume is crucial in geometry. The value of π is an irrational number, often approximated as 3.14 or $\frac{22}{7}$ for calculations. The choice of approximation can slightly affect the final answer, but $\frac{22}{7}$ is commonly used when it helps simplify calculations involving multiples of 7, as seen in this problem with 1386 (which is divisible by 7: $1386 = 7 \times 198$). Always pay attention to the units given in the problem (cm^2 for area, cm for diameter) and ensure your final answer has the correct units.

9. Answer: a

Explanation:

Dimension Line Arrowhead Length Standards in Technical Drawing

In technical drawing and drafting, dimension lines are used to indicate the size or location of features on an object. These lines typically end with arrowheads, which point to the extension lines (or the object outlines themselves) to clearly show the extent of the dimension being measured.

The shape and size of these arrowheads are standardized to ensure clarity and consistency across different drawings and industries. While specific dimensions can vary slightly based on the drafting standard being used (such as ISO or ANSI) and the scale of the drawing, there are common proportions that are generally followed.

A widely accepted proportion for the length-to-width ratio of a standard arrowhead is approximately 3:1. This means the length of the arrowhead is about three times its maximum width.

The question states that the arrowhead is approximately 1 mm wide. Using the common 3:1 length-to-width ratio, we can calculate the approximate length:

$$\text{Length} = \text{Ratio} \times \text{Width}$$

$$\text{Length} = 3 \times 1 \text{ mm}$$

$$\text{Length} = 3 \text{ mm}$$

Therefore, an arrowhead that is 1 mm wide is approximately 3 mm long according to standard drafting practices.

Let's look at the given options based on this understanding:

- Option 1: 3 mm
- Option 2: 5 mm
- Option 3: 1 mm
- Option 4: 1.5 mm

Comparing our calculated length (3 mm) with the options, Option 1 matches the typical standard dimension for an arrowhead with a 1 mm width.

Revision Table: Technical Drawing Standards

Component	Typical Proportion/Size (relative to width)	Purpose
Dimension Line	Thin continuous line	Indicates the measurement
Extension Line	Thin continuous line, slightly gap from object	Extends from object feature to dimension line
Arrowhead	Approx. 3:1 length-to-width ratio	Indicates the ends of a dimension
Dimension Text	Placed above dimension line	Numerical value of the measurement

Additional Information on Arrowheads and Dimensioning

Arrowheads are crucial elements in technical drawing dimensioning. Their shape and size should be consistent throughout a single drawing. Here are some additional points:

- **Types of Arrowheads:** While the filled arrowhead (as implied by the question) is common, other types exist, such as open arrowheads and slashes (obliques), particularly in architectural or structural drawings.
- **Size Consistency:** All arrowheads within a single drawing should generally be the same size, regardless of the size of the dimension they represent. The size is usually chosen based on the drawing scale and sheet size to ensure readability.
- **Placement:** Arrowheads are placed at the intersection of the dimension line and the extension line (or object line). They point inwards for larger dimensions and may be placed outside the extension lines with the dimension text and line extending from them for smaller dimensions where space is limited.
- **Standards Bodies:** Organizations like the International Organization for Standardization (ISO) and the American National Standards Institute (ANSI) publish detailed standards that govern dimensioning practices, including arrowhead specifications.

Understanding these fundamental elements and their standard proportions, such as the typical arrowhead length being three times its width, is essential for creating clear and professional technical drawings.

Your Personal Exams Guide

10. Answer: a

Explanation:

Solving the Analogy: Letter to Envelope, Dagger to What?

The question asks us to find the word related to 'Dagger' in the same way that 'Envelope' is related to 'Letter'. This is a classic analogy question where we need to identify the relationship between the first pair of words and apply that same relationship to the third word to find the fourth word.

Understanding the Relationship: Letter : Envelope

Let's analyse the relationship between 'Letter' and 'Envelope'.

- A Letter is a written communication.
- An Envelope is a flat paper container used to enclose a letter, especially for mailing.

The relationship here is that an Envelope serves as a protective covering or a container for a Letter. It encloses the letter.

Applying the Relationship: Dagger : ?

Now, we need to find something that serves as a protective covering or container for a 'Dagger'. A Dagger is a short, pointed knife used as a weapon.

Let's examine the given options:

1. Sheath
2. Weapon
3. Sword
4. Sharp

Analysing the Options

- **Sheath:** A sheath is a cover or case for a blade, such as a knife, sword, or dagger. It is used to protect the blade and prevent injury. This fits the protective covering/container relationship.
- **Weapon:** A dagger is a type of weapon, but 'Weapon' is a category or classification. It is not a protective covering for a dagger.
- **Sword:** A sword is another type of bladed weapon, often larger than a dagger. While related in function, 'Sword' is not something that protects or contains a dagger.
- **Sharp:** 'Sharp' describes a characteristic of a dagger (its edge or point). It is not a physical object that protects or contains the dagger.

Conclusion: Finding the Related Word

Based on the analysis, the word that shares the same protective covering/container relationship with 'Dagger' as 'Envelope' does with 'Letter' is 'Sheath'. A sheath protects a dagger.

Pair	Relationship
Letter : Envelope	Envelope is a protective cover/container for a Letter.
Dagger : ?	We need a protective cover/container for a Dagger.

Comparing this relationship to the options, 'Sheath' is the correct choice.

Revision Table: Key Concepts

Term	Definition/Relevance
Analogy	A comparison between two things for the purpose of explanation or clarification, often identifying similar relationships.
Letter	Written communication.
Envelope	A container for a letter.
Dagger	A short, pointed knife used as a weapon.
Sheath	A protective cover for a blade (like a dagger).
Relationship	The connection or similarity between the first pair of words that must be applied to the second pair.

Additional Information on Analogies and Word Relationships

Analogies are common in reasoning tests and vocabulary building. They assess your ability to identify relationships between concepts and apply them in different contexts. Common types of relationships include:

- Part to Whole (e.g., Finger : Hand)
- Cause and Effect (e.g., Heat : Expansion)
- Antonyms (e.g., Hot : Cold)
- Synonyms (e.g., Happy : Joyful)
- Worker and Tool (e.g., Carpenter : Hammer)
- Action and Object (e.g., Read : Book)

- Object and Container (e.g., Letter : Envelope, Dagger : Sheath)

Practicing with different types of analogies helps improve logical reasoning and vocabulary skills. When solving analogies, always try to be precise about the relationship between the first pair before looking at the options.

11. Answer: b

Explanation:

Calculating Work Done: Force and Displacement

The question asks us to find the work done when a specific force pushes a cart over a certain distance. Work done in physics is a measure of energy transfer that occurs when an object is moved over a distance by an external force at least part of which is applied in the direction of the displacement. The mass of the cart is provided (30 kg), but in this particular scenario, it is not needed to calculate the work done, as we are directly given the force applied and the displacement.

Understanding Work Done in Physics

Work done (W) is defined as the product of the force (F) applied on an object and the displacement (d) of the object in the direction of the force. Mathematically, this is expressed as:

$$W = F \times d \times \cos(\theta)$$

Where:

- W is the work done
- F is the magnitude of the force
- d is the magnitude of the displacement
- θ is the angle between the force vector and the displacement vector

In this problem, the force of 750 N pushes the cart by 16 m. We can assume that the force is applied in the direction of the displacement, which means the angle θ between the force and displacement is 0 degrees. The cosine of 0 degrees is 1 ($\cos(0^\circ) = 1$).

Therefore, the formula for work done simplifies to:

$$W = F \times d$$

Step-by-Step Work Done Calculation

Let's use the given values to calculate the work done:

- Force (F) = 750 N
- Displacement (d) = 16 m

Using the simplified formula $W = F \times d$:

$$W = 750 \text{ N} \times 16 \text{ m}$$

$$W = 12000 \text{ J}$$

The work done is calculated in Joules (J), as the force is in Newtons (N) and displacement is in meters (m). The standard unit of work (and energy) in the International System of Units (SI) is the Joule.

Converting Work Done from Joules to Kilojoules

The question asks for the work done in kilojoules (kJ). One kilojoule is equal to 1000 Joules. To convert Joules to Kilojoules, we divide the value in Joules by 1000.

$$1 \text{ kJ} = 1000 \text{ J}$$

So, to convert 12000 J to kJ:

$$W \text{ (kJ)} = \frac{12000 \text{ J}}{1000 \text{ J/kJ}}$$

$$W \text{ (kJ)} = 12 \text{ kJ}$$

Thus, the work done is 12 kJ.

Summary of Work Done Calculation

Here is a summary of the values and the calculation steps:

Quantity	Symbol	Value	Unit
Force	F	750	N
Displacement	d	16	m
Angle between F and d	θ	0	degrees

Calculation:

- Work Done (W) in Joules = $F \times d = 750 \text{ N} \times 16 \text{ m} = 12000 \text{ J}$
- Work Done (W) in Kilojoules = $\frac{12000 \text{ J}}{1000} = 12 \text{ kJ}$

Final Answer on Work Done

The work done by the force of 750 N pushing the cart by 16 m is 12 kJ.

Revision Table: Work, Force, and Displacement

Concept	Definition/Formula	Units (SI)	Key Points
Work Done (W)	$W = F \times d \times \cos(\theta)$	Joule (J)	Scalar quantity; involves force and displacement
Force (F)	A push or pull	Newton (N)	Vector quantity; causes acceleration
Displacement (d)	Change in position	Meter (m)	Vector quantity; straight-line distance

Additional Information: Work and Energy Concepts

Work done is closely related to energy. The work-energy theorem states that the net work done on an object is equal to the change in its kinetic energy. When positive work is done on an object, its kinetic energy increases (assuming no other energy transfers). When negative work is done, its kinetic energy decreases.

Different types of forces can do work. For example, friction can do negative work (opposing motion), gravity can do work (positive or negative depending on movement direction), and applied forces like in this problem can do work.

Understanding the angle θ is crucial:

- If $\theta = 0^\circ$, force and displacement are in the same direction, $\cos(0^\circ) = 1$, $W = F \times d$ (Maximum positive work).
- If $\theta = 90^\circ$, force is perpendicular to displacement, $\cos(90^\circ) = 0$, $W = 0$ (No work done by this force).
- If $\theta = 180^\circ$, force and displacement are in opposite directions, $\cos(180^\circ) = -1$, $W = -F \times d$ (Maximum negative work).

In this problem, the applied force is assumed to be doing positive work on the cart, increasing its kinetic energy or overcoming opposing forces like friction, although friction is not mentioned in the problem statement. The question only asks for the work done by the specific 750 N force over the given displacement.

12. Answer: c

Explanation:

All options except 3, have all the letters similar, while in option 3, there is 'O', which is not there in the other options and also, there is 'Q' missing in option '3', which is present in other options.

Hence, 'option 3' is the correct answer.

13. Answer: b

Explanation:

Solving Coding Language Expression with Symbol Substitution

This question requires us to first understand the given code language where standard mathematical operators are replaced by other operators. After substituting the operators based on the provided rules, we need to evaluate the resulting expression using the standard order of operations (BODMAS/PEMDAS).

Understanding the Coding Language Rules

The rules for operator substitution are given as follows:

- '+' represents 'x'
- '÷' represents '+'
- '-' represents '÷'
- 'x' represents '-'

Evaluating the Expression

The given expression is: $12 \div 6 - 2 + 10 \times 5 = ?$

Now, we substitute the operators in the expression according to the coding language rules:

- Replace '÷' with '+'
- Replace '-' with '÷'
- Replace '+' with 'x'
- Replace 'x' with '-'

Applying these substitutions to the expression $12 \div 6 - 2 + 10 \times 5$, we get:

$$12 + 6 \div 2 \times 10 - 5$$

Step-by-Step Calculation using BODMAS/PEMDAS

Now we evaluate the new expression $12 + 6 \div 2 \times 10 - 5$ using the order of operations (BODMAS/PEMDAS):

Brackets

Orders (powers, square roots, etc.)

Division and Multiplication (from left to right)

Addition and Subtraction (from left to right)

Step 1: Division

Perform the division operation first.

$$12 + 6 \div 2 \times 10 - 5$$

Calculate $6 \div 2$:

$$6 \div 2 = 3$$

The expression becomes:

$$12 + 3 \times 10 - 5$$

Step 2: Multiplication

Next, perform the multiplication operation.

$$12 + 3 \times 10 - 5$$

Calculate 3×10 :

$$3 \times 10 = 30$$

The expression becomes:

$$12 + 30 - 5$$

Step 3: Addition and Subtraction

Finally, perform addition and subtraction from left to right.

$$12 + 30 - 5$$

First, perform the addition:

$$12 + 30 = 42$$

The expression becomes:

$$42 - 5$$

Now, perform the subtraction:

$$42 - 5 = 37$$

The final answer is 37.

Revision Table: Operator Substitution

Original Operator	Replaced By
+	$\times$
$\div$	+
-	$\div$
$\times$	-

Additional Information: Order of Operations (BODMAS/PEMDAS)

The order of operations is a set of rules used to clarify which operations come first in an expression with multiple operations. This ensures a unique result for any given expression.

The acronyms BODMAS and PEMDAS represent the same order:

- **BODMAS:** Brackets, Orders (powers, roots), Division and Multiplication, Addition and Subtraction.
- **PEMDAS:** Parentheses, Exponents, Multiplication and Division, Addition and Subtraction.

It is important to remember that Division and Multiplication have equal priority and should be performed from left to right. Similarly, Addition and Subtraction have equal priority and should also be performed from left to right.

14. Answer: a

Explanation:

Evaluating Mathematical Expressions Using BODMAS

This question asks us to evaluate a mathematical expression by following the correct order of operations. The expression is given as:

$$20 - 2[25\% \text{ of } (15 \times 8 \div 6 + 12)]$$

To solve this, we must use the BODMAS rule (or PEMDAS), which dictates the sequence for performing mathematical operations:

- Brackets (or Parentheses)
- Orders (powers, square roots, etc.) or Exponents
- Division and Multiplication (from left to right)
- Addition and Subtraction (from left to right)

Let's break down the expression step by step according to the BODMAS rule.

Step-by-Step Evaluation

We start with the innermost part of the expression, which is inside the parentheses:

$$(15 \times 8 \div 6 + 12)$$

Step 1: Inside the Parentheses – Multiplication and Division

Within the parentheses, we have multiplication and division. We perform these operations from left to right.

- First, multiplication: $15 \times 8 = 120$
- Then, division: $120 \div 6 = 20$

The expression inside the parentheses now becomes:

$$(20 + 12)$$

Step 2: Inside the Parentheses – Addition

Now, perform the addition inside the parentheses:

$$20 + 12 = 32$$

The original expression now simplifies to:

$$20 - 2[25\% \text{ of } 32]$$

Step 3: Calculate the Percentage

Next, we calculate "25% of 32". Remember that 25% can be written as $\frac{25}{100}$ or $\frac{1}{4}$.

$$25\% \text{ of } 32 = \frac{25}{100} \times 32 = \frac{1}{4} \times 32$$

Performing the multiplication:

$$\frac{1}{4} \times 32 = \frac{32}{4} = 8$$

The expression inside the square brackets $[]$ is now 8. The main expression is now:

$$20 - 2[8]$$

Remember that $2[8]$ means 2×8 .

Step 4: Multiplication

Perform the multiplication:

$$2 \times 8 = 16$$

The expression is now:

$$20 - 16$$

Step 5: Subtraction

Finally, perform the subtraction:

$$20 - 16 = 4$$

The value of the expression is 4.

Summary of Steps

1. Solve inside parentheses: $15 \times 8 \div 6 + 12 \implies 120 \div 6 + 12 \implies 20 + 12 = 32$
2. Calculate percentage: $25\% \text{ of } 32 = 8$
3. Perform multiplication: $2 \times 8 = 16$
4. Perform subtraction: $20 - 16 = 4$

Thus, the result of the expression is 4.

Revision Table: BODMAS/PEMDAS Order

Order	Operation	Description
1	Brackets/Parentheses	Evaluate expressions within grouping symbols first.
2	Orders/Exponents	Calculate powers, roots, etc.
3	Division & Multiplication	Perform from left to right.
4	Addition & Subtraction	Perform from left to right.

Additional Information: Percentage Calculation and Order of Operations

Understanding percentages and the strict order of operations is crucial for accurately solving mathematical expressions.

Understanding Percentage

A percentage represents a part of a whole expressed as a fraction of 100. For example, 25% means 25 out of 100, or $\frac{25}{100}$. To find a percentage of a number, you convert the percentage to a decimal or fraction and multiply by the number.

Example: 25% of 32 = $\frac{25}{100} \times 32 = 0.25 \times 32 = 8$.

Importance of Order of Operations

The order of operations ensures that everyone gets the same answer when evaluating an expression. Without a standard order, different people might perform operations in a different sequence, leading to varying results. BODMAS/PEMDAS provides this standard.

For instance, in $15 \times 8 \div 6$, doing division first ($8 \div 6$) and then multiplication would give a different result (and often involve fractions earlier) than doing multiplication first (15×8) and then division, although in this specific case of sequential multiplication and division, the result is the same if done left-to-right as per the rule.

Always remember to work from left to right for operations at the same level (like division and multiplication, or addition and subtraction).

15. Answer: b

Explanation:

Correct answer: Asbestos particles

Asbestos Exposure: Exposure to asbestos often occurs in occupational settings, such as mining, construction, shipbuilding, and manufacturing, but can also occur in older buildings containing asbestos materials. The risk of lung cancer from asbestos exposure is significantly increased for smokers.

How Asbestos Causes Harm: When inhaled, asbestos fibres can become lodged in the lungs. Over time, these fibres irritate the lung tissues, leading to inflammation, scarring, and potentially genetic damage that can result in cancer.

Other Lung Carcinogens: While asbestos is a key one, other substances like radon, certain metals (chromium, nickel), polycyclic aromatic hydrocarbons (PAHs), and diesel exhaust are also known lung carcinogens. Smoking tobacco is the leading cause of lung cancer.

16. Answer: a

Explanation:

Understanding Standard Paper Sizes: A4 Dimensions

The question asks to identify the paper size corresponding to the dimensions 210 mm × 297 mm. Standard paper sizes, especially in most parts of the world, follow the ISO 216 standard. This standard defines the "A" series of paper sizes, among others.

The ISO 216 standard is based on a constant aspect ratio of $\sqrt{2} \approx 1.414$. This means that if you cut a sheet in half parallel to its shorter side, the two resulting smaller sheets will have the same aspect ratio as the original sheet.

The base size in the A series is A0, which has an area of 1 square meter. Each subsequent size (A1, A2, A3, A4, etc.) is obtained by halving the previous size along its longer side. This results in a sequence of dimensions.

Let's look at the dimensions of the common A series paper sizes according to the ISO 216 standard:

Paper Size	Dimensions (width × height in mm)
A0	841 × 1189
A1	594 × 841
A2	420 × 594
A3	297 × 420
A4	210 × 297
A5	148 × 210
A6	105 × 148

Comparing the given dimensions, 210 mm × 297 mm, with the table above, we can see that these dimensions match the standard dimensions for A4 paper.

Therefore, 210 mm × 297 mm are the dimensions of A4 size paper.

Let's review the options provided:

1. A4: Dimensions are 210 mm × 297 mm.
2. A2: Dimensions are 420 mm × 594 mm.
3. A3: Dimensions are 297 mm × 420 mm.
4. A1: Dimensions are 594 mm × 841 mm.

The dimensions 210 mm × 297 mm exactly match the dimensions of A4 paper.

Revision Table: ISO 216 Paper Dimensions

Paper Size	Dimensions (mm)
A0	841 x 1189
A1	594 x 841
A2	420 x 594
A3	297 x 420
A4	210 x 297
A5	148 x 210

Additional Information on ISO 216 Paper Standards

The ISO 216 international standard defines paper sizes, primarily the A, B, and C series. This standard is widely used around the world, except in countries like the United States and Canada.

- **Aspect Ratio:** All sizes in the A, B, and C series have an aspect ratio of $\sqrt{2} : 1$. This property is useful for scaling documents. For example, reducing an A3 document to A4 size is straightforward scaling.
- **Series Relationship:**
 - A series: Starts with A0 (1 sq meter area), each subsequent size is half of the previous one.
 - B series: Dimensions are the geometric mean of the corresponding A series size and the next larger A series size. For example, B1 is the geometric mean of A0 and A1.
 - C series: Dimensions are the geometric mean of the corresponding A series size and the corresponding B series size. Primarily used for envelopes.
- **Practical Use:** A4 is the most common paper size used for general printing, documents, and stationery worldwide. Larger sizes like A3 are often used for drawings and charts, while smaller sizes like A5 are used for notepads and books.

17. Answer: d

Explanation:

Calculating Mass of an Iron Cube using Density

Let's find the mass of the iron cube using the given information. We are provided with the side length of the cube and the density of iron. To find the mass, we first need to calculate the volume of the cube.

The given parameters are:

Parameter	Value
Side length of the iron cube (s)	2 cm
Density of iron (ρ)	7.8 gm/cm ³

Step-by-Step Solution:

Step 1: Calculate the volume of the cube.

The volume (V) of a cube is calculated by cubing its side length (s). The formula is:

Formula for Volume of a Cube

$$V = s^3$$

Substitute the given side length into the formula:

$$V = (2 \text{ cm})^3$$

$$V = 2 \text{ cm} \times 2 \text{ cm} \times 2 \text{ cm}$$

$$V = 8 \text{ cm}^3$$

So, the volume of the iron cube is 8 cm³.

Step 2: Calculate the mass of the cube.

The relationship between mass (m), density (ρ), and volume (V) is given by the formula:

$$\text{Density} = \frac{\text{Mass}}{\text{Volume}}$$

This can be rearranged to find the mass:

$$\text{Mass} = \text{Density} \times \text{Volume}$$

Or, using symbols:

$$m = \rho \times V$$

Substitute the given density and the calculated volume into this formula:

$$m = 7.8 \text{ gm/cm}^3 \times 8 \text{ cm}^3$$

Now, perform the multiplication:

$$m = 62.4 \text{ gm}$$

The mass of the iron cube is 62.4 gm.

Let's quickly review the calculation:

- Side = 2 cm
- Volume = $(2 \text{ cm})^3 = 8 \text{ cm}^3$
- Density = 7.8 gm/cm^3
- Mass = Density $\times$ Volume = $7.8 \text{ gm/cm}^3 \times 8 \text{ cm}^3 = 62.4 \text{ gm}$

The calculated mass matches one of the options provided.

Revision Table: Key Concepts in Mass, Density, and Volume

Concept	Definition	Formula	Units (SI)
Mass (m)	A measure of the amount of matter in an object.	$m = \rho \times V$	kilograms (kg)
Density (ρ)	Mass per unit volume of a substance.	$\rho = \frac{m}{V}$	kg/m ³
Volume (V)	The amount of space an object occupies.	$V = \frac{m}{\rho}$ (for any shape); $V = s^3$ (for cube)	cubic meters (m ³)

Additional Information on Calculating Mass and Density

Understanding the relationship between mass, density, and volume is fundamental in physics and chemistry. Density is an intrinsic property of a substance, meaning it doesn't change regardless of the amount of the substance (as long as temperature and pressure

are constant). For example, a small piece of iron has the same density as a large block of iron.

When solving problems involving mass, density, and volume, always ensure that the units are consistent. In this problem, the side was in cm and density in gm/cm^3 , which allowed direct calculation of volume in cm^3 and mass in gm. If units were mixed (e.g., side in meters, density in gm/cm^3), you would need to convert them to a consistent system (like all SI units or all CGS units) before calculating.

For objects with irregular shapes, volume can be determined using methods like water displacement, provided the object is denser than the liquid and doesn't react with it.

18. Answer: a

Explanation:

Solving Work and Time Problems: A, B, and C Task Completion

This problem involves calculating the time taken by individuals and a group to complete a task. We need to find out how long A and C take to finish the remaining part of the task after A, B, and C work together for a few days.

The key concept here is the relationship between work, time, and efficiency (or work rate). If a person can complete a task in 'n' days, their daily work rate is $\frac{1}{n}$ of the task.

Calculating Individual Work Rates

First, let's determine the daily work rate for each person:

- A can do the task in 12 days, so A's daily work rate is $\frac{1}{12}$.
- B can do the task in 18 days, so B's daily work rate is $\frac{1}{18}$.
- C can do the task in 36 days, so C's daily work rate is $\frac{1}{36}$.

Combined Work Rate of A, B, and C

When A, B, and C work together, their daily work rates add up to find their combined daily work rate.

Combined daily work rate of A, B, and C = A's rate + B's rate + C's rate

$$\text{Combined rate} = \frac{1}{12} + \frac{1}{18} + \frac{1}{36}$$

To add these fractions, we find a common denominator, which is 36.

$$\text{Combined rate} = \frac{3}{36} + \frac{2}{36} + \frac{1}{36} = \frac{3+2+1}{36} = \frac{6}{36} = \frac{1}{6}$$

So, A, B, and C together complete $\frac{1}{6}$ of the task per day.

Work Done by A, B, and C in 2 Days

They all work together for 2 days. The amount of work done in 2 days is calculated by multiplying their combined daily rate by the number of days they worked together.

Work done in 2 days = Combined daily rate $\times$ Number of days

$$\text{Work done in 2 days} = \frac{1}{6} \times 2 = \frac{2}{6} = \frac{1}{3}$$

They completed $\frac{1}{3}$ of the task in the first 2 days.

Calculating the Remaining Work

The total task is considered as 1 unit of work. After completing $\frac{1}{3}$ of the task, the remaining work is the total work minus the work done.

Remaining work = Total work – Work done

$$\text{Remaining work} = 1 - \frac{1}{3} = \frac{3}{3} - \frac{1}{3} = \frac{2}{3}$$

So, $\frac{2}{3}$ of the task remains to be completed.

Combined Work Rate of A and C

After 2 days, B quits. Now, only A and C are working to complete the remaining task. Let's find their combined daily work rate.

Combined daily work rate of A and C = A's rate + C's rate

$$\text{Combined rate of A and C} = \frac{1}{12} + \frac{1}{36}$$

Using the common denominator 36:

$$\text{Combined rate of A and C} = \frac{3}{36} + \frac{1}{36} = \frac{3+1}{36} = \frac{4}{36} = \frac{1}{9}$$

So, A and C together complete $\frac{1}{9}$ of the task per day.

Time Taken by A and C to Complete Remaining Work

The time taken to complete the remaining work is calculated by dividing the remaining work by the combined daily work rate of the people doing the work (A and C).

$$\text{Time taken} = \frac{\text{Remaining work}}{\text{Combined daily rate of A and C}}$$

$$\text{Time taken} = \frac{2/3}{1/9}$$

To divide by a fraction, we multiply by its reciprocal:

$$\text{Time taken} = \frac{2}{3} \times \frac{9}{1} = \frac{2 \times 9}{3 \times 1} = \frac{18}{3} = 6$$

A and C working together will take 6 days to complete the remainder of the task.

Therefore, the correct answer is 6 days.

Revision Table: Work and Time Concepts

Concept	Explanation	Formula
Work Rate	The amount of work done by a person or group in a unit of time (e.g., per day).	If time taken is 'T' days, rate is $\frac{1}{T}$ per day.
Total Work	Often considered as 1 unit.	Work Rate $\times$ Time Taken
Combined Rate	Sum of individual rates when people work together.	$\text{Rate}_{A+B} = \text{Rate}_A + \text{Rate}_B$
Time for Combined Work	Time taken by a group to complete work.	$\text{Time} = \frac{\text{Total Work}}{\text{Combined Rate}}$
Remaining Work	The portion of the task left after some work is done.	Total Work - Work Done

Additional Information on Task Completion

Understanding work and time problems often involves dealing with fractions and rates. The concept is that efficiency is inversely proportional to the time taken to complete a task. A more efficient person takes less time.

When multiple individuals work together, their efficiencies (or work rates) combine to complete the task faster than any single individual could alone. This is why we add their individual daily work rates.

In scenarios where a person leaves or joins, we calculate the work done in different phases and adjust the remaining work and the set of people working accordingly. Always remember that the total work is usually treated as one whole unit.

For instance, if a task involves building a wall, the 'total work' is building the entire wall. If someone builds half the wall, the 'remaining work' is the other half.

19. Answer: b

Explanation:

Understanding Poisonous Gases

The question asks to identify a gas that is highly poisonous, odourless, tasteless, and colourless from the given options. Let's examine each gas to determine which one matches this description.

Properties of Carbon Dioxide (CO_2)

Carbon dioxide is a naturally occurring gas. Here are some of its key properties:

- **Colour:** Colourless
- **Odour:** Odourless
- **Taste:** Tasteless (in gaseous form)
- **Toxicity:** While not typically considered "highly poisonous" in the same way as carbon monoxide, high concentrations of CO_2 can displace oxygen, leading to asphyxiation. It is not primarily known for direct chemical toxicity at low concentrations.

Carbon dioxide fits some of the criteria (colourless, odourless, tasteless), but it is not usually classified as a "highly poisonous" gas in the context of chemical toxicity like carbon monoxide.

Properties of Carbon Monoxide (CO)

Carbon monoxide is a gas produced by incomplete combustion. Let's look at its properties:

- **Colour:** Colourless
- **Odour:** Odourless
- **Taste:** Tasteless
- **Toxicity:** **Highly poisonous.** Carbon monoxide is dangerous because it binds to haemoglobin in red blood cells much more strongly than oxygen, preventing the blood from carrying oxygen to the body's organs and tissues. Even small amounts can be harmful or fatal.

Carbon monoxide perfectly matches all the characteristics mentioned: it is highly poisonous, odourless, tasteless, and colourless.

Properties of Methane (CH_4)

Methane is a primary component of natural gas. Its properties include:

- **Colour:** Colourless
- **Odour:** Odourless (the smell in natural gas is an additive for leak detection)
- **Taste:** Tasteless
- **Toxicity:** Methane is not considered poisonous. It is primarily an asphyxiant at high concentrations because it displaces oxygen. It is also highly flammable.

Methane fits the colourless, odourless, and tasteless criteria, but it is not a highly poisonous gas.

Properties of Nitrogen Dioxide (NO_2)

Nitrogen dioxide is a gas produced from combustion, particularly from vehicles and power plants. Its properties are:

- **Colour:** Reddish-brown
- **Odour:** Pungent, acrid smell
- **Taste:** Acidic taste (when dissolved)
- **Toxicity:** It is toxic and can cause respiratory problems, but it is neither colourless nor odourless.

Nitrogen dioxide does not fit the description of being colourless, odourless, or tasteless.

Comparing the Properties

Gas	Colour	Odour	Taste	Toxicity
Carbon dioxide	Colourless	Odourless	Tasteless	Asphyxiant (high concentrations)
Carbon monoxide	Colourless	Odourless	Tasteless	Highly poisonous
Methane	Colourless	Odourless	Tasteless	Asphyxiant (high concentrations), Not poisonous
Nitrogen dioxide	Reddish-brown	Pungent	Acidic	Toxic (respiratory issues)

Based on the properties, carbon monoxide is the only gas among the options that is described as highly poisonous, odourless, tasteless, and colourless.

Conclusion on Poisonous Gases

Reviewing the characteristics of each gas option confirms that carbon monoxide is the gas that is highly poisonous, odourless, tasteless, and colourless. This makes it particularly dangerous as it can be present without being detected by human senses.

Revision Table: Properties of Gases

Gas Name	Chemical Formula	Colour	Odour	Taste	Toxicity Level
Carbon Dioxide	CO_2	Colourless	Odourless	Tasteless	Low (Asphyxiant)
Carbon Monoxide	CO	Colourless	Odourless	Tasteless	High (Highly Poisonous)
Methane	CH_4	Colourless	Odourless	Tasteless	Very Low (Asphyxiant)
Nitrogen Dioxide	NO_2	Reddish-brown	Pungent	Acidic	Moderate (Toxic)

Additional Information on Dangerous Gases

Understanding the properties of gases is crucial for safety. Here is some additional information:

- **Carbon Monoxide Poisoning:** CO binds to haemoglobin about 200–250 times more tightly than oxygen, forming carboxyhaemoglobin (COHb). This reduces the blood's ability to carry oxygen, leading to symptoms like headache, dizziness, nausea, weakness, and eventually unconsciousness and death in severe cases. It is often called the "silent killer".
- **Sources of Carbon Monoxide:** Incomplete combustion from faulty furnaces, stoves, fireplaces, generators, vehicle exhaust, and tobacco smoke are common sources. Proper ventilation and carbon monoxide detectors are essential safety measures.
- **Other Toxic Gases:** Besides carbon monoxide and nitrogen dioxide, other toxic gases include hydrogen sulfide (H_2S), sulfur dioxide (SO_2), chlorine (Cl_2), and ammonia (NH_3). These gases often have distinct odours, colours, or irritant properties that can serve as warning signs, unlike carbon monoxide.
- **Asphyxiant Gases:** Gases like nitrogen (N_2), argon (Ar), helium (He), methane (CH_4), and high concentrations of carbon dioxide (CO_2) are considered asphyxiants. They reduce the amount of oxygen available to breathe, leading to suffocation without directly poisoning the body's cells.

Identifying gases by their physical and chemical properties helps in understanding their risks and implementing appropriate safety precautions.

20. Answer: d

Explanation:

Understanding Nutrients from Green and Yellow Vegetables

Green and yellow vegetables are important parts of a healthy diet. They are packed with various vitamins and minerals that our bodies need to function properly. The question

asks which element is mostly provided by these types of vegetables among the given options: Copper, Sodium, Zinc, and Potassium.

Key Elements in Vegetables

Let's look at the elements listed and their general presence in green and yellow vegetables:

- **Potassium:** Many green leafy vegetables like spinach and kale, as well as yellow vegetables like squash and sweet potatoes, are excellent sources of Potassium. Potassium is an essential mineral involved in maintaining fluid balance, nerve signals, and muscle contractions.
- **Sodium:** While some vegetables contain a small amount of naturally occurring sodium, they are generally considered low in sodium compared to processed foods. Sodium is crucial for fluid balance and nerve function, but excessive intake is often a concern.
- **Copper:** Copper is an essential trace mineral found in various foods, including some vegetables, nuts, seeds, and legumes. Its concentration in green and yellow vegetables can vary. Copper is needed for iron metabolism and enzyme function.
- **Zinc:** Zinc is another essential trace mineral found in many foods, including some vegetables, but also prominently in meat, fish, legumes, and nuts. Zinc is important for immune function, protein synthesis, and wound healing.

Comparing Element Abundance

When considering the options provided, Potassium is an element that is particularly abundant in many common green and yellow vegetables, and these vegetables are often recommended specifically for their Potassium content as part of a balanced diet aimed at meeting daily requirements.

While Copper, Sodium (in small amounts), and Zinc are also present in various vegetables, Potassium is often highlighted as a primary mineral contribution from a wide range of green and yellow varieties.

Elements and Vegetable Sources

Element	Common Role in Body	Presence in Green/Yellow Vegetables
Potassium	Fluid balance, nerve signals	High (many varieties)
Sodium	Fluid balance, nerve signals	Low (naturally occurring)
Copper	Iron metabolism, enzymes	Moderate (varies by type)
Zinc	Immune function, protein synthesis	Moderate (varies by type)

Based on typical dietary sources and nutritional recommendations, Potassium is the element among the choices that green and yellow vegetables are most notably known for providing in significant amounts.

Conclusion on Vegetable Nutrition

Therefore, green and yellow vegetables in our diet mostly provide us with Potassium as a key element among the options given.

Revision Table: Essential Elements and Diet

Key Dietary Elements

Element	Major Dietary Sources	Why it's Important
Potassium	Bananas, Potatoes, Spinach, Beans, Lentils	Blood pressure regulation, Muscle function
Sodium	Table salt, Processed foods, Some natural sources	Fluid balance, Nerve transmission
Copper	Nuts, Seeds, Legumes, Whole grains, Some vegetables	Energy production, Connective tissue formation
Zinc	Meat, Seafood, Legumes, Nuts, Seeds, Dairy, Some vegetables	Immune function, Growth, Taste and Smell

Additional Information: Role of Potassium in Health

Potassium is an electrolyte, which means it helps conduct electrical charges in the body. This function is vital for numerous processes, including maintaining a normal heart rhythm and ensuring proper function of nerves and muscles. Adequate potassium intake is often associated with lower blood pressure. Green and yellow vegetables offer a great way to increase dietary potassium intake naturally.

21. Answer: b

Explanation:

Sentences Context: The Vine and Johnsy's Condition

The objective of this question is to arrange a series of jumbled sentences, labeled A, B, C, and D, into a specific order that forms a coherent and logical paragraph. The initial sentence provided serves as the background for the narrative: "Outside their window there was a vine with attractive foliage, but it started shedding leaves one by one as summer passed through fall into winter." This sentence establishes the scene, depicting a vine losing its leaves as the seasons change, setting a somber mood.

Sequenced Sentences Analysis

To determine the correct sequence, let's analyze the content and potential role of each sentence within the story:

- **Sentence B:** "Finally, only one last leaf remained. Johnsy told Sue that when that leaf fell, her life would also go." This sentence introduces the critical turning point of the story. It highlights the dire situation where Johnsy, who is ill, has connected her own life to the fate of the last remaining leaf on the vine. This statement establishes her despair and the central conflict.
- **Sentence D:** "A concerned elderly artist friend of the girls, hearing this story, climbed outside the window in the bitterly cold night and painted a life-like leaf at the spot from where the last leaf had just fallen." This sentence describes a significant action taken by a compassionate character, the artist. Hearing about Johnsy's belief, the artist performs a selfless act of painting a leaf to give Johnsy hope, thereby intervening in the situation described in Sentence B.

- **Sentence A:** "Came dawn, when Johnsy looked out of the window from her sick bed, lo and behold, the leaf was still there and it continued to be there on the following days." This sentence depicts the immediate consequence of the artist's action. When Johnsy wakes up the next morning, she sees the leaf still attached to the vine. This sight, crucial for the artist's plan, provides Johnsy with a reason to believe she might survive, directly following the artist's deed.
- **Sentence C:** "The girl regained her will to live and was fully recovered in a few days." This sentence describes the ultimate positive outcome for Johnsy. Influenced by the apparent survival of the last leaf (as described in Sentence A), Johnsy recovers her motivation to live and consequently gets better. This is the resolution of the narrative.

Logical Order Determination

By piecing together the events in a chronological and cause-and-effect manner, we can establish the most logical order for the sentences:

1. The narrative begins with the contextual setup of the vine losing its leaves.
2. **B** introduces the critical problem: Johnsy's belief that her life is tied to the last leaf.
3. **D** presents the solution: the artist's intervention by painting a substitute leaf.
4. **A** shows the immediate result of the intervention: Johnsy observes the painted leaf the next morning.
5. **C** details the final outcome: Johnsy's recovery due to renewed hope.

Therefore, the most logical sequence that constructs a coherent paragraph is **BDAC**.

Step-by-Step Narrative Construction

Let's trace the construction of the coherent paragraph using the sequence BDAC, starting with the initial context:

- **Initial Context:** Outside their window there was a vine with attractive foliage, but it started shedding leaves one by one as summer passed through fall into winter. (Sets the scene of the vine.)
- **B:** Finally, only one last leaf remained. Johnsy told Sue that when that leaf fell, her life would also go. (Johnsy expresses her pessimistic outlook, linking her life to the last leaf.)
- **D:** A concerned elderly artist friend of the girls, hearing this story, climbed outside the window in the bitterly cold night and painted a life-like leaf at the spot from where the

last leaf had just fallen. (The artist performs a compassionate act to save Johnsy's spirit.)

- **A:** Came dawn, when Johnsy looked out of the window from her sick bed, lo and behold, the leaf was still there and it continued to be there on the following days. (Johnsy sees the leaf, which gives her hope, unaware of the artist's efforts.)
- **C:** The girl regained her will to live and was fully recovered in a few days. (This renewed hope leads to Johnsy's eventual recovery.)

This sequence effectively tells the story, moving logically from the initial situation and Johnsy's despair to the artist's intervention and her subsequent recovery, forming a well-structured narrative.

22. Answer: d

Explanation:

Calculating Specific Heat Capacity of a Metal Block

The question asks us to find the specific heat capacity of a block of metal. We are given the mass of the metal block, the amount of heat energy it absorbs, and the resulting change in its temperature. Specific heat capacity is a fundamental 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.

Understanding Specific Heat Capacity

Specific heat capacity (often denoted by c) is defined by the formula that relates the heat energy absorbed or released (Q) by a substance to its mass (m) and the change in its temperature (ΔT). The formula is:

$$Q = mc\Delta T$$

where:

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

A key point to remember is that a change in temperature measured in degrees Celsius (ΔT in $^{\circ}\text{C}$) is numerically equal to the change in temperature measured in Kelvin (ΔT in K). So, $\Delta T_{^{\circ}\text{C}} = \Delta T_{\text{K}}$.

Given Information and Units

Let's list the given values from the question:

- Mass of the metal block, $m = 500 \text{ g}$
- Heat absorbed by the block, $Q = 10 \text{ kJ}$
- Rise in temperature, $\Delta T = 80^{\circ}\text{C}$

The desired unit for specific heat capacity is $\text{Jg}^{-1}\text{K}^{-1}$. This means we need to work with mass in grams (g), heat in Joules (J), and temperature change in Kelvin (K).

Let's convert the units where necessary:

- Mass m : Already in grams (500 g).
- Heat Q : Given in kilojoules (kJ). We need to convert this to Joules (J).

$$1 \text{ kJ} = 1000 \text{ J}$$

So, $Q = 10 \text{ kJ} = 10 \times 1000 \text{ J} = 10000 \text{ J}$.

- Temperature change ΔT : Given in degrees Celsius ($^{\circ}\text{C}$). The change in temperature in $^{\circ}\text{C}$ is the same as the change in temperature in K.

$$\Delta T = 80^{\circ}\text{C} = 80 \text{ K}$$

Quantity	Given Value	Value in Required Units	Unit
Mass (m)	500 g	500	g
Heat absorbed (Q)	10 kJ	10000	J
Temperature change (ΔT)	80 $^{\circ}\text{C}$	80	K

Calculating the Specific Heat Capacity

We use the formula $Q = mc\Delta T$ and rearrange it to solve for c :

$$c = \frac{Q}{m\Delta T}$$

Now, substitute the values with the correct units:

$$c = \frac{10000 \text{ J}}{500 \text{ g} \times 80 \text{ K}}$$

First, calculate the product of mass and temperature change in the denominator:

$$m \times \Delta T = 500 \text{ g} \times 80 \text{ K} = 40000 \text{ g K}$$

Now, substitute this back into the equation for c :

$$c = \frac{10000 \text{ J}}{40000 \text{ g K}}$$

Perform the division:

$$c = \frac{10000}{40000} \text{ Jg}^{-1}\text{K}^{-1}$$

$$c = \frac{1}{4} \text{ Jg}^{-1}\text{K}^{-1}$$

$$c = 0.25 \text{ Jg}^{-1}\text{K}^{-1}$$

So, the specific heat capacity of the metal block is $0.25 \text{ Jg}^{-1}\text{K}^{-1}$.

Checking the Options

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

- Option 1: 0.4
- Option 2: 0.16
- Option 3: 1.56
- Option 4: 0.25

Our calculated value, $0.25 \text{ Jg}^{-1}\text{K}^{-1}$, matches Option 4.

Revision Table: Specific Heat Calculation

Concept/Formula	Details
Specific Heat Capacity (c) Definition	Amount of heat energy required to raise the temperature of 1 unit mass of a substance by 1 degree.
Heat Transfer Formula	$Q = mc\Delta T$
Units	Q in J, m in g or kg, ΔT in $^{\circ}\text{C}$ or K, c in $\text{Jg}^{-1}\text{K}^{-1}$ or $\text{Jkg}^{-1}\text{K}^{-1}$. ΔT in $^{\circ}\text{C} = \Delta T$ in K.
Given Values	$m = 500 \text{ g}$, $Q = 10 \text{ kJ}$, $\Delta T = 80 \text{ }^{\circ}\text{C}$
Unit Conversions	$Q = 10 \text{ kJ} = 10000 \text{ J}$, $\Delta T = 80 \text{ }^{\circ}\text{C} = 80 \text{ K}$
Formula Rearranged	$c = \frac{Q}{m\Delta T}$
Calculation	$c = \frac{10000 \text{ J}}{500 \text{ g} \times 80 \text{ K}} = \frac{10000}{40000} \text{ Jg}^{-1}\text{K}^{-1} = 0.25 \text{ Jg}^{-1}\text{K}^{-1}$

Additional Information: Specific Heat and Thermal Properties

Specific heat capacity is an intensive property, meaning it does not depend on the size or mass of the sample. It is characteristic of the substance itself. Different materials have different specific heat capacities.

- Substances with high specific heat capacity require more energy to change their temperature. Water, for example, has a high specific heat capacity (about $4.18 \text{ Jg}^{-1}\text{K}^{-1}$), which is why it is used as a coolant and plays a significant role in regulating Earth's climate.
- Substances with low specific heat capacity heat up or cool down quickly when heat is added or removed. Metals generally have lower specific heat capacities compared to water.

This property is crucial in various applications, including thermal engineering, climate science, and cooking.

23. Answer: a

Explanation:

Finding $a^2 + b^2$ Given $a - b$ and ab

The problem asks us to find the value of $a^2 + b^2$, given two pieces of information: the difference between a and b is 5, and the product of a and b is 14.

We are given:

- $a - b = 5$
- $ab = 14$

We need to find the value of $a^2 + b^2$.

Using Algebraic Identities to Solve for $a^2 + b^2$

We can use a common algebraic identity that relates the square of a difference to the squares of the individual terms and their product. The identity is:

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

Our goal is to find $a^2 + b^2$. We can rearrange the identity to isolate $a^2 + b^2$ on one side:

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

Substituting Given Values

Now we can substitute the given values into the rearranged identity:

- Substitute $a - b = 5$ into $(a - b)^2$.
- Substitute $ab = 14$ into $2ab$.

So, the expression becomes:

$$a^2 + b^2 = (5)^2 + 2 \times (14)$$

Calculating the Result

Let's perform the calculations:

- $(5)^2 = 5 \times 5 = 25$
- $2 \times 14 = 28$

Now, add these two results together:

$$a^2 + b^2 = 25 + 28$$

$$a^2 + b^2 = 53$$

Final Answer

The value of $a^2 + b^2$ is 53.

Given Information	Required Value	Identity Used	Calculation	Result
$a - b = 5, ab = 14$	$a^2 + b^2$	$(a - b)^2 = a^2 - 2ab + b^2$	$a^2 + b^2 = (a - b)^2 + 2ab = (5)^2 + 2(14) = 25 + 28$	53

Revision Table: Key Concepts

Reviewing the main points from this problem:

- Recognizing the need for an algebraic identity.
- Using the identity $(a - b)^2 = a^2 - 2ab + b^2$.
- Rearranging the identity to solve for the desired term $(a^2 + b^2)$.
- Substituting given numerical values.
- Performing basic arithmetic calculations accurately.

Additional Information: Algebraic Identities

Algebraic identities are equalities that are true for all values of the variables involved. They are powerful tools for simplifying expressions, factoring polynomials, and solving equations. Some other common algebraic identities include:

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

Understanding and remembering these identities can significantly help in solving various algebra problems efficiently.

24. Answer: d

Explanation:

Calculating the Area of an Equilateral Triangle

The problem asks us to find the area of an equilateral triangle when its side length is given. An equilateral triangle is a triangle where all three sides are equal in length, and all three angles are equal (each being 60 degrees).

Understanding the Formula for Equilateral Triangle Area

The area of an equilateral triangle can be calculated using a specific formula that depends only on the length of its side. If the side length of the equilateral triangle is denoted by s , the area A is given by the formula:

$$A = \frac{\sqrt{3}}{4} \times s^2$$

Applying the Formula to Find the Area

We are given that the side length of the equilateral triangle is 8 cm. Let's plug this value into the formula:

$$s = 8 \text{ cm}$$

Substitute $s = 8$ into the area formula:

$$A = \frac{\sqrt{3}}{4} \times (8 \text{ cm})^2$$

First, calculate the square of the side length:

$$8^2 = 8 \times 8 = 64$$

Now substitute this back into the formula:

$$A = \frac{\sqrt{3}}{4} \times 64 \text{ cm}^2$$

Next, simplify the expression:

$$A = \sqrt{3} \times \frac{64}{4} \text{ cm}^2$$

Divide 64 by 4:

$$\frac{64}{4} = 16$$

So, the area of the equilateral triangle is:

$$A = 16\sqrt{3} \text{ cm}^2$$

Comparing with the Options

Let's look at the given options and compare them with our calculated area:

- Option 1: $32\sqrt{3} \text{ cm}^2$
- Option 2: $8\sqrt{3} \text{ cm}^2$
- Option 3: $64\sqrt{3} \text{ cm}^2$
- Option 4: $16\sqrt{3} \text{ cm}^2$

Our calculated area, $16\sqrt{3} \text{ cm}^2$, matches Option 4.

Conclusion on Equilateral Triangle Area

The area of the equilateral triangle with a side length of 8 cm is $16\sqrt{3} \text{ cm}^2$. This calculation used the standard formula for the area of an equilateral triangle based on its side length.

Revision Table: Equilateral Triangle Properties

Property	Formula (side = s)	Example (side = 8 cm)
Perimeter	$P = 3s$	$P = 3 \times 8 = 24 \text{ cm}$
Area	$A = \frac{\sqrt{3}}{4} s^2$	$A = \frac{\sqrt{3}}{4} \times 8^2 = 16\sqrt{3} \text{ cm}^2$
Height	$h = \frac{\sqrt{3}}{2} s$	$h = \frac{\sqrt{3}}{2} \times 8 = 4\sqrt{3} \text{ cm}$

Additional Information on Triangles

Triangles are fundamental shapes in geometry. Here are some key concepts related to triangles:

- **Types of Triangles:** Triangles can be classified by their side lengths (scalene, isosceles, equilateral) or by their angles (acute, right, obtuse).
- **Area of Any Triangle:** The general formula for the area of any triangle is $A = \frac{1}{2} \times \text{base} \times \text{height}$. For an equilateral triangle, the height can be derived in terms of the side, leading to the specific formula used above.
- **Pythagorean Theorem:** For right-angled triangles, the square of the hypotenuse is equal to the sum of the squares of the other two sides ($a^2 + b^2 = c^2$). This theorem is often used in deriving formulas for triangle properties.
- **Units:** When calculating area, the units are always squared (e.g., cm^2 , m^2). Perimeter is measured in linear units (e.g., cm, m).

25. Answer: d

Explanation:

Understanding Physiological Adjustment: Acclimation Explained

The question asks for the term that describes a physiological adjustment made by an organism in response to changes in its environment. This is a fundamental concept in biology and ecology related to how living things cope with varying conditions.

Let's look at the options provided:

- **Bioremediation:** This process involves using living organisms, such as bacteria, fungi, or plants, to remove pollutants or contaminants from soil, water, or air. This is an environmental cleanup technique, not a physiological adjustment by an organism to environmental change.
- **Bioaccumulation:** This refers to the gradual build-up of substances, such as pesticides, heavy metals, or other chemicals, in an organism. The substance is absorbed at a rate faster than that at which it is lost by catabolism and excretion. This is a process related to pollutant uptake over time, not a physiological adjustment to environmental change itself.
- **Cogeneration:** Also known as combined heat and power (CHP), cogeneration is an energy-efficient process that involves the simultaneous production of two or more forms of energy from a single fuel source, typically electricity and useful heat. This is

an engineering or energy concept and has no relevance to biological physiological adjustments in organisms.

- **Acclimation:** This term specifically describes the process where an individual organism adjusts its physiology or behavior in response to a change in its environmental conditions. These adjustments help the organism survive in the new conditions. Examples include adjusting tolerance to temperature, salinity, or light levels. This precisely matches the definition given in the question: physiological adjustment by an organism to environmental change.

Based on the analysis of the options, **Acclimation** is the correct term that defines a physiological adjustment by an organism to environmental change.

Revision Table: Key Terms

Term	Definition	Relation to Physiological Adjustment
Bioremediation	Using organisms to clean up pollutants	No direct relation
Bioaccumulation	Build-up of substances in an organism	Not a direct physiological adjustment to environmental change
Cogeneration	Simultaneous production of electricity and heat	Not a biological term
Acclimation	Physiological adjustment to environmental change	Directly matches the definition

Additional Information: Acclimation vs. Adaptation

It is important to distinguish between acclimation and adaptation.

- **Acclimation:** This is a short-term or medium-term physiological adjustment made by an individual organism during its lifetime in response to changes in its immediate environment. It is typically reversible if the environmental conditions return to the original state. For example, a person moving to a higher altitude will acclimate by increasing red blood cell production.

- **Adaptation:** This is a long-term evolutionary process where a population of organisms undergoes genetic changes over many generations to become better suited to their environment. These changes are heritable and are the result of natural selection acting on variations within the population. For example, the development of thicker fur in mammal populations in cold climates is an adaptation.

While both involve adjusting to the environment, acclimation is a phenotypic response within an individual's lifetime, whereas adaptation is a genetic change over generations in a population.

26. Answer: c

Explanation:

Understanding the Odd One Out Letter Series Question

This question asks us to identify the letter sequence that is different from the others based on a specific pattern or rule. These types of questions test our ability to recognize patterns in sequences of letters.

To solve this, we need to look at the relationship between the letters within each sequence. A common approach is to consider the alphabetical position of each letter.

Analyzing Letter Patterns in Each Option

Let's assign numerical positions to each letter in the English alphabet (A=1, B=2, C=3, ..., Z=26) and analyze the differences between consecutive letters in each option.

Analyzing HJL Pattern

- H is the 8th letter.
- J is the 10th letter.
- L is the 12th letter.

Let's find the differences between consecutive letter positions:

- Position of J - Position of H = $10 - 8 = 2$
- Position of L - Position of J = $12 - 10 = 2$

The pattern for HJL is adding 2 to the position of the previous letter.

Analyzing PRT Pattern

- P is the 16th letter.
- R is the 18th letter.
- T is the 20th letter.

Let's find the differences between consecutive letter positions:

- Position of R - Position of P = $18 - 16 = 2$
- Position of T - Position of R = $20 - 18 = 2$

The pattern for PRT is adding 2 to the position of the previous letter.

Analyzing VTR Pattern

- V is the 22nd letter.
- T is the 20th letter.
- R is the 18th letter.

Let's find the differences between consecutive letter positions:

- Position of T - Position of V = $20 - 22 = -2$
- Position of R - Position of T = $18 - 20 = -2$

The pattern for VTR is subtracting 2 from the position of the previous letter.

Analyzing UWY Pattern

- U is the 21st letter.
- W is the 23rd letter.
- Y is the 25th letter.

Let's find the differences between consecutive letter positions:

- Position of W - Position of U = $23 - 21 = 2$
- Position of Y - Position of W = $25 - 23 = 2$

The pattern for UWY is adding 2 to the position of the previous letter.

Comparing the Patterns to Find the Difference

Let's summarize the patterns we found:

Sequence	Pattern (Difference in position)
HJL	+2, +2
PRT	+2, +2
VTR	-2, -2
UWY	+2, +2

As we can see from the table, the sequences HJL, PRT, and UWY all follow the pattern of adding 2 to the alphabetical position of the preceding letter. The sequence VTR, however, follows a different pattern: subtracting 2 from the alphabetical position of the preceding letter.

Conclusion: Identifying the Different Sequence

Based on the analysis of the letter patterns using their alphabetical positions, the sequence VTR is the odd one out because it follows a decreasing pattern (-2, -2) while the others follow an increasing pattern (+2, +2).

Revision Table: Letter Series Patterns

Understanding letter patterns is key to solving these questions. Common patterns involve:

- Adding or subtracting a constant value to letter positions.
- Adding or subtracting increasing/decreasing values.
- Skipping a fixed number of letters.
- Using reverse alphabetical order.
- Patterns related to vowels and consonants.

Additional Information: Verbal Reasoning and Pattern Recognition

Odd one out questions are a common type in verbal reasoning and logical aptitude tests. They assess your ability to observe, compare, and find the rule that applies to most items

in a group but not to one. Practicing different types of pattern recognition questions, including number series, letter series, and word analogies, can improve your skills in identifying underlying rules and differences.

27. Answer: b

Explanation:

Correct answer: 437

Decimal value = $\backslash(256 + 128 + 0 + 32 + 16 + 0 + 4 + 0 + 1 \backslash)$

Hexadecimal (Base 16): Uses digits 0–9 and letters A–F (where A=10, B=11, ..., F=15). Each position represents a power of 16 ($\backslash(16^0, 16^1, 16^2 \backslash)$, etc.). Hexadecimal is often used in computing as a shorthand for binary because groups of 4 binary digits can be represented by a single hexadecimal digit.

28. Answer: c

Explanation:

Understanding Network Security Software Guide

The question asks about the type of software used to maintain the security of a private network. Let's look at the options provided to determine which one fits this description best.

Analyzing the Options for Network Security

- **Malware:** Malware is a general term for malicious software designed to harm or exploit computer systems or networks. It includes viruses, worms, ransomware, etc. Malware is a threat to network security, not a tool used to maintain it.
- **Encryption:** Encryption is a process of encoding data so that it can only be decoded by authorized parties. It is a crucial component of data security and privacy, often used for securing communications or stored data, but it is not typically described as the primary software for maintaining the overall security of a private network's perimeter against external threats.

- **Firewall:** A firewall is a network security device or software that monitors and controls incoming and outgoing network traffic based on predetermined security rules. It acts as a barrier between a trusted internal network (like a private network) and untrusted external networks (like the internet), preventing unauthorized access and blocking malicious traffic. This definition directly matches the function of maintaining the security of a private network.
- **Clickbait:** Clickbait refers to online content that uses sensational headlines to entice users to click on a link, typically leading to content of low quality or dubious nature, often for advertising revenue. It is a deceptive online practice and has nothing to do with software used for network security.

Conclusion on Network Security Tool

Based on the analysis of the options, the software used to maintain the security of a private network by monitoring and controlling traffic is a firewall.

Term	Function in Network Security
Malware	Threat; harmful software
Encryption	Secures data content
Firewall	Monitors and controls network traffic for security
Clickbait	Deceptive online content practice

Therefore, the correct answer is firewall.

Revision Table: Network Security Terms

Concept	Brief Description	Role in Security
Firewall	Software or hardware filtering network traffic.	Prevents unauthorized access and controls data flow.
Malware	Malicious software (viruses, worms, etc.).	Poses a threat; aims to damage or steal data.
Encryption	Encoding data to protect its contents.	Ensures data confidentiality and integrity.

Additional Information on Firewalls and Private Network Security

Firewalls are a fundamental part of network security architecture. They can be implemented as software running on individual computers (personal firewalls) or as hardware appliances protecting an entire network. Firewalls use various techniques to filter traffic, including:

- Packet filtering: Examining network packets and allowing/denying based on rules (like source/destination IP, port numbers).
- Stateful inspection: Tracking the state of active connections and making decisions based on the context of the traffic.
- Application-level gateway: Examining traffic at the application layer, offering more granular control.

Maintaining a private network's security involves a combination of tools and practices, but the firewall is a core component for establishing a security perimeter.

29. Answer: b

Explanation:

Constitution Drafting Committee Chairman Explained

Understanding the key figures involved in shaping the Indian Constitution is crucial. The Constitution Drafting Committee played a pivotal role in this process. This committee was tasked with the important responsibility of preparing a draft of the Constitution for India.

Role of the Constitution Drafting Committee

The Constituent Assembly formed several committees to handle different aspects of constitution-making. The **Constitution Drafting Committee** was one of the most important ones. Its primary job was to consider the draft constitution prepared by the Constitutional Adviser, Sir B. N. Rau, and to present a revised draft constitution on the basis of the decisions made by the Union Constitution Committee and other committees.

Identifying the Chairman

The question asks specifically about the **chairman** of this crucial **Constitution Drafting Committee**. The committee was established on August 29, 1947.

- **Dr. B. R. Ambedkar** served as the Chairman of the Constitution Drafting Committee. He was a prominent legal expert and statesman, and his contributions were instrumental in drafting the Constitution. He meticulously guided the committee through the complex process of debating, amending, and finalizing the draft.

Other Key Figures in Constitution Making

While other leaders were deeply involved in the Constituent Assembly and the process of making the Indian Constitution, they held different positions:

- **Jawaharlal Nehru**: He was the first Prime Minister of India and played a significant role as a leader of the Constituent Assembly and the chairman of the Union Constitution Committee and the Steering Committee.
- **Dr. Rajendra Prasad**: He was the President of the Constituent Assembly and later the first President of India. He provided leadership to the assembly sessions.
- **Sardar Patel**: He was a key leader, serving as the Deputy Prime Minister and the Home Minister. He chaired several important committees, including the Advisory Committee on Fundamental Rights, Minorities, and Tribal Areas.

However, the specific role of **chairman** for the **Constitution Drafting Committee** belonged solely to **Dr. B. R. Ambedkar**, recognizing his expertise in constitutional law and his

leadership capabilities in navigating the complexities of drafting India's foundational document.

30. Answer: d

Explanation:

Solving Distance, Speed, and Time Problems

This question involves two cars traveling the same distance but at different average speeds, resulting in different travel times. We are given the speeds of both cars and the difference in their travel times. We need to find the distance between the two points.

Understanding the Relationship: Distance, Speed, Time

The fundamental relationship between distance, speed, and time is:

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

This can be rearranged to find time or speed:

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

In this problem, we are interested in finding the distance, and we can use the time formula to set up an equation based on the given information.

Setting up the Equations

Let:

- D be the distance between A and B in km.
- v_X be the average speed of car X, which is 50 km/hr.
- v_Y be the average speed of car Y, which is 75 km/hr.
- t_X be the time taken by car X to travel from A to B in hours.
- t_Y be the time taken by car Y to travel from A to B in hours.

Using the formula $\text{Time} = \frac{\text{Distance}}{\text{Speed}}$, we can write the expressions for t_X and t_Y :

- For car X: $t_X = \frac{D}{v_X} = \frac{D}{50}$
- For car Y: $t_Y = \frac{D}{v_Y} = \frac{D}{75}$

The problem states that car X takes 2 hours more than car Y for the journey. This gives us a relationship between t_X and t_Y :

$$t_X = t_Y + 2$$

Solving for the Distance

Now, we can substitute the expressions for t_X and t_Y into the equation $t_X = t_Y + 2$:

$$\frac{D}{50} = \frac{D}{75} + 2$$

To solve for D , we need to eliminate the denominators. The least common multiple (LCM) of 50 and 75 is 150. Multiply every term in the equation by 150:

$$150 \times \left(\frac{D}{50}\right) = 150 \times \left(\frac{D}{75}\right) + 150 \times 2$$

$$3D = 2D + 300$$

Now, subtract $2D$ from both sides of the equation:

$$3D - 2D = 300$$

$$D = 300$$

So, the distance between A and B is 300 km.

Verification

Let's check if this distance satisfies the given condition:

- If $D = 300$ km, time taken by car X is $t_X = \frac{300 \text{ km}}{50 \text{ km/hr}} = 6$ hours.
- If $D = 300$ km, time taken by car Y is $t_Y = \frac{300 \text{ km}}{75 \text{ km/hr}} = 4$ hours.

The difference in time is $t_X - t_Y = 6 \text{ hours} - 4 \text{ hours} = 2 \text{ hours}$. This matches the information given in the problem (X takes 2 hours more than Y). Therefore, the calculated distance is correct.

Summary of Calculation Steps

Step	Description	Equation/Calculation
1	Define variables and formulas	$t_X = D/50, t_Y = D/75, t_X = t_Y + 2$
2	Substitute time expressions	$D/50 = D/75 + 2$
3	Find LCM of denominators (50, 75)	LCM = 150
4	Multiply by LCM	$150 \times (D/50) = 150 \times (D/75) + 150 \times 2$
5	Simplify the equation	$3D = 2D + 300$
6	Solve for D	$3D - 2D = 300 \implies D = 300$

Revision Table: Distance, Speed, and Time

Here's a quick look at the key components involved in this type of problem:

Concept	Symbol	Units (Common)	Relationship
Distance	D	km, meters, miles	$D = S \times T$
Speed	S or v	km/hr, m/s, mph	$S = D/T$
Time	T or t	hours, seconds, minutes	$T = D/S$

Additional Information: Relative Speed and Time Difference

This problem can also be conceptualized in terms of relative speed, though the algebraic method is straightforward. The time difference between two objects covering the same distance depends on the difference in their speeds.

- When two objects cover the **same distance** D , and their speeds are v_1 and v_2 , the time taken is $t_1 = D/v_1$ and $t_2 = D/v_2$.
- The difference in time is $|t_1 - t_2| = \left| \frac{D}{v_1} - \frac{D}{v_2} \right| = D \left| \frac{1}{v_1} - \frac{1}{v_2} \right|$.
- In our case, $|t_X - t_Y| = 2$ hours, $v_X = 50$ km/hr, $v_Y = 75$ km/hr. Since $v_Y > v_X$, $t_Y < t_X$, so $t_X - t_Y = 2$.
 - $D \left(\frac{1}{50} - \frac{1}{75} \right) = 2$
 - $D \left(\frac{3}{150} - \frac{2}{150} \right) = 2$
 - $D \left(\frac{1}{150} \right) = 2$
 - $D = 2 \times 150 = 300$ km.

This confirms the result obtained through the direct substitution method. Understanding these relationships is crucial for solving various problems involving motion.

31. Answer: b

Explanation:

Understanding the Umaid Bhawan Palace Location

The question asks us to identify the city where the famous Umaid Bhawan Palace is located. This is a question about the geographical location of a prominent landmark in India, specifically within the state of Rajasthan.

Locating the Umaid Bhawan Palace

The Umaid Bhawan Palace is one of the world's largest private residences. It was built by Maharaja Umaid Singh and is a significant historical and architectural site. To answer the question, we need to know which city in Rajasthan is home to this magnificent palace.

Let's consider the options provided:

- Jaipur
- Jodhpur
- Bikaner
- Udaipur

While all these cities are in Rajasthan and are known for their rich history, forts, and palaces, the Umaid Bhawan Palace is specifically located in Jodhpur.

Why Jodhpur is the Correct City

The Umaid Bhawan Palace is situated in the city of Jodhpur, often called the "Blue City" due to the blue-painted houses around the Mehrangarh Fort. The palace is built on Chittar Hill, which is why it is also sometimes referred to as Chittar Palace.

Construction of the Umaid Bhawan Palace began in 1929 and was completed in 1943. It was intended to provide employment to people during a time of famine. Today, part of the

palace is managed by the Taj Hotels group, another part is the residence of the royal family, and there is also a museum.

Examining Other Options

Let's briefly look at why the other options are incorrect:

- **Jaipur:** Jaipur is the capital of Rajasthan and is known for palaces like the Hawa Mahal, City Palace, and Amer Fort. The Umaid Bhawan Palace is not located in Jaipur.
- **Bikaner:** Bikaner is known for Junagarh Fort and Lalgarh Palace. The Umaid Bhawan Palace is not in Bikaner.
- **Udaipur:** Udaipur is famous as the "City of Lakes" and is home to the City Palace, Jag Mandir, and Jag Niwas (Lake Palace). The Umaid Bhawan Palace is not located in Udaipur.

Therefore, based on the historical and geographical facts, the Umaid Bhawan Palace is located in Jodhpur.

Conclusion

The Umaid Bhawan Palace is a significant landmark located in Jodhpur, Rajasthan.

City	Known Palaces/Forts	Umaid Bhawan Palace Location?
Jaipur	Hawa Mahal, City Palace, Amer Fort	No
Jodhpur	Umaid Bhawan Palace, Mehrangarh Fort	Yes
Bikaner	Junagarh Fort, Lalgarh Palace	No
Udaipur	City Palace, Lake Palace, Jag Mandir	No

Revision Table: Key Facts about Umaid Bhawan Palace

Feature	Detail
Location	Jodhpur, Rajasthan
Built by	Maharaja Umaid Singh
Construction Period	1929 - 1943
Purpose	Famine relief project, Royal Residence
Current Use	Hotel, Royal Residence, Museum

Additional Information on Rajasthan's Palaces

Rajasthan is renowned for its numerous historical forts and palaces, testaments to the region's royal past. These structures are major tourist attractions and reflect diverse architectural styles, often blending Rajput, Mughal, and European influences.

- Many palaces have been converted into heritage hotels, offering visitors a chance to experience royal luxury.
- Forts like Mehrangarh in Jodhpur, Amer Fort in Jaipur, and Chittorgarh Fort have played crucial roles in history.
- These historical sites contribute significantly to Rajasthan's economy through tourism.

Understanding the locations of key palaces like Umaid Bhawan Palace is important for general knowledge and competitive exams focusing on Indian history and geography.

32. Answer: c

Explanation:

Understanding the Speed, Distance, and Time Problem

This question is a classic problem involving speed, distance, and time. We have two individuals, A and B, moving between two points, P and Q. The distance between P and Q is given as 9 km. We are given information about their starting times, speeds, and where they meet, and we need to find B's speed.

Step-by-Step Solution to Calculate B's Speed

1. Calculate A's Speed

A's speed is given as 1 kilometre in 10 minutes.

- Distance covered by A = 1 km
- Time taken by A = 10 minutes
- A's speed = $\frac{\text{Distance}}{\text{Time}}$
- A's speed = $\frac{1 \text{ km}}{10 \text{ minutes}} = \frac{1}{10} \text{ km/minute}$

2. Calculate Time A Takes to Reach Point Q

The distance from P to Q is 9 km. A's speed is $\frac{1}{10} \text{ km/minute}$.

- Distance P to Q = 9 km
- A's speed = $\frac{1}{10} \text{ km/minute}$
- Time taken by A to reach Q = $\frac{\text{Distance}}{\text{Speed}}$
- Time taken by A to reach Q = $\frac{9 \text{ km}}{\frac{1}{10} \text{ km/minute}} = 9 \times 10 \text{ minutes} = 90 \text{ minutes}$

3. Determine When and Where A Meets B

A reaches Q in 90 minutes and immediately returns. A meets B after walking 1 km from Q towards P.

- Distance A travelled from Q back towards P = 1 km
- Total distance A covered until meeting B = Distance P to Q + Distance Q back towards P
- Total distance A covered = 9 km + 1 km = 10 km

Now, let's find the total time A walked until meeting B.

- Total distance A covered = 10 km
- A's speed = $\frac{1}{10} \text{ km/minute}$
- Total time A walked = $\frac{\text{Total distance covered by A}}{\text{A's speed}}$
- Total time A walked = $\frac{10 \text{ km}}{\frac{1}{10} \text{ km/minute}} = 10 \times 10 \text{ minutes} = 100 \text{ minutes}$

So, A meets B 100 minutes after A started walking.

4. Determine the Distance B Covered Until Meeting A

The meeting point is 1 km away from Q on the path towards P. Since B started from P (which is 9 km from Q) and walked towards Q, B's distance from P to the meeting point is:

- Meeting point distance from Q = 1 km
- Distance P to Q = 9 km
- Distance B covered = Distance P to Q - Meeting point distance from Q
- Distance B covered = 9 km - 1 km = 8 km

5. Determine the Time B Walked Until Meeting A

B started 4 minutes after A. A walked for 100 minutes until they met.

- Total time elapsed since A started = 100 minutes
- B started late by = 4 minutes
- Time B walked = Total time elapsed - B's start delay
- Time B walked = 100 minutes - 4 minutes = 96 minutes

6. Calculate B's Speed

B covered 8 km in 96 minutes.

- Distance covered by B = 8 km
- Time B walked = 96 minutes
- B's speed = $\frac{\text{Distance covered by B}}{\text{Time B walked}}$
- B's speed = $\frac{8 \text{ km}}{96 \text{ minutes}}$
- B's speed = $\frac{8}{96} \text{ km/minute}$

Simplifying the fraction:

- $\frac{8}{96} = \frac{1 \times 8}{12 \times 8} = \frac{1}{12}$
- B's speed = $\frac{1}{12} \text{ km/minute}$

Final Answer

B's speed is $\frac{1}{12}$ kilometres per minute.

Revision Table: Key Values

Parameter	Value
Distance P to Q	9 km
A's Speed	$\frac{1}{10}$ km/minute
Time A took to reach Q	90 minutes
Total distance A covered until meeting	10 km
Total time elapsed until meeting	100 minutes
B's start delay	4 minutes
Time B walked until meeting	96 minutes
Distance B covered until meeting	8 km
B's Speed	$\frac{1}{12}$ km/minute

Additional Information: Relative Speed Concept

While this problem was solved by calculating the time each person walked and the distance they covered individually, it can also be approached using the concept of relative speed, especially when considering the point where they meet.

- When two objects move towards each other, their relative speed is the sum of their individual speeds. The time taken to meet is the initial distance between them divided by their relative speed.
- When two objects move in the same direction, the faster object's speed relative to the slower one is the difference between their speeds.
- In this problem, A and B are initially moving towards Q (same direction), then A turns back towards B (opposite direction). The meeting point is crucial. The total distance covered by A and B combined from the moment B starts until they meet is equal to the distance from P to Q plus the distance A travels back. This perspective can sometimes simplify problems.

In this specific case, calculating individual times and distances was straightforward because we could easily determine the exact time duration for which each person travelled until the meeting occurred.

33. Answer: a

Explanation:

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

This question asks us to determine the relationship between P and S based on a coded expression representing family ties. We are given specific codes that represent different relationships.

Understanding the Relationship Codes

Let's first break down the meaning of each code provided:

- $A \$ B$ means A is the husband of B.
- $A \# B$ means A is the brother of B.
- $A * B$ means A is the mother of B.

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

Now, let's decode the given expression step by step, moving from left to right:

The expression is $P \$ Q * R \# S$.

Step 1: Analyze $P \$ Q$

According to the given rules, $A \$ B$ means A is the husband of B.

Therefore, $P \$ Q$ means P is the husband of Q.

This implies Q is the wife of P. P is male, and Q is female.

Step 2: Analyze $Q * R$

According to the given rules, $A * B$ means A is the mother of B.

Therefore, $Q * R$ means Q is the mother of R.

From Step 1, we know Q is the wife of P. If Q is the mother of R, and Q is married to P, then P must be the father of R.

So, at this point, we know P is the father of R.

Step 3: Analyze R # S

According to the given rules, A # B means A is the brother of B.

Therefore, R # S means R is the brother of S.

This means R is male, and R and S are siblings (they share the same parents). S could be male or female, as the '# B' rule specifies A's gender (brother) but not B's.

Connecting the Relationships

Let's put the relationships together:

- P is the father of R (from Step 1 and Step 2).
- R is the brother of S (from Step 3).

If P is the father of R, and R is the brother of S, it means R and S share the same parents. Since P is the father of R, P must also be the father of S.

Conclusion: P's Relationship to S

Based on the analysis, the relationship between P and S is that P is the father of S.

Evaluating the Options

Let's compare our conclusion with the given options:

1. P is the father of S
 - This matches our conclusion.
2. P is the brother-in-law of S
 - Incorrect.
3. P is the son of S
 - Incorrect.
4. P is the brother of S

– Incorrect.

Therefore, the relationship P is the father of S is the correct interpretation of the expression $P \$ Q * R \# S$.

Blood Relations Revision Table

Relation Code	Meaning	Example (A, B)
$A \$ B$	A is the husband of B	A is husband, B is wife
$A \# B$	A is the brother of B	A is male, A and B are siblings
$A * B$	A is the mother of B	A is female, B is her child

Additional Information on Blood Relations Concepts

Blood relations questions test your ability to understand and decode relationships within a family tree. They often involve coded language or descriptions of relationships. Key concepts to remember include:

- **Parent:** Father or Mother.
- **Child:** Son or Daughter.
- **Sibling:** Brother or Sister.
- **Spouse:** Husband or Wife.
- **In-laws:** Relatives by marriage (e.g., father-in-law, sister-in-law).
- **Aunt/Uncle:** Sibling of a parent.
- **Cousin:** Child of an aunt or uncle.
- **Grandparent:** Parent of a parent.
- **Grandchild:** Child of a child.

When solving these problems, it is often helpful to draw a family tree or diagram to visualize the relationships as you decode each part of the expression.

34. Answer: a

Explanation:

Understanding Angles Between Clock Hands

The question asks us to find the angle between the hour hand and the minute hand of a clock at half past two, which is 2:30. To solve this, we need to understand how both the hour hand and the minute hand move around the clock face.

Movement of the Minute Hand

The minute hand completes a full circle (360°) in 60 minutes. This means:

- In 1 minute, the minute hand moves: $\frac{360^\circ}{60 \text{ minutes}} = 6^\circ$ per minute.

At 2:30, the minute hand is exactly at the 30-minute mark, which corresponds to the number '6' on the clock face. Its position relative to the '12' (which is 0°) is:

Position of minute hand = 30 minutes $\times 6^\circ/\text{minute} = 180^\circ$.

Movement of the Hour Hand

The hour hand completes a full circle (360°) in 12 hours. This means:

- In 1 hour, the hour hand moves: $\frac{360^\circ}{12 \text{ hours}} = 30^\circ$ per hour.
- Since 1 hour = 60 minutes, in 1 minute, the hour hand moves: $\frac{30^\circ}{60 \text{ minutes}} = 0.5^\circ$ per minute.

At 2:30, the hour hand has moved past the '2' mark. The position of the '2' mark is 2 hours from the '12'. Its initial position at 2:00 would be:

Initial position of hour hand (at 2:00) = 2 hours $\times 30^\circ/\text{hour} = 60^\circ$.

However, at 2:30, the hour hand has also moved for an additional 30 minutes. The distance it moves in 30 minutes is:

Movement in 30 minutes = 30 minutes $\times 0.5^\circ/\text{minute} = 15^\circ$.

So, the total position of the hour hand at 2:30 relative to the '12' is:

Total position of hour hand = Initial position + Movement in 30 minutes = $60^\circ + 15^\circ = 75^\circ$.

Calculating the Angle Between the Hands

To find the angle between the hour hand and the minute hand, we find the absolute difference between their positions relative to the '12'.

$$\text{Angle} = |\text{Position of minute hand} - \text{Position of hour hand}|$$

$$\text{Angle} = |180^\circ - 75^\circ|$$

$$\text{Angle} = |105^\circ| = 105^\circ.$$

The angle between the hands can sometimes be calculated in two ways (clockwise or counter-clockwise, resulting in angles that sum to 360°). We usually refer to the smaller angle. In this case, the calculated angle is 105° , which is less than 180° , so this is the smaller angle.

Summary of Hand Positions at 2:30

Hand	Position from '12' (in degrees)
Minute Hand	180°
Hour Hand	75°

$$\text{Difference} = |180^\circ - 75^\circ| = 105^\circ.$$

Conclusion

The angle between the hour hand and the minute hand at half past two (2:30) is 105° . This aligns with one of the provided options.

Revision Table: Clock Hand Angles

Hand	Movement per Hour	Movement per Minute	Reference Point
Minute Hand	360°	6°	'12' mark (0°)
Hour Hand	30°	0.5°	'12' mark (0°)

Additional Information on Clock Angles

You can use a general formula to find the angle between the hour and minute hands at any given time $H : M$ (where H is hours, M is minutes):

$$\text{Angle} = |30H - 11 \times \frac{M}{2}|$$

Let's apply this formula for 2:30 (H=2, M=30):

$$\text{Angle} = |30 \times 2 - 11 \times \frac{30}{2}|$$

$$\text{Angle} = |60 - 11 \times 15|$$

$$\text{Angle} = |60 - 165|$$

$$\text{Angle} = |-105| = 105^\circ.$$

This formula accounts for both the hour hand's position based on the hour and its additional movement based on the minutes past the hour, and the minute hand's position. Remember that if the formula gives an angle greater than 180° , subtract it from 360° to get the smaller angle.

35. Answer: d

Explanation:

Understanding Shivkumar Sharma and His Musical Instrument

The question asks about the musical instrument associated with the renowned classical musician, Shivkumar Sharma.

Pandit Shivkumar Sharma was a celebrated Indian classical musician who specialized in playing the Santoor. The Santoor is a hammered dulcimer that originated in the Kashmir region. While the instrument has ancient roots, Shivkumar Sharma is widely credited with transforming it from a folk instrument into a respected instrument for Indian classical music on the global stage.

Analysing the Options

Let's look at the given options:

1. Shehnai: The Shehnai is a wind instrument, famous players include Ustad Bismillah Khan. Shivkumar Sharma was not associated with the Shehnai.

2. Tabla: The Tabla is a pair of hand drums, a primary percussion instrument in North Indian classical music. Many great musicians play the Tabla, such as Zakir Hussain. Shivkumar Sharma played a melodic instrument, not percussion like the Tabla.
3. Violin: The Violin is a string instrument played with a bow. While used in Indian classical music (Carnatic and Hindustani), it is not the instrument Shivkumar Sharma is known for.
4. Santoor: The Santoor is a string instrument played with mallets. As mentioned, Shivkumar Sharma is famously associated with the Santoor, pioneering its role in Indian classical music.

Based on his significant contribution and mastery, **Shivkumar Sharma** is synonymous with the **Santoor**.

Famous Musicians and Their Instruments

To further clarify, here is a table listing some famous Indian classical musicians and the instruments they are known for:

Musician	Instrument
Pandit Shivkumar Sharma	Santoor
Ustad Bismillah Khan	Shehnai
Pandit Ravi Shankar	Sitar
Ustad Zakir Hussain	Tabla
Pandit Hariprasad Chaurasia	Flute (Bansuri)

This table clearly shows the association of **Pandit Shivkumar Sharma** with the **Santoor**.

Conclusion on Shivkumar Sharma's Instrument

Therefore, the musical instrument associated with classical musician **Shivkumar Sharma** is the **Santoor**.

Revision Table: Shivkumar Sharma's Instrument

Question Aspect	Key Detail
Musician's Name	Shivkumar Sharma
Associated Instrument	Santoor
Significance	Pioneered Santoor in Indian Classical Music

Additional Information on Santoor and Shivkumar Sharma

The Santoor is a trapezoid-shaped hammered dulcimer, with origins possibly tracing back to Persia or Kashmir. It has a large number of strings stretched over bridges, and it is played by striking the strings with two light wooden mallets. Traditional Santoor from Kashmir were often used in Sufi music.

Pandit Shivkumar Sharma's contribution was crucial in adapting the Santoor for the performance of complex Hindustani classical ragas. He modified the instrument itself to increase its range and sustain, and developed new playing techniques to make it suitable for classical performances, which typically involve elaborate improvisations and intricate melodic patterns. His work elevated the status of the **Santoor** in the world of classical music.

36. Answer: c

Explanation:

Calculating Charge from Work and Potential Difference

This problem requires us to find the amount of charge (Q) moved when a specific amount of work (W) is done against a given potential difference (V). The fundamental relationship connecting these three quantities is:

$$\text{Work Done } (W) = \text{Charge } (Q) \times \text{Potential Difference } (V)$$

The units for these quantities in the SI system are Joules (J) for work, Coulombs (C) for charge, and Volts (V) for potential difference.

We are given:

- Work done, $W = 36 \text{ J}$
- Potential difference, $V = 8 \text{ V}$

We need to find the charge, Q . We can rearrange the formula to solve for Q :

$$Q = \frac{W}{V}$$

Now, substitute the given values into the rearranged formula:

$$Q = \frac{36 \text{ J}}{8 \text{ V}}$$

Performing the division:

$$Q = 4.5 \text{ C}$$

So, the charge moved across the 8 V potential difference, with 36 J of work done, is 4.5 Coulombs.

Let's look at the given options:

- Option 1: 288
- Option 2: 9
- Option 3: 4.5
- Option 4: 16

Our calculated value of 4.5 Coulombs matches Option 3.

Revision Table: Key Concepts

Concept	Symbol	SI Unit	Formula Relation
Work Done	W	Joule (J)	$W = Q \times V$
Charge	Q	Coulomb (C)	$Q = W/V$
Potential Difference	V	Volt (V)	$V = W/Q$

Additional Information: Understanding Work, Charge, and Voltage

Work (W): In the context of electricity, work is done when a force causes a charge to move over a distance against an electric field. It represents the energy transferred.

Charge (Q): Charge is a fundamental property of matter. It is measured in Coulombs (C). A Coulomb is a large unit; electron charge is about 1.602×10^{-19} C.

Potential Difference (V): Also known as voltage, it is the difference in electric potential energy per unit charge between two points in an electric circuit. It represents the energy required to move a unit charge between those points. One Volt is equal to one Joule per Coulomb ($1 \text{ V} = 1 \text{ J/C}$). This unit relationship directly explains why W/V gives the charge Q in Coulombs.

Understanding the relationship $W = QV$ is crucial for solving problems involving energy transfer in electric fields or circuits. This problem is a direct application of this fundamental formula.

37. Answer: d

Explanation:

Understanding the Tank Filling and Emptying Problem

This question involves two pipes, one filling a tank and the other emptying it. To solve this, we need to understand the rate at which each pipe works individually and then calculate their combined effect when working together.

Calculating Individual Pipe Rates

- Pipe A is a filling pipe. It fills the tank in 6 hours. This means in 1 hour, Pipe A fills $\frac{1}{6}$ of the tank. The filling rate of Pipe A is $\frac{1}{6}$ tank per hour.
- Pipe B is an emptying pipe. It empties the tank in 15 hours. This means in 1 hour, Pipe B empties $\frac{1}{15}$ of the tank. The emptying rate of Pipe B is $\frac{1}{15}$ tank per hour. Note that emptying is the opposite of filling, so its contribution is negative in the combined calculation.

Finding the Combined Rate of Filling

When both pipes are opened together, their rates combine. Since Pipe A is filling and Pipe B is emptying, the net rate at which the tank is being filled is the difference between the filling rate and the emptying rate.

Combined rate = (Rate of Pipe A) - (Rate of Pipe B)

Combined rate = $\frac{1}{6} - \frac{1}{15}$ tank per hour.

Calculating the Combined Rate

To subtract these fractions, we need a common denominator. The least common multiple (LCM) of 6 and 15 is 30.

- Convert $\frac{1}{6}$ to a fraction with denominator 30: $\frac{1}{6} \times \frac{5}{5} = \frac{5}{30}$
- Convert $\frac{1}{15}$ to a fraction with denominator 30: $\frac{1}{15} \times \frac{2}{2} = \frac{2}{30}$

Now, subtract the fractions:

Combined rate = $\frac{5}{30} - \frac{2}{30} = \frac{5-2}{30} = \frac{3}{30}$

Simplify the combined rate:

Combined rate = $\frac{3}{30} = \frac{1}{10}$ tank per hour.

This means that when both pipes are open, $\frac{1}{10}$ of the tank is filled every hour.

Determining the Time to Fill the Tank

If $\frac{1}{10}$ of the tank is filled in 1 hour, then the total time taken to fill the entire tank (which is 1 whole tank) is the reciprocal of the combined rate.

Time taken = $\frac{1}{\text{Combined rate}}$

Time taken = $\frac{1}{\frac{1}{10}}$ hours

Time taken = $1 \times \frac{10}{1}$ hours

Time taken = 10 hours.

Therefore, if both pipes are opened together, the tank will be filled in 10 hours.

Revision Table: Key Concepts

Concept	Description	Formula
Individual Rate (Filling)	Fraction of work done by a filling pipe in one unit of time (e.g., 1 hour).	$\frac{1}{\text{Time taken to fill}}$
Individual Rate (Emptying)	Fraction of work done by an emptying pipe in one unit of time. Treated as negative work.	$\frac{-1}{\text{Time taken to empty}}$
Combined Rate	The net fraction of work done by all pipes together in one unit of time. Sum of individual rates.	Sum of individual rates (filling rates are positive, emptying rates are negative)
Time Taken (Combined)	The total time required to complete the work (fill the tank) at the combined rate.	$\frac{1}{\text{Combined rate}}$ (if combined rate is positive, meaning filling is faster than emptying)

Additional Information on Tank Filling Problems

Problems involving pipes and tanks are often solved using the concept of 'work rate'. Here's a bit more detail:

- **Work Rate:** If a job (like filling a tank) can be completed in 'T' units of time, the rate of work is $1/T$ of the job per unit of time.
- **Multiple Workers (Pipes):** When multiple pipes work together, their individual work rates are added to find the combined work rate.
- **Filling vs. Emptying:** Filling pipes contribute positively to the work (filling the tank), while emptying pipes contribute negatively (emptying the tank). So, when both types are involved, you subtract the emptying rate from the filling rate.
- **Total Work:** The total work is typically considered as '1 unit' (representing the full tank).
- **Time = Total Work / Combined Rate:** This fundamental relationship is used to find the time taken when the combined rate is known.

It's important to ensure the combined rate is positive for the tank to actually fill up. If the combined rate were negative (meaning the emptying rate is higher than the filling rate), the tank would be emptied, not filled, and the question might ask for the time taken to empty a full tank instead.

38. Answer: a

Explanation:

Calculating Equilibrium Temperature When Mixing Water

This problem involves mixing two quantities of water at different temperatures and finding the final equilibrium temperature, assuming no heat is lost to the surroundings. This scenario is a classic example of heat exchange or calorimetry.

When substances at different temperatures are mixed, heat flows from the hotter substance to the colder substance until they reach a common final temperature, known as the equilibrium temperature. According to the principle of conservation of energy, if no heat is lost to the surroundings, the heat lost by the hot substance is equal to the heat gained by the cold substance.

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

$$Q = mc\Delta T$$

Where:

- m is the mass of the substance.
- c is the specific heat capacity of the substance.
- ΔT is the change in temperature.

In this problem, we are mixing water with water. Water has a specific heat capacity, denoted by c_w . The density of water (ρ_w) allows us to relate volume (V) to mass (m) using the formula $m = \rho_w V$.

Let's define the properties of the hot and cold water:

- Volume of hot water (V_h) = 0.5 litres
- Initial temperature of hot water (T_h) = 90°C
- Volume of cold water (V_c) = 3.5 litres
- Initial temperature of cold water (T_c) = 10°C
- Final equilibrium temperature (T_f) = ?

We assume the density of water (ρ_w) and specific heat capacity of water (c_w) are constant during this process. The mass of the hot water is $m_h = \rho_w V_h$, and the mass of the

cold water is $m_c = \rho_w V_c$.

Applying the Principle of Heat Exchange

Since no heat is lost to the surroundings, the heat lost by the hot water equals the heat gained by the cold water:

$$Q_{lost,hot} = Q_{gained,cold}$$

The hot water cools down from T_h to T_f , so its temperature change is $\Delta T_h = T_h - T_f$. The cold water heats up from T_c to T_f , so its temperature change is $\Delta T_c = T_f - T_c$.

Using the heat transfer formula:

$$m_h c_w (T_h - T_f) = m_c c_w (T_f - T_c)$$

Substitute $m = \rho_w V$:

$$\rho_w V_h c_w (T_h - T_f) = \rho_w V_c c_w (T_f - T_c)$$

Since ρ_w and c_w are present on both sides and are not zero, we can cancel them out:

$$V_h (T_h - T_f) = V_c (T_f - T_c)$$

Solving for the Final Temperature

Now, substitute the given values into the equation:

$$0.5 \text{ L} \times (90^\circ\text{C} - T_f) = 3.5 \text{ L} \times (T_f - 10^\circ\text{C})$$

Expand both sides of the equation:

$$0.5 \times 90 - 0.5 \times T_f = 3.5 \times T_f - 3.5 \times 10$$

$$45 - 0.5T_f = 3.5T_f - 35$$

Now, rearrange the terms to isolate T_f . Move the terms containing T_f to one side and the constant terms to the other side:

$$45 + 35 = 3.5T_f + 0.5T_f$$

$$80 = 4T_f$$

Finally, solve for T_f :

$$T_f = \frac{80}{4}$$

$$T_f = 20^\circ\text{C}$$

The final equilibrium temperature when mixing half a litre of hot water at 90°C with three and a half litres of cold water at 10°C , assuming no heat loss, is 20°C .

Let's check the options:

Option	Temperature ($^\circ\text{C}$)	Matches Calculation?
1	20	Yes
2	40	No
3	30	No
4	50	No

The calculated temperature of 20°C matches the first option.

Revision Table: Key Concepts for Mixing Water Temperatures

Your Personal Exams Guide

Concept	Description	Formula/Principle
Heat Transfer (Q)	Energy transferred due to temperature difference.	$Q = mc\Delta T$
Specific Heat Capacity (c)	Amount of heat required to raise the temperature of 1 unit mass by 1 degree.	Units: $\text{J}/(\text{kg}\cdot^{\circ}\text{C})$ or $\text{J}/(\text{g}\cdot^{\circ}\text{C})$
Equilibrium Temperature (T_f)	The final common temperature reached by substances in thermal contact.	Intermediate temperature between initial temperatures.
Principle of Heat Exchange	In an isolated system, heat lost by hot objects equals heat gained by cold objects.	$Q_{\text{lost}} = Q_{\text{gained}}$
Density (ρ)	Mass per unit volume. For water, approximately 1 kg/L.	$m = \rho V$

Additional Information: Factors Affecting Equilibrium Temperature

While our calculation assumed an ideal scenario with no heat loss, in reality, the final equilibrium temperature when mixing water can be affected by several factors:

- **Heat Loss to Surroundings:** Some heat will always be transferred to the container, the air, or the environment, making the actual final temperature slightly different from the calculated value.
- **Specific Heat Variation:** The specific heat capacity of water can vary slightly with temperature, though for typical ranges, it's often considered constant for simplicity in introductory problems.
- **Incomplete Mixing:** If the water is not thoroughly mixed, different parts might be at slightly different temperatures.
- **Evaporation:** Evaporation of hot water can cause cooling, reducing the amount of heat available for transfer.

For problems like this in physics or chemistry exams, the assumption of "no heat loss" is standard unless otherwise specified, simplifying the calculation significantly.

39. Answer: a

Explanation:

Number of students who failed in all the three exams = Total students – Total number of students who passed the exam

$$= 60 - (15 + 11 + 9 + 8 + 5 + 3 + 6)$$

$$= 3$$

Hence, '3' is the correct answer.

40. Answer: c

Explanation:

Analyzing Statement and Conclusions on Car Affordability

Let's carefully analyze the given statement and the two conclusions to determine which conclusion, if any, can be definitely drawn.

Understanding the Statement

The statement provides a direct relationship:

- **Statement:** As income increases, cars become more affordable.

This statement tells us that there is a positive correlation between income level and the affordability of cars. Higher income means cars are easier to afford (relative to that income).

It's important to note what the statement **does not** say. It does not mention:

- Technology's impact on car prices.
- Historical car prices (e.g., compared to a generation ago).
- The number of cars.
- Traffic problems.

- Future technological solutions.

Evaluating Conclusion I

Let's look at the first conclusion:

- **Conclusion I:** With greater use of technology, cars are now cheaper than they were a generation ago.

Does the original statement support this conclusion? No. The statement only links current income levels to current car affordability. It says nothing about:

- Whether technology makes cars cheaper.
- Whether cars are cheaper now than in the past.

Conclusion I introduces external factors (technology, historical time comparison) that are not mentioned or implied in the given statement. Therefore, we cannot definitely draw Conclusion I from the statement.

Evaluating Conclusion II

Now let's examine the second conclusion:

- **Conclusion II:** Cars will become so numerous that no future technology will solve the traffic problem.

Does the original statement support this conclusion? Again, no. The statement is solely about the relationship between income and car affordability. It provides no information about:

- The future number of cars.
- Future traffic conditions.
- The capability of future technology to solve problems like traffic.

Conclusion II discusses future scenarios and problems (car numbers, traffic, future technology) that are completely outside the scope of the provided statement. Therefore, we cannot definitely draw Conclusion II from the statement.

Conclusion Analysis Summary

Let's summarise the analysis:

Item	Content	Supported by Statement?	Reason
Statement	As income increases, cars become more affordable.	N/A (Premise)	This is the given information.
Conclusion I	Technology makes cars cheaper now than a generation ago.	No	Statement doesn't mention technology or historical prices.
Conclusion II	Cars will become numerous, traffic unsolvable by future tech.	No	Statement doesn't mention car numbers, traffic, or future technology.

Since neither Conclusion I nor Conclusion II can be logically deduced from the information given in the statement, neither conclusion definitely follows.

Revision Table: Statement and Conclusion Logic

Concept	Description	Key Takeaway for MCQs
Statement	A given piece of information assumed to be true.	Only use the information explicitly provided in the statement.
Conclusion	An inference or deduction made based <i>*only*</i> on the statement.	A valid conclusion must <i>*necessarily*</i> follow from the statement, without adding outside information.
Drawing Conclusions	Process of determining if a conclusion is logically supported by the statement.	Avoid making assumptions based on general knowledge; stick strictly to the statement's facts.

Additional Information: Logical Reasoning Tips

When tackling statement and conclusion questions in logical reasoning, keep these tips in mind:

- **Strict Adherence:** Base your decision **only** on the information given in the statement. Do not use your general knowledge, common sense, or facts from the real world unless the statement explicitly allows it (which is rare in these types of questions).
- **Look for Direct Support:** A valid conclusion must be directly supported by the statement. If the statement talks about 'A causing B', a conclusion about 'C' cannot be drawn unless 'C' is linked to 'A' or 'B' within the statement.
- **Identify Scope:** Understand the scope of the statement. Does it talk about past, present, or future? Does it discuss specific entities or general concepts? Conclusions outside this scope are usually invalid.
- **Negate and Test:** Sometimes, trying to imagine a scenario where the statement is true but the conclusion is false can help. If you can find such a scenario, the conclusion does not definitely follow.
- **Avoid Assumptions:** Do not assume relationships or facts not stated. For instance, if the statement says 'Some A are B', you cannot conclude 'All A are B' or 'No A are B'.

41. Answer: d

Explanation:

Identifying Folk Dances of Indian States

The question asks us to identify which of the given folk dances belongs to the state of Assam. Folk dances are an integral part of India's rich cultural heritage, with each state and region having its unique forms.

Let's examine each option to determine its origin:

- **Giddha:** Giddha is a popular folk dance that originates from the state of Punjab. It is typically performed by women during festive occasions. Therefore, Giddha is not a folk dance from Assam.
- **Lezim:** Lezim is a folk dance form and martial art originating from Maharashtra. It involves the use of a musical instrument called 'Lezim', which consists of two small wooden planks with jingles. Therefore, Lezim is not a folk dance from Assam.
- **Nati:** Nati is the traditional folk dance from the state of Himachal Pradesh. It is recognized by the Guinness Book of World Records as the largest folk dance. Therefore, Nati is not a folk dance from Assam.

- **Bagurumba:** Bagurumba is a traditional folk dance performed by the Bodo people of Assam and Northeast India. It is often referred to as the 'butterfly dance' because of the movements performed by the dancers. This dance is deeply connected with nature and is usually performed during the Bwisagu festival, a major festival of the Bodos.

Based on the origin of each dance form, Bagurumba is the folk dance that comes from Assam.

Here is a summary of the dance forms and their respective states:

Folk Dance	State of Origin
Giddha	Punjab
Lezim	Maharashtra
Nati	Himachal Pradesh
Bagurumba	Assam

Therefore, Bagurumba is the correct answer as it is a widely recognized folk dance from Assam.

Revision Table: Important Indian Folk Dances

Dance Form	State	Key Feature
Bagurumba	Assam	Bodo community dance, 'butterfly dance'
Bihu	Assam	Performed during Bihu festivals
Giddha	Punjab	Vigorous dance by women
Bhangra	Punjab	Energetic male dance
Garba	Gujarat	Performed during Navratri
Lezim	Maharashtra	Uses 'Lezim' instrument
Lavani	Maharashtra	Traditional music and dance
Nati	Himachal Pradesh	Largest folk dance
Kathakali	Kerala	Classical dance-drama
Bharatanatyam	Tamil Nadu	Classical dance form

Additional Information on Folk Dances of Assam

Assam is known for its diverse culture and numerous folk dance forms. Apart from Bagurumba, another very famous folk dance from Assam is Bihu dance. Bihu dance is performed during the Bihu festivals (Rongali Bihu, Kongali Bihu, and Bhogali Bihu) which mark important agricultural cycles. These dances are characterized by lively movements, traditional Assamese music using instruments like Dhol, Taal, Pepa, Gogona, and Bahi, and vibrant traditional attire. Understanding the regional variations of folk dances helps in appreciating the cultural mosaic of India.

42. Answer: b

Explanation:

Calculating Relative Density: Understanding Mass and Immersion

The question asks us to find the relative density of a solid given its mass in air and its weight (or apparent mass) when fully immersed in water. Relative density is a fundamental concept in physics, comparing the density of a substance to the density of a reference substance, usually water for solids and liquids.

Understanding Relative Density and Buoyancy

Relative density is defined as the ratio of the density of a substance to the density of a reference substance (usually water at a specific temperature and pressure).

Mathematically, it is:

$$\text{Relative Density} = \frac{\text{Density of Substance}}{\text{Density of Reference Substance}}$$

When a solid is immersed in water, it experiences an upward force called the buoyant force. This force is equal to the weight of the water displaced by the solid (Archimedes' Principle). The apparent weight of the solid in water is less than its weight in air due to this buoyant force.

$$\text{Weight in Water} = \text{Weight in Air} - \text{Buoyant Force}$$

Since Buoyant Force is the weight of the displaced water, we can also relate this to mass:

$$\text{Apparent Mass in Water} = \text{Mass in Air} - \text{Mass of Displaced Water}$$

From this, we can find the mass of the displaced water:

$$\text{Mass of Displaced Water} = \text{Mass in Air} - \text{Apparent Mass in Water}$$

The volume of the displaced water is equal to the volume of the fully immersed solid. Since the density of water is 1 gm/cm^3 , the mass of the displaced water in grams is numerically equal to the volume of the solid in cm^3 .

An alternative way to calculate relative density, specifically useful in this scenario, is:

$$\text{Relative Density} = \frac{\text{Mass of Solid in Air}}{\text{Mass of Equal Volume of Water}}$$

The mass of an equal volume of water is simply the mass of the water displaced by the solid when fully immersed, which we found earlier:

$$\text{Mass of Equal Volume of Water} = \text{Mass of Displaced Water} = \text{Mass in Air} - \text{Apparent Mass in Water}$$

Step-by-Step Calculation of Relative Density

Let's use the given information:

- Mass of solid in air = 50 gm
- Apparent mass of solid in water = 10 gm

First, calculate the mass of water displaced by the solid:

Mass of Displaced Water = Mass in Air – Apparent Mass in Water

Mass of Displaced Water = 50 gm – 10 gm = 40 gm

Now, we can calculate the relative density using the formula:

Relative Density = $\frac{\text{Mass of Solid in Air}}{\text{Mass of Equal Volume of Water}}$

Relative Density = $\frac{50 \text{ gm}}{40 \text{ gm}}$

Relative Density = $\frac{5}{4}$

Relative Density = 1.25

Final Answer on Relative Density

The relative density of the solid is 1.25. This means the solid is 1.25 times denser than water.

Revision Table: Key Concepts

Concept	Definition/Formula	Application
Relative Density	Ratio of substance density to reference density (usually water). $\frac{\rho_{\text{substance}}}{\rho_{\text{water}}}$	Compares how much denser or lighter a substance is than water.
Archimedes' Principle	Buoyant force = Weight of displaced fluid.	Explains why objects float or sink and apparent weight loss in fluid.
Buoyant Force	Upward force exerted by a fluid on an immersed object.	Calculated as Mass of displaced fluid $\times$ g.
Mass of Displaced Water	Mass in air – Apparent mass in water.	Used to find the volume of the solid (since volume = mass/density for water).

Additional Information: Density vs. Relative Density

It's important to distinguish between density and relative density. Density is an intrinsic property of a substance, measured in units like kg/m^3 or gm/cm^3 . It tells us how much mass is contained in a given volume.

$$\text{Density} = \frac{\text{Mass}}{\text{Volume}}$$

Relative density, on the other hand, is a dimensionless quantity (it has no units) because it is a ratio of two densities. It tells us how dense a substance is compared to water. A relative density greater than 1 means the substance is denser than water and will sink (if fully immersed). A relative density less than 1 means it is less dense than water and will float (partially submerged). A relative density equal to 1 means it has the same density as water and will remain suspended.

In this problem, the density of the solid is:

$$\text{Volume of solid} = \text{Volume of displaced water} = \frac{\text{Mass of Displaced Water}}{\text{Density of Water}} = \frac{40 \text{ gm}}{1 \text{ gm/cm}^3} = 40 \text{ cm}^3$$

$$\text{Density of solid} = \frac{\text{Mass of Solid}}{\text{Volume of Solid}} = \frac{50 \text{ gm}}{40 \text{ cm}^3} = 1.25 \text{ gm/cm}^3$$

$$\text{The relative density is } \frac{1.25 \text{ gm/cm}^3}{1 \text{ gm/cm}^3} = 1.25, \text{ confirming our previous calculation.}$$

43. Answer: b

Explanation:

Understanding Temperature Conversion: Celsius to Fahrenheit

This question asks us to convert a temperature given in degrees Celsius to degrees Fahrenheit. Temperature scales are used to measure how hot or cold something is. Celsius ($^{\circ}\text{C}$) and Fahrenheit ($^{\circ}\text{F}$) are two common scales used worldwide.

The Formula for Celsius to Fahrenheit Conversion

To convert a temperature from Celsius to Fahrenheit, we use a specific mathematical formula. This formula relates a temperature on the Celsius scale to the corresponding

temperature on the Fahrenheit scale.

The formula is:

$$^{\circ}\text{F} = (^{\circ}\text{C} \times \frac{9}{5}) + 32$$

Alternatively, the fraction $\frac{9}{5}$ can be written as 1.8. So, the formula can also be written as:

$$^{\circ}\text{F} = (^{\circ}\text{C} \times 1.8) + 32$$

Step-by-Step Calculation for -100° Celsius

We are given a temperature of -100° Celsius. We will substitute this value into the conversion formula to find the equivalent temperature in Fahrenheit.

1. Start with the formula: $^{\circ}\text{F} = (^{\circ}\text{C} \times 1.8) + 32$
2. Substitute the given Celsius temperature (-100°C) into the formula: $^{\circ}\text{F} = (-100 \times 1.8) + 32$
3. Perform the multiplication: $-100 \times 1.8 = -180$
4. Now, the equation becomes: $^{\circ}\text{F} = -180 + 32$
5. Perform the addition: $-180 + 32 = -148$

So, -100° Celsius is equal to -148° Fahrenheit.

Comparing with Options

Let's look at the given options to see which one matches our calculated result:

- Option 1: -373°
- Option 2: -148°
- Option 3: 173°
- Option 4: -212°

Our calculated value, -148° Fahrenheit, matches Option 2.

Celsius ($^{\circ}\text{C}$)	Fahrenheit ($^{\circ}\text{F}$) Calculation	Fahrenheit ($^{\circ}\text{F}$) Result
-100	$(-100 \times 1.8) + 32$	$-180 + 32 = -148$

Conclusion: -100° Celsius in Fahrenheit

Using the standard formula for converting Celsius to Fahrenheit, we found that -100°C is equivalent to -148°F . This calculation involved a simple multiplication and addition step.

Revision Table: Temperature Conversion

Concept	Description
Celsius Scale	A temperature scale where 0°C is the freezing point of water and 100°C is the boiling point of water at standard atmospheric pressure.
Fahrenheit Scale	A temperature scale where 32°F is the freezing point of water and 212°F is the boiling point of water at standard atmospheric pressure.
Conversion Formula ($^{\circ}\text{C}$ to $^{\circ}\text{F}$)	$^{\circ}\text{F} = (^{\circ}\text{C} \times 1.8) + 32$
Example Calculation	Convert -100°C to $^{\circ}\text{F}$: $^{\circ}\text{F} = (-100 \times 1.8) + 32 = -180 + 32 = -148^{\circ}\text{F}$

Additional Information: Temperature Scales

Besides Celsius and Fahrenheit, another important temperature scale is the Kelvin (K) scale. The Kelvin scale is often used in science because it is an absolute temperature scale, meaning 0 K (absolute zero) is the theoretical point where molecular motion stops. There are also formulas to convert between Kelvin and Celsius or Kelvin and Fahrenheit.

- Celsius to Kelvin: $K = ^{\circ}\text{C} + 273.15$
- Fahrenheit to Celsius: $^{\circ}\text{C} = (^{\circ}\text{F} - 32) \times \frac{5}{9}$
- Fahrenheit to Kelvin: $K = (^{\circ}\text{F} - 32) \times \frac{5}{9} + 273.15$

Understanding how to convert between these scales is fundamental in various scientific and everyday applications.

44. Answer: d

Explanation:

Solving the Number and Percentage Increase Problem

The question asks us to find an original number given that when it is increased by 36, the result is 109% of the original number. Let's break this down using algebra.

Setting Up the Equation

Let the unknown number be represented by the variable x .

According to the problem statement:

- The number is x .
- The increase is 36.
- The increased number is $x + 36$.
- The increased number is also 109% of the original number.
- 109% of x can be written as $\frac{109}{100} \times x$ or $1.09x$.

So, we can set up the equation:

Original Number + Increase = Percentage of Original Number

$$x + 36 = 109\% \text{ of } x$$

$$x + 36 = 1.09x$$

Solving for the Unknown Number (x)

Now we need to solve the equation $x + 36 = 1.09x$ for x .

We want to isolate x on one side of the equation. Let's move the x term from the left side to the right side by subtracting x from both sides:

$$36 = 1.09x - x$$

Combine the terms on the right side:

$$36 = (1.09 - 1)x$$

$$36 = 0.09x$$

To find x , we need to divide both sides of the equation by 0.09:

$$x = \frac{36}{0.09}$$

To make the division easier, we can eliminate the decimal in the denominator by multiplying both the numerator and the denominator by 100:

$$x = \frac{36 \times 100}{0.09 \times 100}$$

$$x = \frac{3600}{9}$$

Now, perform the division:

$$x = 400$$

So, the original number is 400.

Verifying the Solution

Let's check if our answer is correct. If the number is 400 and it is increased by 36, the new number is $400 + 36 = 436$.

Now let's calculate 109% of 400:

$$109\% \text{ of } 400 = \frac{109}{100} \times 400$$

$$= 1.09 \times 400$$

$$= 436$$

Since $400 + 36 = 436$ and 109% of 400 is also 436, our answer is correct.

Comparing with Options

The options provided were:

- 300
- 450
- 360
- 400

Our calculated number, 400, matches the fourth option.

Step	Description	Equation/Calculation
1	Define the number	Let the number be x .
2	Formulate the relationship	$x + 36 = 109\% \text{ of } x$
3	Convert percentage to decimal	$109\% = 1.09$
4	Rewrite the equation	$x + 36 = 1.09x$
5	Solve for x	$36 = 1.09x - x$
6	Simplify and calculate x	$36 = 0.09x$ $x = \frac{36}{0.09} = \frac{3600}{9} = 400$

Revision Table: Percentage Increase Concepts

Concept	Explanation	Formula Example
Percentage	A fraction out of 100. Represented by the symbol %.	$50\% = \frac{50}{100} = 0.5$
Percentage of a number	To find a percentage of a number, convert the percentage to a decimal or fraction and multiply by the number.	$20\% \text{ of } 80 = \frac{20}{100} \times 80 =$ $0.20 \times 80 = 16$
Percentage Increase	The amount of increase expressed as a percentage of the original value.	Percentage Increase $= \frac{\text{Increase Amount}}{\text{Original Amount}} \times 100\%$
Number after Percentage Increase	The original number plus the increase amount. Can also be calculated as $\text{Original Amount} \times (1 + \text{Percentage Increase}/100)$. If a number becomes 109% of itself, it means the original number plus an increase equals 109% of the original number. The increase is $109\% - 100\% = 9\%$ of the original number.	If a number increases by 10%, the new number is $100\% + 10\% = 110\%$ of the original number.

Additional Information: Solving Linear Equations

The problem required solving a simple linear equation. A linear equation in one variable is an equation that can be written in the form $ax + b = c$, where a , b , and c are constants and

$a \neq 0$. The goal is to find the value of the variable (in this case, x) that makes the equation true.

The general strategy involves isolating the variable on one side of the equation using inverse operations:

- If a number is added to the variable term, subtract that number from both sides.
- If a number is subtracted from the variable term, add that number to both sides.
- If the variable term is multiplied by a number, divide both sides by that number.
- If the variable term is divided by a number, multiply both sides by that number.

In our equation, $x + 36 = 1.09x$, we had x terms on both sides. We moved the x term from the left to the right by subtracting x from both sides, resulting in $36 = 0.09x$. Then, since x was multiplied by 0.09, we divided both sides by 0.09 to find x . These are fundamental steps in solving many algebra problems, including those involving percentages and unknown numbers.

45. Answer: d

Explanation:

See-saw Equilibrium Calculation: Finding the Mass M

This problem involves the concept of equilibrium of a see-saw, which is essentially a lever. For a lever to be in rotational equilibrium, the sum of the clockwise moments about the fulcrum must be equal to the sum of the anti-clockwise moments about the fulcrum.

A moment is the turning effect of a force about a pivot (fulcrum). It is calculated as the product of the force and the perpendicular distance from the fulcrum to the line of action of the force. Mathematically, $\text{Moment} = \text{Force} \times \text{Distance}$.

In this scenario, the forces are due to the weights of the children and the boy. Weight is given by $\text{Mass} \times \text{acceleration due to gravity}$ ($W = mg$). Since the acceleration due to gravity (g) is the same for all masses on both sides of the see-saw, we can work directly with $\text{Mass} \times \text{Distance}$ for calculating the moments for equilibrium purposes.

Setting Up the Moments for Equilibrium

Let's define one side as creating anti-clockwise moments and the other side as creating clockwise moments. The two children are sitting on one side, contributing to the total moment on that side. The boy of mass 'M' is sitting on the other side, creating an opposing moment.

- **Side 1 (Children):**

- Child 1: Mass = 30 kg, Distance from fulcrum = 1 m. Moment by Child 1 = $30 \text{ kg} \times 1 \text{ m} = 30 \text{ kg-m}$.
- Child 2: Mass = 15 kg, Distance from fulcrum = 1.2 m. Moment by Child 2 = $15 \text{ kg} \times 1.2 \text{ m} = 18 \text{ kg-m}$.
- Total moment on Side 1 = Moment by Child 1 + Moment by Child 2 = $30 \text{ kg-m} + 18 \text{ kg-m} = 48 \text{ kg-m}$.

- **Side 2 (Boy):**

- Boy: Mass = M kg, Distance from fulcrum = 1.2 m. Moment by Boy = $M \text{ kg} \times 1.2 \text{ m} = 1.2M \text{ kg-m}$.

Applying the Principle of Moments for See-saw Equilibrium

For the see-saw to be in equilibrium, the total moment on one side must equal the total moment on the other side:

Total moment on Side 1 = Total moment on Side 2

$$(30 \text{ kg} \times 1 \text{ m}) + (15 \text{ kg} \times 1.2 \text{ m}) = M \text{ kg} \times 1.2 \text{ m}$$

$$30 \text{ kg-m} + 18 \text{ kg-m} = 1.2M \text{ kg-m}$$

$$48 \text{ kg-m} = 1.2M \text{ kg-m}$$

Solving for M

Now, we need to find the value of M by rearranging the equation:

$$M = \frac{48 \text{ kg-m}}{1.2 \text{ m}}$$

$$M = 40 \text{ kg}$$

So, a boy with a mass of 40 kg needs to sit on the other side at a distance of 1.2 m from the fulcrum for the see-saw to be in equilibrium.

Summary of Calculation

Item	Mass (kg)	Distance (m)	Moment (kg-m)
Child 1	30	1.0	$30 \times 1.0 = 30$
Child 2	15	1.2	$15 \times 1.2 = 18$
Boy	M	1.2	$M \times 1.2 = 1.2M$

Total moment on children's side = $30 + 18 = 48$ kg-m.

Total moment on boy's side = $1.2M$ kg-m.

For equilibrium: $48 = 1.2M$.

Therefore, $M = 48/1.2 = 40$ kg.

Revision Table: See-saw Equilibrium Concepts

Concept	Definition/Principle	Application in Problem
Moment (Torque)	Turning effect of a force about a pivot. Moment = Force $\times$ Perpendicular Distance. (Or Mass $\times$ Distance when g is constant).	Calculated for each child and the boy.
Fulcrum	The pivot point of the lever (see-saw).	Distances are measured from the fulcrum.
Rotational Equilibrium	The state where an object is not rotating or is rotating at a constant angular velocity. For a lever, this means the sum of clockwise moments equals the sum of anti-clockwise moments.	The condition used to set up the equation: $48 \text{ kg-m} = 1.2M \text{ kg-m}$.
Lever	A rigid bar that pivots around a fixed point (fulcrum). See-saw is a type of lever.	The see-saw acts as a lever where moments are balanced.

Additional Information: Understanding Moments and Levers

The problem highlights the practical application of the principle of moments. Moments are crucial in understanding how levers work. Levers are simple machines that help us multiply force or change the direction of force.

There are three classes of levers, classified based on the relative positions of the fulcrum, the effort (applied force), and the load (force being moved or balanced).

- **Class 1 Lever:** The fulcrum is located between the effort and the load. A see-saw is a classic example of a Class 1 lever. Pliers and scissors are other examples. They can be used to multiply force or increase distance/speed depending on where the fulcrum is relative to the effort and load.
- **Class 2 Lever:** The load is located between the fulcrum and the effort. Examples include a wheelbarrow and a nutcracker. Class 2 levers always multiply force (provide mechanical advantage > 1).
- **Class 3 Lever:** The effort is located between the fulcrum and the load. Examples include tweezers, fishing rods, and the human forearm when lifting a weight. Class 3 levers always provide a mechanical advantage < 1 , meaning they don't multiply force but are useful for increasing speed or distance of movement.

In this see-saw problem, the children and the boy represent the loads, and the fulcrum is the pivot. The weights (forces) act downwards, creating moments about the fulcrum that cause rotation. Balancing these moments leads to equilibrium.

Your Personal Exams Guide

46. Answer: a

Explanation:

Understanding Sentence Rearrangement

Sentence rearrangement questions test your ability to understand the logical flow and grammatical structure of a sentence. You are given parts of a sentence and need to arrange them in the correct order to form a coherent and meaningful sentence.

Analysing the Sentence Fragments

Let's look at the given sentence parts:

Sentence beginning: A customer based on reasons

- X: chances are that they actually made the change
- Y: may not come out of the changing room, but
- Z: which is unfortunately inappropriate

Finding the Correct Sequence

We need to connect these parts to the beginning "A customer based on reasons" to form a complete sentence. Let's try linking them:

- Part Y starts with "may not come out of the changing room, but". This seems like a natural continuation after "A customer based on reasons". Reasons can influence why a customer might or might not do something, like come out of a changing room.
- After Y, which ends with "but", we expect a contrasting idea. Part X, "chances are that they actually made the change", provides this contrast. The customer didn't come out, BUT it's likely they did make the change. This connects well after "but".
- Finally, Part Z, "which is unfortunately inappropriate", contains a relative pronoun "which". This pronoun usually refers to the preceding noun or clause. In the sequence Y then X, the most likely thing "which" refers to is "the change" mentioned in part X, or possibly the action of making the change without coming out. Saying that the change (made without coming out) is inappropriate makes sense in a retail context.

So, the sequence Y then X then Z seems to follow logically and grammatically from the sentence beginning.

Let's assemble the complete sentence using the YXZ sequence:

A customer based on reasons + Y + X + Z

A customer based on reasons **may not come out of the changing room, but chances are that they actually made the change which is unfortunately inappropriate.**

Explanation of the Correct Order (YXZ)

Let's break down why the YXZ sequence works:

Part	Content	Connection
Beginning	A customer based on reasons	Sets the subject and introduces the idea of 'reasons' influencing behaviour.
Y	may not come out of the changing room, but	Connects to the beginning, explaining what the customer might do based on reasons. The 'but' signals a coming contrast.
X	chances are that they actually made the change	Provides the contrasting outcome introduced by 'but' in Y. It states the likely action despite not coming out.
Z	which is unfortunately inappropriate	Uses 'which' to refer back to the action or outcome described in X (making the change). It comments on the appropriateness of this action.

The resulting sentence, "A customer based on reasons may not come out of the changing room, but chances are that they actually made the change which is unfortunately inappropriate," is grammatically sound and conveys a clear meaning about a customer's behaviour in a changing room and the perceived inappropriateness of that behaviour.

Revision Table: Key Concepts in Sentence Structure

Your Personal Exams Guide

Concept	Description	Relevance to Question
Subject-Verb Agreement	The verb must agree in number with its subject.	Ensures the main clause is grammatically correct.
Pronoun Reference	Pronouns (like 'which', 'it', 'they') must clearly refer to a specific noun or phrase.	Crucial for understanding what 'which' refers to in part Z.
Connectors (e.g., but)	Words that link ideas, showing relationships like contrast, cause, or effect.	'but' in part Y signals a contrasting idea will follow.
Logical Flow	The sequence of ideas must make sense and follow a rational progression.	Helps determine the most plausible order of the sentence parts.

Additional Information: Tips for Sentence Rearrangement

Tackling sentence rearrangement problems requires a systematic approach. Here are some tips:

- Read all the parts carefully to get a general sense of the topic.
- Look for introductory phrases or clauses that might start a sentence.
- Identify linking words or phrases (like 'but', 'therefore', 'however', 'also') that show connections between ideas.
- Look for pronoun references ('he', 'she', 'it', 'they', 'which', 'this') and determine what noun they are likely referring to. This helps link parts.
- Try to identify the main clause or the core idea of the sentence.
- Mentally (or physically) try arranging the parts in different sequences suggested by the options.
- Read the assembled sentence aloud to check if it flows naturally and makes grammatical sense.
- Pay attention to punctuation (commas, periods) that might give clues about how parts are connected.

47. Answer: d

Explanation:

Understanding Kinetic Energy and Speed Ratios

This problem requires us to find how the kinetic energy (KE) of a car changes when its speed increases. Kinetic energy is the energy associated with an object's motion.

Kinetic Energy Formula

The fundamental formula for calculating kinetic energy is:

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

In this formula, ' m ' represents the mass of the car, and ' v ' represents its speed.

Analyzing the Speed Change

We are given the initial and final speeds of the car:

- Initial Speed (v_1): 54 km/hr
- Final Speed (v_2): 90 km/hr

Calculating the Ratio of Speeds

To determine the ratio of kinetic energies, we first need the ratio of the speeds. We can work directly with the speeds in km/hr because the units will cancel out in the ratio. Let's compute the ratio of the final speed to the initial speed (v_2/v_1):

$$\frac{v_2}{v_1} = \frac{90 \text{ km/hr}}{54 \text{ km/hr}}$$

To simplify this fraction, we can find the greatest common divisor of 90 and 54, which is 18. Dividing both the numerator and the denominator by 18:

$$\frac{90}{54} = \frac{90 \div 18}{54 \div 18} = \frac{5}{3}$$

This calculation shows that the final speed is $\frac{5}{3}$ times the initial speed.

Calculating the Ratio of Kinetic Energies

Since kinetic energy is proportional to the square of the speed ($KE \propto v^2$), the ratio of the final kinetic energy (KE_2) to the initial kinetic energy (KE_1) is equal to the square of the ratio of their speeds:

$$\frac{KE_2}{KE_1} = \left(\frac{v_2}{v_1}\right)^2$$

Substituting the simplified speed ratio we found:

$$\frac{KE_2}{KE_1} = \left(\frac{5}{3}\right)^2 = \frac{5^2}{3^2} = \frac{25}{9}$$

The ratio $\frac{25}{9}$ indicates that the final kinetic energy is $\frac{25}{9}$ times the initial kinetic energy, confirming an increase in energy.

Matching the Options Provided

The question asks for the ratio in which the kinetic energy increases. Our calculated ratio (KE_2/KE_1) is $\frac{25}{9}$. However, observing the multiple-choice options, they are all fractions less than 1. This suggests that the question might be asking for the ratio of the initial kinetic energy (KE_1) to the final kinetic energy (KE_2). Let's calculate this inverse ratio:

$$\frac{KE_1}{KE_2} = \left(\frac{v_1}{v_2}\right)^2$$

We found earlier that $\frac{v_1}{v_2} = \frac{54}{90} = \frac{3}{5}$. Squaring this ratio:

$$\frac{KE_1}{KE_2} = \left(\frac{3}{5}\right)^2 = \frac{3^2}{5^2} = \frac{9}{25}$$

Final Conclusion on Kinetic Energy Ratio

The calculated ratio $\frac{9}{25}$ corresponds to one of the provided options. This value represents the ratio of the initial kinetic energy to the final kinetic energy.

48. Answer: b

Explanation:

Finding Cube Edge Length from Weight and Density

This problem requires us to find the length of the edge of a wooden cube given its weight and density. We'll use the relationships between weight, mass, density, volume, and the dimensions of a cube.

Here's the information provided:

- Weight of the wooden cube (W) = 80 N
- Acceleration due to gravity (g) = 10 m/s²
- Density of wood (ρ) = 1 g/cm³

We need to find the length of the edge (L) of the cube in centimeters (cm).

Step 1: Calculate the Mass of the Cube

Weight is related to mass and gravity by the formula:

$$W = m \times g$$

Where:

- W is the weight
- m is the mass
- g is the acceleration due to gravity

We can rearrange this formula to find the mass:

$$m = \frac{W}{g}$$

Plugging in the given values:

$$m = \frac{80 \text{ N}}{10 \text{ m/s}^2} = 8 \text{ kg}$$

The mass of the wooden cube is 8 kilograms.

Step 2: Convert Mass to Grams

The density is given in grams per cubic centimeter (g/cm³). To use this density value directly, we should convert the mass from kilograms (kg) to grams (g). We know that 1 kg = 1000 g.

$$m = 8 \text{ kg} \times 1000 \text{ g/kg} = 8000 \text{ g}$$

The mass of the wooden cube is 8000 grams.

Step 3: Calculate the Volume of the Cube

Density is defined as mass per unit volume:

$$\rho = \frac{m}{V}$$

Where:

- ρ is the density
- m is the mass
- V is the volume

We can rearrange this formula to find the volume:

$$V = \frac{m}{\rho}$$

Plugging in the calculated mass in grams and the given density in g/cm^3 :

$$V = \frac{8000 \text{ g}}{1 \text{ g/cm}^3} = 8000 \text{ cm}^3$$

The volume of the wooden cube is 8000 cubic centimeters.

Step 4: Calculate the Edge Length of the Cube

For a cube, the volume (V) is related to its edge length (L) by the formula:

$$V = L^3$$

To find the edge length, we need to take the cube root of the volume:

$$L = V^{1/3}$$

Plugging in the calculated volume:

$$L = (8000 \text{ cm}^3)^{1/3}$$

To find the cube root of 8000:

$$8000 = 8 \times 1000$$

The cube root of 8 is 2, and the cube root of 1000 is 10.

$$L = (8 \times 1000)^{1/3} = 8^{1/3} \times 1000^{1/3} = 2 \times 10 = 20$$

So, the edge length of the cube is 20 cm.

The length of the edge of the cube is 20 cm.

Summary of Calculation Steps

Concept	Formula Used	Calculation	Result
Mass from Weight	$m = W/g$	$m = 80 \text{ N}/10 \text{ m/s}^2$	$m = 8 \text{ kg}$
Mass Conversion	$1 \text{ kg} = 1000 \text{ g}$	$m = 8 \text{ kg} \times 1000 \text{ g/kg}$	$m = 8000 \text{ g}$
Volume from Mass & Density	$V = m/\rho$	$V = 8000 \text{ g}/1 \text{ g/cm}^3$	$V = 8000 \text{ cm}^3$
Edge Length from Volume	$L = V^{1/3}$	$L = (8000 \text{ cm}^3)^{1/3}$	$L = 20 \text{ cm}$

Revision Table – Understanding Physics Concepts

Term	Definition	Formula
Weight	The force of gravity on an object's mass. Measured in Newtons (N).	$W = m \times g$
Mass	The amount of matter in an object. Measured in kilograms (kg) or grams (g).	–
Density	Mass per unit volume of a substance. Measured in kg/m^3 or g/cm^3 .	$\rho = m/V$
Volume	The amount of space an object occupies. Measured in m^3 or cm^3 .	For a cube: $V = L^3$
Acceleration due to Gravity (g)	The acceleration experienced by objects due to gravity. Approximately 9.8 m/s^2 on Earth, but given as 10 m/s^2 in this problem.	–

Additional Information – Unit Consistency in Physics Problems

When solving physics problems, it is crucial to ensure that all units are consistent. The SI system (International System of Units) is commonly used, where mass is in kilograms (kg),

length in meters (m), time in seconds (s), and force in Newtons ($N = \text{kg}\cdot\text{m}/\text{s}^2$).

- In this problem, weight is given in Newtons (N), which is an SI unit of force. Gravity is given in m/s^2 , also SI. Dividing N by m/s^2 gives mass in kg.
- However, the density is given in g/cm^3 . This is a common CGS (centimeter-gram-second) unit for density.
- To perform calculations involving density and mass, you can either:
 - Convert mass from kg to g (as done in the solution).
 - Convert density from g/cm^3 to kg/m^3 . ($1 \text{ g}/\text{cm}^3 = 1000 \text{ kg}/\text{m}^3$)
- Since the final answer is requested in cm, working with density in g/cm^3 and calculating volume in cm^3 is convenient, as taking the cube root of cm^3 directly yields cm.

Always check the units of all given quantities and the required units for the final answer before starting calculations. This prevents errors and ensures the result is meaningful.

49. Answer: b

Explanation:

Calculating Profit Percentage

This problem involves calculating the profit percentage on an item when its selling price changes, given the initial selling price and profit percentage. To solve this, we first need to find the cost price (CP) of the item using the information from the first sale. Once we have the CP, we can calculate the profit made in the second sale and then determine the new profit percentage.

Step-by-Step Profit Calculation

Here are the steps to find the new profit percentage:

1. Find the Cost Price (CP):

We know that the shopkeeper sells the item at Rs. 2,700 and makes an 8% profit. The selling price (SP) is the cost price plus the profit. The formula relating SP, CP, and Profit Percentage is:

$$SP = CP \times \left(1 + \frac{\text{Profit \%}}{100}\right)$$

Using the given values:

$$2700 = CP \times \left(1 + \frac{8}{100}\right)$$

$$2700 = CP \times (1 + 0.08)$$

$$2700 = CP \times 1.08\}$$

Now, we can find the Cost Price:

$$CP = \frac{2700}{1.08}$$

$$CP = 2500\}$$

So, the cost price of the item is Rs. 2,500.

2. Calculate the Profit in the Second Sale:

The item is now sold at Rs. 3,000. The profit is the difference between the new selling price and the cost price.

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

$$\text{Profit} = 3000 - 2500\}$$

$$\text{Profit} = 500\}$$

The profit made when selling the item at Rs. 3,000 is Rs. 500.

3. Calculate the New Profit Percentage:

The profit percentage is calculated on the cost price. The formula is:

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

Using the profit from the second sale and the cost price:

$$\text{Profit \%} = \left(\frac{500}{2500}\right) \times 100\}$$

$$\text{Profit \%} = \left(\frac{1}{5}\right) \times 100\}$$

$$\text{Profit \%} = 0.2 \times 100\}$$

$$\text{Profit \%} = 20\%$$

The percentage of profit if the item is sold at Rs. 3,000 is 20%.

Summary of Profit Calculation

Let's summarize the two scenarios:

Scenario	Selling Price (SP)	Cost Price (CP)	Profit	Profit Percentage
Scenario 1 (Given)	Rs. 2,700	Rs. 2,500 (Calculated)	Rs. 2,700 - Rs. 2,500 = Rs. 200	$\left(\frac{200}{2500}\right) \times 100 = 8\%$
Scenario 2 (Questioned)	Rs. 3,000	Rs. 2,500 (Used from Scenario 1)	Rs. 3,000 - Rs. 2,500 = Rs. 500	$\left(\frac{500}{2500}\right) \times 100 = 20\%$

The final profit percentage when selling the item at Rs. 3,000 is 20%.

Revision Table: Profit and Loss Formulas

Concept	Formula
Profit	Selling Price (SP) - Cost Price (CP) (when SP > CP)
Loss	Cost Price (CP) - Selling Price (SP) (when CP > SP)
Profit Percentage	$\left(\frac{\text{Profit}}{\text{CP}}\right) \times 100\%$
Loss Percentage	$\left(\frac{\text{Loss}}{\text{CP}}\right) \times 100\%$
Finding SP given CP and Profit %	$\text{SP} = \text{CP} \times \left(1 + \frac{\text{Profit \%}}{100}\right)$
Finding SP given CP and Loss %	$\text{SP} = \text{CP} \times \left(1 - \frac{\text{Loss \%}}{100}\right)$
Finding CP given SP and Profit %	$\text{CP} = \frac{\text{SP}}{1 + \frac{\text{Profit \%}}{100}}$
Finding CP given SP and Loss %	$\text{CP} = \frac{\text{SP}}{1 - \frac{\text{Loss \%}}{100}}$

Additional Information: Profit and Loss Problems

Profit and loss calculations are fundamental in business mathematics. Understanding the relationship between cost price, selling price, profit, loss, and their respective percentages is crucial. Problems often involve finding an unknown value when others are given, or dealing with successive transactions, discounts, and taxes. Always remember that profit and loss percentages are typically calculated based on the cost price unless otherwise specified.

50. Answer: b

Explanation:

The bullet travels 90 meters in 0.2 seconds

$\therefore$ Bullet speed = $90/0.2 = 450$ m/s

$\therefore$ Speed in km/h = $450 \times 18/5 = 1620$ km/h

51. Answer: d

Explanation:

For two resistors in parallel, $1/24 = 1/R + 1/60$. Therefore, $1/R = 1/24 - 1/60 = 3/120 = 1/40$, giving $R = 40$ ohm.

52. Answer: b

Explanation:

Understanding Antihistamines in First-Aid Boxes

Antihistamines are a common component found in many first-aid kits. Their primary function is to counteract the effects of histamine, a substance produced by the body

during an allergic reaction.

What is Histamine?

Histamine is a chemical involved in the body's immune response. When the body encounters an allergen (like pollen, dust mites, or certain foods), mast cells release histamine. Histamine causes various symptoms associated with allergies, such as:

- Itching
- Sneezing
- Runny nose
- Watery eyes
- Hives or rashes
- Swelling

Antihistamines work by blocking the action of histamine at specific receptors in the body, thus reducing or stopping these symptoms.

Analyzing the Options for Antihistamine Use

Let's look at why antihistamines are used based on the provided options:

1. **To aid in clotting of blood:** This is incorrect. Blood clotting is a complex process involving platelets and clotting factors (like fibrinogen). Antihistamines do not play a role in blood coagulation.
2. **To ease the symptoms of hay fever and other allergies:** This is the primary use of the type of antihistamines typically found in first-aid kits. Hay fever (allergic rhinitis) and other allergic reactions are caused by the release of histamine, and antihistamines are effective in blocking histamine's effects, thereby easing symptoms like sneezing, itching, and runny nose.
3. **To ease the indigestion and heartburn:** Some types of antihistamines, specifically H₂ blockers (like ranitidine or cimetidine), are used to reduce stomach acid production, which can help with indigestion and heartburn. However, the antihistamines commonly found in first-aid boxes for general use (like diphenhydramine or loratadine) are H₁ blockers, primarily used for allergy symptoms and not digestive issues. Therefore, this option is generally incorrect in the context of a standard first-aid box containing allergy medication.
4. **To bring relief from asthma:** Asthma is a chronic respiratory condition characterized by inflammation and narrowing of the airways. While allergies can trigger asthma

symptoms in some individuals, antihistamines are not the main treatment for acute asthma attacks or for managing asthma symptoms. Bronchodilators (like salbutamol) are used to open airways, and corticosteroids are used to reduce inflammation.

Conclusion on Antihistamine Use

Based on the analysis, the most appropriate and common use of antihistamines found in a standard first-aid box is to treat the symptoms of allergic reactions, such as hay fever and other allergies. They work by blocking the histamine released during these reactions.

Condition	Effectiveness of Antihistamines	Typical Treatment
Blood Clotting Issues	No	Anticoagulants (to prevent clotting), Clotting Factors (to aid clotting)
Hay Fever / Allergies	Yes (Primary Use)	Antihistamines, Nasal Sprays, Eye Drops
Indigestion / Heartburn	Limited (H ₂ Blockers, not typical first-aid H ₁ blockers)	Antacids, Proton Pump Inhibitors (PPIs), H ₂ Blockers
Asthma Relief	Limited (Not primary treatment)	Bronchodilators, Inhaled Corticosteroids

Revision Table: Antihistamine Uses

Key Term	Brief Explanation	Relevance to Question
Antihistamines	Drugs that block histamine action.	The subject of the question.
Histamine	Chemical released during allergic reactions.	Blocked by antihistamines to relieve allergy symptoms.
Allergies	Immune response to allergens.	Condition effectively treated by antihistamines.
First-Aid Box	Container of supplies for immediate medical help.	Context where antihistamines are found.
Hay Fever	Seasonal allergic rhinitis.	Specific type of allergy treated by antihistamines.

Additional Information on Antihistamines

There are different generations of antihistamines:

- **First-generation antihistamines:** These can cause drowsiness (e.g., diphenhydramine, chlorpheniramine). They are often found in older first-aid kits.
- **Second-generation antihistamines:** These are less likely to cause drowsiness and are often preferred for daytime use (e.g., loratadine, cetirizine, fexofenadine).

It's important to read the instructions on the specific antihistamine product in a first-aid box, as dosages and specific uses may vary. While antihistamines are generally safe for temporary relief of allergy symptoms, they are not a substitute for professional medical advice, especially for severe allergic reactions (anaphylaxis), which require immediate emergency treatment (like epinephrine).

53. Answer: d

Explanation:

Understanding Direction and Distance Problems

This problem involves tracking a series of movements in different directions to find the final position relative to the starting point. We can solve this by breaking down the movements into components along the North-South axis and the East-West axis.

Step-by-Step Movement Analysis

Let's track the postman's journey from the starting point, the post office.

1. **Starting Point:** Post Office (Let's consider this the origin (0,0)).
2. **First Movement:** Cycles 13 km south. This movement is purely along the South axis.
3. **Second Movement:** Turns west and cycles 10 km. This movement is purely along the West axis.
4. **Third Movement:** Turns north and cycles 8 km. This movement is purely along the North axis.
5. **Fourth Movement:** Turns east and cycles 2 km. This movement is purely along the East axis.
6. **Fifth Movement:** Turns to his left and cycles 5 km. After moving east, a left turn means turning towards North. So, this movement is 5 km north.

Calculating Net Displacement

Now, let's sum up the total distance traveled in each cardinal direction:

- Total South movement: 13 km
- Total North movement: 8 km + 5 km = 13 km
- Total West movement: 10 km
- Total East movement: 2 km

Next, we calculate the net displacement along the North-South axis and the East-West axis.

Net Displacement along North-South Axis:

Net North displacement = Total North - Total South

Net North displacement = 13 km - 13 km = 0 km

This means the postman is at the same latitude (North-South position) as his starting point.

Net Displacement along East–West Axis:

Net West displacement = Total West – Total East

Net West displacement = 10 km – 2 km = 8 km

Since the Net West displacement is positive (meaning West movement was greater than East movement), the postman is to the west of his starting point.

Final Position Relative to Starting Point

Combining the net displacements:

- Net North/South displacement: 0 km
- Net East/West displacement: 8 km West

Therefore, the postman is 8 km west of his starting position at the post office.

We can summarize the movements in a table:

Movement	Direction	Distance (km)	North (+)/South (–)	East (+)/West (–)
Start	–	0	0	0
1	South	13	–13	0
2	West	10	0	–10
3	North	8	+8	0
4	East	2	0	+2
5 (Left after East)	North	5	+5	0
Total Net	–	–	$(-13 + 8 + 5) = 0$	$(-10 + 2) = -8$

A net displacement of –8 in the East/West column corresponds to 8 km West. A net displacement of 0 in the North/South column means no change in the North–South position.

Conclusion

Based on the calculations, the postman's final position is 8 km to the west of the post office.

Revision Table: Direction and Distance Concepts

Concept	Explanation	How it applies here
Starting Point	Reference point for calculating final position.	The Post Office.
Direction	Path of movement (North, South, East, West, Left, Right).	Movements are broken down into cardinal directions. 'Left' after East is North.
Distance	Magnitude of movement in a specific direction.	Summed up for each cardinal direction.
Displacement	Straight-line distance and direction from start to end point. Calculated as net change in North-South and East-West.	Net change in N/S and E/W is calculated to find final displacement.

Additional Information on Relative Position

When solving these types of direction and distance problems, it's helpful to think of movements as changes in coordinates on a grid, where North is positive Y, South is negative Y, East is positive X, and West is negative X. Starting at (0,0):

- 13 km South: $(0, -13)$
- Then 10 km West: $(0 - 10, -13) = (-10, -13)$
- Then 8 km North: $(-10, -13 + 8) = (-10, -5)$
- Then 2 km East: $(-10 + 2, -5) = (-8, -5)$
- Then 5 km Left (North): $(-8, -5 + 5) = (-8, 0)$

The final position is at coordinate $(-8, 0)$. Relative to the start $(0,0)$, this means 8 units in the negative X direction (West) and 0 units in the Y direction (North/South). This confirms the result: 8 km West of the starting position.

54. Answer: b

Explanation:

Understanding Weight and Gravity on Moon vs. Earth

This question asks us to calculate the weight of an astronaut on the moon, given their weight on Earth and the relationship between the gravitational acceleration on the moon and Earth. Weight is the force exerted on an object due to gravity. It is calculated using the formula:

$$W = m \times g$$

where W is weight, m is mass, and g is the acceleration due to gravity.

Mass is a fundamental property of an object and remains constant regardless of location. Weight, however, depends on the acceleration due to gravity, which varies from one celestial body to another.

Analyzing the Given Information

- Astronaut's weight on Earth (W_e) = 90 kgf
- Acceleration due to gravity on Earth (g_e) = 10 m/s^2
- Acceleration due to gravity on Moon (g_m) = $\frac{1}{6} \times g_e$

Understanding the Unit kgf (Kilogram-force)

The unit 'kgf' (kilogram-force) is a unit of force. It is defined as the magnitude of the force exerted by gravity on one kilogram of mass on Earth. With the given value of $g_e = 10 \text{ m/s}^2$, 1 kgf is equal to the weight of a 1 kg mass under this gravity. So, $1 \text{ kgf} = 1 \text{ kg} \times 10 \text{ m/s}^2 = 10 \text{ N}$.

Therefore, the astronaut's weight on Earth in Newtons is:

$$W_e = 90 \text{ kgf} = 90 \times 10 \text{ N} = 900 \text{ N}$$

Step-by-Step Calculation of Astronaut's Weight on Moon

Step 1: Calculate the Astronaut's Mass

We know the weight on Earth (W_e) and the acceleration due to gravity on Earth (g_e). We can use the formula $W_e = m \times g_e$ to find the mass (m).

$$m = \frac{W_e}{g_e}$$

Substituting the values:

$$m = \frac{900 \text{ N}}{10 \text{ m/s}^2} = 90 \text{ kg}$$

The astronaut's mass is 90 kg. This mass will be the same on the moon.

Step 2: Calculate the Acceleration Due to Gravity on the Moon

The question states that the gravity on the moon is 1/6th that on Earth.

$$g_m = \frac{1}{6} \times g_e$$

Substituting the value of g_e :

$$g_m = \frac{1}{6} \times 10 \text{ m/s}^2 = \frac{10}{6} \text{ m/s}^2 = \frac{5}{3} \text{ m/s}^2$$

Step 3: Calculate the Astronaut's Weight on the Moon

Now we can find the astronaut's weight on the moon (W_m) using their mass (m) and the gravity on the moon (g_m):

$$W_m = m \times g_m$$

Substituting the calculated values:

$$W_m = 90 \text{ kg} \times \frac{5}{3} \text{ m/s}^2$$

$$W_m = \left(90 \times \frac{5}{3} \right) \text{ N}$$

$$W_m = \left(\frac{90}{3} \times 5 \right) \text{ N}$$

$$W_m = (30 \times 5) \text{ N}$$

$$W_m = 150 \text{ N}$$

So, the astronaut would weigh 150 N on the moon.

Comparison with Options

Let's check the options provided:

Option	Value
1	15 N
2	150 N
3	90 N
4	9 N

Our calculated weight on the moon is 150 N, which matches Option 2.

Revision Table: Weight and Gravity Concepts

Concept	Definition	Formula	Unit(s)	Varies with Location?
Mass	Amount of matter in an object	Independent	kg, g, etc.	No (usually constant)
Weight	Force of gravity on an object	$W = m \times g$	N, kgf, pounds, etc.	Yes (depends on g)
Acceleration due to gravity (g)	Acceleration experienced by an object due to gravitational force	Varies (depends on mass and radius of celestial body)	m/s^2 , N/kg	Yes (depends on location)

Additional Information: Why Weight Changes but Mass Stays the Same

Weight is a force, specifically the force of attraction between an object and a celestial body (like Earth or the moon). This force depends on the mass of both the object and the celestial body, and the distance between their centers. The acceleration due to gravity (g) encapsulates these factors for a specific location on the surface of a celestial body.

Mass, on the other hand, is an intrinsic property of the object itself – it's a measure of the amount of 'stuff' it contains. Unless the astronaut loses or gains mass (e.g., by eating or removing equipment), their mass remains the same whether they are on Earth, the moon, or in space. Since the moon has significantly less mass than Earth, its gravitational pull is weaker, resulting in a smaller value for g_m compared to g_e . This is why the astronaut's weight is less on the moon, even though their mass hasn't changed.

Understanding the distinction between mass and weight is fundamental in physics. Mass is a scalar quantity, while weight is a vector quantity (having both magnitude and direction – towards the center of the celestial body).

55. Answer: d

Explanation:

Understanding Pump Power and Work

This problem asks us to find the power of a pump. Power is defined as the rate at which work is done or energy is transferred. In this case, the pump does work to lift water against gravity. The work done is converted into the potential energy of the water.

Key Concepts

- **Power (P):** The rate of doing work, calculated as Work done divided by the time taken ($P = \frac{W}{t}$). The unit of power is Watts (W).
- **Work (W):** When a force moves an object over a distance, work is done. In this problem, the pump exerts a force to lift the water against gravity. The work done against gravity to lift an object is given by the potential energy formula ($W = mgh$), where 'm' is mass, 'g' is acceleration due to gravity, and 'h' is the height lifted. The unit of work is Joules (J).
- **Potential Energy (PE):** The energy stored in an object due to its position or state. For an object lifted against gravity, the potential energy gained is $PE = mgh$. Assuming

100% efficiency, the work done by the pump equals the potential energy gained by the water.

Analyzing the Given Information

We are given the following information:

- Mass of water (m) = 1 tonne
- Height the water is lifted (h) = 90 m
- Time taken (t) = 30 minutes
- Acceleration due to gravity (g) = 10 m/s^2
- Efficiency of the pump = 100%

Step-by-Step Power Calculation

To find the power of the pump, we first need to calculate the total work done in lifting the water and then divide it by the time taken.

Step 1: Convert Units

We need to ensure all units are in the standard SI system (meters, kilograms, seconds).

- Mass: 1 tonne = 1000 kilograms (kg)
- Time: 30 minutes. Since 1 minute = 60 seconds, 30 minutes = 30×60 seconds = 1800 seconds (s).

Step 2: Calculate the Work Done

The work done (W) in lifting the water is equal to the potential energy gained by the water. We use the formula $W = mgh$.

$$W = \text{mass} \times \text{gravity} \times \text{height}$$

Substitute the values:

$$W = (1000 \text{ kg}) \times (10 \text{ m/s}^2) \times (90 \text{ m})$$

$$W = 1000 \times 10 \times 90 \text{ J}$$

$$W = 10000 \times 90 \text{ J}$$

$$W = 900000 \text{ J}$$

The work done by the pump is 900,000 Joules.

Step 3: Calculate the Power

Now we can calculate the power (P) using the formula $P = \frac{W}{t}$.

$$P = \frac{\text{Work done}}{\text{Time taken}}$$

Substitute the calculated work and the time in seconds:

$$P = \frac{900000 \text{ J}}{1800 \text{ s}}$$

To simplify the division, we can cancel out zeros:

$$P = \frac{9000}{18} \text{ W}$$

We can divide 9000 by 18. $90 \div 18 = 5$. So, $9000 \div 18 = 500$.

$$P = 500 \text{ W}$$

The power of the pump is 500 Watts.

Comparing with Options

We found the power to be 500 W. Let's compare this with the given options:

Option	Value (W)
1	50
2	250
3	25
4	500

Our calculated value of 500 W matches Option 4.

Revision Table: Pump Power Calculation

Quantity	Symbol	Value	Unit
Mass of water	m	1000	kg (1 tonne)
Height lifted	h	90	m
Time taken	t	1800	s (30 minutes)
Gravity	g	10	m/s ²
Work Done	$W = mgh$	900000	J
Power	$P = W/t$	500	W

Additional Information on Power and Efficiency

The problem assumes 100% efficiency. In reality, pumps are not 100% efficient. Actual efficiency (η) is less than 100%.

Actual Power Output = Work Done / Time Taken (This is the useful power output)

Input Power = Actual Power Output / Efficiency (η)

If the efficiency was, say, 80% (or 0.8), the required input power from the motor driving the pump would be higher than the calculated output power.

$$\text{Input Power} = \frac{500 \text{ W}}{0.8} = 625 \text{ W}$$

This means the motor would need to supply 625 W of power for the pump to deliver 500 W of useful power in lifting the water. The remaining power is typically lost as heat or sound.

Understanding the relationship between work, energy, power, and efficiency is crucial in solving problems related to machinery and energy transfer.

56. Answer: b

Explanation:

Decoding the Artificial Language Puzzle

This question asks us to decode a simple artificial language based on a few given words and their meanings and then find the word for 'overdrive'. We need to look for repeating patterns or parts of words that correspond to parts of the meaning.

Analyzing the Given Words and Meanings

We are given the following:

- unacri means lighthouse
- bancri means lightweight
- banetu means overweight

Step-by-Step Decoding Process

Let's compare the words to find common parts and their meanings:

1. Compare 'bancri' (lightweight) and 'banetu' (overweight).
2. Both words start with 'ban'. The meanings differ in 'light' vs 'over', corresponding to the endings 'cri' vs 'etu'.
3. This suggests that 'ban' might mean 'weight' or something related to the state of being weighted, 'cri' means 'light', and 'etu' means 'over'.
4. Let's test this with the third word: 'unacri' (lighthouse).
5. If 'cri' means 'light', then 'una' must mean 'house' for 'unacri' to mean 'lighthouse'.

So, our potential decoded parts are:

- **ban**: weight
- **cri**: light
- **etu**: over
- **una**: house

Confirming the Decoding

Let's check if these parts fit the original words:

- unacri: una (house) + cri (light) = houselight (close to lighthouse)
- bancri: ban (weight) + cri (light) = lightweight
- banetu: ban (weight) + etu (over) = overweight

The decoding seems consistent with the given words. The structure appears to be combining word parts based on meaning.

Finding the Word for 'overdrive'

We need the word for 'overdrive'. We have already identified the word part for 'over', which is 'etu'. We need to find the word part for 'drive'.

Let's look at the options provided:

1. adiban
2. adietu
3. simuna
4. criteh

We are looking for a word that combines the meaning 'over' ('etu') with 'drive'.

- Option 1: adiban = adi + ban ('weight'). This doesn't match 'overdrive'.
- Option 2: adietu = adi + etu ('over'). This word contains 'etu' (over). If 'adi' means 'drive', then 'adietu' would mean 'drive' + 'over' or 'over' + 'drive'. This looks like a strong candidate.
- Option 3: simuna = simu + una ('house'). This doesn't match 'overdrive'.
- Option 4: criteh = cri ('light') + teh. This doesn't match 'overdrive'.

Based on the options and our decoding, it is highly probable that 'adi' means 'drive' and the word for 'overdrive' is formed by combining 'adi' and 'etu'. The option 'adietu' fits this pattern.

Therefore, combining 'adi' (drive) and 'etu' (over) gives 'adietu', which means 'overdrive'.

Artificial Language Decoding Summary

Word Part	Meaning
una	house
ban	weight
cri	light
etu	over
adi	drive

Based on this analysis, the word for 'overdrive' is constructed from 'adi' (drive) and 'etu' (over), resulting in 'adietu'.

Revision Table: Artificial Language Decoding

Review of Decoded Parts and Combinations

Original Word	Meaning	Parts	Decoded Meaning
unacri	lighthouse	una + cri	house + light
bancri	lightweight	ban + cri	weight + light
banetu	overweight	ban + etu	weight + over
adietu	overdrive	adi + etu	drive + over

Additional Information: Understanding Artificial Languages

Artificial languages, like the one in this puzzle, are constructed languages. They are often created for specific purposes, such as:

- **Fictional Worlds:** Languages like Sindarin and Klingon were created for books or movies to add realism.
- **Communication:** Languages like Esperanto were created to be easy to learn and used as a universal second language.
- **Linguistic Experimentation:** Some languages are created to test hypotheses about language structure or human cognition.

These puzzles help develop logical reasoning and pattern recognition skills by requiring us to act like linguists, analyzing known examples to deduce rules and apply them to new words. The core idea is often that words are built from smaller meaningful units, similar to how words in natural languages are formed from morphemes (like prefixes, suffixes, and roots).

57. Answer: c

Explanation:

Ingredient B of Salad X = 600 gm

After reduction, Ingredient B of Salad X = $(600 - 100)$ gm = 500 gm

Ingredient E of salad Y = 400 gm

After increasing 25%, Ingredient E of salad Y = $(400 + 400 \times 25\%) = 500$ gm

Total weight of Salad X after reduction = $(500 + 500 + 200 + 300 + 100 + 400)$ gm = 2000 gm

Total weight of Salad Y after increasing = $(300 + 100 + 200 + 200 + 500 + 300)$ gm = 1600 gm

$\therefore$ Salad Y would be lighter than salad X = $(2000 - 1600) = 400$ gm

Then salad Y would be lighter than salad X by = $(400/1600) \times 100\% = 25\%$

58. Answer: d

Explanation:

Solving Anu and Dimpy's Age Ratio Problem

This problem involves understanding and solving age-based word problems using ratios and simple algebraic equations. We are given the current ratio of ages and the ratio after a certain number of years and asked to find the age before a certain period.

Step-by-Step Solution for Age Calculation

Let's break down the problem to find Dimpy's age before 2 years.

- 1. Represent the current ages using the given ratio:** The current ratio of Anu's age to Dimpy's age is 5 : 8. Let their current ages be $5x$ years and $8x$ years, respectively, where x is a common factor.
- 2. Express their ages after 2 years:** After 2 years, Anu's age will be $5x + 2$ years, and Dimpy's age will be $8x + 2$ years.
- 3. Set up an equation based on the ratio after 2 years:** The problem states that after 2 years, the ratio of their ages will become 2 : 3. So, we can write the equation:

$$\frac{5x + 2}{8x + 2} = \frac{2}{3}$$

4. **Solve the equation for x :** To solve for x , we will cross-multiply:

$$3 \times (5x + 2) = 2 \times (8x + 2)$$

$$15x + 6 = 16x + 4$$

Now, rearrange the terms to isolate x :

$$6 - 4 = 16x - 15x$$

$$2 = x$$

So, the value of x is 2.

5. **Calculate the current ages:**

- Anu's current age = $5x = 5 \times 2 = 10$ years.
- Dimpy's current age = $8x = 8 \times 2 = 16$ years.

6. **Calculate Dimpy's age before 2 years:** The question asks for Dimpy's age before 2 years.

$$\text{Dimpy's age before 2 years} = \text{Dimpy's current age} - 2$$

$$\text{Dimpy's age before 2 years} = 16 - 2 = 14 \text{ years}$$

Therefore, Dimpy's age before 2 years was 14 years.

Verification of Age Ratio

Let's verify the ages and ratios based on our calculation:

- Current ages: Anu = 10, Dimpy = 16. Ratio = $10 : 16 = 5 : 8$ (Matches the given current ratio).
- Ages after 2 years: Anu = $10 + 2 = 12$, Dimpy = $16 + 2 = 18$. Ratio = $12 : 18 = 2 : 3$ (Matches the given future ratio).

The calculated ages satisfy both conditions in the problem statement.

Understanding Age Ratio Problems

Age ratio problems are common in quantitative aptitude. They typically involve setting up equations based on given ratios at different points in time (present, past, or future).

- **Key Idea:** The difference between the ages of two people remains constant over time. However, the ratio of their ages changes over time.
- **Method:** Use a variable (like x) to represent the common multiple in the ratio. Formulate expressions for ages at different times and set up an equation based on the ratio given for one of those times. Solve the equation to find the variable, and then calculate the required ages.

Revision Table: Key Information

Description	Value/Expression
Current Anu:Dimpy Ratio	5 : 8
Current Ages (Anu, Dimpy)	$5x, 8x$
Ages After 2 Years (Anu, Dimpy)	$5x + 2, 8x + 2$
Ratio After 2 Years	2 : 3
Equation Solved	$\frac{5x+2}{8x+2} = \frac{2}{3}$
Value of x	2
Current Dimpy's Age	16 years
Dimpy's Age Before 2 Years	14 years

Your Personal Exams Guide

Additional Information on Ratio Problems

Ratios provide a way to compare quantities. In age problems, ratios change as time passes, but the difference in age remains constant.

For example, if person A is 10 and person B is 15, the ratio is $10:15 = 2:3$. The difference is 5 years. After 10 years, A is 20 and B is 25. The ratio is $20:25 = 4:5$. The ratio changed, but the age difference is still 5 years ($25 - 20$).

Setting up proportions (like the fraction equation we used) is a standard method to solve problems involving ratios changing over time.

59. Answer: a

Explanation:

Understanding Syllogism Statements and Conclusions

This question requires us to analyze given statements and determine which of the provided conclusions logically follow from them. This type of problem falls under the category of Syllogism or Deductive Reasoning.

Analyzing the Statements

We are given two statements:

1. Statement 1: All flags are banners.
2. Statement 2: Some flags are emblems.

Let's represent these statements using set theory concepts for clarity:

- "All flags are banners" means that the set of 'flags' is a subset of the set of 'banners'. If something is a flag, it must also be a banner.
- "Some flags are emblems" means that there is at least one flag that is also an emblem. The set of 'flags' and the set of 'emblems' have a non-empty intersection.

Evaluating the Conclusions

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

Conclusion I: Some emblems are banners.

From Statement 2, we know there are some flags that are emblems. Let's call one such flag 'X'. So, X is a flag and X is an emblem.

From Statement 1, we know that all flags are banners. Since X is a flag, X must also be a banner.

So, X is an emblem (from Statement 2) and X is a banner (from Statement 1). This means there is at least one entity (X) that is both an emblem and a banner.

Therefore, the conclusion "Some emblems are banners" logically follows from the statements.

Conclusion II: All banners are emblems.

Statement 1 tells us all flags are banners. Statement 2 tells us some flags are emblems. However, the statements do not provide information about all banners. It is possible that there are banners that are not flags. If there are banners that are not flags, the statements give us no information about whether those banners are emblems or not. Therefore, we cannot conclude that ALL banners are emblems.

This conclusion does not logically follow.

Conclusion III: No emblems are banners.

As we determined while evaluating Conclusion I, there must be at least one entity that is both an emblem and a banner (any flag that is also an emblem must be a banner). This directly contradicts the conclusion "No emblems are banners".

Therefore, this conclusion does not logically follow; in fact, it is false based on the statements.

Summary of Conclusions

Based on our analysis:

- Conclusion I: Some emblems are banners – **Follows**
- Conclusion II: All banners are emblems – **Does Not Follow**
- Conclusion III: No emblems are banners – **Does Not Follow**

Only Conclusion I follows from the given statements.

Revision Table: Syllogism Analysis

Statement/Conclusion	Description	Follows?
Statement 1	All flags are banners	Given
Statement 2	Some flags are emblems	Given
Conclusion I	Some emblems are banners	Yes
Conclusion II	All banners are emblems	No
Conclusion III	No emblems are banners	No

Additional Information: Understanding Syllogisms

Syllogism is a form of deductive reasoning where a conclusion is drawn from two or more given or assumed premises (statements). The key is to determine what **MUST** be true if the statements are true, regardless of whether the statements align with real-world facts. Syllogism problems often involve categories and their relationships, which can sometimes be visualized using Venn diagrams.

Common types of statements in syllogism include:

- **All A are B:** Every member of set A is also a member of set B. ($A \subseteq B$)
- **No A are B:** Set A and Set B have no members in common. ($A \cap B = \emptyset$)
- **Some A are B:** There is at least one member that is in both set A and set B. ($A \cap B \neq \emptyset$)
- **Some A are not B:** There is at least one member in set A that is not in set B.

When solving syllogism problems, focus strictly on the information given in the statements and use logical rules to derive conclusions. Avoid using your own knowledge about the real world if it contradicts the statements.

60. Answer: a

Explanation:

Understanding Storage Medium Preparation

The question asks for the term that describes the process of preparing a storage medium, like a disk, so that data can be written to it and read from it. This preparation involves setting up the necessary structure on the disk.

Analyzing the Options for Disk Preparation

Let's look at the given options and determine which one fits the description of preparing a storage medium:

- **format:** This term specifically refers to the process of preparing a data storage device such as a hard disk drive, solid-state drive, floppy disk, or USB flash drive for initial use. Formatting involves setting up a file system on the device, which is a structure that the operating system uses to manage files and folders. It essentially creates the 'address book' for storing data.
- **defrag:** Defragmentation is a process of reorganizing the data stored on a hard drive. Over time, as files are saved, deleted, and modified, parts of files can become scattered (fragmented) across the disk. Defragmenting brings these pieces back together and consolidates free space, which can improve performance. It is a maintenance task performed on a disk that is already in use, not a preparation step for initial use.
- **map:** Mapping generally refers to establishing a correspondence or relationship between different entities, like mapping network drives to local drive letters or mapping logical addresses to physical addresses in memory or storage. It doesn't describe the process of preparing the physical storage medium itself for initial use.
- **boot:** Booting is the process of starting up a computer. It involves loading the operating system into memory so the computer can be used. This process happens after the storage devices (including the one with the operating system) have already been prepared and configured.

Identifying the Correct Term

Based on the analysis:

The term that specifically means to prepare a storage medium, typically a disk, for reading and writing data by creating a file system is **format**.

Formatting is a crucial step before a new hard drive or other storage device can be used by an operating system to store files. It sets up the fundamental structure required for data management.

Conclusion

The correct term for preparing a storage medium like a disk for reading and writing is **format**.

Summary of Storage Operations

Term	Description	Related to Preparation for Use?
Format	Prepares a storage medium for use by creating a file system.	Yes
Defrag	Reorganizes data on an *already used* disk to improve performance.	No
Map	Establishes relationships or assigns addresses (e.g., network drive mapping).	No
Boot	Starts a computer system.	No

Revision Table: Storage Terms

Key Terms in Storage Management

Term	Function
Format	Initial preparation of a disk for reading/writing.
Defragmentation	Optimizing file arrangement on a disk.
File System	Structure on a disk for organizing files.
Storage Medium	Device storing data (e.g., HDD, SSD, USB).

Additional Information on Formatting Storage

Formatting a storage device typically involves two main steps:

1. **Low-Level Formatting (L.L.F):** This is a more fundamental process that defines the basic physical structure of the disk, such as tracks and sectors. Modern hard drives are low-level formatted at the factory, and users typically do not perform this.
2. **High-Level Formatting (H.L.F):** This is what users commonly refer to as "formatting". It involves writing the file system structure onto the disk. This includes creating the boot sector, file allocation tables (like FAT, NTFS, exFAT), and root directories, which allows the operating system to store and retrieve files. When you format a drive in Windows, macOS, or Linux, you are usually performing a high-level format.

Choosing the correct file system (NTFS, FAT32, exFAT, HFS+, APFS, ext4, etc.) depends on the operating system you will use and the intended purpose of the storage device.

61. Answer: c

Explanation:

Identifying the Different Family Relationship

The question asks us to identify the relationship that is different from the others among the given options. We need to analyze the nature of each family connection to find the distinct one.

Analyzing the Family Ties

Let's look at each option and understand the type of relationship it represents:

- **Father:** This is a direct blood relationship. A father is a male parent, part of the immediate family.
- **Son:** This is also a direct blood relationship. A son is a male child, part of the immediate family.
- **Mother-in-law:** This relationship is formed through marriage. A mother-in-law is the mother of one's spouse. It is an affinal (by marriage) relationship, not a direct blood (consanguineous) one.
- **Brother:** This represents a direct blood relationship. A brother is a male sibling, part of the immediate family.

Categorizing the Relationships

We can categorize these relationships based on how they are formed:

- **Blood Relatives (Consanguineous):** These are individuals related by descent from a common ancestor. In this list, Father, Son, and Brother fall into this category. They are part of the nuclear or immediate family structure based on birth.
- **Relatives by Marriage (Affinal):** These are individuals related through marriage. Mother-in-law belongs to this category.

Conclusion on the Different Relationship

Based on the analysis, 'Father', 'Son', and 'Brother' are all direct blood relatives. 'Mother-in-law', however, is a relative by marriage. Therefore, the relationship formed through marriage makes 'Mother-in-law' the one that is different from the rest.

62. Answer: c

Explanation:

Understanding Base Units in Physics

In physics and measurement, units are categorized into base units and derived units. Base units are a set of fundamental units from which all other units can be derived. The International System of Units (SI) defines seven base units. Understanding these base units is crucial for comprehending various physical quantities and their measurements.

Identifying the SI Base Units

The seven fundamental SI base units are:

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

- Candela (cd) for luminous intensity

Analyzing the Given Options

The question asks which of the given options is NOT a base unit. Let's examine each option in relation to the list of SI base units:

- **mole:** The mole (mol) is the SI base unit for the amount of substance.
- **ampere:** The ampere (A) is the SI base unit for electric current.
- **radian:** The radian (rad) is a unit for measuring angles. While fundamental in geometry, it is considered a derived unit in the SI system (specifically, a dimensionless derived unit).
- **candela:** The candela (cd) is the SI base unit for luminous intensity.

Determining the Unit That is NOT a Base Unit

Based on the definition of SI base units and our analysis of the options, the units mole, ampere, and candela are all included in the list of the seven fundamental SI base units. The radian, however, is a unit used for plane angles and is classified as a derived unit within the SI system, even though it is dimensionless.

Therefore, the unit that is NOT a base unit among the given options is the radian.

Unit	Category (SI System)
Mole	Base Unit
Ampere	Base Unit
Radian	Derived Unit (dimensionless)
Candela	Base Unit

Conclusion

The question asks to identify the unit that is NOT a base unit. From the options provided, mole, ampere, and candela are SI base units. Radian is a derived unit for plane angles. Hence, radian is the correct answer.

Revision Table: Key SI Units

Physical 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

Additional Information: Derived Units vs. Base Units

Derived units are units that are combinations of base units. For example, the unit for speed, meters per second (m/s), is derived from the base units of length (meter) and time (second). The unit for force, Newton (N), is derived from kilograms, meters, and seconds (

$$N = kg \cdot m/s^2$$

).

Your Personal Exams Guide

Radian is a special type of derived unit. It is a unit for plane angle, defined as the ratio of the arc length subtending the angle to the radius of the circle. Since it is a ratio of two lengths, it is dimensionless. The steradian (sr) is the SI derived unit for solid angle, and it is also dimensionless. While sometimes referred to as "supplementary units" historically, the SI currently classifies radian and steradian as derived units.

63. Answer: a

Explanation:

Understanding the Invention of Blue Jeans

The question asks about the inventor of blue jeans. Let's look at the options provided to determine the correct answer.

Analyzing the Options

- Levi Strauss: This name is strongly associated with denim clothing and a famous brand of jeans.
- Todd Oldham: An American fashion designer known for colorful and playful designs.
- Calvin Klein: A prominent American fashion designer and founder of the Calvin Klein brand, known for various clothing lines, including jeans.
- Gianni Versace: An Italian fashion designer and founder of the Versace brand, known for luxurious and bold designs.

While Todd Oldham, Calvin Klein, and Gianni Versace are all well-known fashion designers who have certainly designed jeans as part of their collections, they are not credited with inventing the original blue jeans.

The True Story Behind Blue Jeans Invention

The invention of blue jeans is primarily credited to Levi Strauss and Jacob Davis. Jacob Davis, a tailor in Reno, Nevada, was looking for a way to make work pants more durable for his customers, particularly miners and laborers.

He came up with the idea of using metal rivets at points of strain, like pockets and fly, to reinforce the seams. Realizing the potential of this idea but needing a business partner and funding, Davis contacted Levi Strauss, who owned a wholesale dry goods business in San Francisco and was one of Davis's suppliers of denim and canvas fabric.

Strauss and Davis worked together and received a U.S. patent (No. 139,121) for "Improvement in Fastening Pocket-Openings" on May 20, 1873. This patent marked the official birth of riveted denim work pants, which are now known globally as blue jeans. Levi Strauss provided the fabric and business acumen, while Jacob Davis contributed the innovative riveting idea.

Therefore, Levi Strauss, in partnership with Jacob Davis, was instrumental in the invention and popularization of blue jeans.

Conclusion

Based on historical facts and the options provided, Levi Strauss is the correct answer associated with the invention of blue jeans.

Revision Table: Key Facts about Blue Jeans Invention

Key Aspect	Detail
Inventors Credited	Levi Strauss and Jacob Davis
Invention Date (Patent)	May 20, 1873
Key Innovation	Using metal rivets to reinforce stress points
Purpose	Creating durable workwear

Additional Information: Evolution of Blue Jeans

After their invention as durable workwear, blue jeans gradually evolved and became popular casual wear, and eventually a fashion staple worldwide. Key points in their evolution include:

- **Material:** Originally made from durable denim or canvas. Denim became the standard.
- **Color:** The term "blue jeans" comes from the indigo dye used to color the denim.
- **Cultural Impact:** Blue jeans transitioned from work clothes to symbols of rebellion in the mid-20th century, and later became a ubiquitous item of casual clothing.
- **Design Variations:** Over time, various fits, washes, and styles of blue jeans emerged, catering to different fashion trends and preferences.

The legacy of Levi Strauss and Jacob Davis's invention continues today, with blue jeans remaining one of the most popular garments globally.

64. Answer: b

Explanation:

Understanding Resistance and Conductivity in Metal Wires

The electrical properties of a metal wire, such as its resistance, depend on its physical dimensions and the material it is made of. Resistance is a measure of how much a material opposes the flow of electric current. Conductivity is the opposite; it measures how well a material conducts electric current.

Relating Resistance, Length, and Area

For a uniform metal wire, the resistance 'R' is directly proportional to its length 'L' and inversely proportional to its area of cross-section 'A'. This relationship can be written as:

$$R \propto \frac{L}{A}$$

To turn this proportionality into an equation, we introduce a constant of proportionality, which is the resistivity of the material, denoted by ρ . The formula becomes:

$$R = \rho \frac{L}{A}$$

Here, ρ is a property of the material of the wire at a specific temperature.

Connecting Resistivity and Conductivity

Resistivity ρ and conductivity σ are inversely related. Conductivity is the reciprocal of resistivity:

$$\sigma = \frac{1}{\rho}$$

This also means that resistivity is the reciprocal of conductivity:

$$\rho = \frac{1}{\sigma}$$

Deriving the Relation between R, σ , L, and A

Now, we can substitute the expression for resistivity ($\rho = \frac{1}{\sigma}$) into the formula for resistance:

$$R = \left(\frac{1}{\sigma}\right) \frac{L}{A}$$

$$R = \frac{L}{\sigma A}$$

We need to rearrange this equation to match the given options. Let's multiply both sides by σA :

$$R \times \sigma A = \frac{L}{\sigma A} \times \sigma A$$

$$R\sigma A = L$$

This gives us the correct relationship between resistance 'R', conductivity ' σ ', area of cross-section 'A', and length 'L' of the metal wire.

Analyzing the Options

Let's look at the given options and see which one matches our derived relation:

1. $RL = A\sigma$

This can be rearranged to $R = \frac{A\sigma}{L}$. This is not our derived formula $R = \frac{L}{\sigma A}$.

2. $R\sigma A = L$

This is exactly the relation we derived: $R\sigma A = L$.

3. $RA = L\sigma$

This can be rearranged to $R = \frac{L\sigma}{A}$. This is not our derived formula $R = \frac{L}{\sigma A}$.

4. $\sigma = RL/A$

This can be rearranged to $\sigma A = RL$ or $\sigma = R\frac{L}{A}$. This shows that conductivity is proportional to resistance and the ratio of length to area, which is incorrect. In fact, $R = \rho\frac{L}{A}$, so $\frac{R}{L/A} = \rho$, and $\sigma = \frac{1}{\rho} = \frac{1}{R/(L/A)} = \frac{L}{RA}$. So, $\sigma = \frac{L}{RA}$, not $\sigma = \frac{RL}{A}$.

Based on the analysis, the correct relationship is $R\sigma A = L$.

Revision Table: Electrical Properties

Property	Symbol	Definition	Unit (SI)
Resistance	R	Opposition to current flow	Ohm (Ω)
Resistivity	ρ	Intrinsic property of material opposing current flow	Ohm-meter (Ωm)
Conductivity	σ	Intrinsic property of material allowing current flow	Siemens per meter (S/m)
Length	L	Length of the conductor	Meter (m)
Area of Cross-section	A	Area perpendicular to current flow	Square meter (m^2)

Additional Information: Factors Affecting Resistance and Conductivity

The resistance of a wire is not only dependent on its dimensions (length and area) and the material's intrinsic properties (resistivity/conductivity) but also on temperature. For most metals, resistance increases with increasing temperature. This is because the increased thermal vibrations of atoms impede the flow of electrons.

Different materials have vastly different resistivity and conductivity values. Conductors like metals (copper, aluminum, gold) have low resistivity (high conductivity), making them suitable for electrical wiring. Insulators like rubber, glass, and plastic have very high resistivity (low conductivity), preventing current flow.

65. Answer: c

Explanation:

Understanding Word Analogies in Reasoning

Word analogy questions test your ability to recognize the relationship between a pair of words and apply that same relationship to a different word to find a matching word from the options. The core of solving these questions is identifying the exact nature of the connection between the first two words.

Analyzing the Given Analogy: Poem : Stanza :: Book : ?

Let's break down the first pair: **Poem : Stanza**.

- A **Poem** is a piece of writing that often uses rhythm and rhyme.
- A **Stanza** is a division of a poem, consisting of a series of lines arranged together. It's like a paragraph in prose.

So, the relationship here is that a **Stanza** is a structural part or division of a **Poem**.

Applying the Relationship to the Second Pair: Book : ?

Now, we need to find a word that has the same relationship with **Book** as **Stanza** has with **Poem**. We are looking for a structural part or division of a **Book**.

Evaluating the Options

Let's examine each option provided:

- **Option 1: Language**

Language (e.g., English, Spanish) is the medium in which a book is written, not a physical or structural division of the book itself.

- **Option 2: Story**

A story is the narrative content of a book, what the book is about. While a book contains a story, 'story' is not a structural division of the book in the same way a stanza is a division of a poem.

- **Option 3: Page**

A **Page** is one side of a sheet of paper in a book. A book is physically composed of many pages bound together. A page is a fundamental structural unit or division of a book.

- **Option 4: Printing**

Printing is the process used to produce a book. It is not a part or division of the finished book.

Identifying the Correct Analogy Relation

Comparing the relationship 'Stanza is a part/division of a Poem' with the relationship between 'Book' and the options:

- Language is not a part/division of a Book.
- Story is content, not a structural part/division of a Book.
- **Page is a structural part/division of a Book.**
- Printing is not a part/division of a Book.

The relationship between **Book** and **Page** is the same as the relationship between **Poem** and **Stanza** (Part : Whole or Division : Whole).

Therefore, the word that completes the analogy is **Page**.

The completed analogy is: Poem : Stanza :: Book : Page.

Revision Table: Word Analogy Concepts

Term	Definition	Relation in Analogy
Poem	A piece of writing, often with rhythm/rhyme.	Whole (made of stanzas)
Stanza	A unit of lines in a poem.	Part/Division of a Poem
Book	A collection of written or printed pages bound together.	Whole (made of pages)
Page	One side of a sheet of paper in a book.	Part/Division of a Book

Additional Information: Types of Analogies

Word analogies can have various types of relationships. Recognizing these common patterns helps solve analogy questions quickly. Some common types include:

- **Part and Whole:** e.g., Finger : Hand (Finger is part of a Hand)
- **Cause and Effect:** e.g., Rain : Flood (Rain can cause a Flood)
- **Synonyms:** e.g., Happy : Joyful (Both mean the same)
- **Antonyms:** e.g., Hot : Cold (Opposite meanings)
- **Worker and Tool:** e.g., Carpenter : Hammer (Carpenter uses a Hammer)
- **Tool and Object it is used on:** e.g., Knife : Cut (Knife is used to Cut)

- **Action and Object:** e.g., Read : Book (Action performed on the object)
- **Category and Item:** e.g., Fruit : Apple (Apple is an item in the Fruit category)
- **Degree of Intensity:** e.g., Warm : Hot (Hot is a stronger degree of Warm)

For the analogy Poem : Stanza :: Book : ?, the relationship is clearly **Part and Whole**.

66. Answer: b

Explanation:

Evaluate Trigonometric Expression at 30 Degrees

The question asks us to find the value of the expression $5 \sin \theta - 2 \csc \theta$ when $\theta = 30^\circ$. To solve this, we need to know the standard trigonometric values for an angle of 30° and understand the relationship between the sine and cosecant functions.

Understanding Sine and Cosecant

The sine function, denoted as $\sin \theta$, is one of the fundamental trigonometric functions. The cosecant function, denoted as $\csc \theta$, is the reciprocal of the sine function. This means:

$$\csc \theta = \frac{1}{\sin \theta}$$

Provided that $\sin \theta \neq 0$.

Standard Trigonometric Values for 30 Degrees

For standard angles like $0^\circ, 30^\circ, 45^\circ, 60^\circ, 90^\circ$, the trigonometric values are commonly known or can be derived from special right triangles. For $\theta = 30^\circ$:

- The value of $\sin 30^\circ$ is $1/2$.
- Since $\csc \theta$ is the reciprocal of $\sin \theta$, the value of $\csc 30^\circ$ is the reciprocal of $\sin 30^\circ$.

$$\text{So, } \csc 30^\circ = \frac{1}{\sin 30^\circ} = \frac{1}{1/2} = 2.$$

Standard Trigonometric Values for 30°

Trigonometric Function	Value at 30°
$\sin 30^\circ$	$1/2$
$\csc 30^\circ$	2

Evaluating the Expression

Now we substitute the values of $\sin 30^\circ$ and $\csc 30^\circ$ into the given expression $5 \sin \theta - 2 \csc \theta$:

Substitute $\theta = 30^\circ$:

$$5 \sin 30^\circ - 2 \csc 30^\circ$$

Substitute the values $\sin 30^\circ = 1/2$ and $\csc 30^\circ = 2$:

$$5 \left(\frac{1}{2} \right) - 2(2)$$

Perform the multiplication:

$$\frac{5}{2} - 4$$

To subtract 4 from $5/2$, we convert 4 into a fraction with a denominator of 2. Since $4 = 8/2$:

$$\frac{5}{2} - \frac{8}{2}$$

Now subtract the numerators, keeping the common denominator:

$$\frac{5 - 8}{2} = \frac{-3}{2}$$

So, the value of the expression $5 \sin \theta - 2 \csc \theta$ when $\theta = 30^\circ$ is $-3/2$.

Final Answer Calculation Steps

1. Identify the given expression: $5 \sin \theta - 2 \csc \theta$.
2. Identify the value of θ : $\theta = 30^\circ$.
3. Find the value of $\sin 30^\circ$: $\sin 30^\circ = 1/2$.
4. Find the value of $\csc 30^\circ$ using $\csc \theta = 1/\sin \theta$: $\csc 30^\circ = 1/(1/2) = 2$.

5. Substitute these values into the expression: $5(1/2) - 2(2)$.

6. Calculate the result: $5/2 - 4 = 5/2 - 8/2 = -3/2$.

Revision Table: Trigonometry Basics

Key Trigonometric Concepts for Evaluation

Concept	Description	Example ($\theta = 30^\circ$)
Sine ($\sin \theta$)	Ratio of the length of the opposite side to the length of the hypotenuse in a right triangle.	$\sin 30^\circ = 1/2$
Cosecant ($\csc \theta$)	Reciprocal of the sine function. $\csc \theta = 1/\sin \theta$.	$\csc 30^\circ = 1/(1/2) = 2$
Evaluating Expressions	Substituting known values of variables (like θ) into an expression and performing the calculations.	Substitute $\sin 30^\circ = 1/2$ and $\csc 30^\circ = 2$ into $5 \sin \theta - 2 \csc \theta$.

Additional Information: Reciprocal Trigonometric Functions

Besides sine and cosecant, there are other reciprocal trigonometric function pairs:

- **Cosine** ($\cos \theta$) and **Secant** ($\sec \theta$): $\sec \theta = 1/\cos \theta$.
- **Tangent** ($\tan \theta$) and **Cotangent** ($\cot \theta$): $\cot \theta = 1/\tan \theta$.

Understanding these reciprocal relationships is crucial for solving many trigonometric problems and simplifying expressions.

Remembering the standard trigonometric values for common angles ($0^\circ, 30^\circ, 45^\circ, 60^\circ, 90^\circ$) is very helpful for quickly evaluating trigonometric expressions like the one in this problem.

67. Answer: c

Explanation:

Gear Train Calculation: Driver and Follower Turns

This problem involves calculating the relationship between the rotation of two gears in a gear train. We are given the number of teeth on the driver gear and the follower gear, along with the number of turns the driver gear makes. We need to find out how many times the follower gear turns.

Understanding Gear Ratios

In a simple gear train, the ratio of the number of turns made by the follower gear to the number of turns made by the driver gear is equal to the ratio of the number of teeth on the driver gear to the number of teeth on the follower gear. This is a fundamental concept in understanding how gears transmit motion and torque.

The formula representing this relationship is:

$$\frac{\text{Turns of Follower}}{\text{Turns of Driver}} = \frac{\text{Teeth on Driver}}{\text{Teeth on Follower}}$$

Or, using notation:

$$\frac{T_F}{T_D} = \frac{N_D}{N_F}$$

Given Information

Let's list the details provided in the question:

- Number of teeth on the driver gear (N_D): 18
- Number of teeth on the follower gear (N_F): 8
- Number of turns made by the driver gear (T_D): 16

We need to find the number of turns made by the follower gear (T_F).

Parameter	Symbol	Value
Driver Teeth	N_D	18
Follower Teeth	N_F	8
Driver Turns	T_D	16
Follower Turns	T_F	?

Step-by-Step Calculation

We can use the gear ratio formula to find the number of turns for the follower gear. Substitute the given values into the formula:

$$\frac{T_F}{16} = \frac{18}{8}$$

To solve for T_F , we can rearrange the equation:

$$T_F = 16 \times \frac{18}{8}$$

Now, perform the calculation:

First, simplify the fraction $\frac{18}{8}$ or divide 16 by 8:

$$T_F = (16 \div 8) \times 18$$

$$T_F = 2 \times 18$$

$$T_F = 36$$

Conclusion

For every 16 turns of the driver gear, the follower gear turns 36 times. This means the follower gear rotates faster than the driver gear because it has fewer teeth.

Comparing this result with the given options, the correct option is 36.

68. Answer: b

Explanation:

Understanding Units in Physics: Farad per Metre

The question asks for the physical quantity whose unit is Farad per metre. Understanding the units of fundamental electrical quantities is crucial in physics.

What is Permittivity?

Permittivity is a material property that describes how an electric field affects, and is affected by, a dielectric medium. It relates the electric displacement field ($\vec{D}$) to the electric field ($\vec{E}$).

The relationship is given by:

$$\vec{D} = \epsilon \vec{E}$$

where:

- $\vec{D}$ is the electric displacement field
- $\vec{E}$ is the electric field
- ϵ is the permittivity of the medium

Permittivity (ϵ) can be thought of as a measure of how much resistance is encountered when forming an electric field in a medium. In vacuum, this property is represented by the permittivity of free space, ϵ_0 .

Deriving the Unit of Permittivity

The unit of permittivity can be derived from its definition or from formulas involving capacitance. A common formula for the capacitance (C) of a parallel-plate capacitor is:

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

where:

- C is the capacitance
- ϵ is the permittivity of the dielectric material between the plates
- A is the area of one plate
- d is the distance between the plates

Rearranging this formula to solve for permittivity (ϵ):

$$\epsilon = \frac{Cd}{A}$$

Now, let's look at the units of each term:

- Capacitance (C) is measured in Farads (F).
- Distance (d) is measured in metres (m).
- Area (A) is measured in square metres (m²).

Substituting these units into the formula for ϵ :

$$[\epsilon] = \frac{[C][d]}{[A]} = \frac{F \cdot m}{m^2} = \frac{F}{m}$$

Thus, the unit of permittivity is Farad per metre (F/m).

Analyzing Other Options

Let's consider the units of the other given options to confirm they are not Farad per metre.

Quantity	Typical Symbol(s)	Definition/Formula Example	SI Unit	Notes
Watt per steradian		Radiant intensity ($I = P/\Omega$)	W/sr	Unit of radiant intensity or luminous intensity. Not a fundamental property like permittivity.
Permeability	μ	Relates magnetic field (H) to magnetic flux density (B): $\vec{B} = \mu \vec{H}$	Henry per metre (H/m)	Analogous to permittivity, but for magnetic fields.
Electric conductance	G	Reciprocal of resistance ($G = 1/R$)	Siemens (S) or Ω^{-1}	Measures how easily current flows through a material.

Comparing the derived unit for permittivity (F/m) with the units of the other options, it is clear that Farad per metre is the unit specifically associated with permittivity.

Conclusion

Based on the unit derivation from the formula for capacitance and the fundamental definition relating electric displacement and electric field, the unit Farad per metre corresponds to permittivity. The units for watt per steradian, permeability, and electric conductance are distinctly different.

Revision Table: Key Electrical Units

Quantity	SI Unit	Unit Symbol
Capacitance	Farad	F
Permittivity	Farad per metre	F/m
Permeability	Henry per metre	H/m
Electric Resistance	Ohm	Ω
Electric Conductance	Siemens	S
Electric Field Strength	Volt per metre	V/m
Electric Displacement Field	Coulomb per square metre	C/m ²

Additional Information on Permittivity and Units

Permittivity is a key concept in electrostatics, especially when dealing with materials placed in electric fields. It determines how much electrical energy can be stored in a material for a given electric field, which is why it appears in the formula for capacitance. The higher the permittivity, the more charge can be stored for a given voltage and geometry. The unit Farad per metre directly reflects this ability to store charge (related to Farad) over a distance (metre) within a material.

It's important not to confuse permittivity (ϵ) with permeability (μ). While both are material properties, permittivity relates to how a material responds to an electric field, whereas permeability relates to how a material responds to a magnetic field. Their units, F/m and H/m respectively, reflect their distinct roles in electromagnetism.

Understanding the units of physical quantities is fundamental to solving problems in physics and engineering. It helps in checking the consistency of equations and understanding the physical meaning of calculated values. Always pay attention to the units given in a problem statement or derived in a calculation.

69. Answer: c

Explanation:

Understanding Levers and Their Classes

A lever is a simple machine consisting of a beam or rigid rod pivoted at a fixed hinge, or fulcrum. Levers are classified into three types based on the relative positions of the fulcrum, effort (the force applied), and load (the force resisted).

Types of Levers

The three classes of levers are distinguished by the order of the fulcrum, effort, and load along the lever arm:

- **First Class Lever:** The fulcrum is located somewhere between the effort and the load. Think of a seesaw or a crowbar. The effort and load are usually on opposite sides of the fulcrum.
- **Second Class Lever:** The load is located somewhere between the fulcrum and the effort. Think of a wheelbarrow or a nutcracker. The effort and load are on the same side of the fulcrum.
- **Third Class Lever:** The effort is located somewhere between the fulcrum and the load. Think of tweezers or ice tongs. The effort and load are on the same side of the fulcrum.

Analyzing the Given Options

Let's examine each option to determine which class of lever it represents:

- **Ice tongs:** With ice tongs, the pivot point where you hold the tongs together acts as the fulcrum. You apply the effort in the middle by squeezing, and the load (the ice) is at the end. Here, the effort is between the fulcrum and the load. This is an example of a **third-class lever**.
- **Wheelbarrow:** In a wheelbarrow, the wheel acts as the fulcrum. The load (the material in the bin) is placed between the wheel and where you lift the handles (the effort). Here, the load is between the fulcrum and the effort. This is an example of a **second-class lever**.
- **A pair of scissors:** A pair of scissors consists of two levers joined at a pivot. The pivot point where the two blades cross is the fulcrum. The effort is applied on the handles, and the load is where the blades cut (between the fulcrum and the handles). For one blade, the fulcrum is the pivot, the effort is applied on the handle end, and the load is at the cutting point. The fulcrum is between the effort and the load. This is an example of a **first-class lever**.

- **Nut cracker:** In a nutcracker, the hinge is the fulcrum. The nut (the load) is placed between the hinge and where you apply force on the handles (the effort). Here, the load is between the fulcrum and the effort. This is an example of a **second-class lever**.

Based on the analysis, a pair of scissors is an example of a first-class lever.

Revision Table: Lever Classification

Lever Class	Relative Positions	Common Examples
First Class	Fulcrum is between Effort and Load (F-E-L or L-F-E)	Seesaw, Crowbar, Scissors, Pliers
Second Class	Load is between Fulcrum and Effort (F-L-E)	Wheelbarrow, Nutcracker, Bottle Opener
Third Class	Effort is between Fulcrum and Load (F-E-L)	Tweezers, Ice Tongs, Fishing Rod, Human Arm

Additional Information on First Class Levers

First class levers are versatile because the fulcrum can be positioned anywhere between the effort and the load. This allows them to be used for different purposes:

- If the fulcrum is closer to the load, the lever acts as a force multiplier, making it easier to move a heavy load (e.g., using a crowbar to lift a heavy object). In this case, the mechanical advantage is greater than 1.
- If the fulcrum is closer to the effort, the lever is used to increase the distance or speed of the load's movement, although it requires more effort (e.g., using a shovel to throw soil a distance). In this case, the mechanical advantage is less than 1.
- If the fulcrum is exactly in the middle, the effort equals the load (assuming no friction), and the mechanical advantage is approximately 1. Scissors are often designed with the pivot near the center, providing a balance between force and distance of cut.

70. Answer: b

Explanation:

Solving the Equation for the Value of x

The problem asks us to determine the value of the variable x given a specific algebraic equation involving fractions. The equation provided is:

$$\frac{5x}{2} - \frac{1}{4} \left(6x - \frac{5}{3} \right) = \frac{7}{6}$$

We will solve this equation step-by-step to find the precise value of x.

Step 1: Distribute the Outer Fraction

The first step is to simplify the equation by distributing the fraction $-\frac{1}{4}$ to both terms inside the parentheses $(6x - \frac{5}{3})$.

$$\frac{5x}{2} - \left(\frac{1}{4} \times 6x \right) - \left(\frac{1}{4} \times -\frac{5}{3} \right) = \frac{7}{6}$$

Performing the multiplication gives:

$$\frac{5x}{2} - \frac{6x}{4} + \frac{5}{12} = \frac{7}{6}$$

Step 2: Simplify Fractional Coefficients

We can simplify the term $\frac{6x}{4}$. Both the numerator and the denominator are divisible by 2.

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

Substitute this simplified term back into the equation:

$$\frac{5x}{2} - \frac{3x}{2} + \frac{5}{12} = \frac{7}{6}$$

Step 3: Combine Like Terms (x-terms)

Now, combine the terms that contain x, which are $\frac{5x}{2}$ and $-\frac{3x}{2}$. Since they have a common denominator, we can subtract their numerators directly.

$$\left(\frac{5x}{2} - \frac{3x}{2}\right) + \frac{5}{12} = \frac{7}{6}$$

$$\frac{5x - 3x}{2} + \frac{5}{12} = \frac{7}{6}$$

$$\frac{2x}{2} + \frac{5}{12} = \frac{7}{6}$$

Simplify $\frac{2x}{2}$ to x :

$$x + \frac{5}{12} = \frac{7}{6}$$

Step 4: Isolate the Variable x

To find the value of x , we need to get it by itself on one side of the equation. We can do this by subtracting $\frac{5}{12}$ from both sides of the equation.

$$x = \frac{7}{6} - \frac{5}{12}$$

Step 5: Perform Fraction Subtraction

To subtract the fractions $\frac{7}{6}$ and $\frac{5}{12}$, we need a common denominator. The least common denominator for 6 and 12 is 12.

Convert $\frac{7}{6}$ to an equivalent fraction with a denominator of 12:

$$\frac{7}{6} = \frac{7 \times 2}{6 \times 2} = \frac{14}{12}$$

Now substitute this back into the equation for x :

$$x = \frac{14}{12} - \frac{5}{12}$$

Subtract the numerators:

$$x = \frac{14 - 5}{12}$$

$$x = \frac{9}{12}$$

Step 6: Simplify the Final Result

The final step is to simplify the fraction $\frac{9}{12}$. The greatest common divisor of 9 and 12 is 3. Divide both the numerator and the denominator by 3.

$$x = \frac{9 \div 3}{12 \div 3}$$

$$x = \frac{3}{4}$$

Final Value of x

The calculation shows that the value of x satisfying the given equation is $\frac{3}{4}$.

71. Answer: c

Explanation:

Finding the Highest Common Factor of 506 and 782

The highest common factor (HCF), also known as the greatest common divisor (GCD), of two numbers is the largest positive integer that divides both numbers without leaving a remainder.

To find the HCF of 506 and 782, we can use the Euclidean Algorithm, which is an efficient method for computing the HCF of two integers.

Understanding the Euclidean Algorithm

The Euclidean Algorithm works by repeatedly applying the division lemma:

Given two positive integers, say a and b , with $a > b$, we can write $a = bq + r$, where q is the quotient and r is the remainder, such that $0 \leq r < b$. The HCF of a and b is the same as the HCF of b and r . We continue this process until the remainder is 0. The last non-zero remainder is the HCF of the original two numbers.

Applying the Euclidean Algorithm to 506 and 782

Let's find the HCF of 782 and 506 using the steps of the Euclidean Algorithm:

Step 1: Divide 782 (the larger number) by 506 (the smaller number).

$$\begin{equation*} 782 = 506 \times 1 + 276 \end{equation*}$$

The remainder is 276.

Step 2: Now, take the previous divisor (506) and the remainder (276). Divide 506 by 276.

$$\begin{equation*} 506 = 276 \times 1 + 230 \end{equation*}$$

The remainder is 230.

Step 3: Take the previous divisor (276) and the remainder (230). Divide 276 by 230.

$$\begin{equation*} 276 = 230 \times 1 + 46 \end{equation*}$$

The remainder is 46.

Step 4: Take the previous divisor (230) and the remainder (46). Divide 230 by 46.

$$\begin{equation*} 230 = 46 \times 5 + 0 \end{equation*}$$

The remainder is 0.

Since the remainder is now 0, the process stops. The HCF is the last non-zero remainder, which is 46.

Result: HCF of 506 and 782

The highest common factor of 506 and 782 is 46.

Revision Table: Summarizing HCF Calculation Steps

Step	Division Equation ($a = bq + r$)	Divisor (b)	Remainder (r)
1	$782 = 506 \times 1 + 276$	506	276
2	$506 = 276 \times 1 + 230$	276	230
3	$276 = 230 \times 1 + 46$	230	46
4	$230 = 46 \times 5 + 0$	46	0

The last non-zero remainder is 46, confirming it as the HCF.

Additional Information on Finding HCF

Another common method to find the HCF is by using prime factorization.

Prime Factorization Method for HCF

To find the HCF using prime factorization:

- Find the prime factors of each number.
- List all the common prime factors.
- Multiply the common prime factors, using the lowest power that each factor appears in the original numbers' factorizations.

For instance, to find the HCF of 506 and 782 using this method, you would first find their prime factorizations:

- Prime factorization of 506: $2 \times 11 \times 23$
- Prime factorization of 782: $2 \times 17 \times 23$

The common prime factors are 2 and 23. The lowest power of 2 is 2^1 and the lowest power of 23 is 23^1 .

$$\text{HCF} = 2 \times 23 = 46.$$

Both the Euclidean Algorithm and the prime factorization method lead to the same HCF. The Euclidean Algorithm is often more efficient for larger numbers as it avoids the need to find prime factors.

72. Answer: d

Explanation:

5/14

Understanding the Executive's Financial Situation This problem involves calculating an executive's savings fraction after changes in salary and spending.

73. Answer: c

Explanation:

Understanding the Work and Time Problem

This question involves calculating the fraction of a task left after two individuals, A and B, work together for a certain period. We are given the time each person takes to complete the task alone. To solve this, we first determine their individual work rates, then their combined work rate, the amount of work done together, and finally the fraction of the task remaining.

Calculating Individual Work Rates

The work rate of a person is the fraction of the task they can complete in one day. It is the reciprocal of the number of days they take to complete the entire task.

- A completes the task in 35 days. So, A's work rate is $\frac{1}{35}$ of the task per day.
- B completes the task in 14 days. So, B's work rate is $\frac{1}{14}$ of the task per day.

Calculating Combined Work Rate

When A and B work together, their work rates add up to find their combined work rate for one day.

Combined work rate per day = A's work rate + B's work rate

$$\text{Combined work rate per day} = \frac{1}{35} + \frac{1}{14}$$

To add these fractions, we need a common denominator. The least common multiple (LCM) of 35 and 14 is 70.

- $\frac{1}{35} = \frac{1 \times 2}{35 \times 2} = \frac{2}{70}$
- $\frac{1}{14} = \frac{1 \times 5}{14 \times 5} = \frac{5}{70}$

$$\text{Combined work rate per day} = \frac{2}{70} + \frac{5}{70} = \frac{2+5}{70} = \frac{7}{70}$$

Simplifying the fraction, we get:

$$\text{Combined work rate per day} = \frac{7}{70} = \frac{1}{10} \text{ of the task per day.}$$

Calculating Work Done in 5 Days

A and B work together for 5 days. The total work done in 5 days is the combined work rate per day multiplied by the number of days they worked together.

Work done in 5 days = Combined work rate per day $\times$ Number of days worked

$$\text{Work done in 5 days} = \frac{1}{10} \times 5$$

$$\text{Work done in 5 days} = \frac{5}{10} = \frac{1}{2} \text{ of the task.}$$

Calculating Fraction of Task Left

The total task is considered as 1 whole. To find the fraction of the task left, we subtract the work done from the total task.

Fraction of task left = Total task - Work done

$$\text{Fraction of task left} = 1 - \frac{1}{2}$$

$$\text{Fraction of task left} = \frac{2}{2} - \frac{1}{2} = \frac{2-1}{2} = \frac{1}{2}$$

Conclusion

The fraction of the task left after A and B work together for 5 days is $\frac{1}{2}$, which corresponds to "Half".

Parameter	Value
A's time to complete task	35 days
B's time to complete task	14 days
A's daily work rate	$\frac{1}{35}$
B's daily work rate	$\frac{1}{14}$
Combined daily work rate	$\frac{1}{10}$
Days worked together	5 days
Work done in 5 days	$\frac{1}{2}$
Fraction of task left	$\frac{1}{2}$

Revision Table: Work and Time Concepts

Concept	Description	Formula/Relation
Work Rate	The amount of work done per unit of time (e.g., per day).	Work Rate = $\frac{1}{\text{Time taken}}$
Total Work	Usually considered as 1 whole unit.	Work Done + Work Left = 1
Combined Work Rate	Sum of individual work rates when people work together.	$R_{combined} = R_1 + R_2 + \dots$
Work Done	Combined work rate multiplied by the time worked together.	Work Done = $R_{combined} \times \text{Time worked}$

Additional Information: Solving Work and Time Problems

Work and time problems are common in competitive exams. The key is to convert the time taken to complete a task into a work rate (fraction of work per unit time). Here are some additional points:

- If a person completes a task in 'N' days, their work rate is $\frac{1}{N}$ per day.

- If multiple people work together, their individual daily work rates are added to find the combined daily work rate.
- The total work done is the combined daily work rate multiplied by the number of days they work together.
- If someone leaves or joins mid-task, calculate the work done before and after the change separately.
- Efficiency is inversely proportional to the time taken. More efficient means less time taken.

74. Answer: b

Explanation:

The '#' operation returns the average of the two numbers: $a \# b = \frac{a+b}{2}$.

Verify with the given examples:

- $8 \# 12 = \frac{8+12}{2} = 10$
- $5 \# 9 = \frac{5+9}{2} = 7$
- $6 \# 10 = \frac{6+10}{2} = 8$

Applying the rule to $14 \# 4$:

$$14 \# 4 = \frac{14+4}{2} = \frac{18}{2} = 9$$

Hence the value is 9.

75. Answer: a

Explanation:

Understanding the Series Completion Question

This question asks us to identify the pattern in a given series of letters and numbers and find the set of characters that correctly fills the gaps to complete the sequence. The series is presented as:

1_2x_bb_yy_cc_6zzz

There are five gaps in the series, indicated by underscores.

Analyzing the Options for the Series Gaps

We are provided with four options, each containing five characters to fill the five gaps:

1. a345c
2. 3a45b
3. 3a4b5
4. 3a45c

We need to test each option by placing its characters sequentially into the gaps to see which one reveals a consistent pattern in the complete series.

Step-by-Step Solution: Identifying the Pattern

Let's take the first option, a345c, and insert its characters into the gaps in the series:

Original series: 1_2x_bb_yy_cc_6zzz

Inserting a345c: 1a2x3bb4yy5ccc6zzz

The complete series becomes: 1a2x3bb4yy5ccc6zzz

Breaking Down the Series into Segments

Let's look at how the original parts and the inserted characters combine. The original series structure suggests segments separated by the gaps. When filled with a345c, we can observe segments by pairing each original part (except the last one) with the character immediately following it in the combined sequence:

- First original part 1 + first inserted character a = 1a
- Second original part 2x + second inserted character 3 = 2x3
- Third original part bb + third inserted character 4 = bb4
- Fourth original part yy + fourth inserted character 5 = yy5
- Fifth original part cc + fifth inserted character c = cc
- The last original part is 6zzz.

So, the series can be viewed as a sequence of these segments: 1a, 2x3, bb4, yy5, cc, 6zzz.

Identifying the Number Sequence Pattern

Let's look at the numbers present across these segments:

- In 1a, we have the number 1.
- In 2x3, we have the numbers 2 and 3.
- In bb4, we have the number 4.
- In yy5, we have the number 5.
- In ccc, there is no number.
- In 6zzz, we have the number 6.

Arranging the numbers in the order they appear in the series, we get 1, 2, 3, 4, 5, 6. This is a clear sequence of consecutive integers. This strong pattern supports Option 1.

Identifying the Letter Repetition Pattern

Now let's look at the letter parts of the segments:

- 1a contains letter a. Repetition count: 1.
- 2x3 contains letter x. Repetition count: 1.
- bb4 contains letter bb. Repetition count: 2.
- yy5 contains letter yy. Repetition count: 2.
- ccc contains letter ccc. Repetition count: 3. (This is formed by the original cc and the inserted c).
- 6zzz contains letter zzz. Repetition count: 3.

The sequence of letter repetition counts is 1, 1, 2, 2, 3, 3. This is another consistent pattern where repetition counts 1, 2, and 3 each appear twice.

Confirming the Pattern with Option 1

Option 1 (a345c) successfully creates a series where:

- Numbers 1 through 6 appear sequentially.
- Letters appear in groups with repetition counts 1, 1, 2, 2, 3, 3.

Let's quickly check other options. If we used Option 2 (3a45b), the numbers would be inserted as 3, a, 4, 5, b. The series would be 132xabb4yy5ccb6zzz. The number sequence would be 1, 3, 2, 4, 5, 6, which is not sequential. Similarly, other options will also fail to produce the clear sequential number pattern 1, 2, 3, 4, 5, 6.

Thus, Option 1 is the only set of characters that fits the identified patterns.

Revision Table: Series Pattern Analysis

Segment	Numbers Present	Letter Group	Letter Repetition Count
1a	1	a	1
2x3	2, 3	x	1
bb4	4	bb	2
yy5	5	yy	2
ccc	None	ccc	3
6zzz	6	zzz	3

The table shows the number sequence 1, 2, 3, 4, 5, 6 appearing across the segments and the letter repetition counts 1, 1, 2, 2, 3, 3.

Additional Information on Series Patterns

Series completion questions often involve various types of patterns. Recognizing these patterns is key to solving such problems. Some common patterns include:

- **Arithmetic Progressions:** Numbers increasing or decreasing by a constant difference.
- **Geometric Progressions:** Numbers increasing or decreasing by a constant ratio.
- **Fibonacci Series:** Each number is the sum of the two preceding ones.
- **Alphabetical Patterns:** Letters following sequence (A, B, C...), skipping letters (A, C, E...), or moving positions.
- **Combined Patterns:** Involvement of both numbers and letters, sometimes with alternating patterns or patterns based on the position or type of character.
- **Repetition Patterns:** Characters or groups of characters repeating in a predictable way.
- **Positional Patterns:** The position of elements within the series determines the pattern.

Analyzing the given elements carefully and testing potential rules based on number sequences, alphabetical order, repetitions, or positions is a good strategy for solving series

completion questions.

76. Answer: d

Explanation:

Calculating Specific Latent Heat of Vaporisation

The question asks us to find the specific latent heat of vaporisation of nitrogen. We are given the amount of heat released when a certain mass of nitrogen condenses at its boiling point. The specific latent heat of vaporisation (or condensation) is defined as the amount of heat energy absorbed or released per unit mass during a phase change (from liquid to gas or gas to liquid) at a constant temperature (the boiling point).

Understanding Latent Heat of Vaporisation/Condensation

When a substance changes phase from liquid to gas (vaporisation) or from gas to liquid (condensation) at a constant temperature (its boiling point), it either absorbs or releases energy. This energy is called latent heat. Specifically:

- Vaporisation: Energy is absorbed by the substance to overcome the intermolecular forces holding the liquid together and change into a gas. This is called the latent heat of vaporisation.
- Condensation: Energy is released by the substance as the gas molecules come closer together to form a liquid. This is also related to the latent heat of vaporisation (the same amount of energy per unit mass is involved).

The relationship between the heat energy (Q), the mass of the substance (m), and the specific latent heat (L) is given by the formula:

$$Q = mL$$

Where:

- Q is the heat energy absorbed or released (in Joules, J).
- m is the mass of the substance (in grams, g, or kilograms, kg).
- L is the specific latent heat of vaporisation or condensation (in J/g or J/kg).

Applying the Formula to Nitrogen Condensation

We are given the following information:

- Mass of nitrogen, $m = 1.25 \text{ g}$
- Heat released during condensation, $Q = 250 \text{ J}$
- Boiling/condensation point of nitrogen = -196°C (This temperature information confirms the phase change is happening at the boiling point, but is not needed for the calculation itself, only the heat and mass are needed).

We need to find the specific latent heat of vaporisation, L . We can rearrange the formula $Q = mL$ to solve for L :

$$L = \frac{Q}{m}$$

Step-by-Step Calculation

Substitute the given values into the rearranged formula:

$$L = \frac{250 \text{ J}}{1.25 \text{ g}}$$

Now, perform the division:

$$L = \frac{250}{1.25} \text{ J/g}$$

$$L = 200 \text{ J/g}$$

So, the specific latent heat of vaporisation of nitrogen is 200 Jg^{-1} .

Verification with Options

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

- Option 1: 469 Jg^{-1}
- Option 2: 500 Jg^{-1}
- Option 3: 312.5 Jg^{-1}
- Option 4: 200 Jg^{-1}

Our calculated value of 200 Jg^{-1} matches Option 4.

Quantity	Symbol	Value Given
Mass of Nitrogen	m	1.25 g
Heat Released (Condensation)	Q	250 J
Specific Latent Heat	L	?

Calculation summary:

$$L = \frac{Q}{m} = \frac{250 \text{ J}}{1.25 \text{ g}} = 200 \text{ J/g}$$

Revision Table: Specific Latent Heat Concepts

Concept	Definition	Formula	Units
Latent Heat	Energy absorbed or released during a phase change at constant temperature.	Q	Joules (J)
Specific Latent Heat	Latent heat per unit mass of a substance.	$L = \frac{Q}{m}$	J/g or J/kg
Latent Heat of Fusion	Energy per unit mass for melting (solid to liquid) or freezing (liquid to solid).	L_f	J/g or J/kg
Latent Heat of Vaporisation	Energy per unit mass for vaporisation (liquid to gas) or condensation (gas to liquid).	L_v	J/g or J/kg

Additional Information: Phase Changes and Energy

Phase changes, like melting, freezing, boiling, and condensation, involve energy transfer without a change in temperature. This is because the energy is used to change the potential energy between the molecules rather than increasing their kinetic energy (which corresponds to temperature). For example, during boiling, the energy supplied is used to break the bonds holding molecules in the liquid phase, allowing them to escape as a gas, instead of increasing the speed of the molecules.

- **Vaporisation:** Liquid absorbs latent heat to become gas. Molecules gain enough energy to overcome intermolecular forces.
- **Condensation:** Gas releases latent heat to become liquid. Molecules lose energy and intermolecular forces pull them together.
- **Fusion (Melting):** Solid absorbs latent heat to become liquid. Energy breaks down the rigid structure of the solid.
- **Freezing:** Liquid releases latent heat to become solid. Molecules arrange into a more ordered structure.

The specific latent heat is a property of the substance itself and depends on the type of phase change (fusion or vaporisation) and the specific substance (like water, nitrogen, etc.).

77. Answer: b

Explanation:

Understanding Bump Maps in Graphics

In the world of computer graphics and 3D rendering, maps are used to add detail and texture to surfaces without needing to create complex geometric shapes. A specific type of map is designed to give the appearance of raised or indented features on a flat surface.

The question asks about a map that converts color intensity or grayscale information into heights, creating the visual effect that features are lifted above the surface, much like embossed letters. Let's look at the options to identify which type of map fits this description.

Analyzing the Options

- **Shade map:** This term isn't standard in 3D texturing for this specific effect. Shading usually refers to how light interacts with a surface, determining how bright or dark it appears, based on factors like surface angle, light direction, and material properties.
- **Bump map:** A **bump map** is exactly what the question describes. It's a grayscale image where different shades of gray represent different "heights" or displacements from the original surface. Pure black typically represents the lowest point, pure white represents the highest point, and shades of gray represent points in between. During

rendering, this height information isn't used to actually change the geometry, but rather to alter the direction of the surface normal vectors. By perturbing the normals, the lighting calculations are affected, creating the illusion of bumps, wrinkles, or raised text on the surface, giving the appearance that features are raised above the surface.

- **Tone map:** Tone mapping is a technique used in image processing and rendering, especially with High Dynamic Range (HDR) images. It compresses the wide range of luminance values in an HDR image into a range that standard displays can show, preserving perceived contrast and detail. It doesn't relate to adding surface detail based on height information.
- **Tint map:** A tint is a variation of a color produced by adding white to it. A "tint map" isn't a standard term for adding surface detail based on height. Texture maps commonly use color information (like RGB) to define the color or pattern of a surface, but not its apparent height in the way described.

Identifying the Correct Map Type

Based on the analysis, the map that uses grayscale or color intensity to simulate raised features by affecting how light interacts with the surface normals is the **bump map**. This map type is widely used in computer graphics to add surface texture and detail efficiently without increasing the geometric complexity of the 3D model. It effectively converts grayscale information to apparent heights to give the appearance that features are raised above the surface.

Revision Table: Map Types in Graphics

Map Type	Primary Function	How it Works
Bump Map	Simulates surface texture (bumps, indents)	Uses grayscale to perturb surface normals, affecting lighting to create illusion of height.
Normal Map	Simulates surface texture (bumps, indents) - more advanced	Uses RGB color channels to directly store perturbed normal vector directions.
Displacement Map	Modifies actual surface geometry	Uses grayscale to physically move vertices on the surface.
Texture Map (Color/Albedo)	Defines surface color or pattern	Uses RGB color information to color the surface.

Additional Information on Bump Mapping and Related Concepts

Bump maps are a fundamental concept in 3D graphics texturing. While they don't change the actual shape of the object, they are very effective at creating the illusion of detailed surfaces, like the grain of wood, wrinkles on skin, or the texture of fabric. Because they only affect the rendering process (specifically, how light reflects off the surface), they are computationally less expensive than techniques that modify the geometry itself, like displacement mapping.

A related technique is **Normal Mapping**. Like bump mapping, normal mapping also simulates surface details without changing geometry by altering surface normals. However, instead of storing height information in grayscale, a normal map stores the actual direction of the surface normal vector at each point directly in the RGB channels of the image. Normal maps can represent more complex surface variations than simple bump maps because they encode full directional information rather than just a single height value. Both techniques give the appearance that features are raised above the surface or indented, converting some form of image information to a visual height effect through lighting.

Another important map type is the **Displacement Map**. Unlike bump or normal maps, a displacement map actually uses the grayscale information (or color information) to

physically move the vertices of the 3D model. This changes the actual geometry of the object, resulting in true bumps and indents that are visible from silhouettes. Displacement mapping provides the most realistic surface detail but is also the most computationally expensive.

In summary, when a question refers to converting color intensity or grayscale information to apparent heights to give the appearance that features are raised above the surface using lighting effects rather than changing geometry, the term you are looking for is a **bump map**.

78. Answer: c

Explanation:

Calculating the Change in Kinetic Energy for a Truck

This solution explains how to calculate the change in kinetic energy of a truck as it accelerates, using the provided details about its mass and velocities.

Understanding Kinetic Energy

Kinetic energy is the energy an object possesses because it is in motion. The amount of kinetic energy depends on the object's **mass** and its **velocity** (speed). The formula used to calculate kinetic energy (KE) is:

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

where:

- m is the mass of the object (in kilograms, kg)
- v is the velocity of the object (in meters per second, m/s)

The change in kinetic energy is the difference between the final kinetic energy (KE_f) and the initial kinetic energy (KE_i):

$$\Delta KE = KE_f - KE_i$$

Analyzing the Given Information

We are given the following information about the truck:

- **Mass (m):** 5000 kg
- **Initial velocity (v_i):** 25 m/s
- **Final velocity (v_f):** 35 m/s
- We need to find the change in kinetic energy in Megajoules (MJ).

Step-by-Step Calculation

1. Calculate the Initial Kinetic Energy (KE_i)

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

$$KE_i = \frac{1}{2} \times 5000 \text{ kg} \times (25 \text{ m/s})^2$$

First, calculate the square of the initial velocity:

$$(25 \text{ m/s})^2 = 625 \text{ m}^2/\text{s}^2$$

Now, calculate the initial kinetic energy:

$$KE_i = \frac{1}{2} \times 5000 \text{ kg} \times 625 \text{ m}^2/\text{s}^2$$

$$KE_i = 2500 \text{ kg} \times 625 \text{ m}^2/\text{s}^2$$

$$KE_i = 1,562,500 \text{ Joules (J)}$$

2. Calculate the Final Kinetic Energy (KE_f)

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

$$KE_f = \frac{1}{2} \times 5000 \text{ kg} \times (35 \text{ m/s})^2$$

First, calculate the square of the final velocity:

$$(35 \text{ m/s})^2 = 1225 \text{ m}^2/\text{s}^2$$

Now, calculate the final kinetic energy:

$$KE_f = \frac{1}{2} \times 5000 \text{ kg} \times 1225 \text{ m}^2/\text{s}^2$$

$$KE_f = 2500 \text{ kg} \times 1225 \text{ m}^2/\text{s}^2$$

$$KE_f = 3,062,500 \text{ Joules (J)}$$

3. Calculate the Change in Kinetic Energy (ΔKE)

Subtract the initial kinetic energy from the final kinetic energy:

$$\Delta KE = KE_f - KE_i$$

$$\Delta KE = 3,062,500 \text{ J} - 1,562,500 \text{ J}$$

$$\Delta KE = 1,500,000 \text{ J}$$

4. Convert the Change in Kinetic Energy to Megajoules (MJ)

We know that 1 Megajoule (MJ) is equal to 1,000,000 Joules (J). To convert Joules to Megajoules, we divide by 1,000,000:

$$\Delta KE(\text{in MJ}) = \frac{\Delta KE(\text{in J})}{1,000,000 \text{ J/MJ}}$$

$$\Delta KE = \frac{1,500,000 \text{ J}}{1,000,000 \text{ J/MJ}}$$

$$\Delta KE = 1.5 \text{ MJ}$$

Final Answer

The change in the truck's kinetic energy as it accelerates from 25 m/s to 35 m/s is **1.5 MJ**.

79. Answer: a

Explanation:

Understanding Potential Difference Calculation

This problem asks us to find the potential difference across a resistor when we know its resistance and the current flowing through it. This is a classic application of Ohm's Law, a fundamental principle in electrical circuits.

Applying Ohm's Law to Find Voltage

Ohm's Law describes the relationship between voltage (potential difference), current, and resistance in an electrical circuit. It is stated as:

$$V = I \times R$$

Where:

- V is the potential difference (voltage) measured in Volts (V).
- I is the current measured in Amperes (A).
- R is the resistance measured in Ohms (Ω).

Given Values and Units Conversion

We are given the following values:

- Resistance (R) = 2.5 k Ω
- Current (I) = 2 mA

Before we can use Ohm's Law, we need to ensure all units are in their base SI forms (Ohms for resistance, Amperes for current). The given values are in kilohms (k Ω) and milliamperes (mA), so we need to convert them:

- 1 k Ω = $10^3 \Omega$
- 1 mA = 10^{-3} A

Converting the given values:

- Resistance $R = 2.5 \text{ k}\Omega = 2.5 \times 10^3 \Omega$
- Current $I = 2 \text{ mA} = 2 \times 10^{-3} \text{ A}$

Quantity	Given Value	SI Unit	Converted Value (SI Units)
Resistance (R)	2.5 k Ω	Ω	$2.5 \times 10^3 \Omega$
Current (I)	2 mA	A	$2 \times 10^{-3} \text{ A}$

Step-by-Step Potential Difference Calculation

Now that we have the resistance and current in the correct units, we can calculate the potential difference (V) using Ohm's Law:

$$V = I \times R$$

Substitute the converted values into the formula:

$$V = (2 \times 10^{-3} \text{ A}) \times (2.5 \times 10^3 \Omega)$$

Multiply the numerical parts and the powers of ten separately:

$$V = (2 \times 2.5) \times (10^{-3} \times 10^3) \text{ V}$$

$$V = 5 \times 10^{(-3+3)} \text{ V}$$

$$V = 5 \times 10^0 \text{ V}$$

Since $10^0 = 1$:

$$V = 5 \times 1 \text{ V}$$

$$V = 5 \text{ V}$$

The potential difference across the $2.5 \text{ k}\Omega$ resistance is 5 V .

Conclusion: Potential Difference Result

Based on the calculation using Ohm's Law, the potential difference across the resistance is 5 Volts . This matches one of the provided options for the potential difference.

Revision Table: Key Formulas and Units

Concept	Formula	Units
Ohm's Law	$V = I \times R$	Volts (V), Amperes (A), Ohms (Ω)
Potential Difference (Voltage)	V	Volts (V)
Current	I	Amperes (A)
Resistance	R	Ohms (Ω)

Additional Information: Understanding Ohm's Law and Units

Ohm's Law is a linear relationship between voltage and current for many materials, especially conductors like metals, under constant temperature. It essentially states that

the current through a conductor between two points is directly proportional to the voltage across the two points and inversely proportional to the resistance between them.

Understanding unit prefixes like 'kilo' (k) and 'milli' (m) is crucial in physics and engineering calculations. They represent powers of 10:

- kilo (k) means 10^3 (one thousand)
- milli (m) means 10^{-3} (one thousandth)

Always convert values with prefixes to the base unit before performing calculations using formulas like Ohm's Law to avoid errors.

80. Answer: b

Explanation:

Understanding the Living Planet Report

The question asks about the organisation responsible for publishing the Living Planet Report. This report is a major publication in the field of environmental conservation and sustainability.

The Living Planet Report is released periodically, typically every two years. Its primary purpose is to track trends in global biodiversity and the health of the planet. It uses various indicators, most notably the Living Planet Index (LPI), which measures the state of the world's biological diversity based on trends in populations of vertebrate species.

The report also examines humanity's impact on the natural world, often quantified through ecological footprint analysis, and discusses solutions and strategies for conservation and sustainable living.

Identifying the Publisher: World Wide Fund for Nature (WWF)

The Living Planet Report is widely recognised as the flagship publication of a specific international non-governmental organisation dedicated to wilderness preservation, and the reduction of human impact on the environment. This organisation is the World Wide Fund for Nature (WWF).

WWF was founded in 1961 and has grown to become one of the world's largest conservation organisations. Its work spans various areas including climate, food, forests, freshwater, oceans, and wildlife. The Living Planet Report serves as a key tool for WWF to communicate the state of the planet's ecosystems and biodiversity to policymakers, scientists, the public, and other stakeholders, highlighting the urgency of conservation efforts.

Analysing the Options

Let's look at the other options provided:

- **The Nature Conservancy:** This is another major global environmental non-profit organisation. While involved in conservation science and publishing reports, their flagship publication is not the Living Planet Report. They focus heavily on land and water protection projects.
- **Conservation International:** This organisation works on biodiversity conservation and has initiatives like ecosystem profiling and conservation finance. They publish various reports and scientific papers but are not primarily known for the Living Planet Report.
- **Wildlife Conservation Society (WCS):** WCS is an international non-profit organisation based at the Bronx Zoo. They focus on conservation in specific landscapes and seascapes, running field projects worldwide. They produce scientific publications and reports related to their field work, but not the Living Planet Report.

Based on the established knowledge about major environmental reports and the organisations that publish them, the Living Planet Report is definitively a publication of the World Wide Fund for Nature (WWF).

Revision Table: Living Planet Report Key Facts

Feature	Description
Publication Name	Living Planet Report
Publisher	World Wide Fund for Nature (WWF)
Frequency	Typically Every Two Years
Key Metric	Living Planet Index (LPI)
Focus	Global biodiversity trends, ecosystem health, human impact, conservation solutions

Additional Information on Conservation Reports

Many organisations produce reports related to the environment and conservation. Here are a few examples:

- **IPCC Reports:** The Intergovernmental Panel on Climate Change (IPCC) produces comprehensive assessment reports on climate change, its impacts, and future risks.
- **State of the World's Birds:** BirdLife International publishes reports on the status of bird populations globally.
- **IUCN Red List:** The International Union for Conservation of Nature (IUCN) maintains the Red List of Threatened Species, a critical resource for assessing conservation status.

These reports, like the Living Planet Report, are vital for informing conservation strategies, policy decisions, and public awareness regarding the health of our planet and its biodiversity.

81. Answer: d

Explanation:

Decoding the Letter-to-Digit Code for ARID

This problem involves a simple substitution code where specific letters are consistently represented by specific digits. We are given the codes for two words, "DIRTY" and "FOAM", and need to find the code for "ARID" using the same pattern.

Establishing the Letter-to-Digit Mapping

First, let's determine the code for each letter based on the provided examples:

- From "DIRTY = 24759", we can infer the following mappings:
 - D → 2
 - I → 4
 - R → 7
 - T → 5
 - Y → 9
- From "FOAM = 1863", we can infer these mappings:
 - F → 1
 - O → 8
 - A → 6
 - M → 3

We can consolidate these mappings into a single reference. Let's represent this mapping clearly:

Letter	Code
D	2
I	4
R	7
T	5
Y	9
F	1
O	8
A	6
M	3

Coding the Target Word ARID

Now, we use the established letter-to-digit mapping to find the code for the word "ARID". We look up the code for each letter in the word:

1. **A:** According to our mapping, 'A' is coded as 6.
2. **R:** According to our mapping, 'R' is coded as 7.
3. **I:** According to our mapping, 'I' is coded as 4.
4. **D:** According to our mapping, 'D' is coded as 2.

By combining the digits for each letter in order, we get the code for "ARID".

Therefore, ARID = 6742.

Conclusion

Based on the coding pattern derived from "DIRTY" and "FOAM", the code for "ARID" is 6742. This corresponds to the fourth option provided.

82. Answer: d

Explanation:

$1101111_2 = 111$ and $1100101_2 = 101$; their sum is $212 = 11010100_2$. So the answer is **11010100**.

83. Answer: b

Explanation:

Understanding the Question: IPL 2018 Winner

The question asks about the team that emerged victorious in the Indian Premier League (IPL) tournament held in the year 2018.

The Indian Premier League is a professional Twenty20 cricket league contested by ten teams based out of seven Indian cities and three Indian states. The 2018 season was the

eleventh season of the IPL, organized by the BCCI.

Analyzing the Options

Let's look at the teams provided in the options:

1. Sunrisers Hyderabad
2. Chennai Super Kings
3. Royal Challengers Bangalore
4. Kolkata Knight Riders

These are all prominent teams that have participated in the IPL over the years, and some have won the title multiple times. We need to determine which one won the 2018 edition.

Determining the IPL 2018 Champion Team

The 2018 IPL season's final match was played between Chennai Super Kings and Sunrisers Hyderabad on May 27, 2018, at the Wankhede Stadium in Mumbai.

In the final, Chennai Super Kings batted second and chased down the target set by Sunrisers Hyderabad. This victory marked a significant comeback for Chennai Super Kings, who were returning to the tournament after a two-year suspension.

Detailed Result of the IPL 2018 Final

Here is a summary of the final match:

- **Teams:** Chennai Super Kings vs. Sunrisers Hyderabad
- **Venue:** Wankhede Stadium, Mumbai
- **Date:** May 27, 2018
- **Sunrisers Hyderabad Score:** 178/6 in 20 overs
- **Chennai Super Kings Score:** 181/2 in 18.3 overs
- **Result:** Chennai Super Kings won by 8 wickets with 9 balls remaining.

Shane Watson of Chennai Super Kings was the star performer in the final, scoring an unbeaten 117 runs off 57 balls.

Conclusion on the Winner

Based on the outcome of the final match, the team that won the Indian Premier League in 2018 was Chennai Super Kings.

Revision Table: IPL Winners

Year	Winner	Runner-up
2018	Chennai Super Kings	Sunrisers Hyderabad
2017	Mumbai Indians	Rising Pune Supergiant
2019	Mumbai Indians	Chennai Super Kings

Additional Information on IPL 2018

The 2018 IPL season saw intense competition among the eight participating teams. The season featured many memorable performances from players across all teams. Key highlights included:

- Kane Williamson (Sunrisers Hyderabad) was the Orange Cap holder for scoring the most runs (735).
- Andrew Tye (Kings XI Punjab) was the Purple Cap holder for taking the most wickets (24).
- Rishabh Pant (Delhi Daredevils) won the Emerging Player of the Season award.

The tournament structure involved a round-robin league stage followed by playoffs consisting of Qualifiers and Eliminator, culminating in the final match.

84. Answer: a

Explanation:

Understanding Heat Capacity Calculations

The question asks us to determine the heat capacity of a pan given its mass, the amount of heat energy it receives, and the resulting temperature rise. This involves using the fundamental concept of heat capacity in physics.

What is Heat Capacity?

Heat capacity is a physical property that tells us how much heat energy is required to raise the temperature of a substance or an object by one degree Celsius (or one Kelvin). A higher heat capacity means more energy is needed to change the temperature.

The total heat capacity (also sometimes called thermal capacity) of an object depends on its mass and the specific heat capacity of the material it is made from. The formula relating heat, heat capacity, and temperature change is:

$$Q = C\Delta T$$

Where:

- Q is the amount of heat energy transferred (in Joules, J).
- C is the total heat capacity of the object (in Joules per Kelvin, J/K, or Joules per degree Celsius, J/°C).
- ΔT is the change in temperature (in Kelvin, K, or degree Celsius, °C). Note that a change of 1°C is equal to a change of 1K, so the units J/K and J/°C are equivalent for heat capacity.

From this formula, we can rearrange it to find the heat capacity C :

$$C = \frac{Q}{\Delta T}$$

Calculating the Heat Capacity of the Pan

Let's identify the values given in the question:

- Mass of the pan = 200 g
- Heat received by the pan (Q) = 2000 J
- Rise in temperature (ΔT) = 8° C

We need to calculate the heat capacity C of the pan. We can plug the given values for Q and ΔT into the rearranged formula:

$$C = \frac{Q}{\Delta T}$$

$$C = \frac{2000 \text{ J}}{8^\circ \text{C}}$$

Performing the division:

$$C = 250 \text{ J/}^\circ\text{C}$$

Since a change of 1°C is equivalent to a change of 1K , the unit $\text{J/}^\circ\text{C}$ is the same as J/K . Therefore, the heat capacity of the pan is 250 J/K .

It's important to note that the mass of the pan (200 g) is given, but it is not needed to find the *total heat capacity* of the object (the pan). The mass would be required if we were asked to find the *specific heat capacity* of the material the pan is made from (which has units like J/kg/K or $\text{J/g/}^\circ\text{C}$). The question specifically asks for the heat capacity of the pan, which refers to the total heat capacity.

Quantity	Symbol	Value	Units
Heat Received	Q	2000	J
Temperature Change	ΔT	8	$^\circ\text{C}$ (or K)
Heat Capacity (to find)	C	?	J/K

Using the formula $C = \frac{Q}{\Delta T}$:

$$C = \frac{2000 \text{ J}}{8 \text{ K}} = 250 \text{ J/K}$$

Thus, the heat capacity of the pan is 250 JK^{-1} .

Revision Table: Key Concepts

Term	Definition	Formula	Units
Heat Capacity (Total)	Amount of heat needed to raise an object's temp by 1 degree.	$C = Q/\Delta T$	J/K or $\text{J/}^\circ\text{C}$
Specific Heat Capacity	Amount of heat needed to raise 1 unit mass of a substance by 1 degree.	$c = Q/(m\Delta T)$ or $C = mc$	J/kg/K or $\text{J/g/}^\circ\text{C}$

Additional Information: Specific Heat Capacity vs. Heat Capacity

It's important to distinguish between heat capacity (C) and specific heat capacity (c).

- **Total Heat Capacity (C):** This property belongs to a specific object. It depends on both the material and the mass of the object. For example, a large iron pan has a different total heat capacity than a small iron pan, even though they are made of the same material. The unit is typically J/K.
- **Specific Heat Capacity (c):** This is a property of the material itself, independent of the amount of material. It tells us how much heat is needed to raise the temperature of 1 kilogram (or 1 gram) of a substance by 1 Kelvin (or 1°C). For example, all pure iron has the same specific heat capacity. The unit is typically J/kg/K or J/g/°C.

The relationship between total heat capacity and specific heat capacity is $C = m \times c$, where m is the mass of the object.

In this problem, we calculated the total heat capacity C of the pan as an object. If we were asked for the specific heat capacity of the material of the pan, we would use the calculated total heat capacity C and the mass m :

$$c = \frac{C}{m}$$

$$c = \frac{250 \text{ J/K}}{0.2 \text{ kg}} \text{ (converting mass to kg: } 200 \text{ g} = 0.2 \text{ kg)}$$

$$c = 1250 \text{ J/kg/K}$$

However, the question specifically asked for the heat capacity of the pan, which is C , not c . So our initial calculation of 250 J/K is the correct answer based on the question asked.

85. Answer: c

Explanation:

Finding the Ratio A:B:C from the Equation $3A = 6B = 7C$

This problem requires us to find the ratio of three variables A, B, and C, given an equation relating them. The key is to express each variable in terms of a common value and then find the simplest form of their ratio.

Steps to Solve for the Ratio A:B:C

We are given the equation:

$$3A = 6B = 7C$$

To find the ratio $A : B : C$, we can set each part of the equality equal to a constant value, say k . This allows us to express A , B , and C individually in terms of this constant k .

- From $3A = k$, we get $A = \frac{k}{3}$.
- From $6B = k$, we get $B = \frac{k}{6}$.
- From $7C = k$, we get $C = \frac{k}{7}$.

Now, we can write the ratio $A : B : C$ as:

$$A : B : C = \frac{k}{3} : \frac{k}{6} : \frac{k}{7}$$

Since k is a common factor in all terms of the ratio (assuming $k \neq 0$), we can cancel k from each term:

$$A : B : C = \frac{1}{3} : \frac{1}{6} : \frac{1}{7}$$

To express this ratio in its simplest form, where the terms are integers, we need to multiply each fraction by the Least Common Multiple (LCM) of the denominators (3, 6, and 7).

Let's find the LCM of 3, 6, and 7:

- Prime factorization of 3 is 3.
- Prime factorization of 6 is 2×3 .
- Prime factorization of 7 is 7.

$$\text{LCM}(3, 6, 7) = 2 \times 3 \times 7 = 42.$$

Now, multiply each term in the ratio $\frac{1}{3} : \frac{1}{6} : \frac{1}{7}$ by 42:

- First term: $\frac{1}{3} \times 42 = 14$.
- Second term: $\frac{1}{6} \times 42 = 7$.
- Third term: $\frac{1}{7} \times 42 = 6$.

So, the ratio $A : B : C$ is $14 : 7 : 6$.

Verifying the Ratio

We can verify the ratio $14 : 7 : 6$ using the original equation $3A = 6B = 7C$. Let $A = 14x$, $B = 7x$, and $C = 6x$ for some common factor x .

- $3A = 3 \times (14x) = 42x$
- $6B = 6 \times (7x) = 42x$
- $7C = 7 \times (6x) = 42x$

Since $3A = 6B = 7C = 42x$, the ratio $14 : 7 : 6$ is correct.

Variable	Expression in terms of k	Term after multiplying by LCM(42)
A	$\frac{k}{3}$	$\frac{1}{3} \times 42 = 14$
B	$\frac{k}{6}$	$\frac{1}{6} \times 42 = 7$
C	$\frac{k}{7}$	$\frac{1}{7} \times 42 = 6$

Final Ratio Result

The ratio $A : B : C$ is $14 : 7 : 6$.

Revision Table: Key Concepts for Ratio Problems

Concept	Description	Relevance to Question
Ratio	A comparison of two or more quantities of the same kind, usually expressed as $a : b$ or a/b .	Finding the ratio $A : B : C$.
Proportion	An equality between two ratios. While not directly used here, the original equation $3A = 6B = 7C$ implies proportions.	Relates different parts of the equation.
Least Common Multiple (LCM)	The smallest positive integer that is a multiple of two or more numbers.	Used to convert a ratio of fractions to a ratio of integers.
Equating to a Constant	Setting multiple equal expressions to a single variable (like k) to simplify solving.	Crucial step in expressing A , B , and C individually.

Additional Information: Understanding Ratios and Proportions

Ratios and proportions are fundamental concepts in mathematics used to compare quantities. A ratio expresses how much of one quantity there is compared to another. For example, a ratio of 2:3 for boys to girls means for every 2 boys, there are 3 girls.

When you have an equation like $3A = 6B = 7C$, it implies a proportional relationship between the variables. If you increase A, B and C must also increase proportionally to maintain the equality.

The method of setting the expression equal to a constant k is a powerful technique for solving problems where multiple quantities are in a continued equality. It helps in isolating each variable and then forming the ratio easily. Remember that the ratio $A : B : C$ represents the simplest form of the relationship between A, B, and C. If $A : B : C = p : q : r$, it means $A = px$, $B = qx$, and $C = rx$ for some common factor x .

Common mistakes include directly taking the coefficients (3:6:7) or their reciprocals without finding the LCM correctly. Always convert the ratio of fractions to integers by multiplying by the LCM of the denominators.

86. Answer: b

Explanation:

Calculating $(0.1^2 - 0.025^2) \div (0.1 - 0.025)$

The problem asks us to evaluate the expression: $(0.1^2 - 0.025^2) \div (0.1 - 0.025)$. This expression involves squares and division of decimal numbers.

Let's analyze the structure of the expression. The numerator is in the form of a difference of squares, $a^2 - b^2$, where $a = 0.1$ and $b = 0.025$. The denominator is the difference of the same two numbers, $a - b$.

We can use the algebraic identity for the difference of squares, which states:

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

Using this identity, we can rewrite the numerator of the given expression:

$$0.1^2 - 0.025^2 = (0.1 - 0.025)(0.1 + 0.025)$$

Now substitute this back into the original expression:

$$\frac{0.1^2 - 0.025^2}{0.1 - 0.025} = \frac{(0.1 - 0.025)(0.1 + 0.025)}{0.1 - 0.025}$$

Assuming that the denominator is not zero (which is true since $0.1 - 0.025 = 0.075 \neq 0$), we can cancel the term $(0.1 - 0.025)$ from both the numerator and the denominator:

$$\frac{\cancel{(0.1 - 0.025)}(0.1 + 0.025)}{\cancel{(0.1 - 0.025)}} = 0.1 + 0.025$$

Now, we just need to perform the addition:

$$0.1 + 0.025$$

To add these decimal numbers, we can align the decimal points:

$$\begin{array}{r} 0.100 \\ +0.025 \\ \hline 0.125 \end{array}$$

So, the result of the expression is 0.125.

Let's compare this result with the given options:

- Option 1: 0.325
- Option 2: 0.125
- Option 3: 0.625
- Option 4: 0.25

The calculated value 0.125 matches Option 2.

Step-by-Step Calculation Summary

1. Identify the structure of the expression as a difference of squares divided by the difference of the base numbers.
2. Recall or apply the difference of squares formula: $a^2 - b^2 = (a - b)(a + b)$.
3. Substitute the numbers into the formula: $0.1^2 - 0.025^2 = (0.1 - 0.025)(0.1 + 0.025)$.
4. Rewrite the original expression using the factored numerator: $\frac{(0.1 - 0.025)(0.1 + 0.025)}{0.1 - 0.025}$.
5. Cancel out the common term $(0.1 - 0.025)$ from the numerator and the denominator.

6. The expression simplifies to $0.1 + 0.025$.
7. Perform the addition: $0.1 + 0.025 = 0.125$.
8. Match the result with the provided options.

Revision Table: Key Concepts

Concept	Description	Formula/Example
Difference of Squares	An algebraic factorization pattern for an expression that is the difference between two perfect squares.	$a^2 - b^2 = (a - b)(a + b)$
Decimal Addition	Adding numbers that contain decimal points. Requires aligning the decimal points before adding.	$0.1 + 0.025 = 0.125$
Algebraic Simplification	Reducing an algebraic expression to a simpler form, often by factoring and cancelling common terms.	$\frac{(x-y)(x+y)}{(x-y)} = x + y$ (where $x \neq y$)

Additional Information: Understanding Difference of Squares

The difference of squares identity, $a^2 - b^2 = (a - b)(a + b)$, is a fundamental concept in algebra. It is useful for factoring expressions, simplifying fractions like in this problem, and solving equations.

Let's see why the formula $a^2 - b^2 = (a - b)(a + b)$ works:

Consider the product $(a - b)(a + b)$. Using the distributive property (or FOIL method):

- First terms: $a \times a = a^2$
- Outer terms: $a \times b = ab$
- Inner terms: $-b \times a = -ab$
- Last terms: $-b \times b = -b^2$

Adding these terms together:

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

The terms $+ab$ and $-ab$ cancel each other out:

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

This confirms the identity. In our problem, recognizing this pattern allowed us to simplify a complex-looking division into a simple addition.

87. Answer: a

Explanation:

Find the Midpoint of a Line Segment using the Midpoint Formula

This problem asks us to find the coordinates of the midpoint of the line segment connecting two specific points: $(-4, 7)$ and $(2, 3)$. Finding the midpoint is a fundamental concept in coordinate geometry.

Understanding the Midpoint Concept

The midpoint of a line segment is the point exactly halfway between the two endpoints. It divides the segment into two equal parts. In a coordinate plane, we can find the coordinates of this midpoint using a specific formula.

Applying the Midpoint Formula

The midpoint formula is used to find the coordinates (x_m, y_m) of the midpoint of a segment with endpoints (x_1, y_1) and (x_2, y_2) . The formula is:

$$(x_m, y_m) = \left(\frac{x_1 + x_2}{2}, \frac{y_1 + y_2}{2} \right)$$

This means we average the x-coordinates and average the y-coordinates of the two endpoints separately to find the coordinates of the midpoint.

Step-by-Step Midpoint Calculation

Let the given points be $P_1 = (-4, 7)$ and $P_2 = (2, 3)$.

We can assign the coordinates:

- $x_1 = -4$
- $y_1 = 7$
- $x_2 = 2$
- $y_2 = 3$

Now, we apply the midpoint formula:

1. Calculate the x-coordinate of the midpoint (x_m):

$$x_m = \frac{x_1 + x_2}{2} = \frac{-4 + 2}{2}$$

$$x_m = \frac{-2}{2} = -1$$

2. Calculate the y-coordinate of the midpoint (y_m):

$$y_m = \frac{y_1 + y_2}{2} = \frac{7 + 3}{2}$$

$$y_m = \frac{10}{2} = 5$$

So, the coordinates of the midpoint of the segment joining $(-4, 7)$ and $(2, 3)$ are $(-1, 5)$.

Comparing Calculated Midpoint with Options

Let's compare our calculated midpoint coordinates with the given options.

Calculated Midpoint	Option 1	Option 2	Option 3	Option 4
$(-1, 5)$	$(-1, 5)$	$(-2, 3)$	$(1, -5)$	$(2, 4)$

The calculated midpoint $(-1, 5)$ matches the coordinates provided in Option 1.

Revision Table: Essential Coordinate Geometry Formulas

Here is a quick recap of some essential formulas used in coordinate geometry when dealing with points and segments:

Formula	Purpose	Formula Expression
Midpoint Formula	To find the coordinates of the midpoint of a segment between (x_1, y_1) and (x_2, y_2) .	$\left(\frac{x_1+x_2}{2}, \frac{y_1+y_2}{2}\right)$
Distance Formula	To find the distance between two points (x_1, y_1) and (x_2, y_2) .	$\sqrt{(x_2 - x_1)^2 + (y_2 - y_1)^2}$
Slope Formula	To find the slope of a line passing through two points (x_1, y_1) and (x_2, y_2) .	$\frac{y_2 - y_1}{x_2 - x_1}$ (for $x_1 \neq x_2$)

Additional Information: Introduction to Coordinate Geometry

Coordinate geometry, also known as analytical geometry, is a branch of mathematics that uses coordinates to relate algebraic equations and geometric figures. It allows us to solve geometrical problems using algebraic methods.

Key ideas in coordinate geometry include:

- Representing points using ordered pairs (x, y) on a coordinate plane.
- Defining geometric shapes (like lines, circles, parabolas) using algebraic equations.
- Using formulas to calculate properties of figures, such as the distance between points, the midpoint of a segment, the slope of a line, and the area of a polygon.

Understanding the midpoint formula is a crucial step in mastering coordinate geometry problems involving line segments. It's frequently used in various geometry and algebra applications.

88. Answer: a

Explanation:

Understanding the National Skills Development Corporation (NSDC)

The question asks about the government ministry under which the National Skills Development Corporation (NSDC) was formed in India. The NSDC is an important institution focused on skill development across the country.

Formation of NSDC: A Public-Private Partnership

The National Skills Development Corporation (NSDC) was established as a unique public-private partnership (PPP) model. Its primary objective is to promote skill development by creating large, quality, for-profit vocational institutions. It aims to bridge the gap between the demand and supply of skilled workforce across various sectors in India.

Understanding the formation context helps identify the founding ministry. Such large-scale initiatives, especially those involving significant funding and economic impact, are often driven or initiated by ministries responsible for economic policy and financial matters.

Identifying the Founding Ministry of NSDC

Based on the structure and purpose of the NSDC, its formation is linked to a specific government ministry. Let's examine the options provided:

- Ministry of Finance
- Ministry of Electronics & Information Technology
- Ministry of Corporate Affairs
- Ministry of Science & Technology

The establishment of NSDC, aimed at channeling funds and facilitating private sector participation in skill development on a large scale, was initiated under the purview of the ministry responsible for the nation's financial and economic policies.

The National Skills Development Corporation (NSDC) was indeed formed under India's **Ministry of Finance**.

Role and Objectives of NSDC

The core functions of the NSDC include:

- Funding skill development initiatives, often through loans or equity investments in private sector training companies.
- Establishing and maintaining quality standards for training programs.

- Developing sector-specific skill councils to define skill requirements and curricula.
- Acting as a market maker by bringing together various stakeholders in the skill ecosystem.

While other ministries like the Ministry of Skill Development and Entrepreneurship (MSDE) now play a central role in overseeing skill initiatives (and NSDC is administratively under MSDE), its foundational establishment was under the Ministry of Finance to kickstart the initiative with necessary financial backing and policy push.

Analyzing Other Options

- **Ministry of Electronics & Information Technology (MeitY):** This ministry primarily deals with IT policy, electronics manufacturing, and digital services, not the broad spectrum of skill development across all sectors.
- **Ministry of Corporate Affairs (MCA):** This ministry focuses on the administration of the Companies Act and other corporate laws, regulating businesses, not directly on skill development initiatives.
- **Ministry of Science & Technology:** This ministry promotes scientific research, technology development, and innovation, which is different from vocational skill training aimed at various industries.

Therefore, the formation of a large-scale public-private partnership for skill development, involving significant financial mechanisms, aligns with the responsibilities of the Ministry of Finance at the time of its inception.

Your Personal Exams Guide

Aspect	Details
Organisation	National Skills Development Corporation (NSDC)
Type	Public-Private Partnership (PPP)
Formed Under Ministry (Initial)	Ministry of Finance
Current Administrative Ministry	Ministry of Skill Development and Entrepreneurship (MSDE)
Primary Goal	Promote skill development by funding and supporting private sector initiatives.

Revision Table: Key Facts on NSDC Formation

Reviewing the core facts helps solidify understanding of the NSDC's origin:

Key Detail	Description
NSDC Full Form	National Skills Development Corporation
Structure	Public-Private Partnership
Formed By	Government of India
Initial Ministry	Ministry of Finance
Year of Formation	2008

Additional Information on Skill Development in India

Skill development is a crucial focus area for India to leverage its demographic dividend. Several initiatives have been launched over the years:

- **National Skill Development Policy:** Provides a framework for skill development efforts.
- **Pradhan Mantri Kaushal Vikas Yojana (PMKVY):** A flagship scheme for skill training with a focus on youth.
- **Sector Skill Councils (SSCs):** Industry-led bodies under NSDC that define occupational standards and qualifications.
- **Ministry of Skill Development and Entrepreneurship (MSDE):** Created later (in 2014) to provide dedicated focus and coordination for all skill development efforts, consolidating various initiatives previously handled by different ministries, including administrative control over NSDC.

Understanding the evolution of the skill development ecosystem, including the initial role of the Ministry of Finance in establishing NSDC and the later creation of MSDE, provides a complete picture.

89. Answer: b

Explanation:

Understanding Electrical Resistance in a Metal Rod

The electrical resistance of a material, like a metal rod, is a measure of how much it opposes the flow of electric current. It is a fundamental property when studying circuits and electrical conductors.

The resistance of a uniform conductor, such as a metal rod with a constant cross-sectional area, is primarily determined by its physical dimensions and the intrinsic properties of the material it is made from. The well-known formula for resistance is:

$$R = \rho \frac{L}{A}$$

Where:

- R is the resistance of the conductor.
- ρ (rho) is the resistivity of the material.
- L is the length of the conductor.
- A is the cross-sectional area of the conductor.

Factors Directly Affecting Metal Rod Resistance

Let's look at the factors that directly influence the resistance of a metal rod based on the formula and physical principles:

Resistivity (ρ)

Resistivity is an intrinsic property of the material itself. It quantifies how strongly a material resists electric current flow. Different materials have different resistivities (e.g., copper has low resistivity, glass has high resistivity). A material's resistivity also changes with temperature.

Length (L)

From the formula $R = \rho \frac{L}{A}$, we can see that resistance is directly proportional to the length of the rod. This means a longer rod will have higher resistance than a shorter one of the

same material and cross-sectional area. Electrons have to travel a greater distance, encountering more scattering events along the way.

Cross-sectional Area (A)

The formula shows that resistance is inversely proportional to the cross-sectional area. A thicker rod (larger area) will have lower resistance than a thinner one of the same material and length. A larger area provides more space for electrons to flow, reducing the 'bottleneck effect'.

Temperature

For most metals, resistivity (ρ) increases with increasing temperature. As temperature rises, the atoms within the metal vibrate more vigorously, causing more frequent collisions with the flowing electrons. These collisions impede the electron flow, thus increasing the resistance of the metal rod. So, temperature indirectly affects resistance by changing resistivity.

Why Density Does Not Affect Resistance

Density is defined as mass per unit volume ($Density = \frac{Mass}{Volume}$). While the mass and volume of a metal rod depend on its material and dimensions, density itself is not a factor in the fundamental formula for electrical resistance ($R = \rho \frac{L}{A}$). Resistance is about how the material impedes charge flow through its structure and geometry, not its mass-to-volume ratio. While material properties like resistivity are related to the arrangement and type of atoms (which also influence density), density is not the property that directly determines electrical resistance in the way resistivity, length, area, and temperature do.

Analyzing the Given Options

Let's examine each option in the context of what affects the resistance of a metal rod:

- **Resistivity:** Yes, resistance is directly proportional to the resistivity of the material. Different metals have different resistivities.
- **Density:** No, density is not a factor that directly determines electrical resistance according to the standard formula and physical principles.
- **Length:** Yes, resistance is directly proportional to the length of the rod.
- **Temperature:** Yes, resistance of a metal rod typically increases with temperature due to the change in resistivity.

Therefore, the resistance of a metal rod depends on resistivity, length, and temperature, but not directly on density.

Conclusion on Metal Rod Resistance Factors

Based on the formula $R = \rho \frac{L}{A}$ and the relationship between resistivity and temperature, the resistance of a metal rod is dependent on its material's resistivity, its length, its cross-sectional area (implicitly included in the formula), and its temperature. Density is not a direct factor determining the electrical resistance.

Revision Table: Factors Affecting Metal Rod Resistance

Factor	Does it Affect Resistance?	Relationship (for uniform rod)
Resistivity (ρ)	Yes	$R \propto \rho$ (Resistance is proportional to resistivity)
Length (L)	Yes	$R \propto L$ (Resistance is proportional to length)
Cross-sectional Area (A)	Yes	$R \propto \frac{1}{A}$ (Resistance is inversely proportional to area)
Temperature	Yes	For metals, R generally increases with increasing temperature (due to ρ 's dependence on T)
Density	No	Not directly related by the formula for resistance

Additional Information: Electrical Conductivity

Electrical conductivity (σ) is another important property of a material, representing how easily electric current flows through it. Conductivity is the reciprocal of resistivity ($\sigma = \frac{1}{\rho}$). Materials with high conductivity (like copper and aluminum) have low resistivity and are good conductors. Materials with low conductivity (high resistivity) are poor conductors or insulators.

While density is related to the mass of the material per unit volume, conductivity and resistivity are related to how freely charge carriers (electrons in metals) can move through the material's lattice structure. These are distinct physical properties.

90. Answer: d

Explanation:

Correct answer: 8%

Compound Interest (CI): $CI = A - P = P \left[\left(1 + \frac{r}{100} \right)^n - 1 \right]$

If interest is compounded half-yearly, the rate becomes $r/2$ and time becomes $2n$: $A = P \left(1 + \frac{r/2}{100} \right)^{2n}$

If interest is compounded quarterly, the rate becomes $r/4$ and time becomes $4n$: $A = P \left(1 + \frac{r/4}{100} \right)^{4n}$

If interest rates are different for different years $(r_1, r_2, \dots)$: $A = P \left(1 + \frac{r_1}{100} \right) \left(1 + \frac{r_2}{100} \right) \dots$

In our problem, the interest is compounded annually at a constant rate, which allowed us to use the simple relationship between the amounts in consecutive years.

91. Answer: a

Explanation:

Understanding Pulley System Efficiency

This problem asks us to find the efficiency of a pulley system. We are given the mechanical advantage (MA) and the distances moved by the load and the effort. Efficiency tells us how well a machine converts the work input into useful work output. For a pulley system, efficiency depends on factors like friction and the weight of the moving parts.

Key Concepts for Pulley System Efficiency

To calculate the efficiency of a pulley system, we need to understand a few key terms:

- **Mechanical Advantage (MA):** This is the ratio of the load lifted to the effort applied. $MA = \frac{\text{Load}}{\text{Effort}}$. We are given the MA directly in this problem.

- **Velocity Ratio (VR):** This is the ratio of the distance moved by the effort to the distance moved by the load. $VR = \frac{\text{Distance moved by Effort}}{\text{Distance moved by Load}}$. The VR is a theoretical value determined by the geometry of the pulley system, specifically the number of effort segments supporting the load.
- **Efficiency:** This is the ratio of the output work to the input work, usually expressed as a percentage. For a machine, efficiency can also be calculated using the MA and VR: $\text{Efficiency} = \frac{MA}{VR} \times 100\%$.

Calculating Velocity Ratio (VR)

We are given the distances moved by the load and the effort. We can use these values to calculate the Velocity Ratio (VR) of the pulley system.

The formula for Velocity Ratio is:

$$VR = \frac{\text{Distance moved by Effort}}{\text{Distance moved by Load}}$$

Given values:

- Distance moved by Effort (d_E) = 10 m
- Distance moved by Load (d_L) = 2.5 m

Now, let's calculate the VR:

$$VR = \frac{10 \text{ m}}{2.5 \text{ m}}$$

$$VR = 4$$

So, the Velocity Ratio of this pulley system is 4.

Calculating Pulley System Efficiency

Now that we have the Mechanical Advantage (MA) and the Velocity Ratio (VR), we can calculate the efficiency of the pulley system using the formula:

$$\text{Efficiency} = \frac{MA}{VR} \times 100\%$$

Given/Calculated values:

- Mechanical Advantage (MA) = 2.5
- Velocity Ratio (VR) = 4

Substitute these values into the efficiency formula:

$$\text{Efficiency} = \frac{2.5}{4} \times 100\%$$

$$\text{Efficiency} = 0.625 \times 100\%$$

$$\text{Efficiency} = 62.5\%$$

The efficiency of the pulley system is 62.5%.

Final Result

The efficiency of the pulley system is found to be 62.5%.

Revision Table: Pulley System Concepts and Formulas

Concept	Definition	Formula
Mechanical Advantage (MA)	Ratio of Load to Effort	$MA = \frac{\text{Load}}{\text{Effort}}$
Velocity Ratio (VR)	Ratio of distance moved by Effort to distance moved by Load	$VR = \frac{d_E}{d_L}$
Efficiency	Ratio of Useful Work Output to Work Input (or MA to VR)	$\text{Efficiency} = \frac{\text{Output Work}}{\text{Input Work}} \times 100\%$ or $\text{Efficiency} = \frac{MA}{VR} \times 100\%$

Additional Information: Ideal vs Real Pulley Systems Efficiency

In an ideal pulley system, there is no friction, and the ropes and pulleys are massless. In such a theoretical scenario, the Mechanical Advantage (MA) is equal to the Velocity Ratio (VR), and the efficiency is 100%. However, real-world pulley systems always have some energy losses due to:

- Friction in the pulley axles.

- Friction between the rope and the pulleys.
- The weight of the pulleys and the rope itself.

Because of these losses, the actual Mechanical Advantage (MA) in a real pulley system is always less than the theoretical Velocity Ratio (VR). Consequently, the efficiency of a real pulley system is always less than 100%. The difference between MA and VR indicates the extent of these energy losses in the system.

92. Answer: a

Explanation:

Understanding Tessellations: Covering Surfaces with Shapes

The question asks for the term that describes covering a surface with a pattern of flat shapes where there are no gaps or overlaps. This is a fundamental concept in geometry and design.

Defining the Key Term: Tessellation

Let's look at the provided options and see which one accurately fits the description:

- **Tessellation:** A tessellation (or tiling) is the covering of a surface, usually a plane, with one or more geometric shapes, called tiles, with no overlaps and no gaps. The shapes can be regular (like squares or triangles) or irregular. This definition perfectly matches the description given in the question. Famous examples include floor tiles, brick walls, or patterns found in nature like honeycomb.
- **Tracking:** In typography, tracking refers to the consistent adjustment of space between letters across a block of text. This is unrelated to covering a surface with shapes.
- **Gradient:** A gradient is a gradual change in color, shade, or texture across a surface or object. While gradients can be used in design, they do not involve tiling a surface with discrete shapes without gaps or overlaps.
- **Kerning:** Also a typography term, kerning refers to the adjustment of space between specific pairs of characters to improve visual appearance. This is unrelated to covering a surface with shapes.

Based on the definitions, the term that describes covering a surface with a pattern of flat shapes so that there are no overlaps or gaps is Tessellation.

Comparing the Concepts

To further clarify, let's compare the terms:

Term	Description	Related to Covering a Surface with Shapes?
Tessellation	Covering a surface with shapes (tiles) without gaps or overlaps.	Yes
Tracking	Adjusting uniform space between letters in text.	No
Gradient	Gradual change in color or shade.	No
Kerning	Adjusting space between specific letter pairs.	No

Therefore, the correct term is Tessellation.

Revision Table: Key Geometric and Design Terms

Term	Simple Definition	Example/Context
Tessellation	Tiling a surface with shapes, no gaps/overlaps.	Floor tiles, honeycomb
Polygon	A closed 2D shape with straight sides.	Square, triangle, hexagon
Surface	The outer layer of an object; in geometry, often a plane.	Floor, wall, paper

Additional Information: Types of Tessellations

Tessellations can be simple or complex. Here are a few types:

- **Regular Tessellations:** Made up of only one type of regular polygon (all sides and angles equal) that meets at a vertex in the same way. Only three regular polygons can form a regular tessellation: squares, equilateral triangles, and regular hexagons.
- **Semi-regular Tessellations:** Made up of two or more types of regular polygons. The pattern of polygons around each vertex must be the same throughout the tessellation. There are eight such tessellations.
- **Irregular Tessellations:** Made up of irregular polygons or combinations of regular and irregular polygons. These are very common in art and design.

The concept of tessellation is important in various fields, including mathematics, art, architecture, and computer graphics.

93. Answer: c

Explanation:

Finding the Unknown Number Using the Mean

The problem asks us to find the value of 'x' given a set of numbers and their mean. The set of numbers is 10, 4, 1, 15, 15, x, 12, and 14. The given mean is 10.

Understanding the Concept of Mean (Average)

The mean, or average, of a set of numbers is calculated by summing all the numbers in the set and then dividing the sum by the total count of numbers in the set.

Mathematically, the formula for the mean is:

$$\text{Mean} = \frac{\text{Sum of all numbers}}{\text{Total number of values}}$$

Setting Up the Equation to Find x

We are given the numbers: 10, 4, 1, 15, 15, x, 12, and 14.

First, let's count how many numbers are in this set. There are 8 numbers in total.

The sum of these numbers is $10 + 4 + 1 + 15 + 15 + x + 12 + 14$.

The given mean is 10.

Using the formula for the mean, we can write the equation:

$$10 = \frac{10 + 4 + 1 + 15 + 15 + x + 12 + 14}{8}$$

Solving for x Step-by-Step

Now, let's solve the equation to find the value of x.

Step 1: Sum the known numbers in the numerator.

$$10 + 4 + 1 + 15 + 15 + 12 + 14 = 71$$

So the equation becomes:

$$10 = \frac{71 + x}{8}$$

Step 2: Multiply both sides of the equation by 8 to isolate the numerator.

$$10 \times 8 = 71 + x$$

$$80 = 71 + x$$

Step 3: Subtract 71 from both sides of the equation to find x.

$$80 - 71 = x$$

$$x = 9$$

So, the value of x is 9.

Verifying the Solution

Let's check if the mean of the numbers 10, 4, 1, 15, 15, 9, 12, and 14 is indeed 10.

Sum of the numbers: $10 + 4 + 1 + 15 + 15 + 9 + 12 + 14 = 80$.

Total number of values: 8.

Mean =

$$\frac{80}{8} = 10$$

The calculated mean matches the given mean, so our value for x is correct.

Number	Value
Number 1	10
Number 2	4
Number 3	1
Number 4	15
Number 5	15
Number 6	x (which is 9)
Number 7	12
Number 8	14
Total Sum	80
Number of values	8
Mean	$\frac{80}{8} = 10$

Your Personal Exams Guide

Therefore, the value of x is 9.

Revision Table: Mean Calculation

This table summarizes the process of finding the mean.

Concept	Description	Formula
Mean (Average)	Sum of values divided by the count of values.	$\bar{x} = \frac{\sum x_i}{n}$
Sum of values ($\sum x_i$)	Adding up all the numbers in the set.	$x_1 + x_2 + \dots + x_n$
Count of values (n)	The total number of items in the set.	Counting the items: 1, 2, 3, ...

Additional Information on Measures of Central Tendency

The mean is one type of measure of central tendency. Other common measures include the median and the mode.

- **Median:** The middle value in a dataset that is ordered from least to greatest. If there's an even number of values, the median is the average of the two middle values.
- **Mode:** The value that appears most frequently in a dataset. A dataset can have one mode, multiple modes, or no mode.

Understanding the mean, median, and mode helps in analyzing and describing the characteristics of a dataset. Each measure provides different insights into the central value of the data.

94. Answer: d

Explanation:

Understanding Syllogism Statements and Conclusions

This question asks us to analyze two statements and determine which of the two given conclusions logically follow from these statements. We must assume the statements are true, even if they contradict common knowledge.

Analyzing the Given Statements

We have two statements:

- Statement 1: Some wagons are buggies.
- Statement 2: All wagons are carts.

Let's break down what these statements mean. Statement 1 tells us there's at least one wagon that is also a buggy, and at least one buggy that is also a wagon. There might be other wagons that aren't buggies and other buggies that aren't wagons. Statement 2 is stronger; it says every single wagon belongs to the category of carts. There are no wagons outside the group of carts.

Visualizing with Venn Diagrams

We can use Venn diagrams to represent these relationships visually. Let 'W' represent the set of Wagons, 'B' represent the set of Buggies, and 'C' represent the set of Carts.

- Statement 2 (All wagons are carts): The entire circle representing Wagons (W) must be drawn inside the circle representing Carts (C).
- Statement 1 (Some wagons are buggies): There must be an overlapping area between the Wagons circle (W) and the Buggies circle (B). Since the Wagons circle is inside the Carts circle, this overlap between W and B must necessarily fall within the Carts circle (C).

Evaluating the Conclusions

Now let's look at the conclusions based on our understanding and visualization:

- Conclusion I: All buggies are carts.
- Conclusion II: Some carts are buggies.

Checking Conclusion I: All buggies are carts

Does the entire set of Buggies (B) have to be inside the set of Carts (C)?

Based on Statement 1, only *some* buggies are wagons. Statement 2 tells us *all* wagons are carts. This means the portion of buggies that are also wagons *must* be carts. However, the statements don't give us any information about the buggies that are *not* wagons. These buggies could be inside the circle of Carts, or they could be outside the circle of Carts. We cannot be certain that *all* buggies are carts. Therefore, Conclusion I does not necessarily follow from the statements.

Checking Conclusion II: Some carts are buggies

Is there any overlap between the set of Carts (C) and the set of Buggies (B)?

Statement 1 says some wagons are buggies. Let's call this group "Wagons-that-are-buggies".

Statement 2 says all wagons are carts. This means our group "Wagons-that-are-buggies" are also carts, because they are a type of wagon.

So, the group "Wagons-that-are-buggies" are both carts and buggies. This proves that there is an overlap between carts and buggies. Specifically, the carts that are also wagons are also buggies. Therefore, some carts are buggies. Conclusion II logically follows from the statements.

Summary of Conclusions

- Conclusion I: All buggies are carts. (Does not follow)
- Conclusion II: Some carts are buggies. (Follows)

Therefore, only conclusion II follows from the given statements.

Revision Table: Statements and Conclusions Analysis

Statement/Conclusion	Relationship Described	Follows?	Reasoning
Statement 1: Some wagons are buggies.	Partial overlap (W & B)	Given as true	Assumption for the problem
Statement 2: All wagons are carts.	W is subset of C	Given as true	Assumption for the problem
Conclusion I: All buggies are carts.	B is subset of C	No	Statements don't confirm buggies not in W are in C.
Conclusion II: Some carts are buggies.	Partial overlap (C & B)	Yes	Wagons that are buggies (from S1) are also carts (from S2), proving some carts are buggies.

Additional Information on Syllogisms

Syllogisms are a form of logical argument that apply deductive reasoning to arrive at a conclusion based on two or more propositions (statements) that are assumed to be true. They typically involve three parts: a major premise, a minor premise, and a conclusion.

Common types of relationships used in syllogisms include:

- All A are B (Universal Affirmative)
- No A are B (Universal Negative)
- Some A are B (Particular Affirmative)
- Some A are not B (Particular Negative)

To solve syllogism problems, it is often helpful to use Venn diagrams or follow a set of rules derived from logic. Remember that you must strictly adhere to the information given in the statements and not use any outside knowledge.

95. Answer: a

Explanation:

0.0144

Squaring 0.12 involves multiplying 0.12 by itself:

$$X = 0.12 \times 0.12$$

$$X = 0.0144$$

Final Result After performing the calculations step-by-step, we find that the value of X is 0.0144 .

96. Answer: a

Explanation:

Understanding the Dadasaheb Phalke Award

The Dadasaheb Phalke Award is India's highest honour in the field of cinema. It is presented annually at the National Film Awards ceremony by the Directorate of Film Festivals, an organisation set up by the Ministry of Information and Broadcasting. The award is named after Dadasaheb Phalke, who is considered the 'father of Indian cinema' for directing India's first full-length feature film, *Raja Harishchandra* (1913).

The award is given for an individual's "outstanding contribution to the growth and development of Indian cinema". It consists of a Swarna Kamal (Golden Lotus) medallion, a shawl, and a cash prize.

Defining Posthumous Awards

The term "posthumous" refers to something occurring, awarded, or appearing after the death of the person concerned. So, a posthumous award is one that is granted to an individual after they have passed away, in recognition of their achievements during their lifetime.

Analyzing Dadasaheb Phalke Award Recipients

Let's look at the individuals listed in the options and their association with the Dadasaheb Phalke Award:

- **Vinod Khanna:** A highly acclaimed actor, film producer, and politician.
- **Satyajit Ray:** A celebrated filmmaker, screenwriter, music composer, graphic artist, and film critic.
- **Naushad:** A renowned music director, known for his work in Hindi cinema.
- **Durga Khote:** A veteran actress, producer, and writer, active in theatre and cinema.

We need to determine which of these personalities received the Dadasaheb Phalke Award and if the award was conferred after their death.

Personality	Dadasaheb Phalke Award Year	Awarded Posthumously?	Notes
Vinod Khanna	2017	Yes	Award announced and presented after his death in April 2017.
Satyajit Ray	1984	No	Received the award during his lifetime.
Naushad	1981	No	Received the award during his lifetime.
Durga Khote	1983	No	Received the award during her lifetime.

Based on this analysis, Vinod Khanna is the recipient among the given options who was awarded the Dadasaheb Phalke Award posthumously. He passed away in April 2017, and the award for the year 2017 was announced later, making him a posthumous recipient.

Revision Table: Dadasaheb Phalke Awardees

Recipient	Year	Posthumous?
Vinod Khanna	2017	Yes
Satyajit Ray	1984	No
Naushad	1981	No
Durga Khote	1983	No

Additional Information on Dadasaheb Phalke Award

The Dadasaheb Phalke Award was instituted in 1969. The first recipient was Devika Rani. The selection committee includes eminent personalities from the Indian film industry. The award acknowledges the monumental contribution of artists and technicians to the rich

heritage of Indian cinema over the years. It is one of the most prestigious recognitions in the country's cultural landscape.

97. Answer: b

Explanation:

Solving Blood Relation Puzzles

Blood relation questions test your ability to understand and decipher relationships between individuals based on given information. To solve these types of questions, it is helpful to break down the statement into smaller parts and identify the relationships step-by-step.

Analyzing the Relationship Statement

The statement given is: "You are my husband's daughter-in-law's son." A is speaking to B. Let's break this down from A's perspective:

- "My husband": This refers to A's spouse.
- "My husband's daughter-in-law": A husband's daughter-in-law is the wife of his son. Since A is the speaker, "my husband's daughter-in-law" is the wife of A's son.
- "My husband's daughter-in-law's son": This refers to the son of the wife of A's son. The son of A's son is A's grandson.

So, the statement "You are my husband's daughter-in-law's son" means "You are my grandson".

Therefore, B is the grandson of A.

Relationship Part	Meaning (from A's perspective)
My husband	A's spouse
My husband's daughter-in-law	Wife of A's son
My husband's daughter-in-law's son	Son of A's son (A's grandson)

Based on this breakdown, the relationship between B and A is that B is the grandson of A.

Checking the Options

- **B is the father of A:** This is incorrect. If B were A's father, A would be B's daughter or son.
- **B is the grandson of A:** This matches our step-by-step analysis.
- **B is the brother of A:** This is incorrect. B is in a younger generation related through A's son.
- **B is the son of A:** This is incorrect. B is A's son's son, meaning B is A's grandson, not A's son.

The correct relationship is that B is the grandson of A.

Revision Table: Key Blood Relation Terms

Prepp

Your Personal Exams Guide

Term	Relationship
Father's father	Paternal Grandfather
Mother's mother	Maternal Grandmother
Father's mother	Paternal Grandmother
Mother's father	Maternal Grandfather
Son's wife	Daughter-in-law
Daughter's husband	Son-in-law
Husband's or wife's sister	Sister-in-law
Husband's or wife's brother	Brother-in-law
Brother's son	Nephew
Brother's daughter	Niece
Sister's son	Nephew
Sister's daughter	Niece
Son of son	Grandson
Daughter of son	Granddaughter
Son of daughter	Grandson
Daughter of daughter	Granddaughter

Additional Information on Solving Relation Questions

Solving blood relation questions efficiently requires careful reading and understanding of family structures. Here are some tips:

- **Visual Representation:** For complex problems, drawing a family tree can be extremely helpful. Use symbols for different genders and relationships (e.g., squares for males, circles for females, horizontal lines for marriage, vertical lines for parent-child).
- **Break Down Sentences:** Always break down long statements into smaller, manageable parts, starting from the end of the statement and working backward

towards the speaker or the person being referred to.

- **Identify the Subject and Object:** Clearly identify who is speaking (A in this case) and who is being referred to (B).
- **Common Mistakes:** Be careful not to confuse paternal and maternal relations or generations (e.g., son vs. grandson).
- **Practice:** The more you practice different types of blood relation puzzles, the faster and more accurate you will become.

98. Answer: c

Explanation:

Short Trick:

Population density of B in 2000 = $100/100 \text{ sq km} = 1 \text{ per sq km}$

$\therefore$ Area of B = 10 million sq km

Population density of D in 2010 = $300/100 \text{ sq km} = 3 \text{ per sq km}$

$\therefore$ Area of D = $6/3 = 2 \text{ million sq km}$

$\therefore$ Required percentage,

$\Rightarrow (10 - 2)/2 \times 100\% = 400\%$

Detailed Solution:

Let, area of B = $x \text{ sq km}$ and area of D = $y \text{ sq km}$

According to the question,

$\Rightarrow x \times 100/100 = 10 \times 10^6$

$\Rightarrow x = 10 \times 10^6$

$\therefore$ Area of B = 10 million sq km

According to the question,

$\Rightarrow y \times 300/100 = 6 \times 10^6$

$$\Rightarrow x = 2 \times 10^6$$

$\therefore$ Area of D = 2 million sq km

$\therefore$ Required percentage,

$$\Rightarrow (10 - 2)/2 \times 100\% = 400\%$$

99. Answer: b

Explanation:

Understanding Bezier Curves

The question asks to identify a type of curve that is defined using a minimum of three points. Specifically, it mentions two endpoints called **anchor points** and other points that shape the curve, referred to as **handles, tangent points, or nodes**. Let's examine the options based on these characteristics.

Identifying Bezier Curves

Bezier curves are widely used in vector graphics and computer-aided design (CAD). They are defined mathematically using a set of control points. The curve starts at the first control point and ends at the last control point. The points in between, known as control points or handles, influence the curve's shape without necessarily being part of the curve itself.

Components and Structure

The description provided in the question aligns perfectly with the definition of Bezier curves:

- **Anchor Points:** These correspond to the first and last control points of a Bezier curve, defining its start and end.
- **Handles/Tangent Points/Nodes:** These are the intermediate control points that guide the direction and curvature of the line connecting the anchor points.

A quadratic Bezier curve, for instance, is defined by exactly three points: the start anchor point, one control point (handle), and the end anchor point. This directly satisfies the

condition of using "at least three points". Cubic Bezier curves use four points (two anchor points and two handles), which also meets the criteria.

Why Bezier Curves Fit the Description

The fundamental structure of Bezier curves involves anchor points defining the endpoints and handle points dictating the curve's shape. The mention of "at least three points" is key, as the simplest forms (quadratic Bezier) use exactly three, and more complex forms use more. The terminology used (anchor points, handles, nodes) is standard for Bezier curves.

Other options like Bicorn, Deltoid, and Kappa are different types of curves or mathematical shapes with distinct definitions and properties, not typically defined by the anchor/handle point system described.

100. Answer: d

Explanation:

Analyzing Box Weights and Total Combinations

The question asks which of the given total weights cannot be formed by combining boxes that weigh 4 kg, 7 kg, and 10 kg. We are given three specific weights: 4 kg, 7 kg, and 10 kg. The phrase "combination of these boxes" typically implies taking a subset of these distinct boxes. We need to find all possible sums we can get by picking zero, one, or more of these three *specific* boxes.

Finding Possible Total Weights

Let the weights of the three boxes be $w_1 = 4$ kg, $w_2 = 7$ kg, and $w_3 = 10$ kg. We can form combinations by choosing a subset of these boxes and summing their weights. The possible subsets and their corresponding total weights are:

- Choosing no boxes: Total weight = 0 kg (empty combination).
- Choosing one box:
 - Box 1: Total weight = 4 kg.
 - Box 2: Total weight = 7 kg.
 - Box 3: Total weight = 10 kg.

- Choosing two boxes:
 - Box 1 and Box 2: Total weight = $4 + 7 = 11$ kg.
 - Box 1 and Box 3: Total weight = $4 + 10 = 14$ kg.
 - Box 2 and Box 3: Total weight = $7 + 10 = 17$ kg.
- Choosing three boxes:
 - Box 1, Box 2, and Box 3: Total weight = $4 + 7 + 10 = 21$ kg.

The complete set of possible total weights from combinations of these specific three boxes is $\{0, 4, 7, 10, 11, 14, 17, 21\}$ kg.

Possible Total Weights from Combinations

Boxes Included	Calculation	Total Weight (kg)
None	0	0
4 kg	4	4
7 kg	7	7
10 kg	10	10
4 kg, 7 kg	$4 + 7$	11
4 kg, 10 kg	$4 + 10$	14
7 kg, 10 kg	$7 + 10$	17
4 kg, 7 kg, 10 kg	$4 + 7 + 10$	21

Checking the Given Options

Now, we compare the given options for total weight against the list of possible total weights $\{0, 4, 7, 10, 11, 14, 17, 21\}$:

- Option 1: 14 kg. Is 14 in the list? Yes, $14 = 4 + 10$. So, 14 kg **can** be a total weight.
- Option 2: 21 kg. Is 21 in the list? Yes, $21 = 4 + 7 + 10$. So, 21 kg **can** be a total weight.
- Option 3: 17 kg. Is 17 in the list? Yes, $17 = 7 + 10$. So, 17 kg **can** be a total weight.
- Option 4: 18 kg. Is 18 in the list? No, 18 is **not** in the list $\{0, 4, 7, 10, 11, 14, 17, 21\}$. So, 18 kg **cannot** be a total weight from combining these three boxes.

Conclusion

Based on the possible combinations of the three given boxes with weights 4 kg, 7 kg, and 10 kg, the total weight of 18 kg cannot be formed.

Revision Table: Box Weight Combinations

Reviewing Possible Weights

Option (kg)	Is it a Possible Sum?	Combination (if possible)
14	Yes	$4 + 10$
21	Yes	$4 + 7 + 10$
17	Yes	$7 + 10$
18	No	N/A

Additional Information: Understanding Weight Combination Problems

This type of problem is related to finding subset sums. When you have a specific set of items (like boxes with given weights), a "combination" usually means selecting one or more items from that set. If the problem allowed using multiple boxes of the same weight type (e.g., using two 4 kg boxes), the problem would become a variation of the Change-making problem or Frobenius Coin Problem, where you determine which total values can be formed using a given set of coin denominations (or weights) any number of times. However, the phrasing "combination of these boxes" usually points to using each distinct box at most once unless otherwise specified.

In cases where you can use multiple instances of each weight, you would look for solutions to the equation $4a + 7b + 10c = W$, where a, b, c are non-negative integers representing the number of boxes of each weight used, and W is the target total weight. For the given weights 4, 7, and 10, any sufficiently large integer weight can be formed (since the greatest common divisor of 4, 7, and 10 is 1), but smaller weights may not be possible. However, for this specific problem, the simpler subset sum interpretation is the most fitting given the options and standard problem phrasing.

101. Answer: c

Explanation:

Understanding Parallel Screw Ring Gauges and Inspection

Parallel screw ring gauges are used to check external screw threads. They are designed to verify the dimensional accuracy of threaded parts, specifically their functional size and form within specified tolerances. Inspecting these gauges themselves is crucial to ensure they remain accurate and capable of performing their verification task correctly.

Why Comparators are Best for Inspecting Screw Ring Gauges

The question asks for the best gauge or machine to inspect parallel screw ring gauges. While various tools exist for thread measurement, comparators are widely used for the accurate inspection of gauges and precision components. Here's why comparators are considered suitable for this task:

- **Precision Measurement:** Comparators work by comparing a measured dimension to a known standard, amplifying the difference. This allows for very precise measurement of variations in dimensions like major diameter, minor diameter, pitch diameter, and lead of the thread form on the ring gauge.
- **Detailed Analysis:** Unlike simple 'go/no-go' gauges, a comparator provides a quantitative measurement, allowing inspectors to know exactly how much a dimension deviates from the nominal size.
- **Versatility:** With appropriate fixtures and probes, comparators can measure various elements of a screw thread profile, which is essential for a thorough inspection of a screw ring gauge.

While other gauges mentioned have their uses, they are typically not the primary or best tool for the detailed inspection of the thread form and dimensions of a precision gauge like a parallel screw ring gauge.

Analysis of Other Options

Let's look at why the other options are not the best choice for inspecting a parallel screw ring gauge:

- **Thread gauge:** This is a broad term. It could refer to thread plug gauges, ring gauges, or even screw pitch gauges. If it refers to plug or ring gauges, they are typically used for checking the workpiece, not for detailed inspection of the gauge itself. A screw pitch gauge is for visually checking the pitch, not for precise dimensional measurement of a gauge.
- **Plug gauge:** Plug gauges are used to check internal dimensions, specifically holes. A parallel screw ring gauge checks external threads, making a plug gauge irrelevant for its inspection.
- **Ring gauge:** A ring gauge is a type of gauge used to check external dimensions (like the diameter of a shaft). While a parallel screw ring gauge is a specific type of ring gauge, the term 'Ring gauge' here likely refers to a plain ring gauge used for checking plain diameters, not screw threads, or perhaps implies using one ring gauge to check another, which isn't a standard inspection method for dimensional accuracy.

Comparing Inspection Methods

Here is a comparison of the discussed tools in the context of inspecting a parallel screw ring gauge:

Tool	Primary Use	Suitability for Inspecting Parallel Screw Ring Gauge
Comparators	Precision dimensional measurement, comparing to a standard	High - Can accurately measure various thread elements (pitch diameter, etc.) of the gauge.
Thread gauge (general)	Checking workpiece threads (go/no-go) or basic pitch verification	Low to Moderate - Not designed for detailed dimensional inspection of a precision gauge.
Plug gauge	Checking internal holes (go/no-go)	None - Used for internal dimensions.
Ring gauge (plain)	Checking external plain diameters (go/no-go)	None - Used for plain diameters.

Based on their function and precision capabilities, comparators are the most suitable instrument among the given options for conducting a detailed inspection of parallel screw ring gauges to ensure their accuracy.

Conclusion on Parallel Screw Ring Gauge Inspection

Inspecting precision tools like parallel screw ring gauges requires instruments capable of accurate dimensional measurement of complex features like screw threads. Comparators offer the necessary precision and flexibility to measure critical thread parameters effectively, making them the best choice among the given options for this inspection task.

Revision Table: Gauge Inspection Methods

Gauge Type	What it Checks	Typical Inspection Tool
Plain Plug Gauge	Internal plain holes	Comparator (for detailed inspection), Micrometer, Bore Gauge
Plain Ring Gauge	External plain shafts	Comparator (for detailed inspection), Micrometer, Snap Gauge
Thread Plug Gauge	Internal screw threads	Comparator (with thread attachments), Thread Measuring Machine
Parallel Screw Ring Gauge	External screw threads	Comparators (with thread attachments), Thread Measuring Machine

Additional Information: Metrology and Gauge Inspection

Metrology is the scientific study of measurement. In manufacturing, metrology is critical for quality control. Gauges are essential tools in metrology, used for checking dimensions and tolerances of manufactured parts. Regular inspection and calibration of gauges are necessary to maintain their accuracy and ensure the quality of the products being checked. Precision instruments like comparators, coordinate measuring machines (CMMs), and specialized thread measuring machines are used in gauge calibration laboratories for this purpose. The specific method and equipment used for inspecting a gauge depend on the type of gauge, its tolerance class, and the parameters that need verification (e.g., size, form, pitch, lead).

102. Answer: a

Explanation:

Understanding Carbon Content in Steel Composition

Steel is a crucial material used widely in construction, manufacturing, and many other industries. It is primarily an alloy of iron and carbon. The amount of carbon present in the alloy significantly affects the properties of the resulting steel, such as its strength, hardness, and ductility.

What is Steel?

At its core, steel is iron with a small amount of carbon. The carbon atoms fit into the spaces between the iron atoms in the crystal lattice, altering the structure and behavior of the material. Other alloying elements can also be added to create different types of steel with specific properties, but carbon is the primary strengthening agent in plain carbon steel.

Carbon Percentage in Steel

The defining characteristic that differentiates steel from cast iron or pure iron is the percentage of carbon. For a material to be classified as steel, the carbon content must fall within a specific range. This range is generally considered to be between 0.02% and 2.1% by weight. Materials with less than 0.02% carbon are typically considered pure iron, and materials with more than 2.1% carbon are classified as cast iron.

The exact upper limit can vary slightly depending on the definition or standard used, but 2% is a commonly accepted upper boundary for steel.

Analyzing the Options

Let's look at the given options based on our understanding of steel composition:

- **Option 1: Less than 2%**

This option aligns with the typical definition of steel, where the carbon content is significantly less than 2.1% and often kept below 2% for most common types of steel. This seems plausible.

- **Option 2: More than 50%**

A material with more than 50% carbon would not be metal; it would be a form of carbon like graphite. This is far outside the range for any metal alloy, especially steel.

- **Option 3: Greater than 2%**

Materials with carbon content greater than about 2.1% are classified as cast iron, not steel. Cast iron has different properties than steel; it is typically harder, more brittle, and has a lower melting point.

- **Option 4: 100%**

100% carbon is pure carbon (like diamond or graphite). This is not steel, which is an alloy primarily of iron.

Based on the standard definition of steel, the carbon content is relatively low, typically well below 2.1%. Therefore, "Less than 2%" is the most accurate description among the given options for the percentage of carbon in steel.

Material Type	Approximate Carbon Percentage
Pure Iron	Less than 0.02%
Steel	0.02% to 2.1% (commonly less than 2%)
Cast Iron	Greater than 2.1% (typically 2.5% to 4.5%)

Your Personal Exams Guide

Conclusion on Carbon Content in Steel

The percentage of carbon is critical in determining whether an iron alloy is considered steel. Steel is characterized by a relatively low carbon content compared to cast iron. The range is generally considered to be between 0.02% and 2.1%, and many common steels have carbon percentages significantly below 2%.

Revision Table: Carbon in Steel

Concept	Key Point
What is Steel?	Alloy of iron and carbon (primarily).
Role of Carbon	Strengthens iron, affects properties (hardness, ductility).
Carbon Percentage in Steel	Typically 0.02% to 2.1%; commonly < 2%.
> 2.1% Carbon	Classified as Cast Iron, not steel.
< 0.02% Carbon	Considered pure iron.

Additional Information: Types of Steel and Carbon Content

Steel is further categorized based on its carbon content:

- **Low-Carbon Steel (Mild Steel):** Contains 0.02% to 0.3% carbon. It is ductile, easily welded, and formable. Used for car bodies, structural shapes, and pipes.
- **Medium-Carbon Steel:** Contains 0.3% to 0.6% carbon. Harder and stronger than low-carbon steel. Used for machine parts, shafts, and railway tracks.
- **High-Carbon Steel:** Contains 0.6% to 2.1% carbon. Very hard and strong, but less ductile. Used for tools, springs, and high-strength wires.

Even high-carbon steel generally stays within the 0.6% to 1.5% range for most practical applications, rarely reaching the theoretical upper limit of 2.1% unless specifically treated. Therefore, the statement "Less than 2%" is a very accurate general description for the carbon percentage in steel.

103. Answer: c

Explanation:

Understanding the Role of a Jig in Production

In manufacturing and production processes, tools like jigs and fixtures are crucial for ensuring accuracy, repeatability, and efficiency, especially when dealing with multiple

identical parts. These devices help in properly locating and holding the workpiece during machining operations.

What is a Jig?

A jig is a type of tooling used in manufacturing to control the location and motion of a cutting tool relative to a workpiece. While it also holds the workpiece securely in place, its primary and defining function is to guide the tool accurately to the desired location on the workpiece.

Primary Function of a Jig in Manufacturing

The essential feature that distinguishes a jig is its ability to guide the cutting tool (like a drill bit, reamer, or milling cutter). This guidance ensures that the tool operates precisely where needed on the workpiece, leading to consistent dimensions and tolerances across all manufactured parts. Jigs are commonly used for operations such as drilling, boring, tapping, and reaming.

Analyzing the Given Options for Jig Function

Let's evaluate each option based on our understanding of what a jig does in production:

- Option 1: holds the cutting tool and guides the work piece
This statement is incorrect. A jig holds the workpiece and guides the *cutting tool*, not the workpiece itself relative to the tool path.
- Option 2: guides both work piece and cutting tool
This statement is also inaccurate. While the jig fixture holds and locates the workpiece (positioning it correctly), it doesn't actively 'guide' the workpiece's movement during the operation. Its key guidance role is for the cutting tool.
- Option 3: guides the cutting tool
This statement accurately describes the primary function of a jig. The jig incorporates features, such as bushings or templates, that direct the cutting tool's path and depth.
- Option 4: holds the work piece and does not guide the cutting tool
This statement is incorrect. While a jig certainly holds the workpiece, its crucial function, which differentiates it from a simple fixture, is the guidance it provides to the cutting tool.

Conclusion on Jig's Role

Based on the analysis, the most precise description of a jig's function is that it guides the cutting tool relative to the workpiece, in addition to holding the workpiece. This guidance is key to achieving consistent accuracy and repeatability in production.

Revision Table: Jig vs. Fixture Comparison

Feature	Jig	Fixture
Primary Function	Guides the cutting tool	Holds and locates the workpiece
Tool Guidance	Yes, typically includes bushings or templates	No, tool path is controlled by the machine tool
Mobility	Often lighter, may be moved with the workpiece	Usually heavier, fixed to the machine bed
Application	Drilling, reaming, tapping, boring	Milling, turning, grinding, welding

Additional Information on Production Jigs and Tooling

Jigs are essential elements in mass production and batch production environments. They help reduce marking out time, minimize errors caused by human variability, and improve the interchangeability of parts. The design of a jig depends on the specific machining operation, the shape and size of the workpiece, and the required accuracy. Investing in well-designed jigs and tooling is vital for achieving high productivity and consistent quality in manufacturing processes.

104. Answer: c

Explanation:

Understanding Aluminium's Mechanical Properties

Mechanical properties describe how a material behaves when subjected to external forces. These properties are crucial for determining a material's suitability for various

applications.

Analyzing Aluminium's Key Characteristic

Aluminium is a light, versatile metal known for many properties, including corrosion resistance and electrical conductivity. Among its mechanical properties, some stand out. Let's examine the options provided:

- **Strongness:** While aluminium alloys can be very strong, pure aluminium is relatively less strong compared to materials like steel. Its strength-to-weight ratio is excellent, but "strongness" as a standalone main property might be misleading for pure aluminium.
- **Brittleness:** A brittle material breaks easily without significant deformation. Aluminium is known for its ductility and malleability, meaning it can be deformed (stretched or hammered) significantly without breaking. Therefore, brittleness is the opposite of a main property of aluminium.
- **Softness:** Pure aluminium is relatively soft. This means it is easily scratched, indented, and deformed. This inherent softness contributes to its high ductility and malleability, making it easy to process and shape.
- **Toughness:** Toughness is the ability to absorb energy and deform plastically before fracturing. Aluminium alloys can exhibit good toughness, but the fundamental characteristic enabling much of its formability is its relative softness and resulting ductility/malleability.

Why Softness is a Main Property of Aluminium

Considering the fundamental properties of pure aluminium, its relative **Softness**, coupled with high ductility and malleability, is a defining mechanical characteristic. This makes it easy to work with, form into complex shapes, and is the basis for many of its uses, such as in foils and wiring. While alloying significantly increases its strength and hardness, the base metal's softness is a primary trait.

Comparison of Relevant Properties

Property	Description	Aluminium (Pure)
Strength	Resistance to deformation under load	Moderate (Low compared to steel, but good for its weight)
Brittleness	Tendency to fracture with little deformation	Very Low (Highly ductile)
Softness	Resistance to scratching or indentation	Relatively High (Easy to scratch/indent)
Toughness	Ability to absorb energy before fracture	Moderate (Improves significantly with alloying)

Revision Table: Aluminium Mechanical Properties

Property Term	Relevance to Aluminium
Strongness	Primarily achieved through alloying; pure Al is less strong.
Brittleness	Aluminium is the opposite; it is highly ductile and malleable.
Softness	A key characteristic of pure aluminium, facilitating forming processes.
Toughness	Good in alloys, allows energy absorption before fracture.

Additional Information: Alloying and Hardening Aluminium

It is important to understand that while pure aluminium is soft, its mechanical properties, especially strength and hardness (which is related to resistance to softness/indentation), can be greatly enhanced by alloying it with other elements and through heat treatments. For example, duralumin, an alloy of aluminium, copper, magnesium, and manganese, is much stronger and harder than pure aluminium. However, the question asks for a main mechanical property of aluminium in general, and softness (leading to ductility and malleability) is a fundamental trait of the element itself before extensive processing.

Related Mechanical Concepts

- **Ductility:** The ability of a material to undergo significant plastic deformation without fracture when subjected to tensile stress. Aluminium is highly ductile.
- **Malleability:** The ability of a material to undergo significant plastic deformation without fracture when subjected to compressive stress (like hammering or rolling). Aluminium is highly malleable.
- **Hardness:** The resistance of a material to permanent indentation or scratching. Soft materials have lower hardness.

105. Answer: b

Explanation:

Understanding Keys in Shaft and Pulley Connections

Mechanical keys are crucial components used to connect rotating machine elements like pulleys, gears, and sprockets onto a shaft. Their primary function is to prevent relative rotation between the shaft and the component mounted on it, ensuring that torque is transmitted efficiently.

Different types of keys are designed for specific applications, depending on factors such as the amount of torque to be transmitted, the need for easy assembly/disassembly, and the manufacturing complexity.

Analyzing Key Types for Pulley Connections

Let's examine the characteristics of the key types provided in the options to determine which is best suited for connecting a pulley with a shaft:

Key Type	Description & Application	Suitability for Pulley on Shaft
Woodruff key	A semicircular key that fits into a corresponding semicircular keyseat in the shaft. It allows some self-alignment between the key and the mating keyway in the hub. Often used in machine tool and automotive applications.	Less common for large, heavy-duty pulley connections where high torque transmission is required, although used in some smaller applications. The keyseat depth on the shaft can weaken it.
Gib head key	A rectangular or square tapered key with a "gib head" (a prominent head) at one end. This head makes driving the key in and out easier during assembly and disassembly. The taper ensures a tight fit.	Highly suitable. The taper provides a secure fit, and the gib head allows for easy removal, which is often necessary for maintenance or replacement of pulleys. Widely used for general engineering applications involving pulleys, gears, etc.
Flat saddle key	A rectangular key that fits into a keyway in the hub but has a flat surface that rests on the surface of the shaft. It relies solely on friction to transmit torque.	Not suitable. This key relies on friction and is only used for light-duty applications where minimal torque is transmitted. It does not require a keyway on the shaft, but this also means it cannot handle significant loads.
Taper Key (without gib head)	A rectangular or square key that has a taper along its length. It fits into keyways in both the shaft and the hub, with the keyway in the hub also tapered. The taper ensures a tight fit when driven in.	Suitable for pulley connections and transmits torque effectively. However, removing a standard taper key can be difficult, especially after it has been in service for a while, making the gib head key often preferred for components requiring frequent removal like pulleys.

Why the Gib Head Key is Often Preferred for Pulleys

The Gib head key combines the benefits of a tapered key (secure fit and torque transmission) with the added feature of a head for easy extraction. Pulleys sometimes need to be removed for maintenance, belt changes, or replacement. The gib head allows a technician to easily tap or pull the key out, simplifying the disassembly process compared to a standard tapered key which might require a special extractor or can be stubborn to remove.

While a standard taper key can also be used, the convenience offered by the gib head makes the Gib head key a very common choice for connecting pulleys and similar components to shafts.

Conclusion on Pulley–Shaft Key Selection

Based on the function of various keys and their typical applications, the Gib head key is a widely used and practical option for connecting a pulley with a shaft, particularly because of its ease of installation and removal facilitated by the gib head.

Revision Table: Key Types and Applications

Key Type	Keyseat on Shaft?	Torque Capacity	Ease of Removal
Woodruff key	Yes (Semicircular)	Moderate	Moderate
Gib head key	Yes (Rectangular/Square, tapered)	High	Easy (due to head)
Flat saddle key	No (Rests on surface)	Very Low (Friction)	Easy
Taper Key	Yes (Rectangular/Square, tapered)	High	Difficult (without head)

Additional Information on Mechanical Keys

Keys are standardized mechanical components. The fit between the key, keyway (in the shaft), and keyseat (in the hub/pulley) is critical for proper torque transmission and to prevent backlash or vibration. There are different classes of fits (e.g., clearance, transition, interference) depending on the required precision and application.

Other key types exist, such as:

- **Parallel Key:** A rectangular or square key with no taper, relying on an interference fit for torque transmission.
- **Splines:** Multiple keys machined integrally with the shaft, fitting into corresponding grooves in the hub. Used for high torque applications and where axial movement is sometimes required (like gearboxes).
- **Pin Keys:** Dowel pins or tapered pins used for light torque or positional alignment.

The choice of key type also depends on manufacturing cost, required precision, and the operating conditions like speed and load.

106. Answer: c

Explanation:

Understanding the Boring Process in Machining

The question asks about the definition of the machining process called boring. Boring is a fundamental operation performed on machine tools like lathes or boring machines.

Let's analyze the options provided to understand what boring specifically refers to:

- **finishing the hole:** Boring can be used as a finishing operation to achieve a precise diameter and good surface finish in an existing hole. However, finishing is an *application* of boring, not its primary definition.
- **making external threads:** Making external threads is done through processes like turning with a threading tool, die cutting, or thread milling. This is completely different from boring.
- **enlarging a hole:** Boring is primarily defined as the process of enlarging or truing an existing hole. It is typically performed on holes that have been previously drilled, cast, or punched. A single-point cutting tool is used to remove material from the inner surface of the hole.
- **making a hole:** Making a hole from solid material is typically done using a drilling process with a drill bit. Boring starts with an already existing hole.

Based on the standard definitions in machining, the primary purpose and definition of the boring process is to enlarge an existing hole or to improve its accuracy and surface finish.

Let's compare boring with drilling:

Process	Primary Action	Starting Point	Tool Type	Accuracy/Finish
Drilling	Making a hole	Solid material	Multi-point (drill bit)	Generally lower accuracy/finish
Boring	Enlarging/Truing a hole	Existing hole	Single-point	Generally higher accuracy/finish

Therefore, the most accurate description of the boring process among the given options is "enlarging a hole".

Revision Table: Boring Process Basics

Concept	Description
Boring Definition	Enlarging an existing hole using a single-point cutting tool.
Purpose	Increase hole diameter, improve accuracy, achieve better surface finish, correct hole location.
Tooling	Boring bars with single-point inserts.
Machines Used	Lathes, milling machines (with boring head), dedicated boring machines.

Additional Information: Types of Boring

Boring operations can be categorized based on the machine or method used:

- **Through Boring:** Creating a bore that goes completely through the workpiece.
- **Blind Boring:** Creating a bore that does not go completely through the workpiece.
- **Back Boring:** Boring on the back side of a feature, often requiring special tooling.
- **Line Boring:** Boring multiple in-line holes simultaneously using a long boring bar supported at both ends.

- **Bore Finishing:** Using boring for final sizing and surface finish, often after rough boring or drilling.

Understanding the difference between drilling and boring is key to understanding various machining operations.

107. Answer: d

Explanation:

Understanding Nut Shapes in Fasteners

Nuts are fundamental components in fastening systems, designed to be paired with a bolt or screw to mechanically join materials together. They typically have a threaded hole that mates with the threading of the bolt or screw. The outer shape of the nut is crucial for its function, particularly how it is tightened or loosened.

Common Nut Shapes and Their Purpose

While nuts come in various specialized shapes for specific applications, some shapes are far more common than others in general use fasteners. The shape is primarily chosen based on ease of manufacturing, ease of use with tools (like wrenches), and suitability for different load-bearing requirements.

- **Flat Nuts:** This isn't a standard shape for the main body of a common nut. Nuts have thickness and a defined outer profile.
- **Round Nuts:** Some specialized nuts, like coupling nuts or cap nuts, might have a round outer profile, but they often have features like knurling or holes for specific tools, or their function is different from a standard fastening nut tightened by an external tool on its flats. A simple smooth round nut would be difficult to grip and tighten effectively.
- **Square Nuts:** Square nuts exist and are used in certain applications, often in older construction or specific types of channels or assemblies. They provide four points for a wrench to grip.
- **Hexagonal Nuts:** This shape has six sides or "flats". Hexagonal nuts are extremely prevalent across almost all industries where fasteners are used.

Why Hexagonal is the Most Common Shape

The hexagonal shape offers several key advantages that make it the most common:

- **Ease of Manufacturing:** Hexagonal shapes are relatively straightforward to manufacture using common processes like cold forming or machining.
- **Tool Compatibility:** Hex nuts are perfectly suited for use with standard wrenches (open-end, box, or adjustable) and sockets. The six flats provide multiple angles (every 60 degrees) from which a tool can engage the nut, which is particularly helpful in tight spaces where full rotation might not be possible.
- **Grip and Torque:** The six flats provide a secure grip for the tool, allowing significant torque to be applied without slipping or rounding the corners of the nut.

Considering the options and the widespread use of standard tools like wrenches and sockets, the hexagonal shape is overwhelmingly the most common shape for general-purpose nuts.

Revision Table

Nut Shape	Commonality	Typical Use Case	Tooling
Hexagonal	Most Common	General fastening, construction, machinery	Standard wrenches, sockets
Square	Less Common (than Hex)	Specific channels, older applications	Standard wrenches
Round	Specialized	Coupling, cap nuts, specific tools (knurling/holes)	Specialized tools, pliers, or hand tightening
Flat	Not applicable as a main shape	N/A	N/A

Additional Information on Fastener Shapes

While hexagonal is the most common for nuts, bolts heads also commonly feature hexagonal shapes for similar reasons. Other less common nut types include:

- **Wing Nuts:** Designed for hand tightening without tools.
- **Cap Nuts (Acorn Nuts):** Have a dome top, often for aesthetic reasons or to protect the bolt threads. Often hexagonal underneath the cap.
- **Lock Nuts:** Designed with features (like nylon inserts or distorted threads) to resist loosening from vibration. Their outer shape is usually hexagonal.
- **Flange Nuts:** Have a wide flange at the base to distribute the load over a larger area, often hexagonal.

The shape of a nut is primarily determined by its intended application, required torque, and the type of tool needed for installation.

108. Answer: b

Explanation:

Understanding Metal Cutting Coolants

When metal is cut or machined, a lot of heat is generated due to friction between the tool and the workpiece, and the deformation of the metal. This heat can damage the tool, affect the workpiece quality, and even cause thermal distortion. To manage this heat and improve the machining process, coolants are used. Coolants serve multiple purposes, including cooling, lubrication, and chip removal.

Why Soluble Oil is Used with Water for Metal Cutting Coolants

The question asks which substance, when mixed with water, provides a coolant for metal cutting. Among the given options, **soluble oil** is specifically designed for this purpose.

Soluble oil, also known as emulsifiable oil or water-miscible cutting fluid, forms an emulsion when mixed with water. An emulsion is a mixture where tiny droplets of oil are dispersed throughout the water. This mixture combines the cooling properties of water with the lubricating and rust-inhibiting properties of the oil.

Properties of Soluble Oil Emulsions as Coolants:

- **Cooling:** Water is an excellent cooling medium because of its high specific heat capacity. The large amount of water in the emulsion efficiently carries away the heat generated during cutting.
- **Lubrication:** The oil particles in the emulsion provide lubrication between the cutting tool and the workpiece. This reduces friction, which in turn reduces heat generation and tool wear.
- **Rust Prevention:** Soluble oils contain additives that protect both the machine tools and the workpiece from rusting or corrosion, which water alone would cause.
- **Chip Removal:** The fluid helps to wash away the metal chips produced during cutting, keeping the cutting zone clear.

Analysis of Other Options

Let's look at why the other options are typically not used in this specific way for general metal cutting coolants:

- **Vegetable oil and Coconut oil:** These are natural oils. While some natural oils can be used as cutting fluids (sometimes in modified forms or blends), they are often used neat (undiluted) for specific applications or as additives. Simply mixing them with water usually doesn't create a stable emulsion suitable for effective metal cutting cooling and lubrication, and they may lack the necessary additives for rust prevention and stability.
- **Mineral oil:** Mineral oils are petroleum-based oils. Mineral oils can be used as cutting fluids, often neat (straight oils) or as the base for other types of cutting fluids like soluble oils or semi-synthetics. However, simply mixing a regular mineral oil that isn't specifically formulated as a soluble oil with water will not form a stable emulsion and would not function effectively as a water-miscible coolant. Soluble oils are special formulations of mineral oil (or other base oils) with emulsifiers and additives.

Therefore, **soluble oil**, when mixed with water, is the correct substance among the options that provides a coolant widely used for metal cutting operations.

Revision Table: Coolant Types for Metal Cutting

Coolant Type	Base	Typical Use	Notes
Soluble Oil (Emulsifiable Oil)	Mineral oil + Emulsifiers + Water	General machining, turning, milling, drilling	Good cooling, good lubrication, rust protection. Forms milky emulsion.
Straight Oil	Mineral oil or other base oil (used neat)	Heavy-duty cutting, threading, tapping, broaching	Excellent lubrication, poor cooling. No water added.
Semi-Synthetic Fluid	Mineral oil + Synthetic chemicals + Water	Medium-duty machining	Good balance of cooling and lubrication. Forms translucent emulsion.
Synthetic Fluid	Synthetic chemicals + Water (no mineral oil)	Grinding, high- speed machining	Excellent cooling, poor lubrication (additives improve this). Forms clear solution.

Additional Information: Role of Coolants in Machining

Coolants, often called cutting fluids, play a vital role in the efficiency, tool life, and surface finish of machining processes. Beyond cooling and lubrication, they also help to flush away chips, reducing the risk of chip re-cutting and improving the overall process stability.

The concentration of soluble oil in water is crucial. Too little oil reduces lubrication and rust prevention; too much can cause foaming and other issues. The typical concentration for machining is between 3% and 10% oil in water, depending on the specific operation and material.

The choice of coolant depends on several factors, including the material being machined, the type of machining operation, the machine tool capabilities, and environmental considerations.

109. Answer: d

Explanation:

Understanding Cutting Fluids in Machining Operations

Cutting fluids, also known as coolants or lubricants, are essential in many machining operations. They are applied to the cutting tool and workpiece during processes like turning, milling, drilling, and grinding. Their primary functions are to cool the cutting zone, lubricate the interface between the tool and chip/workpiece, flush away chips, and protect surfaces from corrosion.

Why Use Cutting Fluid? Key Functions

Using an appropriate cutting fluid provides several benefits:

- **Cooling:** Removing heat generated by friction and plastic deformation helps maintain tool hardness and prevents thermal damage to the workpiece.
- **Lubrication:** Reducing friction between the tool, workpiece, and chip minimizes wear, lowers cutting forces, and improves surface finish.
- **Chip Removal:** The fluid flow helps wash away chips, preventing them from interfering with the cutting process or damaging the surface.
- **Corrosion Prevention:** Many fluids contain additives to protect the tool and workpiece surfaces from rust and corrosion.

Common Types of Cutting Fluids

Various types of cutting fluids are available, each with different properties and applications. Let's look at the options provided:

- **Lard oil:** This is a natural fatty oil, historically used. While it provides good lubrication, it can become rancid and is expensive. It's not commonly used alone in modern machining.
- **Synthetic Oil:** These are water-based fluids that contain no mineral oil. They offer excellent cooling capabilities and good rust inhibition but typically provide less lubrication than oils.
- **Olive oil:** Another natural oil, primarily known for culinary use. It is not a standard cutting fluid due to cost and suitability for industrial machining environments.
- **Soluble oil:** These are oil-in-water emulsions, meaning small droplets of mineral oil are dispersed in water with the help of emulsifiers. They appear milky when mixed

with water.

Analyzing Soluble Oil as a Common Cutting Fluid

Soluble oil is widely considered the most common type of cutting fluid used in machining operations. Here's why:

- **Balance of Properties:** It provides a good balance of both cooling (due to the high water content) and lubrication (due to the mineral oil).
- **Cost-Effective:** Mixing concentrated soluble oil with water makes it relatively inexpensive compared to straight oils or many synthetic fluids.
- **Versatility:** Suitable for a wide range of machining operations and materials.
- **Ease of Use:** Easily mixed with water to the desired concentration.

While other fluids excel in specific areas (e.g., synthetic fluids for pure cooling, straight oils for heavy-duty lubrication), soluble oil's combination of cost, performance, and versatility makes it the go-to choice for general machining applications.

Conclusion: Identifying the Most Common Cutting Fluid

Considering the properties, cost-effectiveness, and widespread application across various machining tasks, soluble oil stands out as the most common type of cutting fluid used in almost all machining operations.

Revision Table: Cutting Fluid Types

Cutting Fluid Type	Key Characteristics	Commonality
Lard oil	Natural fatty oil, good lubrication, can spoil	Low (historically significant)
Synthetic Oil	Water-based, excellent cooling, less lubrication	Moderate (increasingly used for cooling-intensive tasks)
Olive oil	Natural oil (culinary), not standard industrial use	Very Low
Soluble oil	Oil-in-water emulsion, good balance of cooling & lubrication, cost-effective	High (Most Common)

Additional Information on Cutting Fluids

Choosing the right cutting fluid depends on several factors, including the material being machined, the type of machining operation, the desired surface finish, and environmental considerations. Properly managing cutting fluids, including filtration and maintenance, is crucial for performance, tool life, and operator health and safety. The concentration of soluble oil in water is often adjusted based on the severity of the machining operation; higher concentrations are used for heavier cuts or harder materials to provide more lubrication.

110. Answer: d

Explanation:

Understanding Pipe Cutting Die Construction

A pipe cutting die is a fundamental tool used in plumbing and metalworking to create external screw threads on the end of a pipe. This threading allows the pipe to be connected to other threaded fittings or pipes, forming a secure and leak-proof joint.

The design of a pipe cutting die is crucial for its function. When considering the typical construction, especially for tools used manually or with standard threading machines, the die is built to effectively cut threads onto rounded pipe surfaces.

Typical Configuration: Divided in Two Parts

Pipe cutting dies, particularly those known as 'adjustable dies' or used in common die stocks, are generally constructed by being divided into two separate parts. These two halves are then held together and guided by a die stock or a threading machine head.

This two-part construction offers significant advantages:

- **Adjustability:** Dividing the die into two halves allows the cutting size to be slightly adjusted. The two halves can be moved closer together or further apart within the die stock. This is important for controlling the depth of the thread cut and ensuring a proper fit with mating threads.

- **Ease of Use:** A split die makes it easier to start the threading process on the end of a pipe. The die can be started slightly open and then closed to the correct cutting position once the initial threads begin to form.
- **Manufacturing Simplification:** Manufacturing precise cutting teeth on two smaller pieces can sometimes be simpler than on a single, large, solid piece.

While solid dies (undivided) exist, they cut a fixed size and are less common for general pipe threading where minor adjustments might be needed. Dies divided into more than two parts typically refer to die heads used in industrial machines, which use multiple cutting chasers (often 4 or 5), but the term "pipe cutting die" in a general context usually refers to the two-part adjustable type.

Common Die Constructions for Threading

Construction Type	Description	Typical Use
Undivided (Solid)	Single piece with cutting teeth	Fixed size threading, often for bolts; less common for pipes needing adjustment.
Divided in Two Parts	Two halves held in a stock	Common for manual and machine pipe threading; allows size adjustment.
Divided in Three Parts	Rare for standard dies	Not a typical standard die configuration.
Divided in Four or More Parts	Die head with multiple chasers	Industrial threading machines for high precision and speed.

Conclusion on Pipe Cutting Die Division

Therefore, based on the prevalent design for manual and common machine pipe threading applications, a pipe cutting die is generally divided into two parts to provide adjustability and ease of starting the thread cutting process.

Revision Table: Key Pipe Threading Terms

Summary of Pipe Threading Elements

Term	Explanation	Relevance to Die
Pipe Die Stock	Tool holding the pipe cutting die for manual rotation.	Designed to hold the two halves of the die correctly.
Cutting Oil	Lubricant applied during threading.	Essential for smooth cutting and die longevity.
Threads per Inch (TPI)	Number of threads along one inch.	Dies are specific to the required TPI.
External Thread	Thread cut on the outside surface (e.g., pipe end).	This is what a pipe cutting die creates.

Additional Information: Beyond Basic Pipe Dies

While the two-part die is standard for many applications, threading technology includes variations:

- **Adjustable Dies:** The two-part design is a form of adjustable die, allowing minor changes in the thread diameter.
- **Solid Dies:** Non-adjustable, best for situations where precise, non-variable threads are consistently needed, often on bolts rather than pipes.
- **Receding Dies:** Dies where the chasers move inwards as they cut, often used for tapered pipe threads like NPT, ensuring the thread is cut to the correct taper.
- **Self-Opening Die Heads:** Used in machines, these heads automatically open once the thread is complete, saving time and preventing damage.

Understanding these different types helps appreciate the specific role and design of the general pipe cutting die.

111. Answer: b

Explanation:

Understanding Washer Thickness for Bolt Connections

This question asks about the typical or required thickness of a washer when the nominal diameter of the bolt being used is represented by D . Washers are essential components in bolted connections, providing a bearing surface, distributing load, and sometimes preventing loosening.

Relation Between Washer Thickness and Bolt Diameter

In engineering practice, the dimensions of standard washers are often related to the nominal diameter of the bolt they are intended for. While specific thicknesses can vary based on the type of washer (e.g., plain washer, spring washer) and the specific standard being followed (like ISO, DIN, or national standards), there are common ratios or rules of thumb used for plain washers.

These ratios aim to provide a washer with sufficient strength and bearing area relative to the bolt size without being unnecessarily bulky.

Analyzing the Options

Let's look at the provided options, which suggest a linear relationship between the washer thickness and the bolt nominal diameter D :

- $0.09 D$
- $0.12 D$
- $0.7 D$
- $0.1 D$

These options represent different proportions of the bolt diameter that the washer thickness could be.

Determining Standard Washer Thickness

Based on various engineering standards and common practices for plain washers, the thickness is typically a fraction of the nominal bolt diameter. A widely accepted standard proportion for the thickness of a plain washer is approximately 0.12 times the nominal bolt diameter. This ratio is often found in dimensional standards for general-purpose plain washers used in mechanical and structural applications.

Let's consider why the other options might be less common for a standard plain washer thickness:

- $0.7 D$: A thickness of 0.7 times the diameter would result in a very thick washer, likely unnecessary for most standard applications and adding significant weight and cost. This might be relevant for specialized heavy-duty washers, but not typical.
- $0.09 D$ and $0.1 D$: While closer to the standard, $0.12 D$ is a more frequently cited or standardized value for the typical thickness of plain washers compared to $0.09 D$ or $0.1 D$.

Therefore, the thickness that is typically required or specified for a standard plain washer with a bolt of nominal diameter D is commonly considered to be $0.12 D$.

Revision Table: Key Washer Dimensions

Dimension	Relation to Bolt Diameter (Approx. for Plain Washers)	Typical Ratio
Nominal Bolt Diameter	D	D
Washer Inside Diameter	Slightly larger than D	~ 1.02 to $1.05 D$
Washer Outside Diameter	Larger than inside diameter	$\sim 2 D$ (can vary significantly)
Washer Thickness	Related to D	$\sim 0.12 D$

Your Personal Exams Guide

Additional Information: Types and Functions of Washers

Washers come in various types, each serving specific purposes in bolted connections:

- **Plain Washers:** These are flat rings used to distribute the load of the fastener, prevent damage to the bolted surface, and provide a smooth bearing surface. Their thickness is often standardized based on bolt size.
- **Spring Washers (e.g., Split Lock Washers):** Designed to provide a spring force that helps prevent fasteners from loosening due to vibration or thermal expansion/contraction.
- **Lock Washers (various types):** Include toothed lock washers, conical washers, etc., which create friction or a mechanical interference to resist loosening.

- **Fender Washers:** Plain washers with a much larger outside diameter relative to the inside diameter, used to distribute load over a large area, particularly on soft materials or oversized holes.

The thickness mentioned in the question, $0.12 D$, is a standard value often associated with plain washers.

112. Answer: c

Explanation:

Understanding Knuckle Threads and Their Angle

Let's discuss knuckle threads, a specific type of screw thread, focusing on their characteristic angle. Screw threads are essential components in many engineering applications, used for fastening or transmitting power. Different thread profiles have different angles and shapes tailored for specific purposes.

What is a Knuckle Thread?

A knuckle thread is a thread profile that has a rounded form at both the root and the crest. Unlike sharp V-threads, the rounded shape gives it certain advantages in specific environments.

Angle of Knuckle Threads

The question asks about the angle of a knuckle thread. This angle refers to the angle between the flanks of the thread profile, measured in an axial plane. For a standard knuckle thread, this angle is well-defined.

The standard angle for a knuckle thread is 30° . This angle, combined with the rounded profile, makes knuckle threads suitable for applications where they might encounter rough handling, dirt, or corrosion.

Comparison with Other Thread Types

It's helpful to compare the knuckle thread angle with other common thread types:

- **ISO Metric Thread:** Has a V-profile with a 60° angle. Widely used for general fastening.
- **Whitworth Thread:** Has a V-profile with a 55° angle and rounded roots and crests. Historically significant, still used in some applications.
- **ACME Thread:** Has a trapezoidal profile with a 29° angle. Primarily used for power transmission.
- **Buttress Thread:** Has an asymmetrical profile, typically with one flank perpendicular to the axis and the other at an angle (often 45° or 7°), designed for high axial loads in one direction.

The 30° angle of the knuckle thread is smaller than that of common V-threads (like Metric or Whitworth) and distinct from trapezoidal or buttress threads.

Applications of Knuckle Threads

Due to their rounded profile and 30° angle, knuckle threads are less prone to damage from impacts and are easier to clean. This makes them ideal for applications such as:

- Railway carriage couplings
- Hydrant connectors
- Bulb bases (like Edison screw fittings, although the exact profile can vary)
- Coarse applications where robustness is required

Determining the Correct Angle

Based on engineering standards and common knowledge regarding thread forms, the specific angle associated with the knuckle thread profile is 30° .

The options provided are 55° , 60° , 30° , and 45° . Comparing these to the standard knuckle thread angle, 30° is the correct value.

Thread Type	Profile Shape	Standard Angle	Common Applications
Knuckle Thread	Rounded Crest & Root	30°	Rough handling, dirt exposure (e.g., railway couplings, fire hydrants)
ISO Metric Thread	V-profile	60°	General fastening
Whitworth Thread	V-profile with rounded root & crest	55°	Pipes, some historical machinery
ACME Thread	Trapezoidal	29°	Power transmission (lead screws)
Buttress Thread	Asymmetrical	Typically 45° or 7° (load bearing flank)	High axial loads in one direction

Therefore, the angle of a knuckle thread is 30°.

Revision Table: Common Thread Angles

Thread Type	Thread Angle (θ)
Knuckle Thread	30°
ISO Metric	60°
British Standard Whitworth (BSW)	55°
ACME	29°
Unified (UNC, UNF)	60°

Additional Information: Knuckle Thread Characteristics

Knuckle threads are often used in situations where cleanliness and precision are less critical than robustness and ease of engagement. Their rounded profile makes them more

tolerant of wear and less susceptible to damage than threads with sharp roots and crests. They are typically manufactured by rolling or casting rather than cutting, which is suitable for the often coarse pitch of these threads. While they don't offer the same load-carrying capacity or self-locking properties as some other thread forms, their durability in harsh conditions makes them invaluable for specific heavy-duty applications.

113. Answer: d

Explanation:

Understanding Drainage Pipeline Materials

Drainage pipelines are crucial for carrying waste water away from buildings and areas. The material used for these pipes needs to be durable, resistant to corrosion from various chemicals found in wastewater, and strong enough to withstand external pressures.

Let's examine the suitability of the given options for drainage pipelines:

- **Brass pipe:** Brass is a metal alloy, primarily used for water supply lines, especially in visible fixtures due to its aesthetic appeal and corrosion resistance for clean water. It is generally expensive and not typically used for main drainage lines carrying sewage or greywater.
- **G.I. Pipe:** G.I. stands for Galvanized Iron. These pipes are made of steel coated with zinc to prevent corrosion. While used in older plumbing systems for water supply, they are less common for drainage pipes today due to the potential for internal scaling and corrosion over time, especially with varied wastewater contents.
- **Steel Pipe:** Steel pipes are strong but prone to corrosion if not properly protected. They are used in various industrial applications and sometimes for large-diameter sewers, but require internal lining or coating for corrosion resistance when carrying wastewater. Plain steel is not ideal for typical drainage without protection.
- **Ceramic:** This category includes materials like vitrified clay pipes (VCP). Vitrified clay pipes have been used for gravity drainage and sewer systems for a very long time. They are known for their excellent resistance to chemical attack from acids and alkalis found in sewage, as well as resistance to abrasion. While modern drainage often uses plastic pipes (like PVC or ABS), ceramic materials, particularly vitrified clay, are a recognized material for drainage pipelines, especially in older systems or specific applications.

Considering the options provided and the requirements for drainage pipes, ceramic materials (like vitrified clay) possess properties that make them suitable for drainage pipelines, particularly their chemical resistance which is vital for handling wastewater.

Why Ceramic (Vitrified Clay) is Used for Drainage

Vitrified clay pipes are fired at high temperatures, making them dense, hard, and non-porous. Key advantages for drainage include:

- Excellent chemical resistance: Withstands acids, alkalis, and gases found in sewage.
- Abrasion resistance: Durable against solids transported in wastewater.
- Long service life: Can last for many decades underground.

While other materials like PVC, ABS, and concrete are also widely used for modern drainage systems, among the options listed, Ceramic is the material historically and currently used for drainage pipelines due to its specific properties.

Material Comparison for Drainage Pipes

Material	Suitability for Drainage	Key Property
Brass	Low (expensive, mainly clean water)	Corrosion resistant (clean water)
G.I. Pipe	Moderate (historically used, susceptible to scaling/corrosion)	Galvanized coating (zinc)
Steel Pipe	Moderate (requires lining/coating)	High strength
Ceramic (Vitrified Clay)	High (excellent chemical resistance)	Chemical inertness, durability

Therefore, based on the characteristics required for drainage pipelines and the materials provided in the options, Ceramic is the appropriate choice.

Revision Table: Drainage Pipe Materials

Pipe Material	Typical Use	Drainage Use Suitability
Brass	Clean water supply, fixtures	Not typical
G.I. Pipe	Water supply (older), gas	Less common due to corrosion/scaling
Steel Pipe	Various (requires protection for wastewater)	Requires lining for drainage
Ceramic (Vitrified Clay)	Drainage, sewer systems	High (chemical resistance)

Additional Information: Modern Drainage Pipes

While ceramic pipes (like vitrified clay) are still used and are historically significant for drainage, modern drainage systems frequently utilize plastic pipes such as:

- **PVC (Polyvinyl Chloride)**: Lightweight, easy to install, resistant to corrosion and chemicals, widely used for residential and commercial drainage.
- **ABS (Acrylonitrile Butadiene Styrene)**: Similar properties to PVC, often used for drainage, waste, and vent (DWV) systems.
- **HDPE (High-Density Polyethylene)**: Used for various piping applications, including storm drainage and sewer lines, known for flexibility and durability.

These plastic materials offer advantages in terms of weight, ease of joining, and cost-effectiveness for many drainage applications today.

114. Answer: c

Explanation:

Understanding the Polygon Effect in Drives

The question asks about the Polygon effect and in which type of drive it is observed. The Polygon effect is a phenomenon related to the speed and tension variations in certain types of power transmission systems.

What is the Polygon Effect?

The Polygon effect is a kinematic irregularity that occurs in drives where the flexible element engages with sprockets or toothed wheels. As the engaging element wraps around the sprocket, it doesn't form a perfect circle. Instead, it approximates a polygon shape, with each link or segment acting as one side of the polygon. This polygonal shape causes the effective radius of engagement to change as the sprocket rotates, leading to fluctuations in the velocity of the flexible element.

Why is the Polygon Effect Observed in Chain Drives?

Chain drives are made up of rigid links connected by pins. These links articulate as they wrap around the sprockets. When a chain link engages with a tooth on the sprocket, the center of the link lies on the pitch circle of the sprocket. As the sprocket rotates and the next link engages, the chain forms a chord between the centers of the two engaged pins (or articulation points). The chain thus follows a path that is essentially a series of chords around the sprocket's pitch circle, forming a polygon.

- The number of sides of this polygon is equal to the number of teeth on the sprocket.
- The effective radius of the sprocket, as seen by the chain, fluctuates between the minimum radius (perpendicular distance from the center to a chord) and the maximum radius (distance from the center to a joint).
- This fluctuation in effective radius causes the linear speed of the chain to vary even if the angular speed of the sprocket is constant. This speed variation is the core of the Polygon effect.
- The smaller the number of teeth on the sprocket, the more pronounced the polygonal shape and thus the greater the speed variation and the Polygon effect.

These speed variations can lead to:

- Increased dynamic loads and vibrations in the drive system.
- Noise generation.
- Accelerated wear on the chain and sprockets.

Comparing Different Drive Types

Let's consider the other options provided:

- **Shaft drive:** Shaft drives transmit power through rigid rotating shafts, often connected by couplings or gears. They do not involve flexible elements wrapping around sprockets, so the Polygon effect is not applicable.
- **Gear drive:** Gear drives transmit power through the meshing of teeth on rigid gears. While there are variations in velocity during tooth engagement (known as transmission error), this is a different phenomenon from the Polygon effect seen in chain drives. Gear teeth are designed to minimize such errors.
- **Belt drive:** Belt drives use a flexible belt (like V-belts, flat belts, or timing belts) running on pulleys or sheaves. V-belts and flat belts rely on friction and run smoothly around the pulleys, typically not exhibiting the significant polygonal action seen with rigid links. Timing belts have teeth that engage with grooves on pulleys, but their flexibility and smaller pitch relative to chain links make the Polygon effect less pronounced compared to roller chains, although some level of chordal action exists, especially with smaller pulleys. However, the classic and most significant manifestation of the Polygon effect is associated with the rigid links of roller chains.

Based on the mechanism of power transmission, the Polygon effect is most characteristically and significantly observed in chain drives.

Drive Type	Mechanism	Polygon Effect Presence
Shaft drive	Rigid shafts/couplings	Absent
Gear drive	Meshing rigid teeth	Absent (Velocity variations different)
Chain drive	Rigid links on sprockets	Significantly present
Belt drive	Flexible belt on pulleys	Less significant or absent (depending on belt type)

Conclusion on the Polygon Effect

The kinematic irregularity known as the Polygon effect, causing speed variations and vibrations, is a direct consequence of the polygonal path formed by the rigid links of a chain as it wraps around a sprocket. This phenomenon is characteristic of chain drive systems.

Revision Table: Polygon Effect Key Concepts

- **Definition:** Speed/tension variation due to polygonal chain path on sprockets.
- **Cause:** Rigid links of a chain forming chords around a sprocket's pitch circle.
- **Affected Drive:** Chain drive.
- **Impact:** Vibration, noise, wear, dynamic loads.
- **Mitigation:** Using sprockets with more teeth, specific chain designs, tensioning devices.

Additional Information: Chain Drive Basics

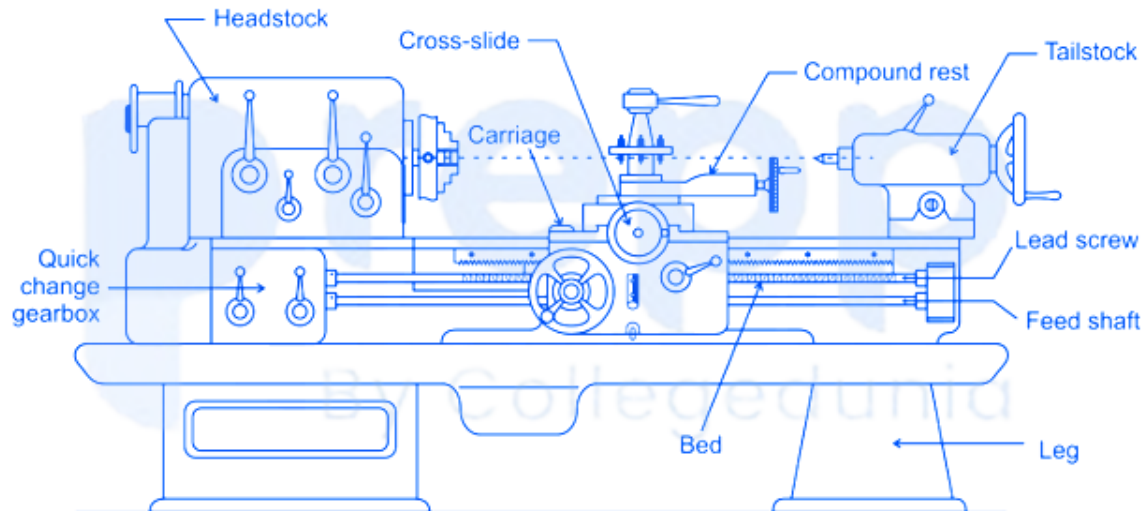
A chain drive system consists of a roller chain or silent chain running over sprockets, which are toothed wheels. It is a positive drive system, meaning there is no slip (under normal operating conditions), maintaining a fixed velocity ratio between the driving and driven shafts. This makes chain drives suitable for applications requiring precise speed synchronization. However, the Polygon effect is an inherent challenge in their design and operation, particularly at higher speeds or with small sprockets.

115. Answer: b

Explanation:

Explanation:

- Headstock is mounted permanently on the inner guide ways at the **left – hand side** of the leg bed.
- The spindle rotates on the two large bearings housed on the headstock casing.
- Tailstock is located on the inner guide ways at the right side of the bed opposite to head stock.
- Carriage is located between the headstock and tailstock on the lathe bed guide ways.
- **The function of Head Stock:**
 - To provide a mean to assemble the work-holding devices.
 - Transmit the drive from the main motor to the work.
 - To accommodate shaft, gears and levers for a wide range of varying work speeds.
 - To ensure arrangement for lubricating the gears, shafts and bearings.
 - The following figure represents various parts of a Lathe machine and their location.



116. Answer: a

Explanation:

Understanding Lubrication: More Than Just Reducing Friction

The primary and most well-known purpose of lubrication in mechanical systems is to reduce friction between moving parts. This helps to minimize wear and tear, reduce heat generation, improve efficiency, and extend the lifespan of components. However, lubrication serves several other important functions beyond just minimizing friction.

Secondary Purposes of Lubrication

While reducing friction is the main goal, lubricants often provide additional benefits. The question asks for a secondary purpose apart from reducing friction. Let's consider the options provided:

- **Reducing corrosion:** Many lubricants contain additives that help protect metal surfaces from corrosion and rust. By forming a protective film and neutralizing corrosive agents, lubricants prevent the degradation of components exposed to moisture, oxygen, or other harsh environments.

- **Limited Motion:** Lubrication facilitates motion by reducing resistance. It does not inherently limit motion; rather, it allows parts to move more freely. Limiting motion is typically achieved through design constraints or other mechanical means, not by the lubricant itself.
- **Maintaining Pressure:** In certain applications, like hydraulic systems, the fluid serves as both a lubricant and a medium for transmitting power under pressure. In such cases, the fluid's viscosity and sealing properties help maintain pressure. However, this is not a universal secondary purpose of lubrication in all mechanical systems (e.g., lubricating a simple bearing).
- **Limited rotation:** Similar to "Limited Motion," lubrication does not typically limit rotation. It makes rotation easier by reducing friction. Limitation of rotation is a function of system design, not the lubricant.

Based on the common functions of lubricants, reducing corrosion is a significant and widely recognized secondary purpose.

Analyzing the Options for Lubrication's Secondary Role

Let's evaluate why 'Reducing corrosion' is the most appropriate secondary purpose among the given options:

Option	Relevance as a Secondary Purpose of Lubrication
Reducing corrosion	Highly relevant. Many lubricants contain anti-corrosion additives and form a protective barrier on metal surfaces.
Limited Motion	Irrelevant. Lubrication facilitates motion, it doesn't limit it.
Maintaining Pressure	Relevant in specific applications (like hydraulics), but not a general secondary purpose of lubrication across all mechanical systems.
Limited rotation	Irrelevant. Lubrication facilitates rotation by reducing friction.

Therefore, apart from reducing friction, reducing corrosion is a key secondary purpose of lubrication.

Revision Table: Lubrication Purposes

Purpose	Description
Primary: Reducing Friction	Minimizes wear, heat generation, and energy loss between moving parts.
Secondary: Reducing Corrosion	Protects metal surfaces from rust and other forms of corrosion.
Secondary: Dissipating Heat	Helps carry away heat generated by friction or external sources.
Secondary: Sealing	Helps prevent the entry of contaminants (dirt, water) or leakage of fluids.
Secondary: Carrying Away Contaminants	Transports wear particles and other debris away from critical areas.
Secondary: Transmitting Power	In hydraulic systems, the fluid acts as a power transmission medium.

Additional Information on Lubricant Functions

Lubrication is a critical aspect of machine design and maintenance, falling under the field of tribology (the science of friction, wear, and lubrication). While friction reduction is primary, the other functions like corrosion protection, heat dissipation, and sealing are vital for system reliability and longevity. Selecting the right lubricant depends on the specific application's requirements, including operating temperature, load, speed, and environmental conditions. Lubricants can be oils, greases, or even solid materials, each offering different properties tailored for various needs.

117. Answer: a

Explanation:

Understanding Tap Drill Size for M10 × 1 mm Metric Threads

When creating an internal thread, like the one for an M10 × 1 mm screw, a hole must first be drilled to a specific size. This size is known as the tap drill size. The tap then cuts the threads into the material of the drilled hole.

For metric threads, the tap drill size is closely related to the nominal diameter and the pitch of the thread.

Calculating Tap Drill Size

The nominal diameter of the M10 × 1 mm thread is 10 mm, and the pitch is 1 mm. The pitch is the distance between adjacent thread crests or roots.

A simplified formula often used to estimate the tap drill size for metric threads is:

$$\text{Tap Drill Size} \approx \text{Nominal Diameter} - \text{Pitch}$$

Using this formula for M10 × 1 mm:

$$\text{Tap Drill Size} \approx 10 \text{ mm} - 1 \text{ mm} = 9 \text{ mm}$$

Why the Tap Drill Size is Often Slightly Larger

While the theoretical minor diameter for an M10 × 1 mm internal thread is 9 mm (Nominal Diameter - Pitch), the standard tap drill size used in practice is typically slightly larger. This is done for several reasons:

- To ensure the tap cuts correctly and doesn't bind.
- To achieve a desired percentage of thread engagement (usually between 75% and 85%). A slightly larger hole results in less than 100% thread engagement, which reduces the torque required for tapping and minimizes the risk of tap breakage, while still providing sufficient thread strength for most applications.

Standard charts for metric tap drill sizes often list specific values that account for these practical considerations.

Standard Tap Drill Size for M10 × 1 mm

Looking up standard tap drill charts, the recommended drill size for an M10 × 1 mm fine pitch metric thread is commonly listed as 9.0 mm or 9.1 mm. The option 9.1 mm aligns with a standard recommended size that provides appropriate thread engagement for this specific thread pitch.

Thread Type	Nominal Diameter (mm)	Pitch (mm)	Standard Tap Drill Size (mm)
M10 × 1 mm	10	1	9.1 (Commonly Used)

Analyzing the Options

- **9.1 mm:** This is a standard and commonly recommended tap drill size for M10 × 1 mm, allowing for correct tapping and appropriate thread engagement.
- **8 mm:** This size is significantly smaller than the theoretical minor diameter (9 mm). Drilling with an 8 mm bit would result in very high thread engagement (>100% theoretically), making tapping extremely difficult, requiring high torque, and greatly increasing the risk of tap breakage.
- **10 mm:** This is the nominal (major) diameter of the thread. Drilling a 10 mm hole would leave no material for the tap to cut threads into; the hole would be too large, and no internal thread would be formed.
- **7.1 mm:** Similar to 8 mm, this size is much too small. It would lead to extremely high thread engagement, making tapping very difficult and prone to tap breakage. This size might be relevant for a different, much coarser pitch thread.

Based on standard engineering practices and tap drill charts, 9.1 mm is the appropriate size for drilling a hole prior to tapping an M10 × 1 mm metric thread.

Revision Table: M10 × 1 mm Tap Drill

Thread Specification	Nominal Diameter	Pitch	Estimated Tap Drill Size (Nominal - Pitch)	Standard Tap Drill Size
M10 × 1 mm	10 mm	1 mm	$10 - 1 = 9 \text{ mm}$	9.1 mm

Additional Information on Threading Concepts

Understanding a few key terms helps in selecting the right drill size:

- **Major Diameter:** The largest diameter of a screw or nut thread. For external threads (screws), it's the outer diameter. For internal threads (nuts or tapped holes), it's the

largest diameter the internal thread will have. For M10, the nominal major diameter is 10 mm.

- **Minor Diameter:** The smallest diameter of a screw or nut thread. For external threads, it's the diameter at the root of the thread. For internal threads, it's the diameter at the root of the internal thread. The tap drill size approximates the minor diameter of the internal thread.
- **Pitch Diameter:** A theoretical diameter between the major and minor diameters where the thread thickness equals the space between threads. It's crucial for thread engagement and strength.
- **Thread Engagement:** This refers to how much the threads of a bolt and a nut/tapped hole overlap. It's usually expressed as a percentage. Higher engagement generally means higher strength, but also requires more tapping torque and increases the risk of tap failure. Tap drill sizes are chosen to achieve a balance, typically 75–85% engagement for general purposes.

118. Answer: c

Explanation:

Understanding Sheet Metal Cutting Tools

The question asks us to identify which of the listed options is NOT a tool typically used for cutting sheet metal. Let's examine each tool:

- **Snip:** A snip is a type of hand shear specifically designed for cutting thin sheet metal. It works by a shearing action, similar to scissors but more robust for metal. Snips are indeed a common sheet metal cutting tool.
- **Power Shear:** A power shear, also known as a shearing machine, is a mechanical tool used to cut sheet metal using a shearing action. It can cut larger sheets and thicker gauges than hand snips. Power shears are definitely used for cutting sheet metal.
- **Milling Cutter:** A milling cutter is a rotary cutting tool used in a milling machine. It removes material from a workpiece by rotating against it and feeding it through a path. While milling can be used to cut profiles or remove sections from sheet metal, its primary action is chip removal through rotation, which is fundamentally different from the shearing action of snips and power shears typically associated with simply cutting sheet metal blanks or strips. Milling is more of a machining process.

- **Mallet:** A mallet is a type of hammer, often with a soft head made of rubber, wood, or plastic. Mallets are used for striking or shaping objects without causing dents or damage to the surface. They are commonly used in sheet metal work for shaping, bending, or seating parts, but they are not cutting tools.

Analyzing the Options for Sheet Metal Cutting

Based on the function of each tool:

Tool	Primary Function	Used for Sheet Metal Cutting?
Snip	Hand shearing/cutting	Yes
Power Shear	Mechanical shearing/cutting	Yes
Milling Cutter	Rotary chip removal/machining	Primarily for machining, not typical shearing/blanking
Mallet	Striking, shaping, bending	No

Both a Milling Cutter and a Mallet are not typical sheet metal cutting tools in the sense of shearing. However, the question asks which is NOT a sheet metal cutting tool among the options. Snips and Power Shears are specifically designed for cutting sheet metal by shearing. A Mallet is used for shaping or striking, not cutting at all. A Milling Cutter cuts material, but through a different rotary chip removal process, and is usually part of a machining operation rather than simple sheet metal cutting like blanking or trimming. Considering the options, the Milling Cutter is the tool among the cutting-related options (excluding the mallet which is clearly not a cutter) that is least specifically associated with the common methods of sheet metal cutting (shearing).

Therefore, the tool that is not a sheet metal cutting tool in the context of typical sheet metal fabrication methods is the Milling Cutter.

Revision Table: Sheet Metal Tools

Tool Type	Function	Example
Cutting Tool	Separates material by shearing, slicing, or chip removal.	Snips, Power Shear, Hacksaw, Plasma Cutter, Laser Cutter, Milling Cutter (for material removal)
Forming Tool	Shapes material without removing it.	Mallet, Hammer, Press Brake, Roller

This table helps clarify the different categories of tools used in metal working, including sheet metal work.

Additional Information on Sheet Metal Processes

Sheet metal work involves various processes beyond just cutting. These include bending, forming, punching, joining (welding, riveting), and finishing. Different tools and machines are used for each process.

- **Shearing:** This is the process used by snips and power shears, where metal is cut between two blades.
- **Blanking:** Cutting a piece out of a sheet (the cut-out piece is the desired part).
- **Punching:** Creating holes in the sheet metal.
- **Forming:** Bending or shaping the sheet metal into desired forms using tools like mallets, press brakes, or rollers.
- **Machining:** Processes like milling or drilling that use cutting tools to remove material and create precise shapes or features. While milling cutters are cutting tools used in machining, they are not the primary tools for the initial cutting of large sheet metal pieces like shears are.

Understanding the specific function of each tool in the context of these processes helps differentiate between sheet metal cutting tools and tools used for other operations like shaping or machining.

119. Answer: d

Explanation:

Understanding Sheet Metal Forming Processes

Sheet metal forming is a manufacturing process where a piece of sheet metal is plastically deformed to create a desired shape without adding or removing material. These processes often involve applying forces that cause the metal to bend, stretch, or shear into the new form.

Let's examine the given options to determine which one fits the description of a sheet metal forming process.

Analysis of Options

We will look at each option and identify the type of manufacturing process it represents:

- **Knurling:** This is a manufacturing process, typically performed on a lathe, to create a pattern on the surface of a part. It involves plastic deformation but is primarily used on solid rods or workpieces, not specifically sheet metal, and is a type of machining or forming depending on the perspective (often considered a cold forming operation for adding texture).
- **Facing:** This is a machining operation performed on a lathe where material is removed from the end surface of a workpiece to create a flat face. It is a material removal process, not a forming process.
- **Drilling:** This is a machining process that uses a rotating cutting tool to create a round hole in a solid material. It is a material removal process, not a forming process.
- **Rolling:** This is a metal forming process where a metal stock is passed through one or more pairs of rollers to reduce the thickness and make it uniform. This process involves plastically deforming the metal. It is widely used in the production of sheet metal and plates from slabs or ingots.

Identifying the Sheet Metal Forming Process

Based on the analysis, we can see that Knurling, Facing, and Drilling are not primarily considered sheet metal forming operations. Facing and Drilling are distinctly machining processes involving material removal. Knurling is a surface forming process but typically applied to solid parts rather than shaping a sheet itself.

Rolling, on the other hand, is a fundamental metal forming process used specifically to produce sheets, plates, and foils by reducing the thickness of the material through plastic

deformation. When applied to produce thin, flat sections, it is indeed a sheet metal forming process, specifically one used in the primary production of sheet metal.

Comparison of Manufacturing Processes

Process	Type	Description	Applies to Sheet Metal Forming?
Knurling	Forming (Surface) / Machining	Creating a pattern on a surface via deformation	Typically applied to solid parts, not shaping the sheet itself.
Facing	Machining (Material Removal)	Creating a flat surface by removing material	No, material removal.
Drilling	Machining (Material Removal)	Creating a hole by removing material	No, material removal.
Rolling	Forming (Bulk/Sheet Metal)	Reducing thickness and shaping by passing through rollers	Yes, a primary process for producing sheet metal.

Therefore, Rolling is the process among the options that is classified as a sheet metal forming process, particularly in the context of producing sheet metal from thicker stock.

Revision Table: Key Manufacturing Processes

Manufacturing Process Types Summary

Process Type	Key Characteristic	Examples
Forming Processes	Shaping material through plastic deformation without material removal.	Rolling, Forging, Extrusion, Bending, Deep Drawing, Stamping
Machining Processes	Removing material to achieve desired shape and size.	Turning, Milling, Drilling, Grinding, Facing

Additional Information: Rolling in Detail

Rolling is a fundamental metalworking process. It can be done hot or cold. Hot rolling is typically used to produce basic shapes and reduce grain size, while cold rolling is used for better surface finish, increased strength, and closer tolerances. The pressure applied by the rollers causes the metal to deform plastically, reducing its thickness and increasing its length. Rolling is a crucial process in manufacturing various metal products, including structural shapes, rails, and especially flat products like sheets and plates which are used in many industries including automotive, construction, and aerospace.

120. Answer: d

Explanation:

Understanding Forces in Spur Gears and Rotation

Spur gears are one of the simplest types of gears, used to transmit motion and power between parallel shafts. When two spur gears mesh, the driving gear exerts a force on the driven gear at the point of contact between the teeth. This force is responsible for causing the driven gear to rotate.

The total force acting between the teeth of meshing spur gears acts along the line of action, which is determined by the pressure angle of the gears. This total force can be resolved into two principal components acting at the point of contact:

- **Radial Component (F_r):** This component acts along the line connecting the center of the gear to the point of contact. It pushes the gear towards or away from its mating gear and results in bearing loads. The radial component does not contribute to the rotation of the gear.
- **Tangential Component (F_t):** This component acts perpendicular to the radius of the gear at the point of contact. It is this component that creates a torque about the center of the gear.

The torque (T) generated on a gear is calculated as the product of the tangential force (F_t) and the effective radius (r) at which the force acts:

$$T = F_t \times r$$

This torque is what causes the gear to rotate. Therefore, the component of the force that assists or causes the rotation on the driven gear is the tangential component.

Let's look at the options provided in the context of spur gear forces and rotation:

- The worm and worm wheel: This is a type of gear drive, not a force component.
- The radial component: As discussed, the radial component causes bearing loads and does not produce torque for rotation.
- The axial component: Spur gears theoretically have no axial thrust component (force acting along the shaft axis). Axial forces are characteristic of helical or bevel gears. Even if present due to misalignment or manufacturing errors, it does not cause the primary rotation on the driven gear.
- The tangential component: This component acts tangentially to the pitch circle (or base circle, depending on the exact force definition and analysis point) and creates the torque necessary for the driven gear to rotate.

Based on the analysis of gear forces and torque generation, the tangential component is the force component that assists the rotation on the driven gear.

Force Component	Direction	Effect on Gear	Assists Rotation?
Radial (F_r)	Towards/Away from center	Bearing load	No
Tangential (F_t)	Perpendicular to radius (tangent)	Torque, Rotation	Yes
Axial (F_a)	Along shaft axis (theoretical for spur gears)	Axial thrust (if present)	No

Revision Table: Key Spur Gear Concepts

Term	Definition/Role
Spur Gears	Gears with straight teeth parallel to the axis of rotation, used for parallel shafts.
Driven Gear	The gear that receives power and motion from another gear (the driving gear).
Force Transmission	The process where one gear exerts force on another, transmitting torque and motion.
Torque	A twisting force that causes rotation. Produced by the tangential force component.
Pressure Angle	The angle between the line of action (direction of total force) and the line tangent to the pitch circles at the point of contact. Determines the ratio of radial to tangential force components.

Additional Information: Gear Force Analysis

The total force (F) between meshing spur gear teeth acts along the line of action. This force can be resolved into its radial (F_r) and tangential (F_t) components using the pressure angle (ϕ).

The relationships are:

- $F_t = F \cos(\phi)$
- $F_r = F \sin(\phi)$

Where F is the total force and ϕ is the pressure angle.

The tangential force (F_t) is directly related to the power transmitted and the speed of the gear. Power (P) transmitted by a gear rotating at angular velocity (ω) with torque (T) is given by $P = T\omega$. Since $T = F_t \times r$, we have $P = F_t \times r \times \omega$. Using the linear velocity ($v = r\omega$), we can also express power as $P = F_t \times v$. This highlights that the tangential force is the component directly involved in the transmission of power and causing rotation.

Understanding these force components is crucial for designing gears and their supporting shafts and bearings to withstand the loads.

121. Answer: b

Explanation:

Understanding Tools for Rectifying Damaged Threads

Threads on bolts, screws, or in nuts can sometimes get damaged due to rust, cross-threading, or impact. When this happens, it's often possible to repair or "rectify" the damaged threads instead of replacing the entire part. Different tools are used for cutting new threads (tapping and threading) versus repairing existing, damaged threads.

Analyzing the Options for Damaged Thread Rectification

Let's look at the given options and determine which one is specifically designed for rectifying damaged or rusted threads:

- **Two-piece die:** This type of die consists of two adjustable halves held in a die stock. It's primarily used for cutting **new external threads** on a rod or bolt, especially when precise adjustment is needed. While they cut threads, they are not the most convenient tool for simply cleaning up or repairing existing damaged threads.
- **Die nut:** A die nut looks like a standard nut but has cutting edges like a die. It's specifically designed to run over existing threads to clean them up, remove rust, burrs, and re-form slightly damaged threads. It's used by screwing it onto the bolt or stud, effectively chasing the threads. This is ideal for rectifying damaged or rusted threads without cutting deeply into the material.
- **Circular split die:** This is a common type of die used in a die stock to cut **new external threads**. The 'split' allows for slight adjustment of the thread size. Like the two-piece die, its primary function is cutting new threads, not repairing existing damaged ones, although it could potentially be used cautiously for cleaning.
- **Die plate:** Also known as a screw plate, a die plate is a flat plate with a series of threaded holes of different sizes. It's an older tool format used for cutting small external threads, particularly for instrument or clockmaking screws. It's for cutting **new threads**, not rectifying damaged ones.

Why Die Nut is Used for Rectifying Damaged Threads

Based on their functions, the die nut is the tool specifically intended for repairing or rectifying damaged and rusted threads. Its design allows it to follow the existing thread

pattern, cleaning and reshaping it without removing excessive material, which is the goal when rectifying damaged threads.

Tool	Primary Use	Suitable for Rectifying Damaged Threads?
Two-piece die	Cutting new external threads	Less suitable; designed for cutting, not just cleaning/repairing
Die nut	Cleaning, repairing, and rectifying existing external threads	Yes, this is its specific purpose
Circular split die	Cutting new external threads	Less suitable; designed for cutting
Die plate	Cutting new external threads (often small sizes)	No; designed for cutting new threads

Therefore, the tool specifically used to rectify damaged or rusted threads is the die nut.

Revision Table: Thread Repair Tools

Term	Definition/Function
Die nut	Tool resembling a nut, used to clean and repair existing external threads.
Die	General term for a tool used to cut external threads (on bolts, rods). Examples include circular split dies and two-piece dies.
Tap	Tool used to cut internal threads (in nuts, holes).
Die stock	Handle or holder for dies (like circular split dies or two-piece dies).

Additional Information on Thread Rectification

When rectifying damaged threads, lubrication is often used to help the die nut cut or clean smoothly. Care must be taken to ensure the die nut is started straight on the thread to avoid further damage. For severely damaged threads, simple rectification might not be sufficient, and re-threading (cutting new threads, potentially of a slightly larger size if material allows) or replacement might be necessary.

122. Answer: d

Explanation:

Understanding V Belt Cross-Section Shape

V belts are a common type of drive belt used for power transmission. They are widely used in various industries and applications due to their efficiency and reliability. A key feature of a V belt is the shape of its cross-section, which is specifically designed to provide effective power transmission.

What is the Shape of a V Belt's Cross-Section?

The defining characteristic of a V belt is its trapezoidal cross-section. This shape allows the belt to wedge into the groove of a pulley, creating a larger contact area and increasing friction compared to a flat belt.

Let's look at why this shape is advantageous and consider the given options:

- **Rectangular:** A rectangular cross-section would not provide the wedging action needed for efficient power transmission in grooved pulleys. It would behave more like a flat belt, relying primarily on tension and surface friction on the top and bottom surfaces, not the sides of the groove.
- **Square:** Similar to a rectangular shape, a square cross-section would not effectively wedge into a V-shaped groove.
- **Circular:** A circular cross-section belt (like a rope) would only have point or line contact in a V-groove, offering very little friction and poor power transmission capability for this application.
- **Trapezoidal:** This is the correct shape. The sloped sides of the trapezoid match the angled sides of the V-groove on the pulley. When the belt is tensioned and enters the groove, it gets slightly wedged, increasing the normal force between the belt

sides and the pulley groove sides. This significantly increases the friction force available for transmitting torque. The wedging action amplifies the effective tension in the belt, improving grip and preventing slip.

Why the Trapezoidal Shape is Important for V Belts

The trapezoidal shape, often described as a "V" shape, is crucial for the functionality of a V belt drive system. The angle of the V allows the belt to seat firmly in the corresponding V-groove of the pulley. As the belt pulls, it is drawn deeper into the groove, creating a wedging effect. This effect provides a mechanical advantage, increasing the grip between the belt and the pulley walls. This enhanced grip ensures more efficient power transfer with less slippage, especially under heavy loads.

Analysis of Options

Based on the characteristic design and function of V belts:

- Option 1 (Rectangular) is incorrect because it lacks the necessary angled sides for wedging.
- Option 2 (Square) is incorrect for the same reason as rectangular.
- Option 3 (Circular) is incorrect as it would not provide sufficient contact area or wedging action in a V-groove.
- Option 4 (Trapezoidal) is correct because this shape is the defining feature that allows V belts to transmit power effectively through wedging action in grooved pulleys.

Therefore, a V belt has a trapezoidal cross-section.

Comparison of Belt Cross-sections

Shape	Effectiveness in V-Groove	Typical Use
Rectangular	Poor (no wedging)	Flat belt drives
Square	Poor (no wedging)	Specialized applications, sometimes ropes
Circular	Poor (point/line contact)	Rope drives, small or light loads
Trapezoidal (V shape)	Excellent (wedging action)	V belt drives (power transmission)

Revision Table: V Belt Basics

Key Properties of V Belts	
Feature	Description
Cross-section Shape	Trapezoidal (V-shaped)
Operating Principle	Wedging action in pulley grooves
Primary Function	Power transmission between parallel shafts
Advantages	Higher power capacity, less slip, shorter center distances possible, quieter operation compared to some drives

Additional Information: Types and Applications of V Belts

V belts come in various standard sizes and types, such as classical V belts (like A, B, C, D sections), narrow V belts (like 3V, 5V, 8V), banded V belts, and cogged V belts. Each type is suited for different power requirements and pulley diameters.

- **Classical V Belts:** Used in many general industrial and agricultural applications.

- **Narrow V Belts:** Offer higher power capacity for their width, suitable for compact drives.
- **Banded V Belts:** Multiple V belts joined together by a band across the top, used to prevent belt whip and vibration on drives with pulsating loads or long center distances.
- **Cogged V Belts:** Have notches or cogs molded into the underside, providing greater flexibility and allowing for smaller pulley diameters and cooler running.

V belts are used in a vast array of machinery, including HVAC systems, industrial pumps, fans, compressors, agricultural equipment, machine tools, and automotive accessories (like alternators and power steering pumps).

123. Answer: b

Explanation:

Understanding Cooling System Tests

Cooling systems are essential for maintaining the optimal operating temperature of an engine. Inspecting these systems helps ensure they are functioning correctly and preventing potential overheating issues. There are different types of tests that can be performed on cooling systems to check their integrity and performance.

Static vs. Dynamic Cooling System Testing

Cooling system tests can generally be categorized into two main types: static and dynamic.

- **Static Test:** A static test is typically performed when the engine is not running and the system is not under normal operating conditions (high temperature and pressure). A common static test is a simple pressure test, where the system is pressurized to a certain level using a hand pump, and the pressure is observed over time to detect leaks. This test is useful for finding leaks in hoses, radiator, heater core, or fittings when the system is cold, but it doesn't fully simulate real operating conditions.
- **Dynamic Test:** A dynamic test is performed while the engine is running and the cooling system is operating under its normal conditions, including temperature,

pressure, and coolant flow. This test simulates the actual stresses and demands placed on the system during vehicle operation.

Why Dynamic Tests are Used for Cooling Systems

Cooling systems experience significant changes in pressure and temperature when the engine is running. Hoses expand, components heat up, and coolant circulates. A static test might not reveal leaks or issues that only appear under these operating conditions. A dynamic test allows mechanics to:

- Observe coolant circulation and flow.
- Check thermostat operation at operating temperature.
- Verify cooling fan activation.
- Monitor system pressure under operating temperature.
- Detect leaks that only manifest when components are hot and under pressure.

Because cooling system performance is critical and needs to be verified under actual working conditions, a dynamic test is the most effective method for a comprehensive inspection of the cooling system's functionality and integrity while the engine is running.

Conclusion on Cooling System Inspection

For inspecting cooling systems, especially to assess their performance and identify issues that occur under normal operating conditions, a **dynamic** test is used. This test simulates the stresses and environment the system experiences when the engine is running, providing a more accurate assessment than a static test alone.

Test Type	Condition	Purpose	Primary Check
Static	Engine off, system cold/ambient	Find leaks in a non-pressurized or manually pressurized system	Component integrity (leaks) at low/no pressure
Dynamic	Engine running, system at operating temperature/pressure	Assess overall system function under load	Leaks under pressure/temperature, flow, thermostat, fan operation

Revision Table: Cooling System Test Types

Feature	Static Test	Dynamic Test
Engine Status	Off	Running
Temperature	Ambient/Cold	Operating Temperature
Pressure	Ambient or Manually Applied (Low)	Operating Pressure (High)
Flow	None	Normal Circulation
What it checks	Leaks (primarily)	Leaks, flow, thermostat, fan, overall system function

Additional Information on Dynamic Testing

A dynamic cooling system inspection involves several steps performed while the engine is running and reaching operating temperature:

- Visually inspect for leaks around hoses, clamps, radiator, water pump, and engine block once the system is hot and pressurized.
- Check that the thermostat opens at the correct temperature, allowing coolant flow through the radiator.
- Verify that the electric cooling fan (if equipped) activates when the coolant reaches the specified temperature.
- Monitor the coolant temperature gauge to ensure the engine does not overheat.
- Use diagnostic tools to read coolant temperature and pressure data if available.
- Listen for unusual noises from the water pump or cooling fans.

Performing a dynamic test gives a complete picture of how the cooling system handles the demands of a running engine, making it a crucial part of cooling system maintenance and diagnostics.

124. Answer: c

Explanation:

Understanding Galvanizing for Corrosion Protection

Galvanizing is a process that involves applying a protective zinc coating to steel or iron. This coating helps prevent rusting and corrosion, significantly extending the lifespan of the metal object.

The Galvanizing Process

The most common method of galvanizing is called hot-dip galvanizing. In this process, the steel or iron object is submerged in a bath of molten zinc. The zinc metallurgically bonds with the iron or steel, forming a series of zinc-iron alloy layers, with a pure zinc layer on the surface.

Why Galvanize? The Role of Zinc

Zinc protects the underlying metal in two main ways:

- **Barrier Protection:** The zinc coating acts as a physical barrier, preventing corrosive substances like moisture and oxygen from reaching the steel or iron surface.
- **Sacrificial Protection:** Zinc is more electrochemically active than iron. If the coating is scratched or damaged, exposing the steel, the zinc will corrode preferentially (sacrificially) instead of the steel, protecting the base metal.

Galvanizing and Material Types

While galvanizing can be applied to various ferrous metals, the base metal composition significantly affects the outcome of the hot-dip galvanizing process, particularly the adhesion and thickness of the zinc coating.

Low Carbon Steel and Galvanizing

Low carbon steel is a very common material for galvanizing. It typically contains less than 0.25% carbon. The composition of low carbon steel, including its silicon content, influences how the zinc reacts with the surface during hot-dip galvanizing. Steels with specific ranges of silicon and phosphorus content tend to react in predictable ways, allowing for effective and adherent zinc coatings. Due to its widespread use in construction,

infrastructure, and manufacturing, and its compatibility with the hot-dip galvanizing process, low carbon steel is frequently galvanized to protect against corrosion.

Considering Other Options

- **Non-metallic substance:** Galvanizing involves coating a metal with another metal (zinc). It is not applicable to non-metallic substances like plastics, wood, or ceramics.
- **Non-ferrous metal:** Non-ferrous metals (metals other than iron or steel, like aluminum, copper, or brass) are not typically galvanized in the same way ferrous metals are. They usually have their own methods for corrosion protection or surface treatment. Applying zinc to a non-ferrous metal wouldn't provide the same sacrificial protection mechanism as it does for iron-based materials.
- **Cast iron:** Cast iron is also a ferrous metal, meaning it contains iron. It can be galvanized. However, its higher carbon content (typically 2-4%) and different microstructure compared to low carbon steel can affect the galvanizing process and the resulting coating. While possible, low carbon steel is a more commonly and predictably galvanized material for many applications compared to cast iron.

Based on typical industrial practices and material properties suitable for the standard hot-dip galvanizing process, low carbon steel is the most common material subjected to galvanizing.

Revision Table: Key Concepts in Galvanizing

Term	Description	Purpose
Galvanizing	Applying a protective zinc coating.	Prevent corrosion of steel/iron.
Hot-Dip Galvanizing	Dipping metal into molten zinc.	Forms durable zinc-iron alloy layers.
Sacrificial Protection	Zinc corrodes instead of steel.	Protects exposed steel areas.

Additional Information on Galvanizing Applications

Galvanizing is a widely used corrosion protection method across numerous industries. Its durability and cost-effectiveness make it suitable for various applications where steel or iron is exposed to corrosive environments. Examples include:

- **Construction:** Steel beams, columns, roofing, fasteners, and rebar (though less common).
- **Infrastructure:** Bridges, guardrails, streetlights, utility poles, and pipelines.
- **Automotive:** Some body panels and chassis components (though electro-galvanizing or other coatings are also common).
- **Agriculture:** Fencing, gates, and equipment.
- **General Fabrication:** Grates, railings, stairs, and hardware.

The effectiveness of galvanizing depends on factors like the thickness of the zinc coating, the environment the galvanized object is exposed to, and the composition of the steel being galvanized.

125. Answer: c

Explanation:

Understanding Valves for Preventing Reverse Flow

Valves are essential components in piping systems, used to control, direct, or regulate the flow of fluids (liquids, gases, or slurries) by opening, closing, or partially obstructing passageways.

The question asks about a specific type of valve used to prevent the product flow in a wrong direction. This function is crucial in many applications to maintain system integrity and prevent damage or contamination caused by backflow.

Analyzing Valve Types and Their Functions

Let's look at the common types of valves mentioned in the options and understand their primary functions:

- **Seat valve:** This term can refer to various valves where a closure element (like a disk or plug) seats against a surface to stop flow. Examples include globe valves, where

the flow path changes direction, offering good throttling capabilities but not primarily used for preventing reverse flow on their own.

- **Half turn valve:** This typically refers to valves that operate by rotating the closure element 90 degrees (a quarter turn). Examples include ball valves and butterfly valves. They are mainly used for isolation (fully open or fully closed) or throttling, not specifically for preventing backflow.
- **Check valve:** Also known as a non-return valve, this valve is specifically designed to allow fluid to flow in only one direction. It automatically closes to prevent flow in the reverse direction. The opening and closing action is usually driven by the differential pressure across the valve.
- **Butterfly valve:** A quarter-turn valve that uses a disc mounted on a shaft to control flow. It's good for isolation and some throttling applications, often used in large pipes, but not designed to prevent reverse flow.

The Role of a Check Valve in Preventing Backflow

A **check valve** is uniquely designed to perform the function described in the question: preventing flow in the wrong direction. When fluid flows in the desired direction, the valve opens. If the flow stops or attempts to reverse direction, the valve automatically closes, blocking the reverse flow.

This automatic operation, without the need for external power or human intervention, makes check valves vital safety devices in systems where backflow could cause problems, such as:

- Preventing contaminated water from entering a clean water supply.
- Protecting pumps from damage due to back-pressure or reverse flow when they shut off.
- Ensuring flow proceeds in the correct sequence in complex piping networks.
- Preventing mixing of fluids from different parts of a system.

Based on their design and function, the **check valve** is the correct answer as it is specifically used to prevent product flow in a wrong direction.

Comparison of Valve Types

Valve Type	Primary Function	Prevents Reverse Flow?
Seat valve (e.g., Globe)	Throttling, Flow Control	No (unless part of a specialized assembly)
Half turn valve (e.g., Ball, Butterfly)	Isolation, Throttling	No
Check valve	Preventing Reverse Flow (Backflow)	Yes
Butterfly valve	Isolation, Throttling	No

Conclusion on Valve Function

The valve specifically designed and utilized to ensure fluid flow in only one direction, thereby preventing flow in the wrong or reverse direction, is the check valve.

Revision Table: Key Valve Functions

Concept	Description	Relevance to Question
Valve Function	Mechanism to control, direct, or regulate fluid flow.	Understanding different valve functions is key to identifying the correct type.
Reverse Flow Prevention	Stopping fluid from moving in the opposite direction of intended flow.	This is the core requirement specified in the question.
Check Valve	A valve that permits flow in one direction only.	This valve type is specifically designed for the function asked.
Other Valve Types	Valves like globe, ball, butterfly have different primary purposes (isolation, throttling).	Understanding their functions helps rule them out for preventing backflow.

Additional Information on Check Valves

Check valves come in various designs, depending on the application, fluid type, and system requirements. Some common types include:

- **Swing Check Valve:** Features a disc that swings on a hinge. Opens with forward flow, closes by gravity and reverse flow.
- **Lift Check Valve:** Has a disc that lifts off its seat with forward flow and drops back onto the seat with reverse flow or gravity. Can be piston type or ball type.
- **Wafer Check Valve:** A compact, flangeless design often placed between flanges, typically of the dual-plate or single-plate swing type.
- **Foot Valve:** A type of check valve with a strainer, typically installed at the bottom of a pump suction line to prevent the pump from losing its prime.

The choice of check valve type depends on factors like line size, flow rate, pressure drop considerations, potential for water hammer, and the nature of the fluid (e.g., presence of solids).

126. Answer: b

Explanation:

Understanding Least Material Condition (LMC) for Shafts

The question asks for the Least Material Condition (LMC) of a shaft with a given size and tolerance. Understanding LMC is crucial in engineering dimensioning and tolerancing (GD&T).

What is Least Material Condition (LMC)?

The Least Material Condition (LMC) refers to the condition of a part where it contains the minimum amount of material within the stated tolerance limits. For external features like shafts, the LMC corresponds to the smallest permissible size. For internal features like holes, the LMC corresponds to the largest permissible size.

In this specific case, we are dealing with a shaft, which is an external feature. Therefore, the LMC will be the minimum allowable size of the shaft.

Calculating LMC for the Shaft

We are given the nominal size and the tolerance of the shaft:

- Nominal Size: ϕ 24.6 mm

- Tolerance: ± 0.5 mm

The tolerance ± 0.5 mm means the actual size of the shaft can be 0.5 mm larger or 0.5 mm smaller than the nominal size. This gives us the upper and lower limits for the shaft's size:

- Upper Limit (Maximum size) = Nominal Size + Tolerance = $24.6 \text{ mm} + 0.5 \text{ mm} = 25.1 \text{ mm}$
- Lower Limit (Minimum size) = Nominal Size - Tolerance = $24.6 \text{ mm} - 0.5 \text{ mm} = 24.1 \text{ mm}$

Since the shaft is an external feature, its LMC is the smallest permissible size, which is the lower limit.

Therefore, the LMC of the shaft is the lower limit calculated above.

LMC = Lower Limit = 24.1 mm

We can summarize the size limits:

Dimension Property	Calculation	Value (mm)
Nominal Size	Given	24.6
Tolerance	Given	± 0.5
Upper Limit (Maximum size)	$24.6 + 0.5$	25.1
Lower Limit (Minimum size)	$24.6 - 0.5$	24.1
Least Material Condition (LMC)	Lower Limit (for shaft)	24.1

Comparing this result with the given options, the LMC is $\phi 24.1$ mm.

Conclusion on Shaft LMC

Based on the calculation, the Least Material Condition (LMC) for the shaft with a nominal size of $\phi 24.6$ mm and a tolerance of ± 0.5 mm is $\phi 24.1$ mm. This represents the smallest acceptable size for the shaft within its specified tolerance.

Revision Table: Key Dimensioning Concepts

Concept	Definition	Shaft (ϕ 24.6 $\pm$ 0.5)	Hole Example (ϕ 24.6 $\pm$ 0.5)
Nominal Size	The theoretical exact size.	24.6	24.6
Upper Limit	Maximum permissible size.	24.6 + 0.5 = 25.1	24.6 + 0.5 = 25.1
Lower Limit	Minimum permissible size.	24.6 - 0.5 = 24.1	24.6 - 0.5 = 24.1
Maximum Material Condition (MMC)	Condition with maximum material (largest shaft, smallest hole).	Upper Limit = 25.1	Lower Limit = 24.1
Least Material Condition (LMC)	Condition with minimum material (smallest shaft, largest hole).	Lower Limit = 24.1	Upper Limit = 25.1

Additional Information on Tolerance and Conditions

Understanding tolerance and material conditions is fundamental in ensuring parts fit and function correctly during assembly. The tolerance defines the acceptable range of variation for a dimension.

- **Tolerance Range:** The difference between the upper limit and the lower limit (25.1 - 24.1 = 1.0 mm). This is also twice the $\pm$ tolerance (2 * 0.5 mm = 1.0 mm).
- **Maximum Material Condition (MMC):** This is the size limit where the part contains the maximum amount of material. For a shaft, this is its largest size (Upper Limit). For a hole, this is its smallest size (Lower Limit). In our shaft example, MMC is ϕ 25.1 mm.
- **Least Material Condition (LMC):** As calculated, this is the size limit where the part contains the minimum amount of material. For a shaft, this is its smallest size (Lower Limit). For a hole, this is its largest size (Upper Limit). In our shaft example, LMC is ϕ 24.1 mm.

These concepts are critical for calculating fits (clearance, interference, transition) between mating parts and are integral to geometric dimensioning and tolerancing (GD&T).

127. Answer: b

Explanation:

Understanding Flat Belt Drive Systems

A flat belt drive is a type of power transmission system that uses a flat belt to connect two or more pulleys. This system is common in various mechanical applications for transferring rotational motion and torque from one shaft to another.

Analyzing the Pulley Configuration Described

The question describes a specific setup for a flat belt drive: two pulleys mounted on the driven shaft and one pulley mounted on the driving shaft. This particular arrangement serves a distinct purpose in controlling the power transmission.

Identifying the Flat Belt Drive Type

Let's examine the characteristics of different flat belt drive types based on their pulley configurations:

- **Cross Belt Drive:** Involves two pulleys, one on the driving shaft and one on the driven shaft. The belt is crossed, causing the driven shaft to rotate in the opposite direction to the driving shaft. This setup typically uses one pulley per shaft.
- **Multiple Belt Drive:** Uses several belts running parallel to each other between pulleys on the driving and driven shafts. This configuration is used to transmit higher power. It involves multiple belts, not necessarily two pulleys on one shaft and one on the other with a single belt.
- **Idle Pulley Drive (or Tension Pulley Drive):** Uses an additional pulley (or pulleys) that are not directly on the driving or driven shafts but are used to maintain belt tension or change the belt path. This setup can involve various pulley arrangements but the key feature is the presence of idle pulleys.
- **Fast and Loose Pulley Drive:** This system specifically uses one pulley on the driving shaft and two pulleys on the driven shaft. The two pulleys on the driven shaft are side-by-side. One pulley, called the **fast pulley**, is keyed or rigidly attached to the driven shaft. The other pulley, called the **loose pulley**, rotates freely on the driven shaft.

How the Fast and Loose Pulley System Works

In a fast and loose pulley drive:

- When the flat belt is on the fast pulley, it drives the driven shaft because the fast pulley is attached to the shaft. Power is transmitted.
- When the flat belt is shifted to the loose pulley, the loose pulley spins freely on the shaft, but the driven shaft does not rotate. Power transmission is disengaged.

This configuration with one driving pulley and two driven pulleys (one fast, one loose) is designed to allow the driven machine or shaft to be started or stopped at will, without stopping the driving shaft.

Based on the description provided in the question – two pulleys on the driven shaft and one pulley on the driving shaft of a flat belt drive – the system matches the characteristics of a fast and loose pulley drive.

Comparing Drive Types

Drive Type	Pulley Configuration (Typical)	Purpose
Cross Pulley Drive	One on driving, one on driven (crossed belt)	Reverse direction of driven shaft
Fast and Loose Pulley Drive	One on driving, Two on driven (Fast & Loose)	Engage/Disengage driven shaft from driving shaft
Multiple Belt Drive	Multiple belts between driving and driven pulleys	Transmit higher power
Idle Pulley Drive	Additional pulley(s) not on main shafts	Maintain tension, change belt path

Conclusion on the Flat Belt Drive Setup

A flat belt drive with one pulley on the driving shaft and two pulleys (one fast and one loose) on the driven shaft is known as a fast and loose pulley drive. This setup allows the driven machinery to be engaged or disengaged from the driving power source easily.

Revision Table: Flat Belt Drive Systems

Flat Belt Drive Type	Pulley Arrangement	Key Functionality
Open Belt Drive	Driving and driven pulleys on parallel shafts, belt runs open	Transmit power, shafts rotate in same direction
Cross Belt Drive	Driving and driven pulleys on parallel shafts, belt crossed	Transmit power, shafts rotate in opposite directions
Fast and Loose Pulley Drive	One driving pulley, Two driven pulleys (Fast & Loose) on same shaft	Start/Stop driven machine without stopping prime mover
Stepped Cone Pulley Drive	Stepped pulleys on driving and driven shafts	Change speed ratio by shifting belt
Jockey/Idle Pulley Drive	Additional pulley(s) used	Maintain belt tension, alter belt path, improve contact angle

Additional Information: Applications of Fast and Loose Pulley Drives

Fast and loose pulley drives were historically very common in workshops and factories, especially before individual electric motors became widespread for each machine. They allowed multiple machines to be driven from a single line shaft (the driving shaft) which was powered by a central engine or motor. Workers could then easily connect or disconnect their specific machine by shifting the belt between the fast and loose pulleys using a belt shifter mechanism.

Examples of older machines that often used fast and loose pulleys include:

- Lathes
- Drilling machines
- Power looms
- Grinding machines

While less common for individual machines with dedicated motors today, the principle of engaging/disengaging driven components is still relevant in various mechanical systems.

128. Answer: c

Explanation:

Understanding the Vernier Height Gauge and Its Use

A Vernier height gauge is a precision measuring instrument primarily used in workshops and manufacturing settings. It is designed to measure vertical dimensions from a reference surface, typically a surface plate. While it can measure height, one of its crucial functions is related to marking or scribing lines on workpieces at specific heights.

Primary Functions of a Vernier Height Gauge

Let's look at the potential uses mentioned in the options:

- **Measuring height of a work piece:** A Vernier height gauge can indeed measure the height of a workpiece. You would typically set the scribe or indicator at the top surface of the workpiece and read the measurement on the main and Vernier scales relative to the base resting on a surface plate.
- **Checking right angular surfaces:** This is not the primary use of a Vernier height gauge. Instruments like try squares or angle plates are designed specifically for checking perpendicularity (right angles). While you might use the vertical column as a reference, the gauge itself isn't built for precise angle checks.
- **Scribing horizontal lines on work piece:** This is a very common and important use of a Vernier height gauge. A scribe is attached to the measuring head. By setting the head to a desired height and moving the entire gauge across the surface plate, the scribe marks a precise horizontal line on the workpiece placed on the same surface plate.
- **Scribing vertical lines on work piece:** A Vernier height gauge is not designed for scribing vertical lines. To scribe vertical lines, you would typically use instruments like a surface gauge with a scribe or a marking out machine, referencing off an edge or other feature.

Analyzing the Options and Key Use

While the Vernier height gauge can measure height, its design, especially with the inclusion of a scribe, makes it exceptionally well-suited for marking. The act of setting a specific height and using the scribe to draw a line at that height is a fundamental marking out operation.

Therefore, among the given options, 'scribing horizontal lines on work piece' accurately describes a core function of the Vernier height gauge that leverages its ability to precisely set a height relative to a reference surface.

Conclusion on Vernier Height Gauge Usage

Based on the functions and how the Vernier height gauge is used in practice, particularly with its scribe attachment, its capability to accurately mark lines at specific heights is a primary application.

Common Uses of Measuring Instruments

Instrument	Typical Use(s)
Vernier Height Gauge	Scribing horizontal lines, Measuring height (often by setting scribe)
Try Square	Checking right angles, Scribing lines perpendicular to an edge
Surface Gauge	Scribing lines parallel to a surface/edge, Checking flatness/parallelism
Vernier Caliper	Measuring external/internal dimensions, depth, step

Revision Table: Vernier Height Gauge Facts

Feature/Aspect	Description for Vernier Height Gauge
Primary Function	Precise vertical measurement and marking/scribing
Key Component	Base, Column, Measuring Head (with Vernier scale), Scribe
Reference Surface	Surface plate is essential for stable measurement/scribing
Typical Use in Marking Out	Scribing horizontal lines at accurate heights

Additional Information on Marking Instruments

Marking out is a crucial step before machining or fabricating a workpiece. It involves transferring design dimensions onto the material using lines, circles, or points. Various tools are used for this process, depending on the required precision and the type of mark needed.

- **Scribers:** Pointed tools used to scratch lines on the workpiece surface.
- **Surface Plate:** A flat, accurate reference surface against which measuring and marking tools like height gauges and surface gauges are used.
- **Centre Punch:** Used to make small indentations (punch marks) along scribed lines or at intersection points to preserve the marking lines after subsequent operations.
- **Vernier Height Gauge:** Excellent for scribing lines parallel to the reference surface (horizontal lines relative to the base).
- **Surface Gauge:** Can also scribe lines parallel to a reference surface or edge, often used freehand or with angle plates.

Understanding the specific application of each marking and measuring tool is key in manufacturing and inspection.

129. Answer: a

Explanation: **Your Personal Exams Guide**

Correct answer: To protect the op-amp from damage.

Option 1: **To protect the op-amp from damage.** As explained above, clamp diodes limit the differential input voltage, thereby protecting the op-amp's input stage from excessive voltage stress. This aligns perfectly with the primary function of clamp diodes in this context.

Based on the analysis, the main purpose of clamp diodes used in comparator circuits is indeed to protect the op-amp from damage caused by large differential input voltages.

130. Answer: b

Explanation:

Understanding Assembly Fits: Hole vs Shaft Size

In engineering and manufacturing, when two components like a hole and a shaft are intended to be assembled, the relationship between their sizes determines the type of fit. This fit describes how tightly or loosely the two parts will join together.

The question describes a specific situation in an assembly where the **Hole size** is smaller than the **Shaft size**. Let's analyze what this means for the assembly process and the resulting fit.

When the nominal size of the hole is designed to be less than the nominal size of the shaft, it implies that there will be an intentional overlap of material when the parts are brought together. This overlap requires force to assemble the parts, and once assembled, they are intended to remain fixed relative to each other due to the pressure between the surfaces.

There are primarily three basic types of fits:

1. Clearance Fit
2. Interference Fit
3. Transition Fit

Let's briefly look at each type:

- **Clearance Fit:** In a clearance fit, the largest possible shaft is always smaller than the smallest possible hole. This means there is always a gap or "clearance" between the shaft and the hole, allowing them to slide or rotate freely.
- **Interference Fit:** In an interference fit, the smallest possible shaft is always larger than the largest possible hole. This results in an "interference" of material, meaning the shaft must be forced into the hole. Once assembled, the parts are held together tightly by the resulting pressure. This fit is used for permanent or semi-permanent assemblies where the parts should not move relative to each other.
- **Transition Fit:** A transition fit falls between a clearance fit and an interference fit. The tolerance zones of the hole and the shaft overlap. This can result in either a small clearance or a slight interference when the parts are assembled, depending on the actual sizes of the manufactured components within their tolerance ranges.

Given the condition in the question, where the **Hole size** is smaller than the **Shaft size**, this directly corresponds to the definition of an **Interference fit**. The shaft is larger than the opening it needs to fit into, creating an interference that requires force for assembly and results in a tight joint.

Consider the relationship using limits (assuming basic hole system for simplicity, where hole minimum size is basic size):

- For a Clearance fit, the upper limit of the shaft is less than the lower limit of the hole. Mathematically, $D_{s_max} < D_{h_min}$.
- For an Interference fit, the lower limit of the shaft is greater than the upper limit of the hole. Mathematically, $D_{s_min} > D_{h_max}$. This implies the smallest shaft is larger than the largest hole.
- For a Transition fit, the tolerance zones overlap. Mathematically, $D_{s_min} < D_{h_max}$ and $D_{s_max} > D_{h_min}$.

In the scenario where the **Hole size is smaller than the Shaft size**, it means the shaft is larger than the hole, leading to an interference fit.

Fit Type	Relationship (Hole vs Shaft)	Characteristics
Clearance Fit	Hole is always larger than Shaft ($D_h > D_s$)	Easy assembly, relative movement possible
Interference Fit	Hole is always smaller than Shaft ($D_h < D_s$)	Requires force for assembly, rigid joint, no relative movement
Transition Fit	Hole and Shaft sizes may overlap (D_h can be $>$, $<$, or $=$ D_s within limits)	Assembly may be easy or require light press/tap, may result in slight clearance or interference

Revision Table: Assembly Fits

Fit Type	Hole Size vs Shaft Size	Assembly	Relative Motion
Clearance	Hole > Shaft	Easy	Possible (sliding, rotation)
Interference	Hole < Shaft	Requires Force (Press Fit)	Prevented (Rigid)
Transition	Overlap (Can be >, <, or =)	Varies (Easy to Light Press)	Limited/Prevented

Additional Information: Tolerances and Fits in Assembly

The type of fit achieved in an assembly depends on the specified tolerances for both the hole and the shaft. Tolerances are the permissible variations in dimensions. Standards like the ISO system of limits and fits define tolerance grades and fundamental deviations that designers use to specify the desired fit.

- **Limits:** The maximum and minimum permissible sizes for a dimension.
- **Tolerance:** The difference between the upper and lower limits.
- **Basic Size:** The theoretical size from which limits are derived.
- **Deviation:** The difference between the actual size and the basic size.
- **Fundamental Deviation:** The deviation closest to the basic size.
- **Tolerance Zone:** The area between the upper and lower limits, defined by the tolerance and the fundamental deviation.

The relationship between the tolerance zones of the hole and the shaft determines whether the fit will be clearance, interference, or transition. When the hole's tolerance zone is entirely below the shaft's tolerance zone, it results in an interference fit, which is the case when the **Hole size is smaller than the Shaft size**.

131. Answer: d

Explanation:

Understanding Power Transmission in Positive Contact Clutches

Positive contact clutches are mechanical devices used to connect two shafts (driving and driven) so that they rotate together at the same speed. Unlike friction clutches, which rely on the force of friction between surfaces to transmit torque, positive contact clutches achieve power transmission through a mechanical interlock.

Let's look at how this interlock works and analyze the given options regarding how power transmission is achieved in these clutches.

Mechanism of Positive Contact Clutches

Positive contact clutches use engaging elements, such as teeth or jaws, on both the driving and driven members. When the clutch is engaged, these teeth or jaws fit together precisely, creating a rigid mechanical connection. Power is transmitted through the direct pressure and shear forces acting on the surfaces of these interlocking elements. When torque is applied to the driving shaft, it is directly transferred to the driven shaft through the mechanical engagement of the teeth or jaws.

Consider the force acting on the engaging teeth. As torque is transmitted, the teeth on one member push against the corresponding teeth on the other member. This pushing force, acting parallel to the surface of engagement, creates shear stress within the material of the teeth and is the primary means by which the torque (and thus power) is transferred from one shaft to the other.

Analyzing the Options

Let's evaluate each option based on our understanding of positive contact clutches:

- **Option 1: Minor part is achieved by friction.** While there might be some incidental friction between the surfaces of the engaging teeth, it is not the primary mechanism for power transmission. The torque capacity of a positive contact clutch is determined by the strength of the interlock, not friction. Thus, friction plays a minor, if any, role in transmitting the main torque.
- **Option 2: Major part is achieved by friction.** This statement is incorrect. If friction were the major part of power transmission, it would behave like a friction clutch, which positive contact clutches are fundamentally different from. Their high torque capacity comes from the rigid mechanical lock.

- **Option 3: It is achieved through friction.** Similar to Option 2, this is incorrect. Positive contact clutches do not rely on friction to transmit power. Friction is the principle behind friction clutches, such as single-plate or multi-plate clutches.
- **Option 4: It is achieved by shear contact.** This statement accurately describes the primary mechanism. The engaging teeth or jaws interlock, and the torque is transmitted through the shear forces acting on the contact surfaces of these interlocking elements. The shear stress in the material of the teeth allows for the transmission of significant torque.

Based on the analysis, power transmission in positive contact clutches is primarily achieved through the mechanical shear contact of the engaging elements.

Comparison: Positive Contact vs. Friction Clutches

Understanding the difference between positive contact and friction clutches helps clarify the mechanism.

Feature	Positive Contact Clutch	Friction Clutch
Power Transmission Mechanism	Mechanical interlock (shear contact)	Friction between surfaces
Engagement	Engaged when teeth/jaws interlock	Engaged when surfaces are pressed together
Slipping during Engagement	No controlled slip; engages only when speeds are close or zero	Allows controlled slip during engagement
Torque Capacity	High, limited by strength of engaging elements	Depends on friction coefficient, normal force, and contact area
Suitability	Where positive, slip-free drive is needed; typically engaged at low speed or rest	Where smooth engagement and speed difference management are needed

As the table highlights, the fundamental difference lies in the power transmission mechanism: mechanical interlock via shear contact versus friction.

Conclusion on Positive Contact Clutch Mechanism

The correct description of power transmission in positive contact clutches is that it is achieved by shear contact. The interlocking teeth or jaws bear the load in shear, enabling the transmission of torque without relying on friction.

Revision Table: Key Concepts

Concept	Description	Relevance to Clutches
Positive Contact Clutch	Clutch type using mechanical interlock (teeth/jaws) for power transmission.	Transmits power via shear contact.
Friction Clutch	Clutch type using friction between surfaces for power transmission.	Contrasts with positive contact clutch mechanism.
Shear Stress	Stress component parallel to a surface, caused by forces acting parallel to the surface.	Primary mechanism for power transmission in positive contact clutches (stress on engaging teeth).
Torque Transmission	Transfer of rotational force from one component to another.	The main function of a clutch.

Additional Information on Positive Contact Clutches

Positive contact clutches are often used in applications where a positive, non-slipping drive is required, and engagement while the shafts are rotating at significantly different speeds is not necessary or desirable. Examples include:

- Gearboxes (to select different gears, often engaged when stopped or moving slowly)
- Certain types of industrial machinery
- Applications requiring high torque transmission in a compact size

Types of positive contact clutches include jaw clutches and spline clutches. Jaw clutches have robust, protruding jaws that interlock, while spline clutches use multiple splines (keys integrated into the shaft) that engage with corresponding slots in the mating hub.

A key characteristic is that these clutches typically cannot be engaged smoothly while the shafts are rotating at high relative speeds. Engaging them under such conditions can cause shock loads and damage the engaging elements. They are best engaged when the shafts are stationary or rotating at very similar speeds.

The term "shear contact" in this context refers to the surfaces of the engaging teeth or jaws transmitting the rotational force through shear loading, as opposed to the normal force used in friction clutches.

132. Answer: c

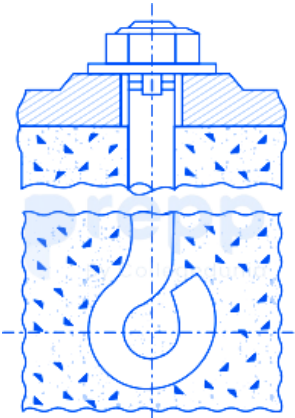
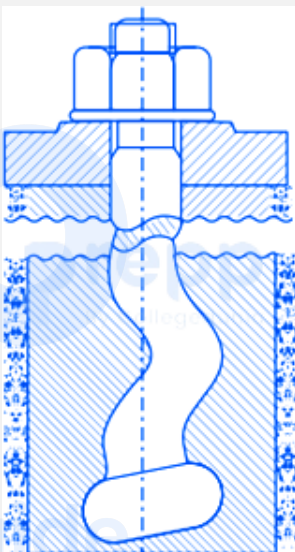
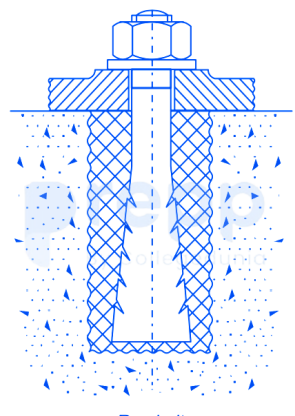
Explanation:

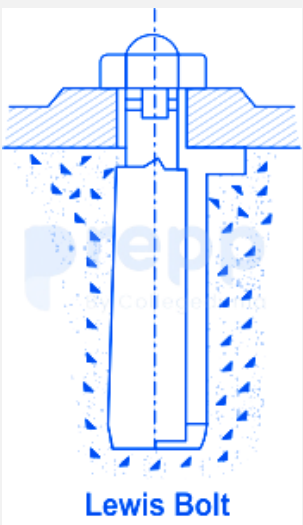
Explanation:

Foundation Bolt

- For some machines, it is very essential to hold down the machines firmly on the foundation to prevent them from moving.
- For this purpose, various types of foundation bolts or anchor bolts are used.
- There are basically four types of Foundation Bolts.

Your Personal Exams Guide

Eye Foundation Bolt	• This is the simplest of all types. • One end of the bolt is forged into an eye and another bar is placed through the eye. • This bolt is used in the machine-shop floor, and then cement concrete is poured around.	 The diagram shows a cross-section of an eye foundation bolt. A vertical bolt passes through a concrete slab. At the top, it has a hexagonal nut and washer. At the bottom, the bolt is forged into a large, circular eye shape. A horizontal bar is inserted through the eye. The entire assembly is surrounded by a layer of concrete. Eye Foundation Bolt
Bent Foundation Bolt	• This bolt can either be cast directly in the cement concrete or in lead. • The latter one is usually employed when machined are to be fixed on stone beds directly.	 The diagram shows a cross-section of a bent foundation bolt. A vertical bolt passes through a concrete slab. At the top, it has a hexagonal nut and washer. The bolt is bent into a wavy, S-shape within the concrete. The entire assembly is surrounded by a layer of concrete.
Rag Foundation Bolt	• This bolt consists of a tapered body, square or rectangular in cross-section, the tapered edges being grooved. • This bolt can be used in similar lines as that of the bent bolt.	 The diagram shows a cross-section of a rag foundation bolt. A vertical bolt passes through a concrete slab. At the top, it has a hexagonal nut and washer. The bolt has a tapered body with grooved edges. The entire assembly is surrounded by a layer of concrete. Rag bolt

Lewis Foundation Bolt	• Lewis bolt is a tapered bolt used as a removable foundation bolt. • This type of Foundation Bolt is not commonly used for mounting machines. • It is used for lifting huge blocks of stones.	 The diagram shows a cross-section of a Lewis Bolt. It is a tapered bolt with a hexagonal head and a threaded section. The bolt is shown passing through a hole in a concrete foundation. The tapered part of the bolt is shown expanding within the hole, which is filled with a material that looks like concrete or grout. The bolt is labeled 'Lewis Bolt' at the bottom.
------------------------------	--	---

133. Answer: d

Explanation:

Understanding Bore Gauges and Measurement

The question asks about the primary measurement performed by a bore gauge. A bore gauge is a precision measuring instrument used specifically for determining the size of bores or holes.

What does a Bore Gauge Measure?

A bore gauge is designed to measure the internal diameter of a hole or cylinder. It typically uses contact points that expand within the bore, and the measurement is read from a dial indicator or digital display. This tool is essential for checking the size and roundness of drilled, reamed, or bored holes with high accuracy.

Analyzing the Options

Let's look at why the other options are incorrect:

- **Option 1: depth of a hole**

This is measured using a depth gauge, which is a different type of measuring tool designed to determine the distance from a reference surface to the bottom of a hole, slot, or step.

- **Option 2: diameter with an accuracy of 0.05 mm**

While a bore gauge does measure diameter, stating a specific accuracy level like 0.05 mm is not the defining function of the tool itself. Many different tools can measure diameter, and accuracy depends on the specific instrument's design and calibration. The primary function is measuring *internal* diameter.

- **Option 3: outside diameter of a shaft**

The outside diameter of cylindrical objects like shafts is typically measured using external measurement tools such as micrometers or calipers, not a bore gauge.

- **Option 4: inside diameter of a hole**

This is the correct and specific function of a bore gauge. It is designed to measure the internal dimension of cylindrical features.

Based on the function and design of the tool, the bore gauge is used for measuring the inside diameter of a hole.

Measurement Tools and Their Uses

Measuring Tool	Typical Measurement
Bore Gauge	Inside diameter of a hole or bore
Depth Gauge	Depth of a hole, slot, or step
Micrometer (External)	Outside diameter of a shaft or object
Caliper	Internal, external, depth, and step measurements (depending on type)

Conclusion on Bore Gauge Function

Therefore, the correct answer is that a bore gauge is used for measuring the inside diameter of a hole.

Revision Table: Key Measurement Tools

Reviewing the functions of various precision measuring instruments is crucial for understanding their applications in manufacturing and inspection.

Summary of Measurement Tools

Tool	Purpose	Measured Dimension
Bore Gauge	Precision measurement of internal dimensions	Inside diameter (of holes, cylinders)
Micrometer	High-precision measurement of small distances	External diameter, length, thickness (depending on type)
Caliper	Versatile measurement for various dimensions	Inside diameter, outside diameter, depth, step
Depth Gauge	Measurement of vertical distance from a reference surface	Depth of holes, slots, etc.

Additional Information on Bore Gauge Types

Bore gauges come in different types, including:

- **Telescoping Bore Gauge:** Used to transfer the inside dimension to an external measuring tool like a micrometer for reading.
- **Small Hole Gauge:** Similar to telescoping gauges but for very small holes.
- **Dial Bore Gauge:** Uses a dial indicator for direct reading of the internal diameter relative to a set master size.
- **Digital Bore Gauge:** Provides a digital display for direct reading.

Regardless of the type, the fundamental purpose of a bore gauge remains the measurement of the inside diameter of a hole.

134. Answer: a

Explanation:

Understanding Split Pins and Nut Locking

The question asks about the type of nut that is typically locked using a split pin. A split pin is a common fastening device used to secure components and prevent them from

loosening under vibration or rotation.

How a Split Pin Works

A split pin, also known as a cotter pin, is a metal fastener with two tines that are bent apart after insertion. It is commonly used to secure a castellated nut (or castle nut) or slotted nut on a bolt or threaded shaft. The bolt or shaft must have a hole drilled through it at the correct position.

Once the nut is tightened, the split pin is inserted through one of the slots in the nut and the hole in the bolt. The tines of the split pin are then spread open, usually bent around the shaft or bolt end. This spreading action prevents the split pin from coming out and, consequently, prevents the nut from rotating backward and loosening.

Types of Nuts and Their Use with Split Pins

Let's look at the given options:

- **Castle Nut:** A castle nut, also called a castellated nut, has slots cut into one end, resembling the battlements of a castle. These slots are specifically designed to accommodate a split pin. The split pin is inserted through the slots of the nut and a corresponding hole in the bolt, providing a positive lock.
- **Square Nut:** A square nut is a standard four-sided nut. It does not have slots for a split pin. While other locking mechanisms like lock washers might be used with square nuts, split pins are not.
- **Flange Nut:** A flange nut has a wide flange at one end that acts as an integrated washer. This distributes the load over a larger area. Like square nuts, they are not designed for use with split pins for locking.
- **Wing Nut:** A wing nut has two large "wings" that allow it to be easily tightened and loosened by hand without tools. They are used in applications requiring frequent adjustment. They do not have features for split pin locking.

Conclusion

Based on the design and function of a split pin, it is specifically designed to be used with nuts that have slots to receive the pin. Among the given options, only the castle nut (or castellated nut) has these slots.

Therefore, the split pin is used for locking a **castle** nut.

Revision Table: Nut Types

Nut Type	Description	Typical Locking Method
Castle Nut (Castellated Nut)	Hex nut with slots on top	Split pin (cotter pin) through slots and bolt hole
Square Nut	Nut with four sides	Washers, double nutting, thread locker
Flange Nut	Nut with an integrated washer (flange)	Built-in load distribution, sometimes thread locker
Wing Nut	Nut with wings for hand tightening	Used for easy adjustment, not typically for critical locking

Additional Information: Other Locking Methods

While split pins with castle nuts are effective for preventing loosening, especially in automotive or mechanical applications, other methods exist for locking nuts and bolts:

- **Lock Washers:** These washers (e.g., split ring, internal/external tooth) provide tension or bite into the nut and surface to resist loosening.
- **Nylon Insert Nuts (Nylock Nuts):** These nuts have a plastic (nylon) ring inside the threads that grips the bolt, creating friction and preventing loosening.
- **Jam Nuts:** Using two nuts tightened against each other on the same bolt thread creates friction that resists loosening.
- **Thread Locking Fluid:** Liquid adhesive applied to threads that hardens to prevent movement between the bolt and nut.
- **Safety Wire:** Wire threaded through holes in bolts or nuts and twisted, so tension in the wire prevents loosening.

135. Answer: b

Explanation:

Understanding Fine Thread Resistance to Unscrewing

The question asks about the resistance offered when a fine thread is being unscrewed. This relates to how easily a threaded fastener, like a screw or bolt, becomes loose or unscrews on its own or under vibration.

Fine threads have more threads per unit length compared to coarse threads. This fundamental difference significantly impacts their resistance to unscrewing.

Why Fine Threads Offer High Resistance to Unscrewing

Here are the key reasons why the resistance of fine thread unscrewing is considered high:

- **Larger Contact Area:** For a given length of engagement, fine threads have a larger total surface area in contact between the male and female threads. This increased contact area leads to higher friction, which opposes unscrewing.
- **Smaller Helix Angle:** Fine threads have a smaller helix angle than coarse threads of the same diameter. The helix angle is the angle of the thread relative to the fastener's axis. A smaller helix angle means it takes more rotations to advance the fastener a certain distance, and conversely, it's harder for axial forces or vibrations to cause rotation and unscrewing.
- **Better Vibration Resistance:** Due to the smaller helix angle and increased friction, fine threads are generally more resistant to loosening under vibration compared to coarse threads. This is a crucial factor in applications where dynamic loads or vibrations are present.
- **Greater Thread Engagement:** Often, fine threads allow for deeper or longer thread engagement in thin materials due to the higher number of threads available over a short length, further increasing resistance.

Comparing fine threads to coarse threads helps clarify the difference in unscrewing resistance.

Feature	Fine Thread	Coarse Thread
Threads per unit length	More	Fewer
Pitch (distance between threads)	Smaller	Larger
Helix Angle	Smaller	Larger
Thread Contact Area (for same engagement length)	Larger	Smaller
Friction during unscrewing	Higher	Lower
Resistance to Unscrewing / Vibration	High	Lower

Based on these characteristics, the resistance of fine thread unscrewing is significantly higher than that of coarse threads.

Revision Table: Fine Thread Characteristics

Let's quickly review the key properties of fine threads that contribute to their performance, including unscrewing resistance.

- Higher thread count per unit length.
- Smaller thread pitch.
- Smaller helix angle.
- Larger contact surface area for a given engagement.
- Greater resistance to vibration and loosening.
- Often used where precise adjustments are needed.
- Less susceptible to stripping in hard materials due to larger contact area distributing load.

Additional Information: Applications and Considerations

Understanding the resistance of fine thread unscrewing is important for selecting the right fastener for an application.

- **Applications:** Fine threads are often used in applications requiring high precision, strong clamping force, and resistance to vibration, such as in aerospace, automotive engines, precision machinery, and adjustment mechanisms.

- **Strength:** While fine threads offer high unscrewing resistance and strength against shear forces (per thread), coarse threads are generally stronger against tensile loads in the bolt itself if the threads are rolled, and are less prone to cross-threading or damage during assembly.
- **Assembly:** Fine threads require more rotations to assemble and disassemble compared to coarse threads, and they are more susceptible to damage from nicks or dirt due to the smaller pitch.

In summary, the physical geometry of fine threads results in a high resistance to unscrewing forces, making them suitable for critical applications where self-loosening must be prevented.

136. Answer: c

Explanation:

Understanding Specialty Files for Metalworking

When working with metals or other materials, selecting the correct file is crucial for achieving the desired shape and finish. Different files have unique shapes, tooth patterns, and purposes. The question asks us to identify a specific type of file characterized by a flat, triangular face with teeth only on the wide face, used for finishing sharp corners.

Characteristics of Key File Types

Let's examine the characteristics of the file types mentioned in the options to determine which one matches the description provided in the question.

Barrette File Features and Uses

A Barrette file is a specialized file designed for filing sharp internal angles and corners. Its key features are:

- **Shape:** It has a flat, triangular cross-section on one side, tapering to a point. The crucial aspect is that one side (usually the back) is smooth or 'safe' (without teeth), while the other sides (the wide face) have teeth.

- **Teeth:** Teeth are typically cut only on the wide, triangular face. This allows the user to file into a sharp corner without damaging the adjacent surface by using the smooth, safe side against it.
- **Use:** It is specifically used for finishing or sharpening internal sharp corners, grooves, and shoulders where precision is needed. The safe edge prevents cutting into the material where it's not intended.

Other File Types

Let's briefly look at the other options:

- **Riffler file:** Riffler files are small, shaped files, often curved or bent, used by sculptors, die sinkers, and jewelers for filing in confined or difficult-to-reach areas and for sculpting intricate shapes. They come in many shapes, not specifically a flat triangular one with teeth only on the wide face.
- **Crossing file:** A crossing file has a shape that is oval on one side and convex on the other, like a double-sided hand file with curved faces. It is used for filing curved surfaces and internal radii. It does not fit the description of having a flat triangular face.
- **Mill saw file:** A mill saw file is typically flat and rectangular with single-cut teeth on the faces and edges, often tapering. It is primarily used for sharpening mill saws or circular saw blades. While flat, it usually has teeth on both faces and sometimes edges, and its shape and primary use don't match the description of finishing sharp internal corners with teeth only on the wide face.

Comparing File Characteristics for Identification

To clearly identify the file that fits the question's description (flat, triangular face, teeth on wide face only, used for finishing sharp corners), let's compare the relevant characteristics:

File Type	Shape	Teeth Location	Primary Use	Matches Description?
Barrette file	Flat, triangular on one side (tapering)	On the wide face only (often with a safe edge)	Finishing sharp internal corners/angles	Yes
Riffler file	Various shapes (curved, bent, etc.)	Often on all working surfaces	Sculpting, confined areas	No
Crossing file	Oval/convex on both sides	On both working surfaces	Filing curved surfaces	No
Mill saw file	Flat, rectangular (often tapering)	Usually on faces and edges	Sharpening saw blades	No

Identifying the Correct File for Sharp Corners

Based on the comparison, the file that precisely matches the description of having a flat, triangular face, teeth only on the wide face, and being used for finishing sharp corners is the Barrette file. The unique combination of its shape and the placement of its teeth makes it ideal for this specific application.

Revision Table: Understanding Different Files

File Name	Key Feature	Common Application
Barrette file	Flat, triangular with teeth on wide face only (safe edge)	Finishing sharp internal corners
Riffler file	Small, various shaped tips	Sculpting, hard-to-reach areas
Crossing file	Double-sided curved profile	Filing curved surfaces and radii
Mill saw file	Flat, single-cut teeth, rectangular/tapering	Sharpening saw teeth

Additional Information on File Types and Uses

Files are fundamental tools in metalworking, woodworking, and other crafts for removing small amounts of material. Beyond shape and teeth location, files are also classified by their cut and grade.

- **Cut:** Refers to how the teeth are cut into the file blank. Common cuts include single-cut (teeth run in one direction), double-cut (two sets of teeth crossing each other), rasp-cut (individual, non-connecting teeth), and curved-cut.
- **Grade:** Refers to the coarseness of the teeth. Grades range from coarse (bastard cut) to medium (second cut) to fine (smooth cut). Finer cuts remove less material but leave a smoother finish.

Understanding these classifications, along with the specific shape and tooth pattern of files like the Barrette file, helps in selecting the right tool for precise tasks like finishing sharp corners or working in confined spaces.

137. Answer: a

Explanation:

Understanding Different Types of Maintenance

Maintenance is crucial for keeping equipment, facilities, and systems running smoothly and efficiently. Different types of maintenance are performed at varying frequencies depending on the task and its purpose. Understanding these types helps in planning and executing maintenance strategies effectively.

Focusing on Daily Maintenance: Routine Maintenance

The question asks about the type of maintenance done on a daily basis. Let's examine the concept of routine maintenance.

- **Routine maintenance** involves tasks that are performed regularly, often daily, weekly, or monthly, to keep equipment or facilities in good working order.
- These are typically simple, repetitive tasks that help prevent minor issues from becoming major problems.

- Examples include daily cleaning, visual inspections, checking fluid levels, lubricating specific parts, or basic operational checks.
- Because these tasks are essential for the immediate operational readiness and longevity of assets, they are often scheduled daily.

Therefore, routine maintenance is the type most commonly associated with daily tasks.

Other Maintenance Types Explained

Let's briefly look at the other options to understand why they are not typically daily activities.

- **Breakdown maintenance:** This is maintenance performed only after equipment has failed. It is reactive rather than proactive. Failures do not necessarily happen daily, so breakdown maintenance is not a daily activity.
- **Ordinary maintenance:** This term can sometimes be used interchangeably with routine maintenance or refer to standard, non-emergency repair work. However, it's not as specific about the frequency as "routine maintenance," which often implies daily or very frequent tasks. If it refers to standard repairs, these are done when needed, not necessarily daily.
- **Special maintenance:** This usually refers to major repairs, overhauls, or unique maintenance tasks that require specialized skills, tools, or significant planning. These are infrequent events, certainly not daily.

Comparing Maintenance Types and Frequency

To further clarify, consider the typical frequency and nature of each type:

Maintenance Type	Typical Frequency	Nature of Work
Routine maintenance	Daily, Weekly, Monthly	Regular checks, cleaning, lubrication, minor adjustments (preventive)
Breakdown maintenance	When failure occurs	Repairing broken equipment (reactive)
Ordinary maintenance	As needed	Standard repairs, general upkeep (mix of preventive/reactive)
Special maintenance	Infrequent (e.g., yearly, multi-yearly)	Major repairs, overhauls, upgrades (planned, extensive)

Based on this comparison, routine maintenance is the type that includes tasks performed on a daily basis.

Routine Maintenance Activities

Common daily routine maintenance activities can include:

- Visual inspection of equipment for any obvious damage or leaks.
- Checking and recording operational parameters (like temperature, pressure, speed).
- Cleaning specific parts to ensure proper functioning.
- Basic lubrication of moving parts if required daily.
- Ensuring safety guards are in place and functional.

These tasks are fundamental to preventive maintenance, aimed at identifying potential issues early and maintaining operational efficiency.

Revision Table – Key Maintenance Concepts

Term	Brief Description	Related Frequency
Routine maintenance	Regular, scheduled tasks	Daily, frequent
Breakdown maintenance	Repair after failure	Unscheduled, reactive
Preventive maintenance	Scheduled tasks to prevent failure	Routine, periodic
Reactive maintenance	Repair after failure	Synonym for Breakdown

Additional Information – Preventive vs Reactive Maintenance

Routine maintenance is a form of preventive maintenance. Preventive maintenance is performed regularly while the equipment is still working to prevent unexpected failure.

In contrast, breakdown maintenance is a form of reactive maintenance, which is performed only after a failure has occurred. Preventive maintenance, including daily routine tasks, is generally preferred over reactive maintenance as it can help avoid costly downtime and extensive repairs.

138. Answer: c

Explanation:

Understanding the Systematic Maintenance Approach

Maintenance is crucial for ensuring that equipment, systems, or facilities operate correctly and efficiently. A systematic approach to maintenance helps in identifying issues, understanding their root causes, and implementing effective solutions. This structured process makes maintenance activities more organized, efficient, and ultimately, more successful in preventing future problems.

Steps in a Systematic Maintenance Process

A typical systematic approach to maintenance follows a logical sequence of steps. Let's break down the key stages involved:

1. **Problem Identification:** The first step is recognizing that a problem exists. This could be a decrease in performance, a complete breakdown, unusual noises, error messages, or any other indicator that something is not working as expected. The problem must be clearly defined.
2. **Cause Analysis:** Once a problem is identified, the next critical step is to determine the underlying cause. What led to the problem? Was it wear and tear, a faulty component, incorrect operation, environmental factors, or something else? Identifying the root cause is essential to prevent recurrence.
3. **Diagnosis:** With a potential cause in mind, diagnosis involves detailed investigation and testing to confirm the cause and assess the extent of the damage or issue. This step uses tools, tests, observations, and historical data to pinpoint the exact nature of the problem and its source. Diagnosis verifies the cause identified in the previous step and provides detailed information needed for planning the repair.
4. **Rectification/Repair:** The final step is taking action to fix the problem and restore the equipment or system to its proper working condition. This involves implementing the necessary repairs, replacing parts, adjusting settings, or carrying out any other corrective actions based on the diagnosis. This step resolves the identified issue.

Analyzing the Given Options

Let's evaluate the sequences provided in the options based on the systematic approach described above:

Option Sequence	Analysis	Systematic?
Problem – Diagnosis – Cause – Rectification	This sequence identifies the problem, then attempts diagnosis. However, diagnosis is typically done to confirm or investigate a *suspected* cause, making 'Cause' coming after 'Diagnosis' less logical in a systematic flow aimed at root cause analysis before detailed investigation.	No
Problem – Measure – Diagnosis – Rectification	'Measure' (presumably meaning measurement or assessment) could be part of Diagnosis, but placing it as a separate step between Problem and Diagnosis doesn't align with the standard flow focusing on finding the underlying cause first or diagnosing the cause.	No
Problem – Cause – Diagnosis – Rectification	This sequence starts with identifying the Problem, then considers potential Causes, followed by Diagnosis (to confirm the cause and extent), and finally Rectification (fixing the problem). This aligns logically with finding out *what* is wrong (Problem), thinking about *why* it happened (Cause), confirming *how* and *where* it happened (Diagnosis), and then *fixing* it (Rectification).	Yes
Problem – Diagnosis – Measure – Rectification	Similar to option 2, this sequence places 'Diagnosis' and 'Measure' before addressing the 'Cause'. A systematic approach often seeks to understand the root cause before extensive diagnosis or measurement for repair planning.	No

Based on this analysis, the sequence "Problem – Cause – Diagnosis – Rectification" represents the most logical and systematic approach to maintenance, emphasizing identifying the problem, understanding the potential reasons (cause), confirming those reasons and assessing the issue (diagnosis), and then fixing it (rectification).

Revision Table: Key Maintenance Steps

Step	Description
Problem	Identifying the symptom or issue.
Cause	Determining the root reason for the problem.
Diagnosis	Investigating and verifying the cause and extent of the issue.
Rectification	Implementing the solution to fix the problem.

Additional Information: Types of Maintenance

While the systematic approach outlines the steps to fix a problem, maintenance itself can be performed in different ways:

- **Corrective Maintenance:** This is performed after a failure has occurred to restore equipment to working order. The systematic approach discussed is a process within corrective maintenance.
- **Preventive Maintenance:** This is performed regularly to lessen the likelihood of equipment failure. It includes activities like inspections, cleaning, lubrication, and minor adjustments, often scheduled based on time or usage.
- **Predictive Maintenance:** This uses data analysis and monitoring (like vibration analysis, thermal imaging) to detect signs of degradation and predict when a failure might occur, allowing maintenance to be scheduled just before failure.
- **Proactive Maintenance:** Focuses on identifying and eliminating the root causes of failures to prevent them from occurring in the first place. This goes beyond just predicting or preventing known failure modes and seeks design or operational improvements.

The systematic approach of Problem – Cause – Diagnosis – Rectification is most directly applicable when responding to an identified 'Problem', which is characteristic of corrective maintenance, but the principles of cause analysis and diagnosis are also valuable in other maintenance types like predictive and proactive approaches.

139. Answer: c

Explanation:

Understanding Recess Height Measurement

Measuring dimensions accurately is crucial in various fields like manufacturing and engineering. Different types of measuring instruments are used for different purposes. The question asks about the specific tool used to measure the height or depth of a recess.

What is Recess Height?

A recess is essentially a cavity or indentation in a surface. The recess height (or often referred to as depth) is the distance from the primary surface down to the bottom of the recess.

Micrometer Types and Their Applications

Micrometers are precision measuring instruments. Let's look at the types mentioned in the options and what they are typically used for:

- **Outside Micrometer:** Used to measure the external dimensions of objects, such as the diameter of shafts, thickness of plates, or length of blocks.
- **Inside Micrometer:** Used to measure internal dimensions, such as the diameter of holes or the distance between two parallel surfaces inside a bore.
- **Depth Micrometer:** Specifically designed to measure the depth of holes, slots, steps, and recesses from a reference surface. It has a base that rests on the reference surface and a measuring rod that extends into the feature being measured.
- **Screw Thread Micrometer:** Used to measure the pitch diameter of screw threads. Its anvils are shaped to fit into the threads.

Measuring Recess Height with a Micrometer

To measure the height (or depth) of a recess, you need an instrument that can bridge the opening of the recess and extend down to its bottom. The depth micrometer is designed precisely for this task. Its base provides a stable reference plane across the opening of the recess, and the measuring rod is lowered until it contacts the bottom of the recess. The distance the rod travels from the base to the bottom is the recess height.

Based on the function of each micrometer type, the instrument best suited for measuring recess height is the depth micrometer.

Why Depth Micrometer is Used for Recess Height

A depth micrometer has a flat base that is placed on the surface from which the depth is to be measured. A rod extends perpendicularly from this base into the recess. The measurement shown on the micrometer thimble and sleeve indicates how far the rod has extended, which directly corresponds to the depth or height of the recess.

Micrometer Type	Primary Application	Suitable for Recess Height?
Outside Micrometer	External dimensions	No
Inside Micrometer	Internal dimensions (e.g., bore diameter)	No
Depth Micrometer	Depth of holes, slots, recesses	Yes
Screw Thread Micrometer	Screw thread pitch diameter	No

Therefore, to accurately measure the recess height, a depth micrometer is the correct tool.

Revision Table: Micrometer Applications

Micrometer Type	Measures
Outside	External length, diameter, thickness
Inside	Internal length, diameter
Depth	Depth of holes, slots, recesses, steps
Screw Thread	Pitch diameter of threads

Additional Information on Depth Micrometer Use

When using a depth micrometer to measure recess height:

- Ensure the base is resting firmly and squarely on the reference surface around the recess.

- Lower the measuring rod gently until it makes contact with the bottom of the recess. Avoid excessive force, which can lead to inaccurate readings or damage.
- Lock the thimble before removing the micrometer to preserve the reading.
- Read the measurement from the barrel and thimble scales, similar to an outside micrometer but typically measuring the distance the rod extends.

Proper handling and technique are essential for obtaining accurate recess height measurements with a depth micrometer.

140. Answer: c

Explanation:

Understanding Bolts in Marine Shaft Couplings

Marine shaft couplings are critical components that connect sections of a ship's propulsion shafting or connect the engine/gearbox to the shaft. These couplings must reliably transmit significant torque and maintain precise alignment, often in demanding marine environments. The type of bolt used in these couplings is vital for ensuring a secure and durable connection.

Identifying the Correct Bolt Type for Marine Couplings

Let's examine the types of bolts commonly used in mechanical applications and consider their suitability for marine shaft couplings:

- **Ellen bolt:** This typically refers to a bolt with a square head, often used in structural or timber applications. While strong, it does not provide the precision fit required for high-torque, high-alignment shaft couplings.
- **Eye bolt:** An eye bolt has a loop or "eye" at one end, used for lifting, rigging, or securing cables. It is not designed for joining coupling halves to transmit rotational power.
- **Headless tapered bolt:** This bolt has a conical or tapered shank and no head. It is designed to be driven into precisely reamed holes, creating a tight, interference fit between the joined components. The tapered design helps in drawing the coupling halves together and ensuring accurate centering and alignment. The interference fit prevents slippage and distributes the load effectively, making it ideal for

transmitting torque in shaft couplings, particularly in marine applications where reliability and alignment are paramount. Being headless often allows for flush mounting within the coupling flange.

- **Hexagonal bolt:** A standard hexagonal bolt is a versatile fastener used in countless applications. While strong, a standard hex bolt in a clearance hole does not provide the necessary interference fit and alignment precision required for critical marine shaft couplings that experience high stress and torque variations.

Based on the requirements for a secure, high-torque, and precisely aligned connection in marine shaft couplings, the headless tapered bolt is the most appropriate choice among the given options.

Why Headless Tapered Bolts are Preferred for Marine Shaft Couplings

The primary reasons headless tapered bolts are commonly used in marine shaft couplings include:

- **Precision Fit:** The tapered shape allows the bolt to create an interference fit when driven into a corresponding tapered or reamed hole in the coupling halves. This ensures accurate alignment and concentricity between the joined shafts.
- **Torque Transmission:** The tight fit significantly increases the friction and shear resistance, enabling the coupling to transmit high levels of torque without relative movement or slip between the flanges.
- **Load Distribution:** The tapered design helps distribute the stress evenly across the bolt and the mating surfaces of the coupling.
- **Security:** The interference fit provides a vibration-resistant connection, crucial in the dynamic environment of a ship.
- **Ease of Assembly/Disassembly (relative):** While requiring careful installation, the taper facilitates drawing the components together. Specialized tools are often used for precise tightening and removal.

Therefore, for the critical application of joining sections of a marine propulsion shaft, a bolt that provides a reliable, high-precision, and high-torque connection is essential. The headless tapered bolt fulfills these requirements effectively.

Bolt Type	Common Application	Suitable for Marine Shaft Coupling?	Reason
Ellen bolt	Structural, Timber	No	Lacks precision fit for high torque/alignment
Eye bolt	Lifting, Securing	No	Not designed for transmitting rotational power
Headless tapered bolt	Shaft Couplings (Marine, Industrial), Precision Assembly	Yes	Provides precise interference fit, high torque transmission, alignment
Hexagonal bolt	General Fastening	No (in standard form)	Standard fit lacks precision/interference needed

Revision Table: Key Concepts

Term	Explanation	Relevance to Marine Couplings
Marine Shaft Coupling	Connects engine/gearbox to propeller shaft or sections of shafting.	Transmits power, maintains alignment.
Tapered Bolt	Bolt with a conical shank.	Creates interference fit for precision and torque transmission.
Interference Fit	A fit where one part is slightly larger than the hole it goes into, requiring force to assemble.	Prevents slip, ensures alignment, crucial for torque.
Torque Transmission	The ability to transfer rotational force from one component to another.	Primary function of a shaft coupling.

Additional Information on Marine Shaft Couplings and Bolts

Marine shaft couplings come in various types, including flanged couplings, rigid couplings, and flexible couplings. Flanged couplings, which are very common, often use bolts to join the two flanges. The type of bolt used is specified based on the power rating, speed, and criticality of the shaft system.

The installation of tapered bolts requires careful preparation of the holes and controlled tightening to achieve the desired interference fit. Improper installation can lead to misalignment, fretting, or failure of the coupling.

Other types of precision bolts, sometimes with heads but still tapered, might also be used depending on the specific coupling design and manufacturer specifications. However, the principle of using a tapered shank for a precise, interference fit remains key for critical shaft connections.

Maintaining the integrity of marine shaft couplings and their bolts through regular inspection and maintenance is vital for the safe and efficient operation of a vessel's propulsion system.

141. Answer: d

Explanation:

Understanding Abrasives for Lapping Soft Metals

Lapping is a precision finishing process used to improve the surface finish, flatness, and dimensional accuracy of a workpiece. It typically involves rubbing the workpiece against a lap using a fine abrasive suspended in a liquid vehicle (slurry). The choice of abrasive is crucial for achieving the desired results on specific materials.

Why Abrasive Choice Matters in Lapping

The abrasive particles do the cutting in the lapping process. Different materials require different types and sizes of abrasives. For softer materials like soft ferrous and non-ferrous metals, the abrasive needs to be hard enough to cut the material effectively but not so hard or coarse that it causes excessive scratching or damage. The abrasive should also break down slowly and evenly during the process.

Common Lapping Abrasives

Several types of abrasives are used in lapping, each with properties suitable for different applications:

- Aluminium Oxide (Al_2O_3): A common, versatile abrasive. It is tough and fractured, making it suitable for many materials, including steel, cast iron, and non-ferrous metals.
- Silicon Carbide (SiC): Harder and sharper than aluminium oxide, used for lapping hard materials like cemented carbides, ceramics, and some hard steels. It breaks down more quickly.
- Boron Carbide (B_4C): Extremely hard, used for very hard materials and situations requiring high stock removal rates.
- Diamond: The hardest abrasive, used for lapping superhard materials like ceramics, carbides, and superalloys.

Selecting Abrasives for Soft Ferrous and Non-Ferrous Metals

Soft ferrous metals (like low-carbon steel) and non-ferrous metals (like aluminium, copper, brass) are relatively easy to cut. While silicon carbide is harder, aluminium oxide is often preferred for lapping these softer materials because:

- It provides a good balance of cutting ability and surface finish.
- It is less likely to cause excessive surface damage or embedment in softer materials compared to harder abrasives like silicon carbide or boron carbide.
- It is cost-effective and widely available.

Aluminium oxide effectively removes material from these softer surfaces, leaving a smooth, uniform finish without deep scratches.

Analyzing the Given Options

Let's consider the provided options:

1. copper oxide: Copper oxide is not typically used as a primary abrasive for lapping metals. It is much softer than traditional abrasives like aluminium oxide or silicon carbide.
2. Fly Ash: Fly ash is a byproduct of coal combustion and consists of fine particles. While it has some abrasive properties and can be used in certain polishing or mild abrasive

applications, it is not a standard or effective abrasive for precision metal lapping, especially compared to engineered abrasives.

3. sodium chloride: Sodium chloride (common salt) is a crystalline compound but is not an abrasive material used in precision machining or lapping applications for metals. It is relatively soft and soluble.
4. Aluminium oxide: As discussed, aluminium oxide is a widely used and effective abrasive for lapping a variety of materials, including soft ferrous and non-ferrous metals. Its properties make it suitable for achieving good surface finishes on these materials.

Based on the characteristics and typical applications of these materials in lapping, Aluminium oxide is the most suitable abrasive among the given options for lapping soft ferrous and non-ferrous metals.

Abrasive Type	Hardness (Mohs Scale, Approx.)	Typical Applications
Aluminium Oxide (Al_2O_3)	9	Steel, Cast Iron, Non-ferrous metals (Al, Cu, Brass), Ceramics
Silicon Carbide (SiC)	9.5	Hard steels, Cemented Carbides, Ceramics, Glass
Boron Carbide (B_4C)	9.75	Very hard materials, High stock removal
Diamond	10	Ceramics, Carbides, Superalloys, Superhard materials
Copper Oxide	~3-4	Not a standard lapping abrasive for metals
Fly Ash	Varies (low)	Mild abrasive/filler, not precision metal lapping
Sodium Chloride	~2.5	Not an abrasive for metal lapping

Conclusion

Considering the properties and common uses of abrasives in precision finishing processes like lapping, Aluminium oxide is the appropriate choice for working with soft ferrous and non-ferrous metals from the given options.

Revision Table: Lapping Abrasives

Key Term	Brief Description	Relevance to Question
Lapping	Precision finishing process for surface flatness and finish.	Process where the abrasive is used.
Abrasive	Hard material particles that remove material from a surface.	The subject of the question - which abrasive is needed.
Soft Ferrous Metals	Iron-based metals with low carbon content, relatively soft.	One type of material being lapped.
Non-ferrous Metals	Metals without significant iron content (Al, Cu, Brass etc.).	Another type of material being lapped.
Aluminium Oxide	Common, versatile abrasive (Al_2O_3).	The identified suitable abrasive.

Additional Information: Lapping Process Details

The lapping process involves several factors that influence the outcome:

- **Lap Material:** The lap itself can be made of various materials like cast iron, copper, tin, or ceramic. The workpiece rubs against the lap with the abrasive slurry between them.
- **Vehicle:** The abrasive is mixed with a liquid vehicle (like oil, water, or grease) to form a slurry. The vehicle helps distribute the abrasive, cool the surface, and carry away swarf (removed material).
- **Abrasive Grit Size:** The size of the abrasive particles affects the material removal rate and the final surface finish. Coarser grits remove material faster but leave a rougher finish; finer grits remove material slower but produce a smoother finish.
- **Pressure and Speed:** The pressure applied and the relative speed between the lap and the workpiece also affect the lapping efficiency and surface quality.

Proper selection of the abrasive, vehicle, grit size, and process parameters is essential for successful lapping of different materials, including soft ferrous and non-ferrous metals.

142. Answer: d

Explanation:

Understanding DC Voltage in Pneumatic Solenoid Systems

Pneumatic systems often use electrical control signals to operate valves that direct airflow. These valves are commonly controlled by solenoids, which are essentially electromagnetic switches. When an electric current passes through the coil of the solenoid, it generates a magnetic field that moves a piston or spool, changing the state of the valve (e.g., opening or closing a port).

Solenoids in Pneumatic Automation

Solenoids are key components in automated pneumatic systems. They interface the electrical control system (like a PLC or control panel) with the pneumatic power system. The voltage and current supplied to the solenoid coil determine if it activates or deactivates the valve.

The question asks about the typical DC voltage range used for these solenoids in pneumatic systems. While solenoids can be designed for various AC and DC voltages, specific ranges are standard in industrial automation and control systems.

Typical DC Voltage Ranges for Solenoids

Control systems, particularly those involving electronics and lower power consumption for control signals, often utilize low voltage DC. This offers benefits in terms of safety, ease of wiring, and compatibility with electronic control circuits.

Common voltage ranges encountered in industrial control, including for pneumatic solenoids, include both AC and DC options. However, for DC applications in control systems, a specific lower voltage range is very prevalent.

- Higher AC voltages (like 110 V, 220 V, 230 V, 440 V) are common for powering larger machinery or main distribution but are less frequently used for the fine control

signals of individual valve solenoids, especially in DC systems.

- Lower DC voltages are preferred for control circuits due to safety and compatibility with electronic components.

Analyzing the given options for DC voltage ranges:

- 110 V to 220 V: This range typically involves higher AC voltages, not standard low-voltage DC control.
- 220 V to 250 V: This also represents higher voltage AC power, not typical low-voltage DC control.
- 230 V to 440 V: This covers common industrial three-phase or higher single-phase AC voltages, not low-voltage DC control.
- 12 V to 24 V: This range, specifically 24 V DC, is a widely accepted standard for control voltage in industrial automation systems globally. 12 V DC is also used in some applications, particularly mobile or vehicle-based systems, or less complex controls.

Therefore, the most common and standard DC range for solenoids in pneumatic control systems is 12 V to 24 V.

Benefits of Using Low Voltage DC (12V to 24V)

Using low DC voltages like 12V or 24V for pneumatic solenoids offers several advantages:

- **Safety:** Lower voltages significantly reduce the risk of electrical shock compared to higher AC voltages.
- **Compatibility:** This voltage range is standard for PLCs, sensors, and other control system components, simplifying integration.
- **Wiring:** Lower voltage wiring often has less stringent insulation requirements, potentially simplifying installation (though proper wiring practices are always essential).
- **Battery Backup:** Systems can often be designed with battery backup using 12V or 24V batteries to maintain operation during power outages.
- **Reduced Electromagnetic Interference (EMI):** DC systems can sometimes generate less EMI than comparable AC systems.

Conclusion

Based on standard practices in industrial automation and the safety and compatibility benefits for control circuits, the typical DC voltage range for solenoids in pneumatic

systems is 12 V to 24 V.

Common Solenoid Voltages

Type	Typical Voltage Ranges	Application Context
DC Solenoids (Control)	12 V, 24 V	Industrial automation, vehicles, battery-powered systems, electronic controls
AC Solenoids (Control)	24 V AC, 110 V AC, 230 V AC	Industrial automation (less common for new control systems compared to DC, but still widely used)
Higher Voltage Applications	110 V – 480 V AC	Motor control, main power distribution (not typically for individual valve solenoids)

Revision Table: Pneumatic Solenoid DC Voltage

Concept	Key Point
Solenoid Function	Electromagnetic switch controlling pneumatic valves.
Typical DC Range	12 V to 24 V DC.
Why Low DC?	Safety, compatibility with control systems (PLCs, sensors), ease of wiring.
Other Voltages	AC options (24V, 110V, 230V) are also used, but low DC is standard for many modern control circuits.

Additional Information: Pneumatic System Components

Beyond solenoids and valves, a typical pneumatic system includes several other essential components:

- **Air Compressor:** Generates compressed air.

- **Air Treatment Unit (FRL – Filter, Regulator, Lubricator):** Cleans the air (filter), controls pressure (regulator), and sometimes adds lubrication (lubricator) before it enters the system.
- **Actuators:** Devices that convert pneumatic energy into mechanical motion, such as cylinders (linear motion) or motors (rotary motion).
- **Fittings and Hoses:** Connect components and route compressed air throughout the system.
- **Sensors:** Provide feedback about system status, such as pressure sensors, position sensors for cylinders, etc. These also typically operate on low DC voltage (like 24V DC).

Understanding the standard operating voltages of components like solenoids and sensors is crucial for designing, troubleshooting, and maintaining pneumatic control systems.

143. Answer: a

Explanation:

Understanding Tap Types and Chamfer Lengths

Threading is a common process in manufacturing and repair, and taps are essential tools used to cut internal screw threads. Taps come in different types, each designed for a specific purpose, particularly concerning the starting and finishing of threads, especially in blind holes (holes that do not go all the way through a workpiece).

The key difference between common tap types lies in the length of their chamfer. The chamfer is the tapered cutting section at the leading end of the tap. This taper helps the tap to enter the hole easily and distribute the cutting action over several teeth.

Let's look at the three main types of taps and their typical chamfer lengths:

- **Taper Tap:** This tap has the longest chamfer, typically extending over 7 to 10 threads. The long taper allows the tap to start cutting easily and gradually, making it suitable for starting new threads, especially in through holes or where starting is difficult.
- **Plug Tap:** This tap has a medium-length chamfer, usually spanning 4 to 6 threads. It's the most common general-purpose tap, used after a taper tap (if needed) or as

the first tap in many applications where starting isn't too difficult. It cuts threads deeper than a taper tap.

- **Bottoming Tap:** This tap has the shortest chamfer, covering only 1 to 2 threads. It is designed to cut threads almost to the very bottom of a blind hole, typically used after a taper tap and/or plug tap have already cut the majority of the thread depth.

Here is a summary:

Tap Type	Purpose	Typical Chamfer Length (Threads/Teeth)
Taper Tap	Starting threads, through holes	7 to 10
Plug Tap	General purpose, deeper threads	4 to 6
Bottoming Tap	Finishing threads in blind holes	1 to 2

The question asks about the number of chamfered teeth of a "taper tap bottom". The phrasing "taper tap bottom" can be slightly confusing. Given the options, and specifically the option "1 to 2", it strongly suggests that the question is referring to the chamfer length characteristic of a **bottoming tap**, which has the shortest chamfer designed for the bottom of a blind hole.

Considering the standard classifications and the provided options, the range of 1 to 2 chamfered teeth is characteristic of a bottoming tap. While the term "taper tap" usually refers to the tap type with the longest chamfer (7-10 teeth), the combination with "bottom" and the presence of the 1-2 option indicates the intended answer relates to the tool used to finish threads at the bottom of a hole, which is the bottoming tap.

Let's evaluate the options based on standard tap types:

- 1 to 2 teeth: This is characteristic of a **bottoming tap**.
- 4 to 6 teeth: This is characteristic of a **plug tap**.
- 20 to 30 teeth: This range is significantly larger than typical tap chamfers. Tap chamfers are usually measured in a small number of threads/teeth.
- 7 to 10 teeth: This is characteristic of a **taper tap**.

Therefore, aligning the options with standard tap terminology, the number of chamfered teeth that corresponds to the range 1 to 2 is found on a bottoming tap. Assuming the question intends to refer to the tap type used for threading to the bottom of a hole, or perhaps uses "taper tap bottom" loosely to imply this function, the correct range from the options is 1 to 2 chamfered teeth.

Revision Table: Tap Chamfer Details

Tap Type	Chamfer Purpose	Chamfer Length (Approx. Threads)	When to Use
Taper Tap	Easy starting, guide	7-10	To start new threads, through holes, hard materials
Plug Tap	General threading	4-6	After taper tap, or as first tap in many cases
Bottoming Tap	Thread to bottom of blind hole	1-2	After plug/taper tap to finish threads in blind holes

Additional Information on Threading Taps

Threading taps are typically made from high-speed steel (HSS) or carbon steel. HSS taps are more durable and suitable for production work, while carbon steel taps are generally used for hand tapping in softer materials.

Taps can be used for hand tapping (using a tap wrench) or machine tapping. Machine tapping requires precise alignment and speed control.

Lubrication is crucial when tapping to reduce friction, heat, and prevent tap breakage. Different materials require specific cutting fluids.

Taps come in various thread forms (e.g., UNC, UNF, Metric) and sizes. Selecting the correct tap for the desired thread is essential.

When tapping a blind hole, it's common to use a set of taps: first the taper tap, then the plug tap, and finally the bottoming tap to ensure full thread depth is achieved.

144. Answer: a

Explanation:

Understanding High Viscosity Fluids

Viscosity is a fundamental property of fluids that describes their resistance to flow. Think of it as the internal friction within the fluid. Fluids with low viscosity, like water or alcohol, flow easily. Fluids with high viscosity, like honey or tar, resist flow strongly.

Let's look at the options provided and see how they relate to the characteristics of a high viscosity fluid.

Analyzing High Viscosity Characteristics Options

We will examine each option to determine which one is a characteristic of a high viscosity fluid.

- **Semi-liquid state:** High viscosity fluids flow very slowly. Consider substances like tar or very thick syrup. At room temperature, they pour extremely slowly and might even appear to hold their shape for a while before eventually flattening out. This slow, resistant flow can give the impression of a state that is thicker and less free-flowing than a typical liquid, almost like a state between a solid and a liquid. While "semi-liquid" isn't a precise scientific state of matter, it describes the consistency and behavior of many high viscosity substances as perceived in everyday life.
- **Less power consumption:** Moving or pumping fluids requires overcoming their internal resistance to flow (viscosity). For a given flow rate, a higher viscosity fluid creates more resistance, meaning more energy (and thus more power) is needed to pump it compared to a low viscosity fluid. Therefore, high viscosity typically leads to **more** power consumption, not less. This option is incorrect.
- **Slow operation:** When high viscosity fluids are involved in processes (like flow through pipes, mixing, or pumping), the operations tend to be slower compared to using low viscosity fluids under similar conditions. This is a direct result of the fluid's high resistance to flow. The fluid moves sluggishly.
- **Pressure drop:** As a fluid flows through a system (like a pipe), there is a drop in pressure due to friction between the fluid layers and between the fluid and the pipe walls. This pressure drop is influenced by viscosity. For a given flow rate and geometry, a high viscosity fluid will experience a **larger** pressure drop compared to a

low viscosity fluid because of the increased internal friction. While pressure drop occurs with any flowing fluid, a significantly larger pressure drop is characteristic of high viscosity fluids. However, "pressure drop" by itself is not uniquely characteristic; the *magnitude* of the pressure drop is.

Comparing Characteristics of High Viscosity Fluids

Based on the analysis, both "Semi-liquid state" (describing the appearance/consistency) and "Slow operation" (describing the flow behavior in processes) and increased "Pressure drop" are related to high viscosity. However, among the given options, "Semi-liquid state" often captures the commonly perceived nature of extremely viscous substances.

Characteristics vs. Viscosity Level

Characteristic	Low Viscosity Fluid	High Viscosity Fluid
Flow Rate (under same pressure)	Fast	Slow
Resistance to Flow	Low	High
Pumping Power	Low	High
Pressure Drop (for same flow rate)	Low	High
Consistency/Appearance	Runny	Thick / Semi-liquid like

Conclusion on High Viscosity Fluid Characteristics

Considering the options, the description "Semi-liquid state" aligns with how many extremely high viscosity fluids appear and behave, flowing very slowly and resisting changes in shape strongly, thus resembling a state between a typical liquid and a solid.

Revision Table: Key Concepts in Fluid Viscosity

Fluid Viscosity Terms and Definitions

Term	Definition	Relevance to High Viscosity Fluids
Viscosity (μ)	A measure of a fluid's internal resistance to flow.	High value means the fluid resists flow strongly.
Shear Stress (τ)	The force per unit area required to make the fluid flow.	Higher for high viscosity fluids at a given shear rate.
Shear Rate ($\frac{du}{dy}$)	The rate at which parallel layers of fluid move past each other.	Viscosity relates shear stress to shear rate ($\tau = \mu \frac{du}{dy}$ for Newtonian fluids).
Newtonian Fluid	Fluid where viscosity is constant regardless of shear rate and time.	Many common fluids (water, air, simple oils) are approximately Newtonian. Behavior is predictable based on μ .
Non-Newtonian Fluid	Fluid where viscosity changes with shear rate or time.	Examples: ketchup, paint, blood, some polymers. High viscosity fluids can be Newtonian or non-Newtonian. Their behavior is more complex.

Additional Information on Viscosity and Flow

Viscosity is highly dependent on temperature. For most liquids, viscosity decreases as temperature increases (they become more "runny"). For most gases, viscosity increases as temperature increases.

Understanding fluid viscosity is crucial in many engineering applications, such as:

- Designing piping systems (predicting pressure drop and flow rates).
- Selecting pumps (determining required power).
- Designing lubrication systems (ensuring proper film thickness).
- Analyzing flow in natural systems (rivers, blood circulation).

High viscosity fluids can pose challenges in industrial processes, often requiring specialized equipment or heating to facilitate flow and processing.

145. Answer: b

Explanation:

Understanding Striking Tools in Workshop Practice

Striking tools are hand tools used to apply force to an object, typically to drive fasteners, shape materials, or break objects. The most common type of striking tool is the hammer, but other tools like mallets, punches, and chisels (when struck by a hammer) also fall under this category.

Let's examine the given options to determine which one is a striking tool:

- **V-Block:** A V-block is a metal block with a V-shaped groove, used to hold cylindrical workpieces for operations like drilling, milling, or inspection. It is a work-holding or supporting tool, not a striking tool.
- **Ball peen hammer:** A ball peen hammer is a type of hammer commonly used in metalworking. It has a hardened steel head with two faces: a flat face used for striking and driving, and a rounded or "ball" peen face used for shaping metal, riveting, and peening surfaces. This tool is explicitly designed for striking.
- **File:** A file is a hand tool used to cut, shape, and smooth material by removing small amounts from a workpiece. It works by pushing or pulling it across the surface. It is an abrasive cutting tool, not a striking tool.
- **Scribing block:** A scribing block, also known as a surface gauge, is used for marking parallel lines on a workpiece from a surface plate, or for checking heights. It is a marking and measuring tool, not a striking tool.

Based on the functions of these tools, the **Ball peen hammer** is clearly designed for striking and applying impact force.

Identifying the Striking Tool

Comparing the descriptions:

- V-Block: Holds work.
- Ball peen hammer: Applies force through impact (striking).
- File: Removes material by abrasion.
- Scribing block: Marks lines or checks height.

Therefore, the tool that fits the definition of a striking tool among the options is the Ball peen hammer.

Tool	Primary Function	Is it a Striking Tool?
V-Block	Work holding/support	No
Ball peen hammer	Striking, shaping, riveting	Yes
File	Cutting/smoothing by abrasion	No
Scribing block	Marking/measuring	No

The Ball peen hammer's primary use involves hitting objects or materials, making it a fundamental striking tool in mechanical and metalworking trades.

Revision Table: Workshop Hand Tools

Tool Type	Examples	Common Uses
Striking Tools	Hammers (Ball peen, Claw), Mallets, Sledgehammers	Driving nails, shaping metal, demolition, assembling
Cutting Tools	Files, Saws, Chisels, Snips	Removing material, shaping, cutting sheet metal
Measuring Tools	Rulers, Calipers, Micrometers, Scribing blocks	Taking dimensions, marking out
Holding/Supporting Tools	Vices, Clamps, V-Blocks	Securing workpieces
Fastening Tools	Screwdrivers, Wrenches	Tightening/loosening fasteners

Additional Information: Types of Hammers

Hammers come in many varieties, each designed for specific tasks. Some common types include:

- **Claw Hammer:** Used for driving nails and pulling them out. Common in woodworking.
- **Ball Peen Hammer:** Used in metalworking for shaping metal, striking punches and chisels, and riveting.
- **Sledgehammer:** A large, heavy hammer used for heavy-duty striking, such as breaking concrete or driving large stakes.
- **Rubber Mallet:** Has a rubber head and is used for striking surfaces without causing damage, often for assembling furniture or working with sheet metal.
- **Dead Blow Hammer:** Contains lead or sand in the head to minimize rebound when striking. Used for striking surfaces without marring and reducing bounce-back.

Understanding the specific function of each tool is crucial for safe and effective work in any trade or workshop environment.

146. Answer: a

Explanation:

Understanding Metal Joining Processes

The question asks to identify a process for joining similar metals using heat, with or without pressure, and with or without filler material. Let's examine the given options to find the best fit.

Analyzing the Options for Joining Metals

- **Welding:** This is a fabrication process that joins materials, usually metals or thermoplastics, by causing coalescence. This is often done by melting the workpieces and adding a filler material to form a pool of molten material (the weld pool) that cools to become a strong joint. Pressure may also be used in conjunction with heat, or by itself. Welding is widely used to join similar metals, and it can be done with or without filler material depending on the specific welding process and application.
- **Nut and bolt joint:** This is a mechanical fastening method. It involves using separate components like nuts and bolts to hold two or more pieces together. It does not involve the application of heat to melt and join the materials at a molecular level.

- **Soldering:** This is a joining process used to join different types of metals together by melting a filler metal (solder) into the joint. The key characteristic of soldering is that the filler metal has a much lower melting point than the workpieces, and the workpieces themselves are not melted. While it uses heat and filler, it's primarily a low-temperature process and often used for electrical connections or non-structural joints, and the base metals don't fuse.
- **Forming:** This is a manufacturing process where metal is shaped into a desired geometry by plastic deformation. Examples include bending, stamping, and forging. Forming does not involve joining two separate pieces of material; it only changes the shape of a single piece or blank.

Detailed Examination of Welding

Based on the definitions, welding perfectly matches the description provided in the question:

- **Joining similar metals:** Welding is commonly used for this purpose (e.g., steel to steel, aluminum to aluminum).
- **Application of heat:** Most welding processes use heat to melt the base metals and/or filler material.
- **With or without application of pressure:** Some welding processes, like resistance welding or forge welding, use pressure. Others, like arc welding or gas welding, primarily rely on heat and might not use significant external pressure.
- **Addition of filler material:** Many welding processes use filler material to add volume and strength to the joint. However, some processes, like autogenous welding, join the base metals directly without adding filler material.

Therefore, welding is the comprehensive term that covers the process described.

Process	Uses Heat?	Uses Pressure?	Uses Filler Material?	Joins Similar Metals?	Base Metal Melts?	Fits Question?
Welding	Yes (usually)	Yes (sometimes)	Yes (sometimes)	Yes	Yes (usually)	Yes
Nut and bolt joint	No	Yes	N/A (mechanical)	Yes	No	No
Soldering	Yes	No (typically)	Yes	Yes (can be)	No	No (base metal doesn't melt)
Forming	Yes (sometimes, e.g., hot forming)	Yes	No	N/A (shapes one piece)	No	No

Conclusion on Metal Joining Process

Comparing the options against the criteria given in the question, welding is the only process that involves joining similar metals by applying heat, with or without pressure, and with or without filler material, typically resulting in the melting and fusion of the base materials.

Revision Table: Key Concepts in Metal Joining

Term	Definition	Key Characteristics
Welding	Joining materials by causing coalescence, usually through melting the workpieces.	Heat, often pressure, often filler, base metal fuses.
Soldering	Joining metals using a low-melting-point filler metal; base metals do not melt.	Heat, filler (low MP), no base metal melting.
Brazing	Similar to soldering but uses a higher-melting-point filler metal; base metals do not melt.	Heat, filler (higher MP than solder), no base metal melting.
Mechanical Fastening	Joining parts using mechanical devices like bolts, rivets, screws.	No heat or melting of base materials.

Additional Information on Welding Processes

Welding encompasses a wide range of specific techniques. These processes are classified based on the energy source used, such as electrical arc, gas flame, laser, electron beam, friction, and ultrasound. Some common welding processes include:

- Shielded Metal Arc Welding (SMAW)
- Gas Metal Arc Welding (GMAW / MIG)
- Gas Tungsten Arc Welding (GTAW / TIG)
- Flux-Cored Arc Welding (FCAW)
- Submerged Arc Welding (SAW)
- Resistance Welding (Spot, Seam, etc.)
- Oxy-Acetylene Welding (OAW)
- Laser Beam Welding (LBW)
- Electron Beam Welding (EBW)

Each process has its own advantages, disadvantages, and typical applications, but they all fundamentally involve joining materials, often similar metals, using heat and sometimes pressure and/or filler material, leading to material fusion.

147. Answer: c

Explanation:

Understanding Lubricants from Petroleum Refining

Lubricants are substances used to reduce friction between moving surfaces. They can come from various sources, including petroleum, plants, animals, or they can be synthesized chemically. The question asks about a specific type of lubricant obtained through fractional distillation of petroleum.

Fractional Distillation of Petroleum

Petroleum, or crude oil, is a complex mixture of hydrocarbons. Fractional distillation is a key process used in petroleum refineries to separate this mixture into different components (fractions) based on their boiling points. The crude oil is heated, and the vapors rise up a fractionating column. As the vapors cool at different levels, they condense into liquids, which are collected as different fractions. These fractions range from light gases (like LPG) at the top to heavy residues (like bitumen) at the bottom. Various products like gasoline, kerosene, diesel, and lubricating oils are obtained through this process.

Analyzing the Options

Let's look at the given options to identify which lubricant type is derived from the fractional distillation of petroleum:

- **Coolant:** Coolants are primarily used to regulate temperature, usually by carrying away heat. While some coolants might have lubricating properties or be mixed with lubricants (like antifreeze in engine coolant), a coolant itself is not typically classified as a lubricant obtained directly by the fractional distillation of petroleum.
- **Synthetic oils:** Synthetic oils are lubricants produced artificially from chemical compounds rather than being extracted directly from crude petroleum through physical processes like distillation. They are synthesized to have specific properties that may not be found in natural oils.
- **Mineral Oils:** Mineral oils are derived from crude petroleum. They are obtained from the heavier fractions produced during the fractional distillation of petroleum, after further refining processes like solvent extraction, dewaxing, and hydrotreating. These

processes remove impurities and undesirable components from the distilled fractions to produce base oils, which are the primary component of many lubricating oils. Thus, mineral oils are indeed obtained as a product of petroleum processing, starting with fractional distillation.

- **Fatty oils:** Fatty oils, also known as natural oils or vegetable oils/animal fats, are derived from plants (e.g., castor oil, palm oil) or animals (e.g., lard, tallow). These are esters of glycerol with fatty acids and are chemically different from petroleum-based lubricants. They are not obtained from the fractional distillation of petroleum.

Based on the analysis, mineral oils are the type of lubricant obtained by the fractional distillation and subsequent refining of petroleum.

Sources of Different Lubricant Types

Lubricant Type	Primary Source	Obtained from Petroleum Distillation?
Coolant	Various (e.g., water, glycol, synthetic compounds)	Indirectly or not typically classified this way
Synthetic Oils	Chemical Synthesis	No (Synthesized)
Mineral Oils	Crude Petroleum	Yes (via fractional distillation & refining)
Fatty Oils	Plants or Animals	No

Revision Table: Lubricants and Sources

Let's quickly review the key points about the sources of different lubricants related to petroleum.

- Fractional distillation separates crude petroleum into various fractions.
- Lubricating oil fractions are obtained from this process.
- Further refining of these fractions yields mineral base oils.
- Mineral oils are petroleum-based lubricants.

Additional Information: Petroleum Refining Process

Petroleum refining is a complex process that goes beyond simple distillation. After fractional distillation separates crude oil into broad categories based on boiling point, various other processes are used to improve the quality and yield of specific products. These include:

- **Cracking:** Breaking down heavier hydrocarbon molecules into lighter, more valuable ones.
- **Reforming:** Reshaping hydrocarbon molecules to improve properties like octane number (for gasoline).
- **Alkylation and Polymerization:** Combining smaller molecules to create larger ones.
- **Treating:** Removing impurities like sulfur, nitrogen, and metals.
- **Lubricant Processing:** Specific processes like solvent extraction, dewaxing, and hydrofinishing are used on the lubricant fractions from distillation to produce high-quality base oils suitable for various lubricant applications.

Mineral oils are a direct result of these refining steps applied to the heavier fractions obtained from the initial fractional distillation of petroleum.

148. Answer: a

Explanation:

Recognizing Drills for Soft Metals

Identifying the correct drill bit for a specific material like soft metal is crucial for efficient and safe drilling. Drill bits come in various designs optimized for different materials, and one key feature that distinguishes them is the helix angle.

Understanding Drill Bit Helix Angle

The helix angle of a drill bit is the angle formed by the cutting edge (or flute) relative to the axis of the drill bit. This angle significantly impacts how chips are removed and the cutting action itself. Different helix angles are suited for different materials:

- **Small Helix Angle:** Typically ranges from 10° to 20° . These drills have a strong cutting edge and are designed for hard, brittle materials like cast iron or hardened steel. The chips tend to be short and break easily.

- **Normal Helix Angle:** Usually around 25° to 30° . This is a general-purpose design suitable for a wide range of materials, including mild steel and various plastics.
- **Large Helix Angle:** Generally falls between 30° and 45° (sometimes called "fast helix"). These drills have a sharper cutting edge and flutes that promote rapid chip evacuation. They are ideal for drilling soft, gummy, and ductile materials, such as soft metals like aluminum, copper, brass, and soft plastics, where chips tend to be long and continuous.

The large helix angle in drills for soft metals helps prevent chips from clogging the flutes, which is a common problem when drilling ductile materials. The faster spiral angle lifts the continuous chips away from the cutting area efficiently.

Analyzing the Options

Let's evaluate the given options based on our understanding of drill bit design:

- **Option 1: By the large helix angle**

As discussed, a large helix angle is specifically designed for soft and ductile materials to facilitate chip removal and provide a keen cutting edge suitable for such materials. This aligns with our knowledge.

- **Option 2: Based on cutting parameters**

Cutting parameters like speed and feed rate are adjustments made during the drilling operation based on the material and the chosen drill bit. They are not features of the drill bit itself used for identification.

- **Option 3: By the small helix angle**

A small helix angle is suitable for hard and brittle materials, not soft metals. So, this is incorrect.

- **Option 4: Based on required hole diameter**

The required hole diameter determines the size of the drill bit needed, not the specific design features like helix angle that make it suitable for a particular material.

Therefore, the primary way to recognize a drill bit designed for soft metals is by observing its large helix angle.

Comparison of Drill Helix Angles for Different Materials

Helix Angle Type	Typical Angle Range	Suitable Materials	Chip Characteristics
Small Helix	10° – 20°	Hard, brittle materials (Cast Iron, Hard Steel)	Short, broken chips
Normal Helix	25° – 30°	General Purpose (Mild Steel, Plastics)	Varies, typically coiled or segmented
Large Helix	30° – 45°	Soft, ductile materials (Aluminum, Copper, Brass, Soft Plastics)	Long, continuous chips

Conclusion

A drill bit designed for soft metals is recognized by its large helix angle, which helps manage the continuous chips produced when drilling ductile materials.

Revision Table: Drill Bits for Soft Metals

Key features to remember when selecting drills for soft metals:

- Feature: Helix Angle
- Value for Soft Metals: Large (30° – 45°)
- Reason: Promotes rapid chip evacuation
- Benefit: Prevents flute clogging

Additional Information: Drill Bit Terminology

Beyond the helix angle, other parts of a drill bit also play a role in its performance:

- **Flutes:** Grooves that spiral around the body of the drill. They allow chips to escape and coolant to reach the cutting edges. The shape and size are affected by the helix angle.
- **Lands:** The raised areas between the flutes that provide body clearance and support.
- **Point Angle:** The angle formed at the tip of the drill bit by the cutting edges. A common angle is 118° for general purpose, but angles vary for different materials

(e.g., 135° split point for harder materials, sharper angles for softer materials).

- **Web:** The thin portion of the drill bit body between the flutes, running down the center.

Understanding these features helps in selecting the most appropriate drilling tool for any given application and material.

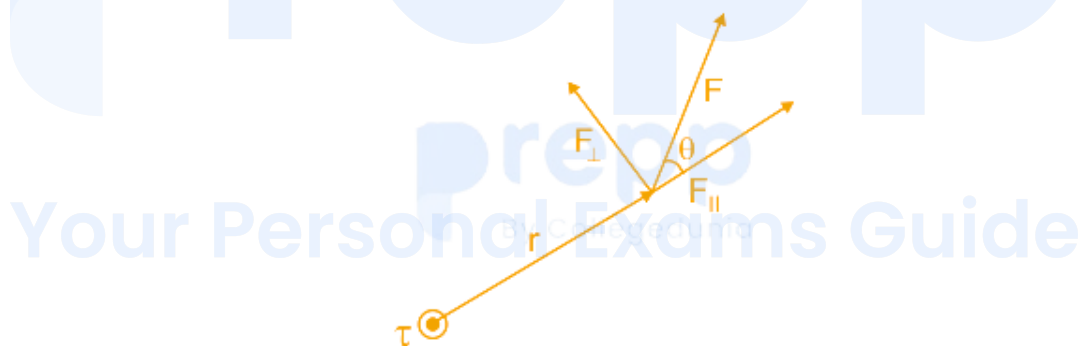
149. Answer: c

Explanation:

CONCEPT:

Torque (τ):

- The physical quantity, similar as force that **causes the rotational motion**. It is the **cross product of the force with the perpendicular distance** between the axis of rotation and the point of application of the force with the force.



- Torque (τ) = $r \times F = r F \sin \theta$

[Where, r = distance from point of application of force (in metre), f = force (in Newton) and θ = angle between $r \rightarrow$ and $F \rightarrow$]

- Also, Torque (τ) = $r_{\perp} F = r F_{\perp}$

Where, $r_{\perp}$ = component of the distance in the direction perpendicular to the $F \rightarrow$

$F_{\perp}$ = component of the force in the direction perpendicular to $r \rightarrow$

EXPLANATION:

We know that,

$$\text{Torque } (\tau) = r F \sin \theta$$

r = in metre and F = in Newton and $\sin \theta$ = no unit

$\therefore$ The SI unit of Torque (τ) is Newton-metre or N-m, Hence Option 3 is correct

★ **Important Points**

Some Basic Dimensions are-

Prepp

Your Personal Exams Guide

Sl. No.	Quantity	Common Symbol	SI Unit	Dimension
1	Velocity	v, u	ms^{-1}	LT^{-1}
2	Acceleration	a	ms^{-2}	LT^{-2}
3	Force	F	Newton (N)	M L T^{-2}
4	Momentum	p	Kg-ms^{-1}	M L T^{-1}
5	Gravitational constant	G	$\text{N-m}^2\text{Kg}^{-2}$	$\text{L}^3 \text{M}^{-1} \text{T}^{-2}$
6	Torque	τ	N-m	$\text{M L}^2 \text{T}^{-2}$
7	Bulk Modulus	B	Nm^2	$\text{M L}^{-1} \text{T}^{-2}$
8	Energy	E, U, K	joule (J)	$\text{M L}^2 \text{T}^{-2}$
9	Heat	Q	joule (J)	$\text{M L}^2 \text{T}^{-2}$
10	Pressure	P	Nm^{-2} (Pa)	$\text{M L}^{-1} \text{T}^{-2}$
11	Electric Field	E	$\text{Vm}^{-1}, \text{NC}^{-1}$	$\text{M L I}^{-1} \text{T}^{-3}$
12	Potential (voltage)	V	V, JC^{-1}	$\text{M L}^2 \text{I}^{-1} \text{T}^{-3}$

13	Magnetic Field	B	Tesla (T), Wb m ⁻¹	M I ⁻¹ T ⁻²
14	Magnetic Flux	Φ _B	Wb	M L ² I ⁻¹ T ⁻²
15	Resistance	R	Ohm (Ω)	M L ² I ⁻² T ⁻³
16	Electromotive Force	E	Volt (V)	M L ² I ⁻¹ T ⁻³

150. Answer: c

Explanation:

Understanding Metal Stock Removal Processes

The question asks to identify the process specifically known for removing stock from a metal surface. Removing stock means taking away material from the workpiece.

Analyzing the Options for Stock Removal

Let's examine each option to see if it involves removing material from a metal surface:

- **Soldering:** This is a metal joining process. It involves melting a filler metal (solder) to join two pieces of metal. Solder is added to the joint; material is not removed from the surfaces being joined, although surface cleaning might be done beforehand.
- **Brazing:** Similar to soldering, brazing is also a metal joining process. It uses a filler metal that melts at a higher temperature than solder but below the melting point of the base metals. Like soldering, material is added as a filler; stock is not removed from the base metal surfaces during the joining process itself.
- **Honing:** Honing is an abrasive machining process. It uses abrasive stones (hones) to enlarge and refine a surface, typically a cylindrical bore. While its main goal is to improve surface finish and geometry (like straightness and roundness), it inherently

involves removing a small amount of material (stock) through abrasion to achieve the desired precision and surface characteristics.

- **Coating:** Coating is the process of applying a layer of material onto a surface. Examples include painting, plating, or powder coating. This process adds material to the surface; it does not remove stock.

Identifying the Correct Process

Based on the analysis, honing is the only process among the options that involves the removal of material, or stock, from a metal surface. Although the amount of stock removed during honing is usually small compared to processes like turning or milling, it is fundamentally a stock removal process aimed at achieving high precision and excellent surface finish.

Conclusion on Stock Removal

Therefore, the process of removing stock from a metal surface among the given options is honing.

Revision Table: Comparing Metal Processes

Process	Primary Function	Involves Stock Removal?	Notes
Soldering	Joining metals	No (adds filler material)	Low melting point filler
Brazing	Joining metals	No (adds filler material)	Higher melting point filler than solder
Honing	Surface finishing & geometry improvement (typically bores)	Yes (removes small amount abrasively)	Precision abrasive machining
Coating	Applying a surface layer	No (adds material)	Protection, aesthetics, function

Additional Information on Honing and Machining

Honing is part of the family of abrasive machining processes. These processes use abrasive particles to remove material from a workpiece. Other abrasive processes include grinding, lapping, and polishing. Each process serves a specific purpose, often related to achieving tight tolerances or a specific surface finish. Honing is particularly effective for finishing internal cylindrical surfaces like engine cylinders.

Understanding different manufacturing processes, including how they modify materials (adding or removing stock), is crucial in engineering and manufacturing fields. Stock removal processes are fundamental to shaping raw materials into finished parts with the required dimensions and surface properties.

151. Answer: d

Explanation:

Understanding Footstep Bearings and Load Types

A footstep bearing is a specific type of bearing commonly used in mechanical systems. It is typically situated at the bottom of a vertical shaft. Its primary function is to support the weight or axial load acting along the shaft.

What is a Footstep Bearing?

A footstep bearing is essentially a pivot bearing designed to support vertical shafts. Think of it as the base upon which a spinning vertical rod rests. It takes the downward force exerted by the weight of the shaft and any attached components.

Types of Loads on Bearings

Bearings are designed to handle different types of mechanical loads:

- **Static Load:** A load that does not change over time. While a footstep bearing might experience a constant static load from the shaft's weight, this term describes the nature of the load, not the direction it acts in relative to the bearing.
- **Torsional Load:** A load that causes twisting or rotation. This is usually resisted by shafts and couplings, not typically the primary load supported by a footstep bearing, although some torque might be transmitted through it.

- **Tangential Load:** A force acting tangentially (at a tangent) to a rotating surface. This is more relevant to bearings supporting radial forces or transmitting torque in a tangential direction.
- **Thrust Load:** A load acting axially, along the axis of rotation or movement. This force pushes or pulls along the shaft's length.

Footstep Bearings and Thrust Loads

Given its placement at the base of a vertical shaft, the main force a footstep bearing must support is the weight of the shaft, which acts vertically downwards, along the axis of the shaft. This type of load, acting parallel to the axis of rotation, is known as a thrust load.

Therefore, a footstep bearing is specifically designed to carry significant axial or thrust loads. It prevents the vertical shaft from moving downwards.

Analyzing the Options

Let's look at the provided options in the context of a footstep bearing:

- Option 1: Static load bearing – Describes the time-dependence of the load, not its direction relative to the bearing.
- Option 2: Torsional load bearing – Primarily resists twisting forces, which is not the main function of a footstep bearing.
- Option 3: Tangential load bearing – Resists forces acting tangential to rotation, related more to radial or torque transmission.
- Option 4: Thrust load bearing – Specifically designed to support loads acting along the axis of the shaft, which is exactly what a footstep bearing does.

Based on the function of a footstep bearing, it primarily supports thrust loads.

Conclusion

A footstep bearing is a type of bearing whose main purpose is to support axial loads, also known as thrust loads. It is designed to withstand forces acting along the axis of the shaft it supports.

Bearing Type	Primary Load Supported	Direction of Load
Radial Bearing	Radial Load	Perpendicular to shaft axis
Thrust Bearing	Thrust Load (Axial Load)	Parallel to shaft axis
Footstep Bearing	Thrust Load (Axial Load)	Along the vertical shaft axis

Revision Table: Key Concepts

Term	Definition	Relevance to Footstep Bearing
Footstep Bearing	Bearing supporting the end of a vertical shaft	The specific bearing in question
Thrust Load	Load acting along the axis of a shaft	The primary load supported by a footstep bearing
Axial Load	Same as Thrust Load	Alternative term for the load type
Vertical Shaft	Shaft oriented vertically	Footstep bearings are typically used with these

Additional Information: Types of Thrust Bearings

Footstep bearings are a specific design within the broader category of thrust bearings. Other types of thrust bearings include:

- Ball thrust bearings
- Roller thrust bearings (cylindrical, spherical, tapered)
- Hydrostatic thrust bearings
- Magnetic thrust bearings

Each type is suited for different load capacities, speeds, and operating conditions. Footstep bearings are typically plain thrust bearings or use simple rolling elements at the base.

152. Answer: b

Explanation:

Understanding the Die Nut Tool

The question asks for the definition of a die nut. Let's break down what this specific tool is used for in threading applications.

A die nut is a specialized tool designed primarily for thread repair and cleaning, rather than cutting new threads from scratch like a standard die. Its unique characteristic is its form.

- **Standard Die:** Typically round or hexagonal, used with a die stock to cut new external threads on a cylindrical rod or bolt.
- **Die Nut:** Shaped like a standard nut, often hexagonal, allowing it to be driven with a wrench or socket.

The key distinction lies in the shape and primary function. While a standard die is used for cutting new threads, a die nut is used to chase, clean, or repair existing threads that might be damaged, rusty, or clogged with debris. Because of its shape, it can be easily screwed onto an existing bolt or stud, much like a regular nut would be tightened.

Let's look at the options provided:

- Option 1: A die which is used for shaping. While dies are shaping tools in a broad sense (shaping threads), this is a general description and doesn't capture the specific form or primary use of a die nut.
- Option 2: A die in the shape of a nut. This precisely describes the die nut. It performs the function of a die (acting on threads) but is manufactured in the familiar shape of a hexagonal nut.
- Option 3: A forged nut. A forged nut is a standard fastener, typically made by forging. A die nut is a cutting/repairing tool, not a fastener, although it shares the nut's shape.
- Option 4: A nut for tightening a die. Tools for tightening dies are typically die stocks or wrenches, not a type of nut itself.

Therefore, the most accurate description of a die nut is that it is a die tool specifically made in the shape of a conventional nut.

Comparing Die Nut vs. Standard Die

Here's a quick comparison to highlight the differences:

Feature	Die Nut	Standard Die
Shape	Hexagonal (like a nut)	Typically round or hexagonal (often requiring a die stock)
Primary Use	Cleaning, repairing, chasing existing threads	Cutting new external threads
Driving Method	Wrench, socket	Die stock

Revision Table: Key Terms for Threading Tools

Term	Description
Die Nut	A die shaped like a nut, used for repairing/cleaning external threads.
Standard Die	A tool used to cut new external threads on rods or bolts.
Tap	A tool used to cut new internal threads (like inside a nut or hole).
Die Stock	A handle used to hold and turn a standard die for cutting threads.
Threading	The process of cutting screw threads onto fasteners or holes.

Additional Information: Applications of Die Nuts

Die nuts are indispensable tools in various applications, particularly in automotive repair, plumbing, and general maintenance. They are excellent for:

- Restoring threads on studs or bolts that have been damaged by cross-threading, rust, or impact.
- Cleaning paint, rust, or dirt out of existing threads to allow nuts to screw on smoothly.
- Chamfering the end of a bolt slightly to help start a nut.

Because they are driven by standard wrenches, they can be easier to use in confined spaces compared to a die stock and standard die.

153. Answer: c

Explanation:

Understanding Gauges for Checking Inside Taper Holes

Gauges are essential tools used in manufacturing and inspection to check if a manufactured part conforms to the required dimensions and tolerances. They do not measure a specific value like a micrometer or vernier caliper, but rather indicate whether a feature is within or outside the acceptable range. Different types of gauges are designed for checking specific features, such as diameters, lengths, angles, or tapers.

The question asks specifically about checking an **inside taper hole**. An inside taper hole is a bore or hole that gradually changes in diameter along its length. Checking such a feature requires a gauge that is designed to verify both the taper angle and the size at specific points along the taper.

Analyzing the Options for Checking Inside Taper Holes

- **Thread gauge:** A thread gauge is used to check the pitch and form of screw threads, both external and internal. It is not suitable for checking the taper or size of a plain taper hole.
- **Rectangular gauge:** The term "rectangular gauge" is not a standard classification for a specific type of dimensional gauge used in precision measurement for features like taper holes. It does not apply here.
- **Taper plug gauge:** A taper plug gauge is a precision inspection tool specifically designed to check internal taper holes. It is a plug-shaped gauge with the exact specified taper angle. It typically has 'Go' and 'No-Go' size limits marked by lines or steps along its length. When the gauge is inserted into the taper hole, the position where the 'Go' line/step aligns with the edge of the hole indicates the minimum acceptable size, and the 'No-Go' line/step indicates the maximum acceptable size. It checks both the taper angle and the size of the hole.
- **Square gauge:** Similar to a rectangular gauge, "square gauge" is not a standard type of precision gauge used for checking features like taper holes. It might refer to a

squareness check but not a taper.

Why Taper Plug Gauges Check Inside Taper Holes

The **taper plug gauge** is the correct tool for checking an inside taper hole because its form is designed to match the desired taper. By inserting the gauge into the hole, an inspector can verify:

- Whether the taper angle of the hole matches the taper angle of the gauge. If the contact is not uniform along the length, the angle is incorrect.
- Whether the size of the hole at a specified depth (indicated by lines or steps on the gauge) is within the acceptable tolerance limits ('Go' and 'No-Go').

This makes the **taper plug gauge** the precise and appropriate tool for verifying the dimensions and form of an inside taper hole.

Revision Table: Gauges and Their Applications

Gauge Type	Primary Application
Thread Gauge	Checking screw threads (pitch, form)
Plain Plug Gauge	Checking straight cylindrical holes (Go/No-Go diameters)
Plain Ring Gauge	Checking straight cylindrical shafts (Go/No-Go diameters)
Snap Gauge	Checking external dimensions (e.g., shaft diameter)
Taper Plug Gauge	Checking inside taper holes (taper angle and size)
Taper Ring Gauge	Checking external taper shafts (taper angle and size)

Additional Information on Taper Plug Gauges

Taper plug gauges are crucial for quality control in industries like manufacturing, where precision taper holes are common, such as in machine tool spindles, collets, and fittings. Common standards for tapers include Morse taper, Jacobs taper, and Brown & Sharpe taper. A specific **taper plug gauge** would be made to match one of these standard tapers or a custom taper specification.

The accuracy of the inspection depends heavily on the accuracy of the **taper plug gauge** itself, which must be manufactured to very tight tolerances. Proper usage involves inserting the gauge into the clean, dry hole and checking the position of the 'Go' and 'No-Go' marks relative to the face of the hole.

154. Answer: c

Explanation:

Understanding the Process of Joining Parts

The question asks for the term that describes a series of processes where various individual parts are brought together and joined to create a complete, finished product. Let's look at the options provided to determine which one accurately fits this description.

Analyzing the Options

We are given four options:

1. storing
2. queuing
3. assembling
4. disassembling

Evaluation of Each Option:

- **Storing:** This refers to the act of keeping goods or materials in a warehouse or other storage location. It is a process related to inventory management, not the act of creating a product by joining parts. Therefore, 'storing' is not the correct answer.
- **Queuing:** This term describes waiting in a line or sequence. It is relevant in processes where items or tasks wait for their turn, but it does not involve the physical act of joining parts to form a product. Therefore, 'queuing' is not the correct answer.
- **Assembling:** This process involves bringing together individual components or parts and fitting them together in a specific order to create a complete functional product. This directly matches the description given in the question: a series of processes to joint various parts to get a complete product. This is the correct term for this manufacturing activity.

- **Disassembling:** This is the opposite of assembling. It involves taking apart a finished product into its individual components. While related to products, it is not the process of creating a product by joining parts. Therefore, 'disassembling' is not the correct answer.

Why Assembling is the Correct Process

Assembling is a fundamental process in manufacturing across many industries, including automotive, electronics, furniture, and machinery. It can range from simple manual tasks to complex automated production lines. The core idea remains the same: combining disparate components to form a unified whole that serves a specific purpose as a complete product.

Consider manufacturing a car. Thousands of individual parts like the engine, chassis, wheels, seats, and electronic components are brought together and joined in a specific sequence to form the final vehicle. This entire process is known as automobile assembly.

Based on the analysis, the term that correctly describes the series of processes to joint various parts to get a complete product is assembling.

Comparison of Processes

Process	Description	Relevant to Joining Parts for Product?			
Storing	Keeping items in a location	No			
Queuing	Waiting in line or sequence	No	Assembling	Joining parts to form a whole product	Yes
Disassembling	Taking a product apart	No			

Conclusion on Joining Parts Process

The process that involves joining various parts together sequentially to form a complete product is universally known as assembling. The other options describe different activities

unrelated to the creation of a product from its components.

Revision Table: Manufacturing Processes

Term	Meaning	Example
Assembly	Combining components to create a product.	Building a computer from CPU, motherboard, etc.
Manufacturing	The overall process of making goods.	Car production, electronics fabrication.
Fabrication	Making parts, often by cutting or shaping material.	Cutting metal sheets, molding plastic.
Disassembly	Taking apart a product.	Dismantling old electronics for recycling.

Additional Information: Assembly Line Concepts

A common way to perform assembly is using an assembly line. In an assembly line, the product moves along a line, and different workers or machines add parts or perform assembly steps at each station. This method is highly efficient for mass production because it allows for specialization of tasks and smooth workflow. Automation plays a significant role in modern assembly lines, with robots performing precise and repetitive joining tasks.

Key aspects of assembly include:

- Component readiness: Ensuring all necessary parts are available.
- Sequence of operations: Following a specific order for joining parts.
- Tools and equipment: Using appropriate tools for fastening, welding, or other joining methods.
- Quality control: Inspecting the assembled product to ensure it meets standards.

155. Answer: c

Explanation:

Understanding Forklift Arms and Components

A forklift is a powerful industrial truck used to lift and move materials over short distances. It's equipped with a power-operated forked platform that can be raised and lowered directly in front of the load.

Let's look at the main part responsible for lifting the loads.

The question asks about the two arms of the forklift. These are the long, prong-like extensions at the front of the machine that slide under loads, typically placed on pallets, to lift them.

Identifying the Correct Term for Forklift Arms

Based on their appearance and function, these two arms have a specific and common name:

- They are shaped somewhat like the tines of a large agricultural fork, used for lifting hay or similar materials.
- Their primary function is to 'fork' or lift the load from below.

Considering this, the most appropriate term for the two lifting arms of a forklift is 'fork'.

Analyzing the Options

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

- **Tail:** The 'tail' is usually the rear part of the forklift, often containing the counterweight. It is not the lifting arms.
- **Gibb:** This is not a standard term used to describe a part of a forklift.
- **Head:** The 'head' could refer to the overhead guard or parts of the mast assembly, but not specifically the lifting arms.
- **Fork:** This term accurately describes the two lifting arms.

Summary Table

Part Name	Description
Fork (Forks)	The two horizontal arms at the front used for lifting loads, typically inserted under pallets.
Tail	The rear section of the forklift, often housing the counterweight.
Gibb	Not a standard forklift part name.
Head	Could refer to various upper parts, not specifically the lifting arms.

Therefore, the correct name for the two arms of a forklift is 'fork'.

Revision Table: Key Forklift Terminology

Term	Description Related to Forklift
Fork / Forks	The lifting arms
Mast	The vertical assembly that raises and lowers the forks
Carriage	The part that moves up and down the mast, to which the forks are attached
Counterweight	A weight in the back to balance the load being lifted

Additional Information: Forklift Parts Explained

Understanding the different parts of a forklift is important for safe and efficient operation. Here are a few more key components:

- **Mast:** This is the upright structure at the front that allows the forks to be raised and lowered. It can have multiple stages to allow for different lifting heights.
- **Carriage:** The forks are mounted onto the carriage, which is the component that travels up and down the mast.
- **Counterweight:** Located at the rear of the forklift, the counterweight balances the weight of the load being carried on the forks, preventing the forklift from tipping forward.

- **Overhead Guard:** This is a sturdy frame above the operator's compartment designed to protect the operator from falling objects.
- **Tires:** Forklifts use different types of tires depending on the work environment (e.g., cushion tires for indoor use on smooth surfaces, pneumatic tires for outdoor or rough surfaces).

The primary function of the forklift is achieved through the interaction of the forks, carriage, and mast, enabling the machine to lift and transport goods effectively.

156. Answer: b

Explanation:

Understanding Lathe Components for Depth of Cut

A lathe is a machine tool used for shaping metal, wood, or other materials by rotating the workpiece while a cutting tool is fed into it. Different parts of the lathe carriage are responsible for moving the cutting tool in various directions to perform operations like turning, facing, threading, etc. The question asks specifically which part gives depth to the cut during a machining operation.

Analysis of Lathe Parts and Their Functions

Let's examine the function of each component listed in the options to understand how they contribute to the machining process:

1. **Lead screw:** The lead screw is a long, threaded shaft that runs parallel to the lathe bed. It is primarily used for cutting screw threads or for providing power feed to the carriage during operations like turning or facing. It controls the longitudinal movement of the carriage when engaged, but it does not directly control the depth of the cut.
2. **Cross slide:** The cross slide is mounted on top of the saddle and moves the cutting tool perpendicular to the lathe's axis of rotation. This movement is crucial for setting the depth of the cut in turning operations, as the tool is fed into the workpiece radially. It is also used for facing operations.
3. **Tailstock:** The tailstock is located at the opposite end of the lathe bed from the headstock. Its main function is to support long workpieces with a center or to hold

tools like drills, reamers, or taps for operations performed on the end of the workpiece. It does not control the movement of the cutting tool relative to the workpiece's diameter.

4. **Compound slide:** The compound slide, or compound rest, is mounted on top of the cross slide. It has its own base that can be swiveled to any angle in the horizontal plane. This component provides angular tool movement and also allows for fine adjustments of the tool position, particularly useful for cutting tapers or angles. While it contributes to tool positioning, the primary control for setting the overall depth of cut in standard turning is the cross slide.

Identifying the Component for Depth of Cut

Based on the functions of these lathe components, the part responsible for moving the cutting tool perpendicularly towards or away from the workpiece axis, thereby determining how deep the tool penetrates the material, is the cross slide.

- The cross slide carriage movement directly impacts the radius reduction on the workpiece, which translates to the depth of the cut.
- Adjusting the cross slide inward increases the depth of the cut, while moving it outward decreases it or positions the tool away from the work.

Conclusion: Lathe Part for Depth Control

The portion of a lathe that is primarily used to give depth to the cut during operations like turning or facing is the cross slide.

Revision Table: Lathe Components and Uses

Component	Primary Function Related to Tool/Workpiece Position
Headstock	Holds and rotates the workpiece
Tailstock	Supports workpiece; holds tools (drills, reamers)
Carriage (Saddle & Apron)	Moves along the bed (longitudinal feed)
Cross Slide	Moves tool perpendicular to axis (controls depth/facing)
Compound Slide	Swivels for angular cuts; fine tool adjustment
Lead Screw	Used for threading or automatic power feed

Additional Information: Lathe Operations and Parameters

In lathe operations, depth of cut is one of the key machining parameters, alongside cutting speed and feed rate.

- **Depth of Cut:** This is the distance the cutting tool penetrates the workpiece surface in a single pass. It is controlled by the cross slide. It directly affects the amount of material removed and the cutting forces.
- **Feed Rate:** This is the distance the cutting tool advances along the workpiece per revolution (in turning) or per minute. It is controlled by the carriage movement (longitudinal feed) or the cross slide movement (facing feed), often driven by the lead screw or a feed rod.
- **Cutting Speed:** This is the speed at which the surface of the workpiece passes the cutting edge of the tool. It is determined by the spindle speed (RPM) and the workpiece diameter.

Understanding how each lathe component functions and relates to machining parameters like depth of cut is essential for effective and accurate machining.

157. Answer: b

Explanation:

Understanding Measuring Instrument Principles

The question asks about a specific measuring instrument that operates based on the principle of a rack and pinion mechanism. Let's examine the working principles of the given options to identify the correct one.

Working Principles of Measuring Instruments

Different measuring instruments employ various mechanical principles to achieve accurate measurements. Understanding these principles is key to identifying the instrument in question.

- **Vernier Bevel Protractor:** This instrument is used for measuring angles. It typically uses a main scale and a vernier scale on a circular disc to measure angles with high precision. Its operation relies on the vernier principle for reading fractional divisions of the main scale. It does not use a rack and pinion for its primary measurement function.
- **Dial Test Indicator:** Also known as a lever indicator or DTI, or more broadly, a Dial Indicator (which includes plunger type), this instrument is used to measure small distances and displacements. The pointer on the dial magnifies the linear movement of a contact point or plunger. In many dial indicators (especially plunger types), the linear movement of the plunger is converted into rotational movement of the pointer via a rack and pinion gear system. A rack attached to the plunger engages with a pinion gear, which is connected to the pointer shaft, thereby amplifying the small linear displacement into a larger rotational movement on the dial.
- **Vernier Height Gauge:** This instrument is used to measure vertical heights from a reference surface. It consists of a heavy base, a vertical column with a main scale, and a sliding jaw with a vernier scale. Measurements are taken by adjusting the jaw and reading the height using the vernier principle against the main scale. Its movement is typically controlled by a screw or a friction drive, not a rack and pinion for the measurement principle itself.
- **Micrometer:** A micrometer is used for precise measurement of dimensions (like outer diameter, inner diameter, depth). Its working principle is based on the accurate screw thread, specifically the micrometer screw principle, which converts a small rotation of the screw into a precise linear movement of the spindle. It does not use a rack and pinion mechanism for measurement.

Analyzing the Rack & Pinion Principle

The rack and pinion is a type of linear actuator that comprises a circular gear (the pinion) engaging a linear gear (the rack). Rotational motion of the pinion causes the rack to move linearly, and vice versa. This mechanism is often used for amplification or conversion of motion types.

In the context of measuring instruments, a rack and pinion is particularly useful for converting a small linear displacement into a larger, easily readable rotational displacement on a dial. This is exactly how many dial indicators work to magnify the movement of the contact point.

Instrument	Primary Working Principle	Uses Rack & Pinion?
Vernier Bevel Protractor	Vernier Scale for Angle Measurement	No
Dial Test Indicator / Dial Indicator	Rack & Pinion / Lever Mechanism for Amplification	Yes (common in plunger type Dial Indicators)
Vernier Height Gauge	Vernier Scale for Height Measurement	No
Micrometer	Micrometer Screw Principle	No

Based on the analysis of the working principles, the instrument that commonly employs the rack and pinion principle for its operation is the Dial Test Indicator (or more broadly, a Dial Indicator). The rack and pinion amplifies the small linear movement of the contact point or plunger, allowing it to be displayed as a larger, more easily read rotation on the dial.

Revision Table: Measuring Instrument Principles

Instrument	Measurement Type	Mechanism for Measurement/Amplification
Vernier Bevel Protractor	Angle	Vernier scale
Dial Test Indicator / Dial Indicator	Small linear displacement	Rack and pinion or Lever (depending on type)
Vernier Height Gauge	Height	Vernier scale
Micrometer	Linear dimension	Micrometer screw thread

Additional Information: Rack & Pinion in Metrology

While the rack and pinion principle is fundamental to many dial indicators for motion amplification, it's important to note variations exist, such as lever-type dial test indicators that use a different lever-based mechanism for amplification. However, the rack and pinion is a classic and widely used method in plunger-type dial indicators to convert and amplify linear movement into rotational movement for the dial display.

The primary function of the rack and pinion in a dial indicator is to magnify small physical movements. If the plunger moves by a small amount, say Δx , the rack attached to it also moves by Δx . This linear movement drives the pinion gear, causing it to rotate. The rotation of the pinion is transferred through a gear train to the pointer on the dial. The gear train is designed to provide a significant amplification ratio, so that a very small Δx results in a noticeable rotation of the pointer, allowing for precise measurement of small displacements.

158. Answer: d

Explanation:

When creating internal threads in a material, such as in a nut or a tapped hole, a process called tapping is used. Tapping involves using a tool called a tap, which is essentially a screw-shaped cutting tool with flutes for chip removal.

Understanding the Tapping Hole and Tap Size Relationship

Before a tap can be used to cut threads into a material, a hole must first be drilled through the material. This initial hole is known as the tapping hole or tap drill hole. The size of this hole is crucial for successful tapping and for creating threads with the correct strength and engagement.

The purpose of the tapping hole is to provide the space for the tap to cut the threads. However, if the hole were the same size as the major diameter of the tap (which is the largest diameter of the screw thread), there would be no material left for the tap to cut and form the internal threads.

Why the Tapping Hole Must Be Smaller

The tapping hole must be **smaller than** the tap size. Here's why:

- The tap cuts material from the wall of the pre-drilled hole.
- This removed material forms the grooves (or roots) of the internal thread.
- The material that remains between these grooves forms the ridges (or crests) of the internal thread.
- If the tapping hole were equal to or larger than the tap's major diameter, there would be insufficient material left in the hole wall to form the thread crests, resulting in little to no thread engagement.

The specific size of the tapping hole is called the **tap drill size**. This size is calculated or looked up in charts based on the nominal size of the tap and the thread pitch (the distance between adjacent threads). The tap drill size is slightly larger than the minor diameter of the thread but smaller than the major diameter, allowing the tap to cut the full thread form.

Summary of Tapping Hole Size

In summary, for a tap to successfully cut internal threads into a material, the preparatory hole, known as the tapping hole or tap drill hole, must be smaller than the nominal size of the tap itself. This ensures there is enough material available for the tap to cut and form the complete thread profile.

Revision Table: Tapping Concepts

Term	Description	Size Relation to Tap
Tap	Tool used to cut internal threads.	Reference size.
Tapping Hole (Tap Drill Hole)	Hole drilled before tapping.	Smaller than the tap size.
Threads	Helical ridges (external) or grooves (internal) allowing connection.	Formed by the tap in the tapping hole.

Additional Information: Tap Drill Size Calculation

While charts are commonly used, the tap drill size can be approximated using formulas. A common formula for UNC/UNF threads is:

$$\text{Tap Drill Size} \approx \text{Major Diameter} - \frac{1}{\text{Threads Per Inch}}$$

For metric threads, a common approximation is:

$$\text{Tap Drill Size} \approx \text{Nominal Diameter} - \text{Pitch}$$

These formulas highlight that the tap drill size is always less than the nominal or major diameter of the tap, confirming that the tapping hole is indeed smaller than the tap size.

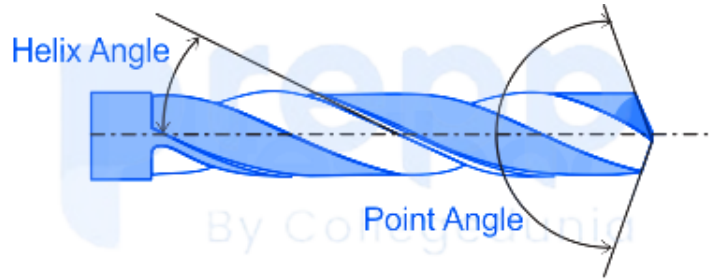
Your Personal Exams Guide

159. Answer: b

Explanation:

Concept:

- Drilling is a cutting process in which a hole is originated by means of a multipoint, fluted, end cutting tool. As the drill is rotated and advanced into the workpiece, material is removed in the form of chips that move along the fluted shank of the twist drill.



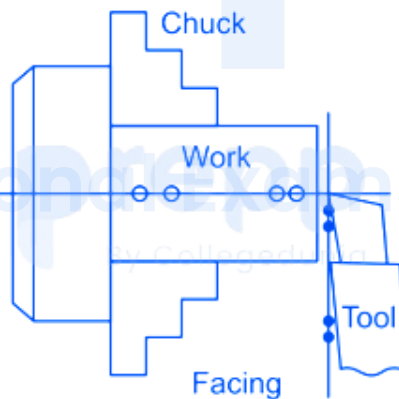
- Tool geometry:
- The figure shown the three most common angles,
 - Point angle
 - Helix angle
 - Lip relief angle

The standard point angle of 118° is most common and the standard helix angle is 24° .

★ Additional Information

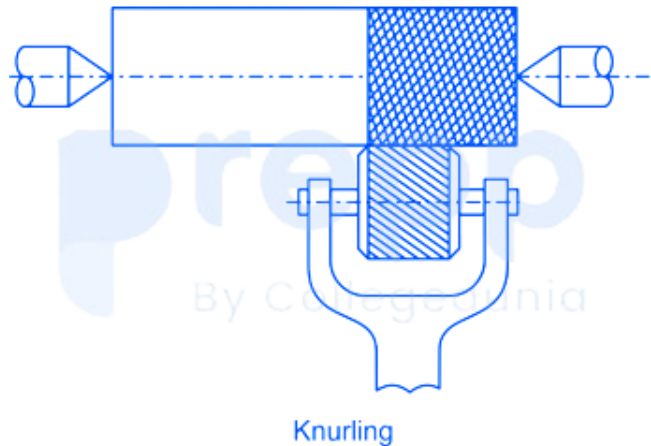
Facing:

- Facing is a machining operation by which the end surface of the workpiece is made flat by removing metal from it.



Knurling

- Knurling is the operation of producing a straight-lined, diamond-shaped pattern or cross lined pattern on a cylindrical external surface by pressing a tool called knurling tool. Knurling is not a cutting operation but it is a forming operation.
- A lathe is used for many operations such as turning, threading, facing, grooving, ; Knurling, Chamfering, center drilling



160. Answer: a

Explanation:

Understanding Alloy Steels and Alloying Elements

Steel is primarily an alloy of iron and carbon. However, to achieve specific properties like increased strength, hardness, corrosion resistance, or heat resistance, other elements are added. Steel with intentionally added alloying elements (besides carbon) is known as alloy steel.

These alloying elements alter the crystal structure, modify phase transformations, and form carbides or other compounds within the steel, leading to desirable changes in its mechanical, physical, and chemical properties.

Common Alloying Elements in Steel

Many elements can be added to steel to improve its performance. Some of the most common alloying elements include:

- **Chromium (Cr):** Enhances hardness, strength, and corrosion resistance (essential for stainless steel).
- **Nickel (Ni):** Increases toughness, strength, and corrosion resistance.
- **Molybdenum (Mo):** Improves strength, hardness, and creep resistance at high temperatures.

- **Vanadium (V):** Refines grain size and increases strength and toughness by forming carbides.
- **Tungsten (W):** Improves hardness, wear resistance, and strength at high temperatures (used in high-speed steels).
- **Manganese (Mn):** Acts as a deoxidizer and desulfurizer; increases hardness and strength, especially impact strength.
- **Silicon (Si):** Acts as a deoxidizer; increases strength and elasticity (used in spring steels); also used in electrical steels.
- **Copper (Cu):** Improves corrosion resistance in some environments.

Analyzing the Given Options

Let's consider each material listed in the options:

- **Sodium (Na):** Sodium is a highly reactive alkali metal. It is soft, silvery-white, and reacts vigorously with water and oxygen. Due to its extreme reactivity and low melting point, it is not suitable as an alloying element to enhance the structural or chemical properties of steel. Adding sodium to molten steel would be highly dangerous and would not result in a stable, usable alloy.
- **Magnesium (Mg):** While magnesium is used in some metallic alloys (like aluminum alloys), it is not commonly used as a primary alloying element in steel to modify its properties significantly like chromium or nickel do. It can be used in metallurgical processes related to steelmaking, such as desulfurization or in the production of nodular cast iron, but it's not a standard element added to create alloy steel as discussed in the context of modifying the steel matrix.
- **Silicon (Si):** Silicon is a common alloying element in steel. It is added as a deoxidizer during steelmaking and is also used in specific types of alloy steels, such as silicon steels (electrical steels) where it increases electrical resistivity, or in spring steels where it enhances elastic limits.
- **Chromium (Cr):** Chromium is one of the most important alloying elements in steel, particularly for creating stainless steels (which contain at least 10.5% Chromium). It significantly improves corrosion resistance, hardness, and strength.

Identifying the Material Not Used

Based on the properties and common uses of these materials in metallurgy, Sodium is clearly not used as an alloying element in steel due to its high reactivity and unsuitability for the required properties.

Magnesium is also not a standard alloying element in the same way Silicon or Chromium are for modifying the bulk properties of steel alloys intended for structural or high-performance applications, but Sodium is definitively not used at all.

Conclusion

The material from the given options that is not used in alloy steels is Sodium.

Revision Table: Alloying Elements vs. Non-Alloying Material

Material	Used in Alloy Steels?	Typical Role/Reason
Sodium (Na)	No	Highly reactive, not suitable as a structural alloying element.
Magnesium (Mg)	Generally No (as primary alloying element)	Not a standard element for modifying steel matrix properties; used in related processes or other alloys.
Silicon (Si)	Yes	Deoxidizer, increases strength and elasticity, used in electrical steels.
Chromium (Cr)	Yes	Increases hardness, strength, and corrosion resistance (especially in stainless steel).

Additional Information: Why Alloying is Important

Adding alloying elements allows engineers to tailor the properties of steel for specific applications, such as:

- Increasing strength and hardness for tools and structural components.
- Improving toughness and ductility for applications requiring impact resistance.
- Enhancing corrosion resistance for use in harsh environments (e.g., stainless steels).
- Improving heat resistance and creep strength for high-temperature applications.
- Modifying magnetic and electrical properties for electrical and electronic uses.

The specific combination and percentage of alloying elements determine the type of alloy steel and its resulting properties.

161. Answer: a

Explanation:

Understanding Why CO₂ Fire Extinguisher Nozzles Get Cold

When you use a carbon dioxide (CO₂) fire extinguisher, the CO₂ is stored under high pressure as a liquid inside the cylinder. When released, this liquid CO₂ rapidly expands into a gas as it exits the nozzle. This rapid expansion is a physical process that requires energy.

The energy needed for this expansion is drawn from the surroundings, specifically from the metal of the nozzle and the escaping gas itself. This causes a significant drop in temperature, often below the freezing point of water and even CO₂ (which freezes into 'dry ice' at around -78.5°C). This phenomenon is related to the Joule-Thomson effect, where a gas cools upon expansion if it's below its inversion temperature.

Safety Implications of a Cold CO₂ Extinguisher Nozzle

The question asks what the extremely cold nozzle indicates. The primary consequence of the nozzle becoming extremely cold is a safety hazard. Let's look at the options provided in the context of this extreme cold:

- It is safe to continue.
- You should not use it.
- The material inside is finished.
- The materials inside are running low.

An extremely cold nozzle presents direct risks:

- **Frostbite Risk:** Touching the cold nozzle or horn can cause severe frostbite almost instantly, leading to tissue damage.
- **Material Integrity:** While less common with modern materials, extreme cold can potentially make some materials brittle, though this is less likely to be the primary safety concern mentioned in typical fire extinguisher training.

Therefore, the extreme cold is a warning sign indicating a dangerous condition.

Analyzing the Options Based on Nozzle Temperature

Let's evaluate each option:

Option 1: You should not use it.

- This directly addresses the safety risk. An extremely cold nozzle means there's a high risk of frostbite if handled improperly or if contact is made with skin. While you might still be able to operate the extinguisher by holding a protected handle, the extreme temperature itself is a strong indicator that handling requires extreme caution, implying a state where continued casual use is unsafe without proper protective measures (like gloves, which aren't always available in an emergency). In many safety contexts, extreme conditions like this indicate you should cease operation if you cannot do so safely.

Option 2: The material inside is finished.

- The nozzle gets cold precisely **because** the CO₂ is being discharged. If the extinguisher were finished, there would be no high-pressure release, and thus no significant cooling effect. So, a cold nozzle indicates the material is being used, not finished.

Option 3: It is safe to continue.

- As discussed, the extreme cold poses a direct hazard (frostbite). Therefore, it is not necessarily safe to continue using it without extreme caution and awareness of the temperature, making this option incorrect as a general statement. The cold indicates a danger, not safety.

Option 4: The materials inside are running low.

- The cooling effect is most pronounced during the discharge process when the pressure is high enough to expel liquid CO₂ which then expands. While the discharge might continue until the material is low, the cold nozzle is a symptom of the current discharge mechanism, not necessarily an indicator of the remaining quantity. A nozzle could become very cold even if the extinguisher is still relatively full, as long as it's discharging rapidly.

Considering the immediate safety hazard posed by the extremely low temperature of the nozzle, the most appropriate indication is related to the safety of continued use or handling.

Symptom	Cause	Indication
Extremely cold nozzle	Rapid expansion of high-pressure CO ₂ gas	Immediate safety hazard (frostbite) requiring caution or cessation of use if safety cannot be ensured.

The extreme cold is a direct consequence of the phase change and expansion of the CO₂ during discharge. This cold poses a significant safety risk, primarily frostbite. Therefore, the presence of an extremely cold nozzle indicates that continuing to use the extinguisher requires extreme caution due to the temperature hazard, or potentially, that you should cease using it if you cannot avoid contact with the cold parts safely.

Revision Table: CO2 Extinguisher Operation

Component	Function	Safety Note
Cylinder	Stores liquid CO ₂ under high pressure.	Handle with care, avoid damage.
Valve & Handle	Controls the release of CO ₂ .	Pull pin, squeeze handle to operate.
Discharge Horn/Nozzle	Directs the flow of CO ₂ towards the fire.	Becomes extremely cold during discharge; DO NOT TOUCH bare skin to the horn or nozzle.
CO ₂ Agent	Displaces oxygen and cools the fire surface (minor effect). Works primarily by smothering.	Can cause asphyxiation in confined spaces.

Additional Information on CO2 Fire Extinguisher Safety

Carbon dioxide fire extinguishers are effective on Class B (flammable liquids) and Class C (electrical) fires. They work by displacing the oxygen around the fire, suffocating it. CO₂ is a clean agent, leaving no residue, which is why they are preferred for electrical equipment.

Important safety points when using a CO₂ extinguisher:

- Always ensure ventilation when using CO₂ in enclosed spaces, as it displaces oxygen and can cause asphyxiation.
- Hold the extinguisher by the insulated handle; avoid touching the metal cylinder or the horn/nozzle during or immediately after use due to the extreme cold.
- Aim the horn at the base of the fire.
- Be aware of the limited range of CO₂ extinguishers.
- Ensure the fire is completely out, as there is a risk of re-ignition if the heat source is not sufficiently cooled (which CO₂ is not very good at).

The extreme cold is a significant indicator of the operational state and a critical safety hazard associated with CO₂ extinguishers. Recognizing this indicates the need for immediate safety precautions.

162. Answer: c

Explanation:

Measuring Small Hole Diameters: Choosing the Right Gauge

Accurately measuring the diameter of very small holes, such as one measuring only 0.6 mm, requires specialized tools. Different measuring gauges are designed for different purposes and size ranges. Let's examine the options provided to determine the most suitable tool for measuring a 0.6 mm hole diameter.

Analyzing the Measuring Gauge Options

- **Screw Gauge (Micrometer Screw Gauge):** This instrument is primarily used for measuring external dimensions like the diameter of wires, thickness of sheets, etc. While very precise, it is not typically designed for measuring the internal diameter of small holes directly.
- **Ring Gauge:** A ring gauge is a type of 'go/no-go' gauge used to check the external diameter of cylindrical parts. The part is passed through or stopped by the ring gauge to verify if its diameter is within tolerance limits. It is used for external measurements, not internal hole diameters.

- **Pin Gauge:** Pin gauges are precision cylindrical rods of exact, known diameters. They are specifically used for measuring the diameter of small holes, or checking if a hole is within a specified size tolerance (using go/no-go sets). Pin gauges are available in very small diameters, making them ideal for measuring holes like 0.6 mm. A pin gauge of 0.6 mm would be used to check if the hole is at least that size, or a set of pin gauges can be used to determine the actual diameter or size range.
- **Internal Micrometer:** Internal micrometers are used to measure the internal diameter of bores or holes. They are available in different types (e.g., two-point or three-point contact). However, standard internal micrometers are typically designed for measuring larger diameters, usually starting from a few millimeters upwards. They are not suitable or practical for accurately measuring a hole as small as 0.6 mm.

Selecting the Appropriate Gauge for a 0.6 mm Hole

Based on the function and typical size range of each measuring gauge, the pin gauge is the most appropriate tool for measuring or checking the diameter of a 0.6 mm hole. Pin gauges are specifically designed for this purpose and are available in the required small sizes.

Therefore, a pin gauge is used to measure the diameter of a 0.6 mm hole.

Revision Table: Comparison of Measuring Gauges

Gauge Type	Primary Use	Suitable for 0.6 mm Hole?
Screw Gauge	External dimensions, thickness	No
Ring Gauge	Checking external diameters (go/no-go)	No
Pin Gauge	Measuring/checking small internal hole diameters	Yes
Internal Micrometer	Measuring larger internal hole diameters	No

Additional Information: Precision Measurement Tools

Precision measurement is crucial in manufacturing and engineering. Different tools offer varying levels of accuracy and are designed for specific measurement tasks.

- **Micrometers:** These are high-precision tools used for measuring dimensions with accuracy typically up to 0.01 mm or even 0.001 mm. Types include outside, inside, and depth micrometers.
- **Calipers:** Versatile tools for measuring external dimensions, internal dimensions, and depths. Vernier calipers and digital calipers are common types, offering precision generally up to 0.02 mm or 0.01 mm.
- **Gauges:** These tools are often used for checking if a dimension is within a specified tolerance rather than measuring the exact size. Examples include go/no-go gauges (like pin and ring gauges used in this way), feeler gauges, radius gauges, etc.
- **Coordinate Measuring Machine (CMM):** A highly sophisticated device used for measuring physical geometric characteristics of an object. It provides very high accuracy and can measure complex shapes.

Choosing the right tool depends on the required accuracy, the size and shape of the feature being measured, and the specific measurement task (e.g., checking tolerance vs. determining exact size).

163. Answer: a

Explanation:

Understanding Zinc: Metal Classification

Let's analyze the question asking about the classification of Zinc. Zinc is a chemical element with the symbol Zn and atomic number 30. To classify it correctly based on the options provided, we need to understand the definitions of the terms used in the options.

Defining Metal Categories

Metals are generally divided into different categories based on their properties and composition. The options refer to classifications like ferrous metal, non-ferrous metal, non-metal, and types of steel.

- **Ferrous Metals:** These are metals that contain iron as their primary component. Iron itself, steel (which is an alloy of iron and carbon), and cast iron are common examples of ferrous metals. Ferrous metals are often magnetic and are prone to rusting (oxidation).
- **Non-ferrous Metals:** These are metals or alloys that do not contain significant amounts of iron. Examples include copper, aluminum, lead, zinc, tin, gold, silver, etc. Non-ferrous metals are generally more resistant to corrosion than ferrous metals and are often lighter.
- **Non-metals:** These are elements that lack the characteristics of metals. They are typically poor conductors of heat and electricity, are not malleable or ductile, and can be found in solid, liquid, or gaseous states at room temperature. Examples include oxygen, nitrogen, carbon, sulfur, chlorine, etc.
- **Steel:** Steel is specifically an alloy, primarily composed of iron and carbon, often with other elements added to improve certain properties. It is a type of ferrous metal.

Analyzing Zinc's Properties

Zinc is a metallic element. It has properties typical of metals, such as being a good conductor of electricity and heat, being malleable (can be hammered into shape), and ductile (can be drawn into wires). It is not a non-metal.

Now, let's consider whether Zinc is a ferrous or non-ferrous metal. Zinc is not an alloy of iron, nor does it contain iron as its main component. Therefore, it does not fit the definition of a ferrous metal or a type of steel.

Based on the definition, Zinc is a metal that does not contain significant iron. This places it in the category of non-ferrous metals.

Evaluating the Options

Let's examine each option in light of our understanding:

- Option 1: **non-ferrous metal** – Zinc is a metal and its composition does not primarily include iron. This definition aligns with Zinc's classification.
- Option 2: **ferrous metal** – Ferrous metals are primarily iron-based. Zinc is not iron-based. This option is incorrect.
- Option 3: **non-metal** – Zinc exhibits metallic properties (conductivity, malleability, ductility). It is not a non-metal. This option is incorrect.

- Option 4: **type of steel** – Steel is an alloy of iron and carbon. Zinc is a pure metallic element, not an alloy of iron, and certainly not a type of steel. This option is incorrect.

Therefore, the correct classification for Zinc among the given options is a non-ferrous metal.

Classification	Key Characteristic	Zinc
Ferrous Metal	Contains significant Iron	Does not contain significant Iron
Non-ferrous Metal	Does not contain significant Iron	Does not contain significant Iron
Non-metal	Lacks metallic properties	Has metallic properties
Type of Steel	Alloy of Iron and Carbon	Pure metallic element

The table above summarises why Zinc fits the definition of a non-ferrous metal and not the other categories.

Revision Table: Understanding Materials

Your Personal Exams Guide

Term	Definition	Example
Metal	Typically solid, shiny, malleable, ductile, good conductors of heat/electricity.	Copper, Aluminium, Iron, Zinc
Ferrous Metal	Metal or alloy containing significant iron.	Iron, Steel, Cast Iron
Non-ferrous Metal	Metal or alloy not containing significant iron.	Copper, Aluminium, Zinc, Lead, Gold
Non-metal	Element lacking metallic properties; poor conductor.	Oxygen, Carbon, Sulfur
Alloy	Mixture of two or more elements, at least one of which is a metal.	Steel (Iron + Carbon), Brass (Copper + Zinc), Bronze (Copper + Tin)
Steel	Alloy of iron and carbon.	Stainless Steel, Carbon Steel

Additional Information: Uses of Zinc

Zinc is a versatile non-ferrous metal with many important applications. Some of its common uses include:

- **Galvanizing:** Coating iron or steel with a layer of zinc to prevent rusting. This is a major use of zinc.
- **Alloys:** Zinc is used to create various alloys, such as brass (copper and zinc) and bronze (copper and tin, sometimes with zinc).
- **Batteries:** Zinc is used in batteries, including alkaline batteries and zinc-carbon batteries.
- **Die Casting:** Zinc alloys are used for die casting parts due to their low melting point and strength.
- **Chemicals:** Zinc compounds are used in pigments, medicines, cosmetics, and agriculture.
- **Nutrition:** Zinc is an essential trace element for human health.

Understanding the classification of elements and materials like Zinc into categories such as ferrous and non-ferrous metals is fundamental in materials science and engineering.

164. Answer: c

Explanation:

Understanding Belt Drives and Speed Differences

Belt and pulley systems are common mechanical power transmission devices. They transfer power from one shaft to another using a continuous belt running over pulleys. Ideally, the belt moves at the same surface speed as the pulley. However, in reality, there is often a difference in speed.

What is Belt Slip?

In a belt drive system, the term that describes the actual difference between the surface speed of the belt and the surface speed of the pulley is known as **slip**.

Slip occurs when the friction between the belt and the pulley surface is not sufficient to maintain a perfect grip. This causes the belt to move slightly slower than the pulley's surface speed. It's like when your shoes 'slip' on a smooth floor; there's relative motion despite contact.

Mathematically, if the surface speed of the pulley is V_{pulley} and the speed of the belt is V_{belt} , the speed difference due to slip is $V_{pulley} - V_{belt}$. Slip is often expressed as a percentage:

$$\text{Percentage Slip} = \frac{V_{pulley} - V_{belt}}{V_{pulley}} \times 100$$

Significant slip can lead to power loss, reduced efficiency, and increased wear on the belt and pulleys, potentially generating heat.

Why 'Slip' is the Correct Answer for Speed Difference

The question specifically asks for the term describing the "actual difference between the surface speed of the belt and pulley". Based on the definition, **slip** is precisely this phenomenon – the relative motion or speed difference between the belt and the pulley surface due to friction limitations.

Other Belt Drive Terms Explained

Let's look at the other options provided to understand why they are not the correct answer for the speed difference:

- **Dressing:** This refers to a substance applied to the surface of a belt to improve its condition, often to increase friction (and thus reduce slip) or improve flexibility and lifespan. It is a treatment, not a speed difference.
- **Torque:** Torque is a rotational force that causes rotation. In a belt drive, torque is transmitted from the driving pulley to the belt and then to the driven pulley, enabling power transmission. While related to the forces involved, it is not the term for the speed difference itself.
- **Creep:** Creep is another type of speed difference in belt drives, caused by the elastic stretching of the belt as it moves from the slack side (lower tension) to the tight side (higher tension) and contracting as it moves back. Creep is usually much smaller than slip and occurs even with sufficient friction. However, the term **slip** is the standard term for the relative movement between belt and pulley surface when friction is insufficient, which is the more general phenomenon resulting in the 'actual difference' in speeds as implied by the question, especially when friction limits are met. In many contexts, 'slip' is used broadly to include frictional slip and sometimes creep, but primarily refers to frictional slip. Since 'slip' is an option and accurately describes the speed difference due to relative motion at the surface, it is the most appropriate answer here.

Therefore, out of the given options, **slip** is the term used to describe the actual difference between the surface speed of the belt and the pulley.

Revision Table: Belt Drive Concepts

Term	Definition in Belt Drives	Relation to Speed Difference
Slip	Relative motion between belt and pulley surface due to insufficient friction.	Directly represents the speed difference between belt and pulley surfaces.
Dressing	Substance applied to belt surface.	Can affect slip (often reduces it), but is not a speed difference itself.
Torque	Rotational force transmitted.	Enables power transmission, but is not the speed difference.
Creep	Elastic stretching/contracting of belt causing slight relative motion.	Also causes a speed difference, but 'slip' is the general term for speed difference due to surface relative motion, especially frictional.

Additional Information on Belt Drive Efficiency

Understanding concepts like **slip** and **creep** is crucial for analyzing the efficiency of belt drive systems. While essential for transmitting power, belt drives are subject to energy losses.

- Losses due to **slip** and **creep** reduce the output speed of the driven pulley compared to the theoretical speed, leading to a reduction in power transmission efficiency.
- Other losses include bending losses (energy absorbed as the belt bends around the pulleys), windage losses (air resistance), and bearing friction in the pulleys.
- Minimizing **slip** (e.g., by ensuring proper belt tension and using appropriate belt dressing or types) is key to maintaining efficiency and extending the life of the belt and pulleys. Excessive **slip** generates heat, which can degrade the belt material.

Proper design, installation, and maintenance are necessary to manage **slip** and optimize the performance of belt drive systems.

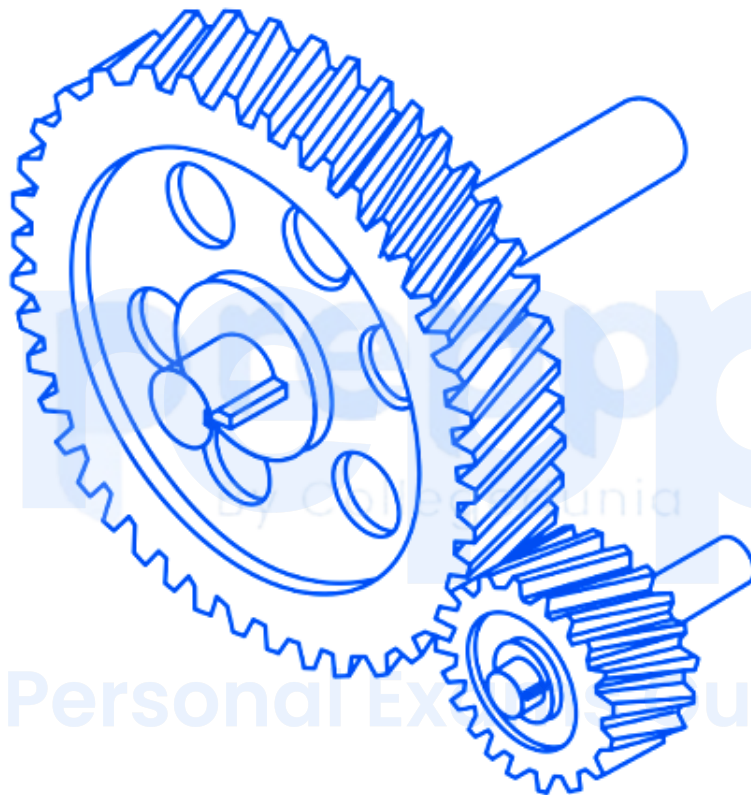
165. Answer: c

Explanation:

Explanation:

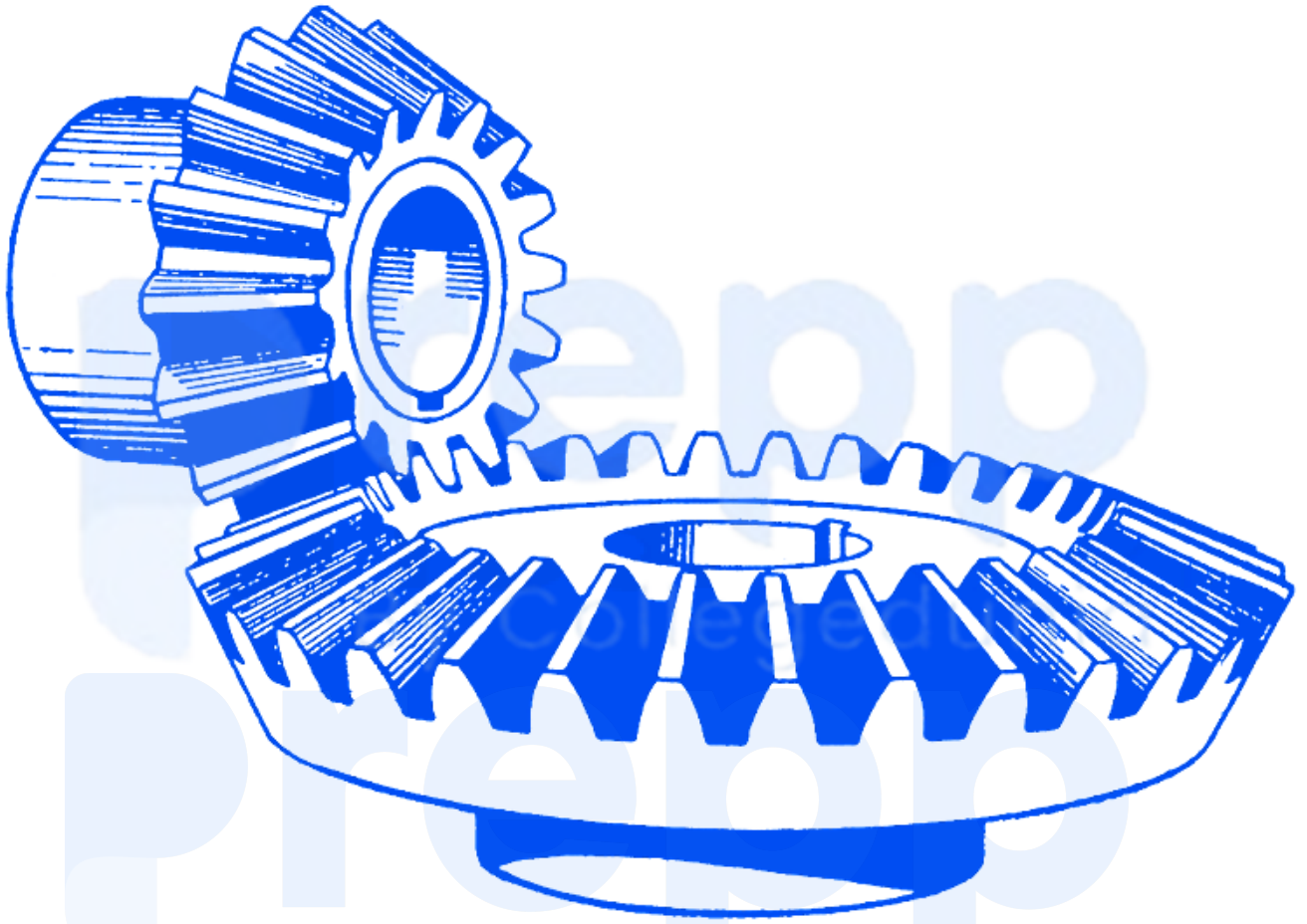
Helical Gear:

- In a helical gear, the teeth are cut at an angle to the axis of rotation.
- It is used to transmit power between two parallel shafts.
- The end thrust is exerted by the driving and driven gears in the case of helical gears and the thrust may be eliminated by using double helical gears. These gears are called herring-bone gears.



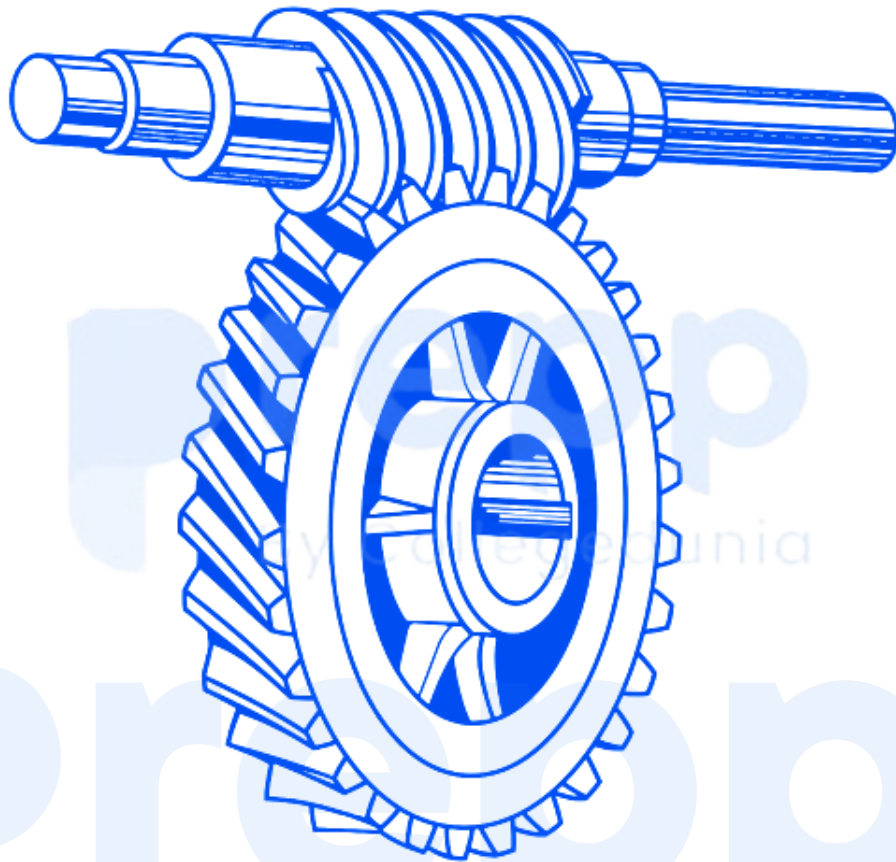
Bevel Gear:

- Bevel gears are used to transmit motion between shafts at various angles to each other.
- The teeth profile may be straight or spiral.



Worm shaft and worm gear:

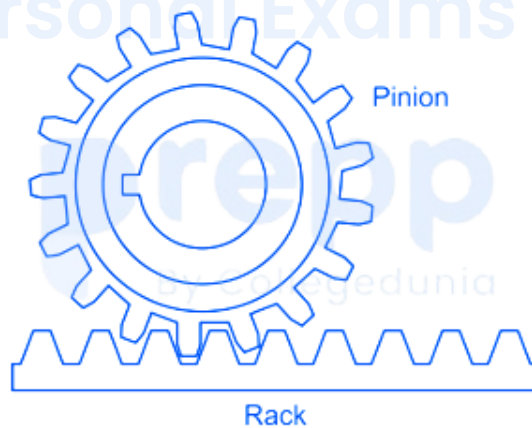
- The worm shaft has spiral teeth cut on the shaft and worm wheel is a special form of gear teeth cut to mesh with the worm shaft.
- These are widely used for speed reduction purpose.



Rack and pinion

- The rack and pinion can change rotary into linear movement and vice versa.

Your Personal Exams Guide



166. Answer: b

Explanation:

Understanding Jigs for Machining Operations

Jigs and fixtures are essential tools used in manufacturing to accurately locate and hold a workpiece for machining or assembly operations. They ensure interchangeability and accuracy of the manufactured parts by guiding the cutting tool to the correct position relative to the workpiece.

The question asks which type of jig allows for machining on multiple surfaces or planes of a workpiece in a single setup. Let's look at the characteristics of different types of jigs.

Types of Jigs and Their Applications

Jigs come in various forms, each designed for specific types of operations and workpiece geometries:

- **Template Jig:** This is the simplest type. It's a plate that rests on the workpiece, with holes acting as guides for the drill or other tools. It's mainly used for simple drilling operations on flat surfaces and doesn't offer support for multi-plane machining in one setup.
- **Plate Jig:** Similar to a template jig but may include pins to locate the workpiece more accurately. It's also generally used for operations on a single surface.
- **Open Type Jig:** These jigs are open at the top or sides, allowing for easy loading and unloading of the workpiece. They are often used for operations where the workpiece requires support from underneath or the sides, but they may not easily facilitate indexing for machining on multiple planes without repositioning the workpiece.
- **Box Type Jig:** Also known as a tumble jig, this type encloses the workpiece like a box. It has feet or locating pads on multiple sides. This design allows the jig to be rotated or "tumbled" onto different faces, presenting different planes of the workpiece to the cutting tool while the workpiece remains fixed within the jig. This capability makes it ideal for machining operations (like drilling, reaming, tapping) on multiple surfaces in a single clamping setup.

Why Box Type Jigs Excel in Multi-Plane Machining

The key feature of a box type jig that enables machining in more than one plane by a single setting is its enclosed structure with locating surfaces (feet) on multiple sides. Once the workpiece is securely clamped inside the box jig, the operator can simply rotate

the jig onto a different set of feet to expose a new surface for machining without unclamping and reclamping the workpiece. This maintains the relative position between the machined features on different planes.

Jig Type	Primary Use	Multi-Plane Machining (Single Setting)
Template Jig	Simple drilling on flat surfaces	Generally No
Plate Jig	Drilling with locating pins on flat surfaces	Generally No
Open Type Jig	Operations on one or two accessible surfaces	Limited / Difficult
Box Type Jig	Operations on multiple faces/planes	Yes, by rotating the jig

Therefore, the box type jig is specifically designed to handle operations that require machining on different planes of a workpiece efficiently in a single setup, significantly reducing setup time and potential errors from reclamping.

Revision Table: Key Jig Types

Jig Type	Description	Suitability for Multi-Plane
Template Jig	Flat plate guide, rests on workpiece	Low
Plate Jig	Plate guide with locators	Low
Open Jig	Accessible design, open sides/top	Medium (depends on design)
Box Jig	Enclosed structure with feet on multiple faces	High (designed for it)

Additional Information: Advantages of Box Jigs

- **Efficiency:** Allows multiple operations on different planes in one setup, saving time.
- **Accuracy:** Maintains the workpiece's position relative to all machining planes, improving accuracy of feature locations.
- **Versatility:** Can be designed for complex parts requiring operations on many sides.
- **Setup Reduction:** Single clamping reduces overall setup time compared to reclamping for each face.

167. Answer: b

Explanation:

Understanding Machine Part Fixing and Alignment

When working with machines, ensuring that all parts are correctly assembled and positioned is crucial for their proper function, efficiency, and longevity. The question asks for the specific term used for fixing different parts of a machine accurately according to the required specifications.

Analyzing the Options

Let's look at the given options and understand what each term means in the context of machine assembly and maintenance:

- **Unpacking:** This refers to removing the machine or its parts from its packaging. It is the first step when receiving equipment but doesn't describe the fixing or positioning of parts.
- **Alignment:** This involves positioning different parts relative to each other in a precise manner so that they are in a correct line, plane, or orientation. Proper alignment ensures smooth operation and prevents unnecessary wear and tear. For instance, aligning shafts or couplings in a motor system is critical.
- **Dismantling:** This is the process of taking a machine or assembly apart into its component parts. It is the opposite of fixing or assembling parts.
- **Checking:** This is a general term for verifying the condition, function, or position of something. While checking might be part of the fixing process (e.g., checking if parts are aligned correctly), it is not the term for the act of fixing them in position.

Defining Alignment in Machine Assembly

Alignment in machine assembly specifically means bringing machine components into a precise, intended relationship with each other. This often involves ensuring shafts are coaxial, surfaces are parallel or perpendicular as required, and components are positioned to minimize stress and vibration during operation. Correct **alignment** is essential for preventing premature failure of bearings, seals, couplings, and other components, and for achieving the machine's designed performance.

Therefore, the fixing of different parts of a machine properly as per the requirement is accurately described by the term **alignment**.

Revision Table: Machine Operations Terminology

Term	Description	Relevance to Fixing Parts
Unpacking	Removing from packaging	Not relevant
Alignment	Precisely positioning parts relative to each other	Directly relevant (The act of proper fixing)
Dismantling	Taking apart a machine	Opposite process
Checking	Verifying condition or position	A step within the process, not the process itself

Additional Information: Importance of Proper Alignment

Proper machine **alignment** is not just about putting parts together; it's about putting them together correctly to ensure peak performance and longevity. Misalignment can lead to several problems, including:

- Increased vibration
- Excessive wear on components like bearings, seals, and couplings
- Higher energy consumption
- Reduced lifespan of the machine
- Increased risk of catastrophic failure

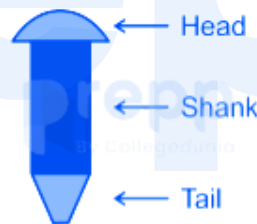
Techniques for achieving proper **alignment** vary depending on the machine and components involved. Common methods include using straight edges, dial indicators, or laser alignment tools for more precise applications like rotating machinery.

168. **Answer: c**

Explanation:

Explanation:

- Rivets are cylindrical rods having heads of various shapes.
- Rivets are used to join together two or more sheets of metal permanently.
- Threads are not present on the rivet.
- In sheet metal work riveting is done where: brazing is not suitable, the structure changes owing to welding heat, the distortion due to welding cannot be easily removed etc.



Riveting

- Riveting is one of the methods of making permanent joints of two or more pieces (Metal strips).
- It is customary to use rivets of the same metal as that of the parts that are being joined together.
- Therefore, For riveting aluminium sheets, the material of the rivet should be aluminium
- Uses
 - Rivets are used for joining metal sheets and plates in fabrication work, such as bridges, ship building, games, structural steel work.
- Types of rivets :
 - Tinman's rivet
 - Flat head rivet
 - Round head rivet

- Countersunk head rivet

169. Answer: a

Explanation:

Understanding Bench Vice Body Material

A bench vice is a mechanical tool used to secure an object to allow work to be performed on it with tools such as saws, planes, drills, grinders, files, etc. It is commonly used in workshops for woodworking, metalworking, and other tasks. The main body of a bench vice needs to be strong and rigid enough to hold the workpiece firmly without bending or breaking under the forces applied during various operations.

Materials Used for Bench Vice Bodies

Different materials have different properties that make them suitable or unsuitable for specific applications like the body of a bench vice. Let's examine the common materials mentioned in the options:

- **Cast Iron:** This is an alloy of iron that is typically strong in compression and has good rigidity. It can be easily cast into complex shapes, which is advantageous for manufacturing the intricate parts of a vice body. Grey cast iron, in particular, is commonly used due to its good machinability and damping properties (ability to absorb vibrations).
- **Forged Steel:** Steel is an alloy of iron and carbon. Forging is a manufacturing process that involves shaping metal using localized compressive forces. Forged steel is generally stronger and tougher than cast iron. However, forging complex shapes like a vice body can be more difficult and expensive than casting.
- **Aluminium Alloy:** Aluminium alloys are known for their lightness and corrosion resistance. While strong for their weight, they are generally not as strong or rigid as cast iron or steel for the same volume, and they are significantly more expensive than cast iron. They might be used for lighter-duty vices or specific components, but typically not the main body of a standard workshop bench vice.
- **Steel:** This is a broad category. General steel can be strong and tough. Like forged steel, manufacturing a vice body from general steel might involve processes that

are more complex or costly than casting iron, especially for the main structural components that require mass and rigidity.

Why Cast Iron is Preferred for Bench Vice Bodies

The body of a bench vice is subjected to significant clamping forces and potential impacts. It needs to be:

- **Strong:** To hold the workpiece securely.
- **Rigid:** To prevent deflection under load, ensuring accuracy.
- **Durable:** To withstand years of use and potential abuse.
- **Economical to Produce:** To make the tool affordable.

Cast iron meets these requirements effectively. Its compressive strength and rigidity are excellent for resisting clamping forces. The casting process allows for efficient production of the complex shapes needed for the vice body, including guide ways, jaws, and mounting points, at a relatively low cost compared to machining or forging steel into the same shape. While forged steel is stronger, the benefits in this application are often outweighed by the increased manufacturing difficulty and cost for the main body component.

Your Personal Exams Guide

Material Properties Comparison for Bench Vice Body

Material	Strength	Rigidity	Ease of Shaping (Body)	Cost	Suitability for Bench Vice Body
Cast Iron	Good (especially compression)	Good	Excellent (casting)	Low	High (Standard choice)
Forged Steel	Excellent	Excellent	Difficult (for complex body)	High	Lower (for main body)
Aluminium Alloy	Moderate	Moderate	Good (casting/machining)	High	Low (for standard workshop vice)
Steel	Good to Excellent	Good to Excellent	Moderate to Difficult	Moderate to High	Moderate (depends on type/process)

Based on the typical requirements for a sturdy and economical workshop bench vice, **cast iron** is the widely adopted material for the main body structure.

Revision Table: Bench Vice Components

Bench Vice Components and Materials

Component	Typical Material(s)
Body (Fixed & Sliding Jaw)	Cast Iron (Grey Cast Iron is common)
Jaws (Replaceable inserts)	Hardened Steel (often serrated)
Lead Screw	Steel
Handle	Steel
Nut (for Lead Screw)	Cast Iron or Bronze

Additional Information on Bench Vices

While the body is commonly made of cast iron, other parts of the bench vice use different materials optimized for their function:

- The screw and handle are typically made of steel for strength and toughness.
- The replaceable jaw inserts are usually made of hardened steel to resist wear and provide a strong grip.
- The nut that the lead screw turns in might also be made of cast iron or bronze for wear resistance.

Understanding the materials used in tools like a bench vice helps explain their characteristics and why certain materials are chosen for specific parts.

170. Answer: b

Explanation:

Understanding Pneumatic Systems and Their Power Source

A pneumatic system is a type of fluid power system that uses compressed air to transmit energy and perform work. These systems are widely used in various industrial and automotive applications due to their simplicity, cleanliness, and reliability.

Every fluid power system, whether pneumatic or hydraulic, requires a source to generate the pressurized fluid that drives the system. In pneumatic systems, this fluid is compressed air.

Identifying Key Components of a Pneumatic System

Let's look at the typical components found in a pneumatic system and understand their function:

- **Compressor:** This device takes atmospheric air and compresses it, increasing its pressure and energy. This compressed air is the working medium that powers the entire system. Therefore, the compressor is responsible for generating the pressurized fluid.

- **Air Receiver:** This is a storage tank for the compressed air produced by the compressor. It helps to store a reserve of compressed air, smooth out pressure fluctuations, and allow the compressor to cycle on and off efficiently. While it stores energy, it doesn't generate the initial pressure.
- **Valve:** Valves are control devices used to direct, regulate, or control the flow, pressure, or direction of the compressed air within the system. They act as switches or regulators, but they do not generate the power.
- **Muffler:** A muffler, or silencer, is used at the exhaust ports of components (like cylinders or valves) to reduce the noise generated when compressed air is vented to the atmosphere. It deals with the exhausted air but is not involved in generating or transmitting power.

Determining the Power Source

Based on the functions of these components, the device that introduces energy into the system by compressing air and creating pressure is the power source.

The compressor is the component that performs this crucial task. It converts mechanical or electrical energy into potential energy stored in the form of compressed air under pressure. This pressurized air is then used to actuate cylinders, motors, and other pneumatic devices, thereby doing work.

Component	Primary Function	Role as Power Source
Compressor	Compresses air, increases pressure	Generates the pressurized fluid (power)
Air Receiver	Stores compressed air	Stores power, does not generate
Valve	Controls flow/pressure/direction	Controls power, does not generate
Muffler	Reduces exhaust noise	Handles exhausted fluid, does not generate power

Therefore, the component that serves as the power source in a pneumatic system is the compressor.

Revision Table: Pneumatic System Essentials

Let's quickly review the main point about the power source in pneumatic systems:

- The power in a pneumatic system comes from compressed air.
- The device responsible for generating this compressed air at high pressure is the compressor.
- Other components like air receivers, valves, and mufflers play important roles in storage, control, and noise reduction, but they do not generate the primary power.

Additional Information on Pneumatic Power Sources

While the compressor is the primary power source, the entire process involves several steps to make the air usable for the system:

- **Air Intake:** Air is drawn from the atmosphere, often passed through filters.
- **Compression:** The compressor increases the pressure of the air. Different types of compressors exist, such as piston, rotary screw, or centrifugal compressors, each suited for different pressure and flow requirements.
- **Air Treatment:** Compressed air often contains moisture, oil, and particles. It is typically treated using filters, regulators, and lubricators (often combined into an FRL unit) before being sent to the control valves and actuators. This treatment ensures the air is clean and at the correct pressure for system components.
- **Distribution:** The treated, compressed air is distributed through pipes and hoses to various parts of the system.

Understanding the role of the compressor as the generator of compressed air is fundamental to understanding how pneumatic systems operate and deliver power.

171. Answer: d

Explanation:

Understanding Dial Gauge Measurement Range

A dial gauge, also known as a dial indicator, is a precision measuring instrument used to accurately measure small linear distances or displacements. It converts the linear

movement of a spindle into the rotational movement of a pointer on a circular dial. This makes it easy to read very small variations in size or position.

Dial gauges are commonly used in workshops and laboratories for tasks such as:

- Checking the runout of shafts or surfaces.
- Measuring variations in thickness of materials.
- Comparing dimensions against a master component.
- Setting up machine tools.

Typical Measurement Range of Dial Gauges

The measurement range of a dial gauge refers to the total distance the spindle can travel and be accurately indicated on the dial. This range is also sometimes called the "travel" or "stroke" of the indicator.

Dial gauges are designed for measuring relatively small displacements with high precision, often down to fractions of a millimeter or an inch. Consequently, their typical measurement range is limited compared to tools like measuring tapes or large calipers.

Let's consider the options provided for the typical measurement range:

- 600 mm to 700 mm: This is a very large range, much greater than the usual travel of a standard dial gauge. Tools like large calipers or height gauges might operate in this range.
- 400 to 500 mm: Similar to the above, this range is too large for a typical dial gauge, which is designed for precise measurement over short distances.
- 800 to 900 mm: Again, this represents a range far exceeding the capabilities of standard dial indicators.
- 0.25 to 300 mm: This option provides a range that includes common dial gauge travels. While a single standard dial gauge usually has a travel between a few millimeters (e.g., 1 mm, 10 mm, 25 mm, 50 mm), the option suggests a *spectrum* of available dial gauges. Some specialized dial gauges, particularly electronic ones or those designed for specific industrial applications, can indeed have a total measurement range up to 300 mm or more, starting from very small values like 0.25 mm which could represent a starting point or minimum measurable displacement in certain setups or with specific types like test indicators. The lower end might also imply the least count or starting point from a reference. Given the options, this range from a very small displacement (0.25 mm) up to a considerable travel (300 mm)

encompasses the capabilities of various types of dial gauges and is the most plausible range among the choices for the *typical spectrum* available.

Based on standard manufacturing and applications, dial gauges typically have a relatively short stroke (travel) for high precision. Common mechanical dial gauges have ranges like 1 mm, 5 mm, 10 mm, 25 mm, or 50 mm. However, considering the possibility of larger or electronic versions, and the nature of the options given as broad ranges, the range represented by 0.25 to 300 mm covers the spectrum from very sensitive indicators to those with larger total travel, making it the most appropriate answer among the choices.

Common Dial Gauge Ranges (Examples)

Type	Typical Ranges
Mechanical Dial Indicator	1 mm, 5 mm, 10 mm, 25 mm, 50 mm
Electronic Dial Indicator	Up to 50 mm, 100 mm, or even larger
Dial Test Indicator	0.8 mm, 1 mm, 2 mm

The range 0.25 to 300 mm provided in the option is a broad classification that includes the capabilities of various types of dial gauges, from standard ones with small travel to potentially larger or electronic ones.

Revision Table: Dial Gauge Measurement Range

Concept	Explanation
Dial Gauge	Measures small linear displacements precisely.
Measurement Range (Travel)	Total distance the spindle can move and be measured.
Typical Range Spectrum	0.25 mm up to around 300 mm, covering various types.

Additional Information on Dial Gauges and Metrology

Dial gauges are fundamental tools in the field of metrology (the science of measurement). Their accuracy and precision are crucial for quality control and manufacturing processes.

Key Features of Dial Gauges:

- **Least Count:** The smallest measurement increment the gauge can display (e.g., 0.01 mm, 0.001 mm).
- **Accuracy:** How close the measured value is to the true value.
- **Repeatability:** The ability to get the same measurement value under the same conditions.
- **Dial Face:** Can be balanced (reading in both directions from zero) or continuous (reading in one direction up to the full range).

Comparison with Other Measuring Tools:

While dial gauges are excellent for comparative measurements and measuring small displacements with high precision, they are different from tools used for larger dimensions:

- **Calipers:** Used for measuring external, internal, step, and depth dimensions over larger ranges (e.g., 150 mm, 300 mm, up to 1000 mm or more). Less precise than dial gauges for small variations.
- **Micrometers:** Used for highly accurate measurement of external, internal, or depth dimensions, typically over smaller ranges (e.g., 0–25 mm, 25–50 mm) with very high precision (e.g., 0.01 mm, 0.001 mm).

The 0.25 to 300 mm range for dial gauges represents the overall variety available, from very sensitive instruments with limited travel (like 0.25 mm least count or starting point) to larger industrial dial gauges with significant total travel (up to 300 mm).

172. Answer: c

Explanation:

Understanding the Preparation of Alumina

Alumina, scientifically known as aluminium oxide, has the chemical formula Al_2O_3 . It is a very important chemical compound used extensively in various industries, including as the primary material for producing aluminium metal.

Alumina is not typically found in a pure state in nature. It is usually extracted from naturally occurring ores that contain aluminium compounds. Let's look at the given options to identify the main source material for preparing alumina.

- **Aluminium:** Aluminium is the metal produced from alumina through processes like the Hall-Hérault process. Therefore, aluminium metal is not the source material for alumina; rather, alumina is a precursor to aluminium metal.
- **Copper:** Copper is a completely different metal, primarily extracted from ores like chalcopyrite (CuFeS_2) or bornite (Cu_5FeS_4). Copper ores do not contain significant amounts of aluminium compounds suitable for commercial alumina production.
- **Bauxite:** Bauxite is the principal ore from which aluminium and its compounds, including alumina, are commercially extracted. Bauxite is a sedimentary rock that is rich in aluminium oxides and hydroxides, such as gibbsite ($\text{Al}(\text{OH})_3$), boehmite ($\text{AlO}(\text{OH})$), and diaspore ($\text{AlO}(\text{OH})$). The process used to refine bauxite into pure alumina is most commonly the Bayer process.
- **Galena:** Galena is the primary ore of lead, with the chemical formula PbS (lead sulfide). Like copper ores, galena is not a source material for preparing alumina.

Based on the composition and industrial processes, bauxite is the raw material used to prepare alumina. The Bayer process purifies the bauxite ore to obtain pure aluminium oxide (alumina), which is then used to produce aluminium metal or for other applications like ceramics and abrasives.

Revision Table: Ores and Their Products

Ore	Primary Metal Extracted / Compound Prepared	Key Component
Bauxite	Aluminium / Alumina (Al_2O_3)	Aluminium oxides and hydroxides
Chalcopyrite	Copper	Copper iron sulfide (CuFeS_2)
Galena	Lead	Lead sulfide (PbS)
Hematite	Iron	Iron(III) oxide (Fe_2O_3)

Additional Information on Alumina and Bauxite

Alumina (Al_2O_3) is a hard, chemically inert material with a high melting point. Besides being the precursor for aluminium production, it is used in:

- Ceramics (e.g., spark plug insulators, cutting tools)
- Abrasives (e.g., sandpaper, grinding wheels)
- Refractory materials (materials resistant to high temperatures)
- Catalyst supports

Bauxite deposits are found in various parts of the world, typically in tropical and subtropical regions. The exact composition of bauxite varies depending on the location, but it always contains a significant amount of aluminium-bearing minerals along with impurities like iron oxides, silica, and titanium dioxide. The Bayer process involves dissolving the aluminium compounds in a hot concentrated solution of sodium hydroxide, separating the liquid from the solid impurities, and then precipitating pure aluminium hydroxide, which is finally heated (calcined) to produce alumina (Al_2O_3).

173. Answer: a

Explanation:

Understanding Multiple V-Belt Drives and Maintenance

Multiple V-belt drives are common in power transmission systems, using several belts running parallel in corresponding grooves on pulleys. This setup allows for greater power transmission capacity compared to a single belt and provides redundancy – if one belt fails, the others can often continue transmitting some power, though usually at reduced capacity.

Addressing a Broken Belt in a Multiple V-Belt Drive

When one belt in a multiple V-belt drive system breaks, it indicates a failure point within the drive. While it might seem economical to replace only the broken belt, this is generally not the recommended practice for maintaining the efficiency and longevity of the system. Here's why:

- **Unequal Wear:** All belts in a multiple V-belt drive are typically installed and operate under similar conditions. Therefore, if one belt has failed, it's highly probable that the

other belts have also undergone significant wear and tear, even if they haven't broken yet.

- **Tension Imbalance:** A new belt will have different elasticity and length characteristics compared to the older, stretched belts. If a new belt is installed alongside old belts, the new belt will initially be tighter and carry a disproportionately higher share of the load.
- **Premature Failure:** The increased load on the new belt can lead to its premature failure. Simultaneously, the older belts, already weakened by wear, will experience additional stress due to the load imbalance, accelerating their failure as well.
- **Reduced Efficiency:** Uneven tension and load distribution among the belts reduce the overall efficiency of the power transmission system and can cause excessive vibration and noise.

Why Replacing All Belts is the Best Practice

To ensure optimal performance, reliability, and lifespan of a multiple V-belt drive, the standard recommendation is to replace all belts as a complete set if any one belt fails. Replacing all belts ensures:

- **Uniform Tension:** All belts will have similar tension characteristics, leading to balanced load sharing across the drive.
- **Even Load Distribution:** The power transmitted will be distributed evenly among all belts, preventing overloading of any single belt.
- **Extended System Life:** By starting with a fresh set of belts, the entire system operates more smoothly and efficiently, extending the life of the belts, pulleys, and bearings.
- **Reduced Downtime:** Replacing all belts at once reduces the likelihood of frequent, unpredictable failures of individual older belts, minimizing costly downtime for repairs.

Therefore, when one belt of a multiple V-belt drive is broken, replacing all the belts is the most effective solution.

Comparison: Replacing All Belts vs. Replacing Only Broken Belt

Feature	Replace All Belts	Replace Only Broken Belt
Load Distribution	Uniform and balanced	Uneven, new belt carries more load
Belt Lifespan	Maximized for the set	New belt fails prematurely, old belts fail sooner
System Efficiency	Optimized	Reduced due to imbalance
Risk of Future Failure	Minimized in the short term	High risk of subsequent failures
Maintenance Cost	Higher initial cost	Lower initial cost, higher long-term cost due to frequent replacements

Revision Table: Key Takeaways on Multiple V-Belt Maintenance

Concept	Explanation
Multiple V-Belts	Used for higher power and redundancy.
Broken Belt	Indicates wear across the system.
Recommended Action	Replace all belts simultaneously.
Reason	Ensure uniform tension, even load distribution, extended life.
Consequence of Partial Replacement	Uneven wear, premature failure, reduced efficiency.

Additional Information: V-Belt Maintenance Best Practices

- Always ensure proper belt tension during installation.
- Check belt and pulley alignment regularly.
- Inspect belts for signs of wear, cracking, or damage periodically.

- Use belts from the same manufacturer and production batch if possible.
- Store belts in a cool, dry place away from direct sunlight and chemicals.
- Never force belts onto pulleys; use proper installation tools.

174. Answer: a

Explanation:

Understanding Bolts for Machine Table T-Slots

Machine tables, often found on milling machines, drilling machines, and other machine tools, have special slots called T-slots. These T-slots are designed for securing workpieces, fixtures, vises, and other accessories to the table using clamping systems. Choosing the correct type of bolt is crucial for a secure and safe setup.

Let's look at the common types of bolts mentioned and their suitability for machine table T-slots:

- **T-headed bolt:** This type of bolt is specifically designed with a 'T'-shaped head. This head fits perfectly into the widened bottom section of the T-slot. When the nut is tightened on the threaded end, the T-head wedges against the underside of the slot, providing a strong clamping force that holds the attached item securely to the table.
- **Flange bolt:** A flange bolt has a washer-like flange integrated into the head. This spreads the load over a larger area. While useful in many applications, the head shape is typically hexagonal or similar and would not fit properly into the restrictive profile of a T-slot.
- **Hook bolt:** A hook bolt has a hooked end instead of a standard head. The hook is designed to catch onto an edge or through a hole. While variations might be used in some clamping setups, a standard hook bolt is not the primary or most effective fastener for directly fitting into and utilizing the T-slot profile itself for clamping workpieces.
- **Hexagon-headed bolt:** This is the most common type of bolt with a six-sided head. A standard hexagon-headed bolt requires a washer and nut, but its head is too large and the wrong shape to slide into or lock within a T-slot. Special adapters or T-slot nuts are sometimes used with hexagon-headed bolts, but the bolt itself doesn't directly engage the slot's T-shape.

Based on their design and how they interact with the machine table T-slot profile, the **T-headed bolt** is the fastener specifically made for this application.

The T-head slides along the slot and rotates slightly to lock into place before tightening, ensuring a positive mechanical connection within the T-slot structure.

Why T-headed Bolts Fit T-Slots

The design of a T-slot and a T-headed bolt is complementary:

- The T-slot has a narrow opening at the top and a wider section underneath, forming the 'T' shape when viewed in cross-section.
- The T-headed bolt has a rectangular or square head with dimensions designed to pass through the narrow top opening and then rotate 90 degrees to sit within the wider bottom section of the T-slot.

This design allows the bolt to be inserted anywhere along the length of the slot and provides a robust anchoring point for clamping.

Bolt Types vs. T-Slot Compatibility			
Bolt Type	Head Shape	Fits T-Slot?	Primary Use Example
T-headed bolt	T-shaped (rectangular/square base)	Yes	Securing items to machine table T-slots
Flange bolt	Hexagon with integrated washer	No (head is wrong shape)	Connecting components, reducing need for separate washer
Hook bolt	Hook	No (designed to hook edges)	Specific clamping or lifting applications
Hexagon-headed bolt	Hexagon	No (head is too large/wrong shape)	General fastening (requires T-slot nut for T-slots)

Conclusion on Machine Table Bolts

For direct attachment and clamping using the inherent design of the T-slots on a machine table, the **T-headed bolt** is the appropriate fastener. Its unique head shape is specifically engineered to lock into the T-slot profile, providing a secure and stable anchor point for holding workpieces or fixtures during machining operations.

Revision Table: Machine Shop Fasteners

Reviewing key points about fasteners used in machine shops:

- **T-headed bolts:** Essential for machine table T-slots.
- **Standard bolts (Hex, Cap Screws):** Used with nuts; require T-nuts or adapters for T-slots.
- **Clamps:** Used with T-slot bolts (like T-headed bolts) to hold workpieces.
- **T-nuts:** Inserted into T-slots to provide a threaded hole for standard bolts.

Additional Information: T-Slot Applications

Beyond machine tables, T-slots and T-headed bolts (or T-nuts with standard bolts) are used in various applications requiring adjustable fastening along a linear track:

- Assembly lines and automation setups
- Framing systems (e.g., aluminum profiles)
- Fixture building
- CNC router tables
- Jigs and custom clamping setups

Understanding the function of T-slots and the fasteners designed for them is fundamental in machining and manufacturing.

175. Answer: a

Explanation:

Understanding GO and NO GO Limit Gauges

'GO' and 'NO GO' gauges are essential tools used in manufacturing and quality control to quickly and efficiently check if a manufactured part's dimension is within the specified

tolerance limits.

These gauges operate on the principle of attribute gauging, meaning they don't measure the exact size but rather check if a dimension falls within a permissible range. This range is defined by upper and lower limits, which are derived from the nominal size and the specified tolerance.

What is a Limit Gauge?

A limit gauge is a measuring tool used to check whether the dimensions of a part lie within specified upper and lower limits. These limits are often referred to as the maximum material condition (MMC) and the minimum material condition (LMC or MMC based on feature type – internal/external). Limit gauges are designed based on Taylor's principle, which guides their design for both the GO and NO GO sides.

The Function of GO and NO GO Gauges

- **GO Gauge:** The GO side of the gauge is designed to check the maximum material condition (MMC). For an external feature (like a shaft), the GO gauge checks if the shaft is not too large. For an internal feature (like a hole), the GO gauge checks if the hole is not too small. If the GO gauge fits or enters the feature, it means the part is at or below the MMC for an external feature or at or above the MMC for an internal feature.
- **NO GO Gauge:** The NO GO side of the gauge is designed to check the minimum material condition (LMC for external, MMC for internal, depending on context but generally related to the 'other' limit). For an external feature, the NO GO gauge checks if the shaft is not too small. For an internal feature, the NO GO gauge checks if the hole is not too large. If the NO GO gauge does **not** fit or enter the feature, it means the part is at or outside the specified limit checked by the NO GO gauge, which is the desired outcome for a good part.

For a part to be acceptable, the GO gauge must fit, and the NO GO gauge must not fit.

Why are they Limit Gauges?

The name 'Limit gauge' comes directly from their purpose: checking if a dimension is within the specified upper and lower size limits (tolerances) without measuring the actual size. 'GO' and 'NO GO' gauges are the most common type of limit gauge.

Analysing the Options

- **Limit gauge:** This category includes gauges designed to check if a dimension is within upper and lower limits. 'GO' and 'NO GO' gauges perfectly fit this description.
- **Ring gauge:** A ring gauge is specifically used to check the outside diameter of cylindrical parts. While a ring gauge can be a type of limit gauge (e.g., a GO ring gauge and a NO GO ring gauge), the term 'Limit gauge' is the broader category that includes 'GO' and 'NO GO' functionality for various feature types.
- **Slip gauge:** Slip gauges (also known as gauge blocks) are precision length standards used for calibrating other measuring instruments, setting up machine tools, or verifying dimensions. They are not used as 'GO' and 'NO GO' checks on manufactured parts in the same way.
- **Plug gauge:** A plug gauge is specifically used to check the diameter of holes or internal features. Similar to ring gauges, a plug gauge can be a type of limit gauge (e.g., a GO plug gauge and a NO GO plug gauge), but 'Limit gauge' is the encompassing term.

Therefore, 'GO' and 'NO GO' gauges are fundamentally a type of limit gauge because their design and function are based on checking the dimensional limits of a feature.

Revision Table: Key Gauge Types

Gauge Type	Primary Function	Example Use
Limit Gauge (GO/NO GO)	Check if dimension is within specified limits (pass/fail).	Checking bore diameter tolerance, shaft diameter tolerance.
Ring Gauge	Check external diameter.	Checking outside diameter of a pin or shaft.
Slip Gauge (Gauge Block)	Precision length standard for calibration/setting.	Calibrating micrometers, setting milling machine depth.
Plug Gauge	Check internal diameter (hole size).	Checking diameter of a drilled or reamed hole.

Additional Information on Dimensional Tolerances and Gauging

In manufacturing, it's impossible to produce parts with theoretically perfect dimensions. Therefore, engineers specify tolerances, which are the permissible variations in size or form.

- **Tolerance:** The total permissible variation in a dimension. It is the difference between the upper limit and the lower limit of size.
- **Limits of Size:** The maximum and minimum permissible sizes of a feature.
- **Fit:** The relationship between two mating parts (e.g., a shaft and a hole) concerning the difference in their sizes before assembly. Types of fits include clearance fit, interference fit, and transition fit.

'GO' and 'NO GO' gauges are designed precisely based on these limits of size to ensure that parts will assemble correctly and function as intended according to the specified fit.

Prepp

Your Personal Exams Guide