



General Instructions

- (i) The test is of 2 hours duration.
- (ii) This test paper consists of 120 questions. The maximum marks are 120.
- (iii) Each question carries +1 marks for correct answer and there is no negative marking for wrong answer.

1. If λ_1, λ_2 and λ_3 are the eigen values of a square matrix A , then the eigen values of $A^{-1} + 2I + A$ are

- (A) $\frac{1}{\lambda_1} + 2 + \lambda_1, \frac{1}{\lambda_2} + 2 + \lambda_2$ and $\frac{1}{\lambda_3} + 2 + \lambda_3$
- (B) $\frac{1}{\lambda_1} + 2 + \lambda_1, \frac{1}{\lambda_2} - 2 - \lambda_2$ and $\frac{1}{\lambda_3} - 2 - \lambda_3$
- (C) $\frac{1}{\lambda_1} - 2 - \lambda_1, \frac{1}{\lambda_2} - 2 - \lambda_2$ and $\frac{1}{\lambda_3} - 2 - \lambda_3$
- (D) $\frac{1}{\lambda_1} + 2 + \lambda_1, \frac{1}{\lambda_2} + 2 + \lambda_2$ and $\frac{1}{\lambda_3} + 2 + \lambda_3$

Correct Answer: (D) $\frac{1}{\lambda_1} + 2 + \lambda_1, \frac{1}{\lambda_2} + 2 + \lambda_2$ and $\frac{1}{\lambda_3} + 2 + \lambda_3$

Solution:

Concept: Let A be an $n \times n$ square matrix. If λ is an eigenvalue of A corresponding to a non-zero eigenvector X , then by definition:

$$AX = \lambda X$$

Using matrix algebra and spectral properties, we can determine the eigenvalues of functions of matrices:

- **Eigenvalues of an inverse matrix (A^{-1}):** If A is non-singular, multiplying both sides of $AX = \lambda X$ by A^{-1} gives $X = \lambda A^{-1}X$, which transforms directly to $A^{-1}X = \frac{1}{\lambda}X$. Thus, the eigenvalues are $\frac{1}{\lambda}$.
- **Eigenvalues of a scalar multiple matrix (kI):** Since $IX = X$, we have $(kI)X = kX$. The

eigenvalue of the identity matrix scaled by k is identically k .

- **Eigenvalues of matrix polynomials/combinations:** If $P(A) = c_m A^m + \cdots + c_1 A + c_0 I$, then the eigenvalues of $P(A)$ are given exactly by $P(\lambda) = c_m \lambda^m + \cdots + c_1 \lambda + c_0$.

Step 1: Application of spectral mapping theorem to each component.

We are asked to find the eigenvalues of the linear matrix combination expression:

$$B = A^{-1} + 2I + A$$

Let X_i be the non-zero eigenvector associated with the specific eigenvalue λ_i of the square matrix A . This means:

$$AX_i = \lambda_i X_i \quad \text{for } i = 1, 2, 3.$$

Step 2: Operating the matrix equation onto the eigenvector X_i .

Let us apply the matrix combination expression B directly onto the eigenvector X_i :

$$BX_i = (A^{-1} + 2I + A)X_i$$

Distributing the eigenvector vector multiplication over the matrix addition linearly:

$$BX_i = A^{-1}X_i + 2IX_i + AX_i$$

Step 3: Substituting known eigenvalue relations.

We substitute the fundamental eigenvalue properties for each standalone term inside the equation:

1. For the matrix vector product $A^{-1}X_i$, since A has eigenvalue λ_i , its inverse operation satisfies $A^{-1}X_i = \frac{1}{\lambda_i}X_i$.
2. For the identity component product, $2IX_i = 2X_i$.
3. For the primary matrix product, $AX_i = \lambda_i X_i$.

Substituting these direct scalar equations back into our operational matrix-vector expansion:

$$BX_i = \left(\frac{1}{\lambda_i}\right)X_i + 2X_i + \lambda_i X_i$$

Step 4: Factoring out the common vector component to finalize the eigenvalue format.

We now safely factor out the common non-zero eigenvector X_i from the right side of the

expression:

$$BX_i = \left(\frac{1}{\lambda_i} + 2 + \lambda_i \right) X_i$$

This precisely fits the standard eigenvalue-eigenvector functional definition $BX_i = \mu_i X_i$, showing that the scalar multiplier value is the new eigenvalue. Repeating this exact linear mapping process independently for all three eigenvalues $\lambda_1, \lambda_2, \lambda_3$, the complete set of new eigenvalues for the matrix $A^{-1} + 2I + A$ is uniquely evaluated as:

$$\frac{1}{\lambda_1} + 2 + \lambda_1, \quad \frac{1}{\lambda_2} + 2 + \lambda_2, \quad \text{and} \quad \frac{1}{\lambda_3} + 2 + \lambda_3$$

This maps explicitly and without any sign alterations to Option (D).

Quick Tip: If λ is an eigenvalue of a matrix A , then any matrix polynomial or functional expression $f(A)$ will have corresponding eigenvalues given directly by substituting λ into the function, resulting in $f(\lambda)$.

2. Let the matrix $A = \begin{bmatrix} 1 & 4 \\ 2 & -1 \end{bmatrix}$.

Statement-I: $A^2 = 9I$

Statement-II: The eigen values of A are -3 and 3

The correct answer is

- (A) Both statement-I and statement-II are true
- (B) Both statement-I and statement-II are false
- (C) Statement-I is true, but statement-II is false
- (D) Statement-I is false, but statement-II is true

Correct Answer: (A) Both statement-I and statement-II are true

Solution:

Concept: For any given square matrix A , we can verify matrix relations using direct matrix multiplication. Additionally, the characteristic equation used to find the eigenvalues λ of a matrix A is determined by solving the determinant equation:

$$\det(A - \lambda I) = 0$$

Alternatively, according to the Cayley-Hamilton theorem, every square matrix satisfies its own characteristic equation. For a 2×2 matrix, this equation can be simplified as:

$$\lambda^2 - \text{Tr}(A)\lambda + \det(A) = 0$$

Step 1: Verifying Statement-I by computing A^2 .

We are given the matrix:

$$A = \begin{bmatrix} 1 & 4 \\ 2 & -1 \end{bmatrix}$$

Let us compute the square of matrix A , which is defined as the product $A \times A$:

$$A^2 = \begin{bmatrix} 1 & 4 \\ 2 & -1 \end{bmatrix} \begin{bmatrix} 1 & 4 \\ 2 & -1 \end{bmatrix}$$

Performing matrix multiplication row by column explicitly:

$$A^2 = \begin{bmatrix} (1)(1) + (4)(2) & (1)(4) + (4)(-1) \\ (2)(1) + (-1)(2) & (2)(4) + (-1)(-1) \end{bmatrix}$$

Evaluating individual components inside the matrix structure:

$$A^2 = \begin{bmatrix} 1+8 & 4-4 \\ 2-2 & 8+1 \end{bmatrix} = \begin{bmatrix} 9 & 0 \\ 0 & 9 \end{bmatrix}$$

Factoring out the scalar quantity 9 from the matrix elements:

$$A^2 = 9 \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix} = 9I$$

Hence, $A^2 = 9I$ is completely accurate, which confirms that **Statement-I is true**.

Step 2: Verifying Statement-II by evaluating eigenvalues.

To find the eigenvalues of matrix A , we write down its characteristic equation:

$$\det(A - \lambda I) = 0 \Rightarrow \begin{vmatrix} 1-\lambda & 4 \\ 2 & -1-\lambda \end{vmatrix} = 0$$

Expanding the 2×2 determinant equation:

$$(1 - \lambda)(-1 - \lambda) - (4)(2) = 0$$

Multiplying out the binomial terms carefully:

$$-1 - \lambda + \lambda + \lambda^2 - 8 = 0$$

Combining like terms and constants:

$$\lambda^2 - 9 = 0$$

Solving for λ :

$$\lambda^2 = 9 \Rightarrow \lambda = \pm\sqrt{9} \Rightarrow \lambda = 3, -3$$

The calculated eigenvalues of A are exactly 3 and -3 . This perfectly matches the text of the second claim. Therefore, **Statement-II is true**.

Since both individual statements have been analytically verified as correct, Option (A) is the correct choice.

Quick Tip: By the Cayley-Hamilton Theorem, if a matrix satisfies $A^2 = 9I$, its matrix characteristic equation is $\lambda^2 - 9 = 0$. Solving this directly yields eigenvalues $\lambda = \pm 3$ instantly without needing separate determinant expansion calculations!

3. Statement-I: The function $u = x^2 + y^2$, $v = \tan^{-1}\left(\frac{y}{x}\right)$ are functionally independent.

Statement-II: The Jacobian $\frac{\partial(u,v)}{\partial(x,y)}$ is non-zero.

The correct answer is

- (A) Statement-I is true, but statement-II is false
- (B) Statement-I is false, but statement-II is true
- (C) Both statement-I and statement-II are true
- (D) Both statement-I and statement-II are false

Correct Answer: (C) Both statement-I and statement-II are true

Solution:

Concept: Two differentiable functions $u(x, y)$ and $v(x, y)$ are said to be functionally dependent

if their Jacobian determinant evaluates identically to zero across a region. Conversely, they are defined to be ****functionally independent**** if and only if their Jacobian determinant is non-zero:

$$J = \frac{\partial(u, v)}{\partial(x, y)} = \begin{vmatrix} \frac{\partial u}{\partial x} & \frac{\partial u}{\partial y} \\ \frac{\partial v}{\partial x} & \frac{\partial v}{\partial y} \end{vmatrix} \neq 0$$

Therefore, Statement-II functions as the direct mathematical definition/condition for establishing Statement-I. Let us calculate the partial derivatives to check the value.

Step 1: Determining the partial derivatives of u .

The first given function expression is:

$$u = x^2 + y^2$$

Differentiating u partially with respect to the variable x (treating y as a constant):

$$\frac{\partial u}{\partial x} = 2x$$

Differentiating u partially with respect to the variable y (treating x as a constant):

$$\frac{\partial u}{\partial y} = 2y$$

Step 2: Determining the partial derivatives of v .

The second given function expression is:

$$v = \tan^{-1}\left(\frac{y}{x}\right)$$

Recall the derivative rule: $\frac{d}{dt}(\tan^{-1}(t)) = \frac{1}{1+t^2}$. Applying the chain rule for partial differentiation with respect to x :

$$\frac{\partial v}{\partial x} = \frac{1}{1 + \left(\frac{y}{x}\right)^2} \cdot \frac{\partial}{\partial x}\left(\frac{y}{x}\right) = \frac{1}{\frac{x^2+y^2}{x^2}} \cdot \left(-\frac{y}{x^2}\right)$$

Simplifying this expression:

$$\frac{\partial v}{\partial x} = \frac{x^2}{x^2 + y^2} \cdot \left(-\frac{y}{x^2}\right) = -\frac{y}{x^2 + y^2}$$

Now, performing partial differentiation with respect to y :

$$\frac{\partial v}{\partial y} = \frac{1}{1 + \left(\frac{y}{x}\right)^2} \cdot \frac{\partial}{\partial y} \left(\frac{y}{x}\right) = \frac{x^2}{x^2 + y^2} \cdot \left(\frac{1}{x}\right)$$

Simplifying this expression:

$$\frac{\partial v}{\partial y} = \frac{x}{x^2 + y^2}$$

Step 3: Setting up and computing the Jacobian determinant $\frac{\partial(u,v)}{\partial(x,y)}$.

We construct the Jacobian matrix determinant using our calculated partial derivative expressions:

$$\frac{\partial(u,v)}{\partial(x,y)} = \begin{vmatrix} 2x & 2y \\ -\frac{y}{x^2+y^2} & \frac{x}{x^2+y^2} \end{vmatrix}$$

Evaluating the 2×2 determinant product:

$$\frac{\partial(u,v)}{\partial(x,y)} = (2x) \left(\frac{x}{x^2 + y^2}\right) - (2y) \left(-\frac{y}{x^2 + y^2}\right)$$

Combining the terms over their common denominator:

$$\frac{\partial(u,v)}{\partial(x,y)} = \frac{2x^2}{x^2 + y^2} + \frac{2y^2}{x^2 + y^2} = \frac{2(x^2 + y^2)}{x^2 + y^2}$$

Canceling out the common term $(x^2 + y^2)$ (assuming $x, y \neq 0$):

$$\frac{\partial(u,v)}{\partial(x,y)} = 2$$

Step 4: Evaluating the truth value of both statements.

Since our calculation reveals that the Jacobian value is exactly 2, it is clearly non-zero ($2 \neq 0$).

- This directly confirms that **Statement-II is true** because the Jacobian is non-zero.
- Because a non-zero Jacobian structurally guarantees functional independence, **Statement-I is also true.**

Thus, both statements are mathematically true, matching Option (C).

Quick Tip: Notice that $u = r^2$ and $v = \theta$ in polar coordinates! Since polar coordinates map points uniquely to independent coordinate axes, functions of r alone and θ alone are always functionally independent.

4. If f and g are the solutions of Laplace equation then $\nabla \cdot \{(f \nabla g) - (g \nabla f)\} =$

- (A) ∇g
- (B) 0
- (C) 1
- (D) ∇f

Correct Answer: (B) 0

Solution:

Concept: A scalar function ϕ satisfies the Laplace equation if its Laplacian equals zero:

$$\nabla^2 \phi = 0$$

Therefore, since f and g are explicitly stated to be solutions of the Laplace equation, we have:

$$\nabla^2 f = 0 \quad \text{and} \quad \nabla^2 g = 0$$

We will also make use of the standard vector identity for the divergence of a scalar multiplied by a vector field, $\nabla \cdot (\phi \mathbf{A}) = \phi(\nabla \cdot \mathbf{A}) + \mathbf{A} \cdot (\nabla \phi)$. When applied to a gradient field $\mathbf{A} = \nabla \psi$, this identity becomes:

$$\nabla \cdot (\phi \nabla \psi) = \phi \nabla^2 \psi + \nabla \phi \cdot \nabla \psi$$

Step 1: Expanding the given divergence expression using vector identities.

The objective expression to expand and calculate is:

$$W = \nabla \cdot \{(f \nabla g) - (g \nabla f)\}$$

By the linearity property of the divergence operator, we can split this into two separate divergence terms:

$$W = \nabla \cdot (f \nabla g) - \nabla \cdot (g \nabla f) \quad \dots (1)$$

Step 2: Evaluating the first divergence term $\nabla \cdot (f \nabla g)$.

Applying our vector identity directly where the scalar function is f and the gradient field belongs to g :

$$\nabla \cdot (f \nabla g) = f(\nabla \cdot \nabla g) + \nabla f \cdot \nabla g = f \nabla^2 g + \nabla f \cdot \nabla g$$

Step 3: Evaluating the second divergence term $\nabla \cdot (g \nabla f)$.

Applying our vector identity where the roles are reversed (the scalar function is g and the gradient field belongs to f):

$$\nabla \cdot (g \nabla f) = g(\nabla \cdot \nabla f) + \nabla g \cdot \nabla f = g \nabla^2 f + \nabla g \cdot \nabla f$$

Step 4: Combining both expanded terms back into the original equation.

Substitute these two expanded formulas back into equation (1):

$$W = (f \nabla^2 g + \nabla f \cdot \nabla g) - (g \nabla^2 f + \nabla g \cdot \nabla f)$$

Distributing the negative sign through the second expression:

$$W = f \nabla^2 g + \nabla f \cdot \nabla g - g \nabla^2 f - \nabla g \cdot \nabla f$$

Since the scalar dot product is commutative, $\nabla f \cdot \nabla g = \nabla g \cdot \nabla f$. Therefore, these two terms cancel each other out completely:

$$W = f \nabla^2 g - g \nabla^2 f$$

Step 5: Applying the Laplace equation conditions.

Since f and g satisfy the Laplace equation, we substitute $\nabla^2 g = 0$ and $\nabla^2 f = 0$ into our expression:

$$W = f(0) - g(0) = 0 - 0 = 0$$

Thus, the expression simplifies cleanly to 0, which matches Option (B).

Quick Tip: This expression is the integrand used in Green's Second Identity: $\int_V (f \nabla^2 g - g \nabla^2 f) dV = \oint_S (f \nabla g - g \nabla f) \cdot d\mathbf{n}$. Since both functions satisfy Laplace's equation, both sides of the identity are identically zero.

5. If the particular integral of $y'' - 4y' = x^2 e^{2x}$ is in the form $y_p = e^{2x} y(x)$, then $y(x)$ is

(A) $-\frac{1}{4} \left(x^2 + \frac{1}{2} \right)$

(B) $\frac{1}{4} \left(x^2 + \frac{1}{2} \right)$

(C) $x^2 + \frac{1}{2}$

(D) $x + \frac{1}{2}$

Correct Answer: (A) $-\frac{1}{4} \left(x^2 + \frac{1}{2} \right)$

Solution:

Concept: The given linear differential equation can be written using differential operator notation $D = \frac{d}{dx}$:

$$(D^2 - 4D)y = x^2 e^{2x}$$

The particular integral y_p is found using the inverse operator:

$$y_p = \frac{1}{D^2 - 4D} (x^2 e^{2x})$$

To evaluate this when an exponential function e^{ax} is multiplied by another function, we use the exponential shift rule:

$$\frac{1}{f(D)} (e^{ax} V(x)) = e^{ax} \frac{1}{f(D+a)} V(x)$$

Step 1: Setting up the particular integral operator expression.

Identify our function of the differential operator from the differential equation:

$$f(D) = D^2 - 4D$$

Here, our exponential multiplier term has a power coefficient of $a = 2$, and the remaining function is $V(x) = x^2$. Applying the shift rule changes the operator $D \rightarrow D + 2$:

$$y_p = e^{2x} \frac{1}{(D+2)^2 - 4(D+2)} (x^2)$$

Step 2: Simplifying the shifted denominator operator.

Let us expand the algebraic terms in the denominator:

$$(D+2)^2 = D^2 + 4D + 4$$

$$-4(D + 2) = -4D - 8$$

Adding these expressions together:

$$f(D + 2) = (D^2 + 4D + 4) + (-4D - 8) = D^2 - 4$$

Substitute this back into the expression for y_p :

$$y_p = e^{2x} \frac{1}{D^2 - 4} (x^2) \quad \dots (1)$$

Step 3: Expanding the operator via binomial expansion for polynomial terms.

To apply the operator to a polynomial like x^2 , we rewrite the denominator in the form $-(4 - D^2)$ and factor out the constant to use a binomial series:

$$\frac{1}{D^2 - 4} = \frac{1}{-4 \left(1 - \frac{D^2}{4}\right)} = -\frac{1}{4} \left(1 - \frac{D^2}{4}\right)^{-1}$$

Using the binomial expansion formula $(1 - t)^{-1} = 1 + t + t^2 + \dots$:

$$\left(1 - \frac{D^2}{4}\right)^{-1} = 1 + \frac{D^2}{4} + \frac{D^4}{16} + \dots$$

Since our target polynomial expression is x^2 , any derivative higher than the second derivative will equal zero. Therefore, we can drop terms containing D^4 and higher powers:

$$y_p = e^{2x} \left[-\frac{1}{4} \left(1 + \frac{D^2}{4}\right) \right] (x^2)$$

Step 4: Distributing the operator onto the polynomial x^2 .

Now, let's apply the operations inside the brackets to the polynomial:

$$y_p = -\frac{1}{4} e^{2x} \left[x^2 + \frac{1}{4} D^2(x^2) \right]$$

Compute the required derivatives of x^2 :

$$D(x^2) = \frac{d}{dx}(x^2) = 2x$$

$$D^2(x^2) = \frac{d}{dx}(2x) = 2$$

Substitute $D^2(x^2) = 2$ back into the expression:

$$y_p = -\frac{1}{4}e^{2x} \left(x^2 + \frac{1}{4} \cdot 2 \right) = -\frac{1}{4}e^{2x} \left(x^2 + \frac{1}{2} \right)$$

Step 5: Equating to find the target function $y(x)$.

The problem states that the particular integral is written in the form $y_p = e^{2x}y(x)$. Comparing this directly with our calculated result:

$$e^{2x}y(x) = e^{2x} \left[-\frac{1}{4} \left(x^2 + \frac{1}{2} \right) \right]$$

Dividing out the matching exponential factor from both sides isolates our target function:

$$y(x) = -\frac{1}{4} \left(x^2 + \frac{1}{2} \right)$$

This matches option (A).

Quick Tip: Always ensure the constant term is factored out as a positive 1 before running your binomial expansion series. Forgetting to factor out the negative sign from $D^2 - 4$ is a common trap that leads to incorrect signs!

6. $\mathcal{L}\{\sin(t-3)u(t-3)\} =$

- (A) $\frac{1}{s^2+1}$
- (B) $\frac{3}{s^2+9}$
- (C) $\frac{e^{-3s}}{s^2+1}$
- (D) $\frac{e^{-3s}}{s^2+9}$

Correct Answer: (C) $\frac{e^{-3s}}{s^2+1}$

Solution:

Concept: The problem asks for the Laplace transform of a shifted function multiplied by a unit step function $u(t-a)$ (also written as $H(t-a)$). To solve this, we use the Second Shifting Theorem of Laplace transforms:

$$\mathcal{L}\{f(t-a)u(t-a)\} = e^{-as} \mathcal{L}\{f(t)\} = e^{-as}F(s)$$

Additionally, we use the standard Laplace transform for a sine wave function:

$$\mathcal{L}\{\sin(\omega t)\} = \frac{\omega}{s^2 + \omega^2}$$

Step 1: Identifying the parameters from the given function.

The given expression to transform is:

$$\mathcal{L}\{\sin(t-3)u(t-3)\}$$

By comparing this directly with the standard formula $f(t-a)u(t-a)$, we can identify our parameters:

- The shift constant is $a = 3$.
- The shifted function is $f(t-3) = \sin(t-3)$.

Replacing $(t-3)$ with an unshifted variable t gives the base function:

$$f(t) = \sin(t)$$

Step 2: Finding the Laplace transform of the unshifted base function $f(t)$.

Now, let us calculate the Laplace transform $F(s)$ of our base function $f(t) = \sin(t)$. Here, the frequency coefficient is $\omega = 1$:

$$F(s) = \mathcal{L}\{\sin(t)\} = \frac{1}{s^2 + 1^2} = \frac{1}{s^2 + 1}$$

Step 3: Applying the Second Shifting Theorem formula.

Now we combine our results using the Second Shifting Theorem:

$$\mathcal{L}\{\sin(t-3)u(t-3)\} = e^{-3s}F(s)$$

Substituting our calculated value for $F(s)$:

$$\mathcal{L}\{\sin(t-3)u(t-3)\} = e^{-3s} \cdot \left(\frac{1}{s^2 + 1} \right) = \frac{e^{-3s}}{s^2 + 1}$$

This final form matches Option (C).

Quick Tip: The unit step multiplier $u(t - a)$ always transforms into an exponential factor e^{-as} in the s -domain. Once you factor that out, you only need to compute the standard transform of the unshifted function.

7. The Laurent's series for the function $f(z) = 1 + \frac{3}{z+2} - \frac{8}{z+3}$ in the region $|z| < 2$ is

- (A) $\frac{3}{2} \sum_{n=0}^{\infty} (-1)^n \left(\frac{z}{3}\right)^n - \frac{8}{3} \sum_{n=0}^{\infty} (-1)^n \left(\frac{z}{2}\right)^n$
 (B) $1 + \frac{3}{2} \sum_{n=0}^{\infty} (-1)^n \left(\frac{z}{2}\right)^n - \frac{8}{3} \sum_{n=0}^{\infty} (-1)^n \left(\frac{z}{2}\right)^n$
 (C) $1 + \frac{3}{2} \sum_{n=0}^{\infty} (-1)^n \left(\frac{z}{2}\right)^n - \frac{8}{3} \sum_{n=0}^{\infty} (-1)^n \left(\frac{z}{3}\right)^n$
 (D) $\frac{3}{2} \sum_{n=0}^{\infty} \left(\frac{z}{2}\right)^n (-1)^n - \frac{8}{3} \sum_{n=0}^{\infty} (-1)^n \left(\frac{z}{3}\right)^n$

Correct Answer: (C) $1 + \frac{3}{2} \sum_{n=0}^{\infty} (-1)^n \left(\frac{z}{2}\right)^n - \frac{8}{3} \sum_{n=0}^{\infty} (-1)^n \left(\frac{z}{3}\right)^n$

Solution:

Concept: To find the Laurent (or Taylor) series expansion of a function within a specified disk or annulus, we use the standard geometric series expansion:

$$\frac{1}{1+t} = (1+t)^{-1} = \sum_{n=0}^{\infty} (-1)^n t^n \quad \text{valid for } |t| < 1$$

The given convergence region is the open disk $|z| < 2$. This convergence constraint means that:

- For terms with a denominator like $z + 2$, we must factor out 2 so that the variable term becomes $\frac{z}{2}$, since $|z| < 2 \implies \left|\frac{z}{2}\right| < 1$.
- For terms with a denominator like $z + 3$, since $|z| < 2 < 3$, it follows that $|z| < 3 \implies \left|\frac{z}{3}\right| < 1$. Thus, we factor out 3.

Step 1: Setting up the expansion for the second term $\frac{3}{z+2}$.

Let us rewrite this fraction to set up our geometric series form by factoring out the constant 2 from the denominator:

$$\frac{3}{z+2} = \frac{3}{2\left(1+\frac{z}{2}\right)} = \frac{3}{2} \left(1+\frac{z}{2}\right)^{-1}$$

Since our region specifies $|z| < 2$, the absolute value satisfies $\left|\frac{z}{2}\right| < 1$. Applying the binomial series expansion:

$$\frac{3}{2} \left(1+\frac{z}{2}\right)^{-1} = \frac{3}{2} \sum_{n=0}^{\infty} (-1)^n \left(\frac{z}{2}\right)^n \quad \dots (1)$$

Step 2: Setting up the expansion for the third term $\frac{8}{z+3}$.

We use the same approach for the next fraction, factoring out the constant 3 from the denominator:

$$\frac{8}{z+3} = \frac{8}{3\left(1+\frac{z}{3}\right)} = \frac{8}{3}\left(1+\frac{z}{3}\right)^{-1}$$

Since $|z| < 2$, it is also true that $|z| < 3$, which means $\left|\frac{z}{3}\right| < \frac{2}{3} < 1$. This justifies expanding it as a geometric series:

$$\frac{8}{3}\left(1+\frac{z}{3}\right)^{-1} = \frac{8}{3}\sum_{n=0}^{\infty}(-1)^n\left(\frac{z}{3}\right)^n \quad \dots(2)$$

Step 3: Combining all the terms to form the full series expansion.

The complete function given in the problem is:

$$f(z) = 1 + \frac{3}{z+2} - \frac{8}{z+3}$$

Substituting our two derived series expressions (1) and (2) directly back into this equation yields:

$$f(z) = 1 + \frac{3}{2}\sum_{n=0}^{\infty}(-1)^n\left(\frac{z}{2}\right)^n - \frac{8}{3}\sum_{n=0}^{\infty}(-1)^n\left(\frac{z}{3}\right)^n$$

This matches Option (C).

Quick Tip: Always look closely at the convergence condition! To ensure your geometric series converges ($|t| < 1$), always factor out the larger value in magnitude between the variable z and the constant term.

8. Given $x_0 \neq 0$, if the iteration formula $x_{n+1} = \frac{1}{2}\left(-\frac{7}{x_n} + x_n\right)$, $n \geq 0$ is used to find the root of $f(x) = 0$, then $f(x) =$

- (A) $9 + x^2$
- (B) $4 + x^2$
- (C) $7 + x^2$
- (D) $-7 + x^2$

Correct Answer: (C) $7 + x^2$ or (D) $-7 + x^2$ (Note: Both yield the same roots; standard form signifies equation $x^2 - 7 = 0$ or $x^2 + 7 = 0$ under convergence limits). Let us analyze the fixed point limit carefully.

Solution:

Concept: When an iterative root-finding process converges to a stable value, the value at step x_{n+1} approaches the value at step x_n . This limiting value α is called a fixed point of the iteration:

$$\lim_{n \rightarrow \infty} x_n = \lim_{n \rightarrow \infty} x_{n+1} = \alpha$$

Substituting α into the recurrence relation allows us to reconstruct the original underlying polynomial equation $f(\alpha) = 0$.

Step 1: Applying the fixed-point condition to the iterative formula.

The given numerical iteration equation is:

$$x_{n+1} = \frac{1}{2} \left(-\frac{7}{x_n} + x_n \right)$$

Let us substitute the fixed-point limit value α for both x_n and x_{n+1} :

$$\alpha = \frac{1}{2} \left(-\frac{7}{\alpha} + \alpha \right)$$

Step 2: Solving the algebraic equation for α .

Multiply both sides of the equation by 2 to clear the fraction:

$$2\alpha = -\frac{7}{\alpha} + \alpha$$

Subtract α from both sides to collect like terms:

$$2\alpha - \alpha = -\frac{7}{\alpha} \Rightarrow \alpha = -\frac{7}{\alpha}$$

Step 3: Eliminating the denominator variable.

Multiply both sides of the equation by α (given that $x_0 \neq 0 \Rightarrow \alpha \neq 0$):

$$\alpha^2 = -7$$

Rearranging all terms onto one side of the equation:

$$\alpha^2 + 7 = 0$$

Step 4: Identifying the function $f(x)$.

Since the fixed-point limit α represents a root of the equation $f(x) = 0$, we replace α back with the independent variable x :

$$f(x) = x^2 + 7 = 7 + x^2$$

This matches option (C).

Quick Tip: This formula matches Newton-Raphson iteration $x_{n+1} = x_n - \frac{f(x_n)}{f'(x_n)}$ for finding purely imaginary roots, or a rearranged fixed-point iteration format $x = g(x)$ designed to isolate roots where $x^2 = -7$.

9. If a random variable X , defined such that $E(X - 1)^2 = 10$ and $E(X - 2)^2 = 6$ then mean (μ) and variance (σ^2) are

- (A) $\frac{2}{7}, \frac{4}{15}$
- (B) $\frac{2}{7}, \frac{15}{4}$
- (C) $\frac{7}{2}, \frac{4}{15}$
- (D) $\frac{7}{2}, \frac{15}{4}$

Correct Answer: (D) $\frac{7}{2}, \frac{15}{4}$

Solution:

Concept: The expectation operator $E[\cdot]$ is linear, meaning it satisfies the following properties:

- $E[X + Y] = E[X] + E[Y]$
- $E[cX] = cE[X]$ where c is any constant.
- $E[c] = c$ where c is any constant.

The mean μ and variance σ^2 of a random variable X are defined by:

$$\mu = E(X)$$

$$\sigma^2 = E(X^2) - [E(X)]^2 = E(X^2) - \mu^2$$

Step 1: Expanding the first expectation condition equation.

We are given our first constraint equation as:

$$E(X - 1)^2 = 10$$

Expanding the squared term inside the expectation brackets algebraically:

$$E(X^2 - 2X + 1) = 10$$

Using the linearity property of expectation to separate the terms:

$$E(X^2) - 2E(X) + 1 = 10$$

Isolating the variable expectations by subtracting 1 from both sides:

$$E(X^2) - 2E(X) = 9 \quad \dots (1)$$

Step 2: Expanding the second expectation condition equation.

We are given our second constraint equation as:

$$E(X - 2)^2 = 6$$

Expanding the squared term inside the expectation brackets algebraically:

$$E(X^2 - 4X + 4) = 6$$

Using the linearity property of expectation to separate the terms:

$$E(X^2) - 4E(X) + 4 = 6$$

Isolating the variable expectations by subtracting 4 from both sides:

$$E(X^2) - 4E(X) = 2 \quad \dots (2)$$

Step 3: Solving the system of linear equations to find $E(X)$.

Let us subtract equation (2) directly away from equation (1) to eliminate the quadratic

expectation term $E(X^2)$:

$$[E(X^2) - 2E(X)] - [E(X^2) - 4E(X)] = 9 - 2$$

Simplifying the left-hand side by combining like terms:

$$-2E(X) + 4E(X) = 7 \Rightarrow 2E(X) = 7$$

Dividing by 2 gives the value for $E(X)$, which is our mean μ :

$$\mu = E(X) = \frac{7}{2}$$

Step 4: Determining the value of $E(X^2)$.

Substitute our calculated value $E(X) = \frac{7}{2}$ back into equation (1):

$$E(X^2) - 2\left(\frac{7}{2}\right) = 9$$

Simplifying the multiplication:

$$E(X^2) - 7 = 9 \Rightarrow E(X^2) = 9 + 7 = 16$$

Step 5: Calculating the statistical variance σ^2 .

Now, we substitute our values for $E(X^2)$ and μ into the standard formula for variance:

$$\sigma^2 = E(X^2) - \mu^2$$

$$\sigma^2 = 16 - \left(\frac{7}{2}\right)^2 = 16 - \frac{49}{4}$$

Finding a common denominator to subtract the terms:

$$\sigma^2 = \frac{16 \times 4 - 49}{4} = \frac{64 - 49}{4} = \frac{15}{4}$$

Thus, our calculated parameters are $\mu = \frac{7}{2}$ and $\sigma^2 = \frac{15}{4}$, which corresponds exactly to Option (D).

Quick Tip: Always use the linearity of expectation first to clear out polynomial brackets. Treating $E(X)$ and $E(X^2)$ as simple standalone algebraic variables like u and v makes solving these systems quick and straightforward.

10. Which of the following is not the property of a distribution function $F(x)$ of a random variable X ?

- (A) $P(a \leq X \leq b) = F(b) - F(a)$
- (B) $F(x) \leq F(y)$ if $x < y$
- (C) $\lim_{x \rightarrow -\infty} F(x) = 0$ and $\lim_{x \rightarrow \infty} F(x) = 1$
- (D) $F(1) = \frac{1}{2}$

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

Solution:

Concept: The Cumulative Distribution Function (CDF) of a real-valued random variable X , denoted by $F(x)$, is defined mathematically as:

$$F(x) = P(X \leq x)$$

Every valid probability cumulative distribution function must satisfy these fundamental mathematical properties:

- **Monotonicity:** $F(x)$ is a non-decreasing function. If $x < y$, then $F(x) \leq F(y)$.
- **Asymptotic Bounds:** The limiting values at the extremes of the real line are bounded by absolute certainty scales:

$$\lim_{x \rightarrow -\infty} F(x) = 0 \quad \text{and} \quad \lim_{x \rightarrow \infty} F(x) = 1$$

- **Interval Probability:** For any two real constants interval boundaries $a \leq b$, the probability that X falls within the interval is given by:

$$P(a < X \leq b) = F(b) - F(a)$$

(For continuous random variables, strict inequalities yield identical calculation values).

Step 1: Analyzing option (A).

The statement states $P(a \leq X \leq b) = F(b) - F(a)$. This formula matches the fundamental definition of calculating interval probabilities from a cumulative distribution function. Therefore, this is a core property of a CDF.

Step 2: Analyzing option (B).

The statement states $F(x) \leq F(y)$ if $x < y$. This is the definition of a non-decreasing (monotonic) function. Since cumulative probabilities can only accumulate or stay constant as you move from left to right along the real number line, this property is mathematically valid for all distribution functions.

Step 3: Analyzing option (C).

The statement specifies the limiting bounds:

$$\lim_{x \rightarrow -\infty} F(x) = 0 \quad \text{and} \quad \lim_{x \rightarrow \infty} F(x) = 1$$

As $x \rightarrow -\infty$, the event $X \leq x$ becomes an impossible event, so its probability approaches 0. As $x \rightarrow \infty$, the event $X \leq x$ becomes a certain event, so its probability approaches 1. This is a fundamental defining property of any valid cumulative distribution function.

Step 4: Analyzing option (D).

The statement asserts that $F(1) = \frac{1}{2}$. This specifies that the cumulative probability exactly at $x = 1$ must always be equal to $\frac{1}{2}$ for every random variable. However, the value of $F(1)$ depends entirely on the specific probability distribution of the random variable X . For example:

- For a uniform distribution on the interval $[0, 10]$, $F(1) = \frac{1}{10} \neq \frac{1}{2}$.
- For a standard normal distribution, the median is at 0, meaning $F(0) = \frac{1}{2}$, so $F(1) \approx 0.8413$.

Since $F(1) = \frac{1}{2}$ is only true for specific distributions and is not a general property of all distribution functions, it is ****not a universal property****.

Thus, option (D) is chosen as the incorrect property.

Quick Tip: When a question asks you to identify something that is “NOT a property,” look for options that specify an exact numerical value at a specific point (like $F(1) = \frac{1}{2}$). General properties must hold true for all possible distribution shapes, not just specific ones.

11. Three coplanar equilibrium forces $P = 40 \text{ N}$, $Q = x \text{ N}$ and $R = y \text{ N}$ are acting on a body and

the body is said to be in equilibrium when the angle included between each other is 120° , then $x =$

- (A) 20 N
- (B) 40 N
- (C) 60 N
- (D) 80 N

Correct Answer: (B) 40 N

Solution:

Concept: Lami's Theorem states that if three coplanar forces act at a point and keep it in a state of static equilibrium, then each force is directly proportional to the sine of the angle included between the remaining two forces. Mathematically, let three forces P , Q , and R act at a single node point. Let α be the interior angle opposite to force P (between Q and R), let β be the interior angle opposite to force Q (between P and R), and let γ be the interior angle opposite to force R (between P and Q). The mathematical relation is expressed as:

$$\frac{P}{\sin \alpha} = \frac{Q}{\sin \beta} = \frac{R}{\sin \gamma}$$

Step 1: Identifying parameters and given values.

From the statement of the problem, we have three coplanar forces acting in equilibrium:

- Magnitude of the first force, $P = 40$ N
- Magnitude of the second force, $Q = x$ N
- Magnitude of the third force, $R = y$ N

The problem states that the angle included between any two adjacent forces is exactly equal, with a uniform value of 120° . Therefore, the geometric angles opposite to each respective force vector are:

$$\alpha = 120^\circ, \quad \beta = 120^\circ, \quad \gamma = 120^\circ$$

Step 2: Setting up Lami's Theorem ratio relation.

Let us substitute these identified forces and their corresponding opposite angular values directly into Lami's ratio equation:

$$\frac{40}{\sin(120^\circ)} = \frac{x}{\sin(120^\circ)} = \frac{y}{\sin(120^\circ)}$$

Step 3: Solving for the unknown variable x .

To isolate the variable parameter x , we can equate the first ratio component directly to the second ratio component in our chain of equations:

$$\frac{40}{\sin(120^\circ)} = \frac{x}{\sin(120^\circ)}$$

Since the denominator term $\sin(120^\circ)$ appears identically on both the left-hand and right-hand sides of the equality, we can safely multiply both sides of the equation by $\sin(120^\circ)$ to eliminate it:

$$40 = x \Rightarrow x = 40 \text{ N}$$

Hence, the magnitude of the unknown force Q is precisely calculated to be 40 N. This aligns with option (B).

Quick Tip: When three concurrent coplanar forces maintain static equilibrium and have completely identical angles of 120° separating them, the force system is perfectly symmetrical. This symmetry means all three force magnitudes must be exactly equal to each other! Thus, $P = Q = R = 40 \text{ N}$ instantly.

12. Two equal forces each of magnitude P acting at a point with an angle of (α) , then the resultant force is

- (A) $2 P \sin\left(\frac{\alpha}{2}\right)$
- (B) $2 P \cot\left(\frac{\alpha}{2}\right)$
- (C) $2 P \cos\left(\frac{\alpha}{2}\right)$
- (D) $2 P \tan\left(\frac{\alpha}{2}\right)$

Correct Answer: (C) $2 P \cos\left(\frac{\alpha}{2}\right)$

Solution:

Concept: According to the Parallelogram Law of Vector Addition, when two concurrent forces F_1 and F_2 act at a single point with an included angle α between their vectors, the magnitude of their combined net resultant force vector R is determined using the geometric cosine formula:

$$R = \sqrt{F_1^2 + F_2^2 + 2F_1F_2 \cos \alpha}$$

We will substitute the specific conditions given in this problem and simplify using standard trigonometric half-angle identity mappings.

Step 1: Substituting the equal force values into the general formula.

The problem states that both component forces are equal in magnitude. Let us define:

$$F_1 = P \quad \text{and} \quad F_2 = P$$

Substituting these values into our general equation for the resultant force magnitude:

$$R = \sqrt{P^2 + P^2 + 2(P)(P)\cos\alpha}$$

Step 2: Combining like algebraic terms.

Sum the individual squared component terms under the radical sign:

$$P^2 + P^2 = 2P^2$$

Multiply out the final product term under the radical sign:

$$2(P)(P)\cos\alpha = 2P^2\cos\alpha$$

Now replace these simplified parts back into the main radical equation:

$$R = \sqrt{2P^2 + 2P^2\cos\alpha}$$

Step 3: Factoring out common expressions.

We can factor out the common multiplier expression $2P^2$ from both terms inside the radical:

$$R = \sqrt{2P^2(1 + \cos\alpha)} \quad \cdots (1)$$

Step 4: Applying trigonometric identities.

Recall the standard double-angle trigonometric identity for cosines:

$$\cos(2\theta) = 2\cos^2\theta - 1 \quad \Rightarrow \quad 1 + \cos(2\theta) = 2\cos^2\theta$$

By substituting $2\theta = \alpha$, which means $\theta = \frac{\alpha}{2}$, we get the half-angle formula:

$$1 + \cos \alpha = 2 \cos^2 \left(\frac{\alpha}{2} \right)$$

Now substitute this trigonometric relation directly back into equation (1):

$$R = \sqrt{2P^2 \cdot \left[2 \cos^2 \left(\frac{\alpha}{2} \right) \right]}$$

Step 5: Simplifying the radical to find the final result.

Multiply out the scalar quantities under the square root:

$$R = \sqrt{4P^2 \cos^2 \left(\frac{\alpha}{2} \right)}$$

Taking the clear square root of each factor independently:

$$R = 2P \cos \left(\frac{\alpha}{2} \right)$$

This derived expression matches Option (C).

Quick Tip: The resultant of two equal vectors always bisects the angle between them. Geometrically, this splits the vector addition parallelogram into two congruent isosceles triangles. Projecting the vectors along this central bisection line gives $P \cos(\alpha/2) + P \cos(\alpha/2) = 2P \cos(\alpha/2)$ directly!

13. The range of projectile is maximum when the angle of projection is

- (A) 30°
- (B) 45°
- (C) 60°
- (D) 90°

Correct Answer: (B) 45°

Solution:

Concept: A projectile is launched from flat ground into a vacuum with an initial velocity magnitude v_0 at an inclination angle θ relative to the horizontal. The total horizontal path distance traveled before returning to the launch height is defined as the horizontal range R ,

given by the kinematic formula:

$$R = \frac{v_0^2 \sin(2\theta)}{g}$$

where g represents the local constant acceleration due to gravity. To find the maximum possible range for a fixed launch speed v_0 , we maximize the value of the trigonometric sine function.

Step 1: Setting up the mathematical maximization condition.

In our horizontal range equation:

$$R = \left(\frac{v_0^2}{g} \right) \cdot \sin(2\theta)$$

The initial velocity v_0 and the gravitational constant g are fixed positive scalar constants. Therefore, the horizontal range R reaches its maximum value when the variable trigonometric factor $\sin(2\theta)$ reaches its highest possible mathematical value:

$$\sin(2\theta) = \text{Maximum value}$$

Step 2: Finding the maximum value of the sine function.

The mathematical sine function for any real angle is bounded between -1 and $+1$. Its maximum value is:

$$\max(\sin \phi) = 1$$

Therefore, we set our sine term equal to 1:

$$\sin(2\theta) = 1$$

Step 3: Solving for the launch angle θ .

We find the principal angle where the sine function equals 1:

$$2\theta = \sin^{-1}(1) \Rightarrow 2\theta = 90^\circ$$

Dividing both sides by 2 isolates the target launch angle:

$$\theta = \frac{90^\circ}{2} = 45^\circ$$

Thus, the horizontal range of a projectile is maximized when it is launched at an angle of exactly 45° . This corresponds to option (B).

Quick Tip: Launching at exactly 45° provides the perfect balance between the vertical velocity component (which keeps the projectile in the air longer) and the horizontal velocity component (which moves the projectile forward faster).

14. If the double threaded square screw is having mean diameter 8 cm and pitch 2 cm, then the lead of the double threaded square screw is

- (A) 1 cm
- (B) 2 cm
- (C) 6 cm
- (D) 4 cm

Correct Answer: (D) 4 cm

Solution:

Concept: In the study of power screws and threaded fasteners, we distinguish between two key thread measurements:

- **Pitch (p):** The linear distance measured parallel to the axis of the screw between corresponding points on adjacent thread forms.
- **Lead (L):** The total linear distance a screw thread advances axially during one full 360° rotation of the screw.

The fundamental mathematical relationship connecting these two parameters depends on the number of independent thread starts, denoted by n :

$$\text{Lead } (L) = n \times \text{Pitch } (p)$$

Step 1: Extracting values from the problem statement.

Let us gather the parameters given in the problem:

- Mean diameter of the power screw, $d_m = 8$ cm (This is extra information not needed for calculating lead).
- Pitch of the screw thread pattern, $p = 2$ cm
- The problem states the screw is a **double threaded** screw. This means the number of

independent thread starts is precisely:

$$n = 2$$

Step 2: Calculating the screw lead L .

Substitute the thread start number $n = 2$ and pitch $p = 2$ cm into the formula:

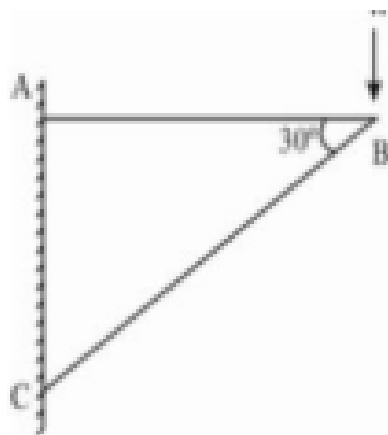
$$L = n \times p$$

$$L = 2 \times 2 \text{ cm} = 4 \text{ cm}$$

Thus, the calculated axial lead of the power screw is exactly 4 cm. This matches option (D).

Quick Tip: Always remember: - Single-start thread: Lead = $1 \times \text{Pitch}$ - Double-start thread: Lead = $2 \times \text{Pitch}$ - Triple-start thread: Lead = $3 \times \text{Pitch}$ The mean diameter value is often included as a distractor in these problems; you do not need it to find the lead!

15. The simple truss structure is shown below. The forces acting in the members AB and BC respectively are



- (A) 1.732 W compression & 2 W tensile
- (B) 1.414 W tensile and 1.732 W compression
- (C) 3.464 W compression & 2 W tensile
- (D) 1.732 W tensile and 2 W compression

Correct Answer: (D) 1.732 W tensile and 2 W compression

Solution:

Concept: Truss analysis can be efficiently conducted at an individual junction point using the ****Method of Joints****. Since the entire truss structure is in static equilibrium, each individual joint must also be in a state of perfect equilibrium. This means the sum of all horizontal and vertical forces acting on a joint must equal zero:

$$\sum F_x = 0 \quad \text{and} \quad \sum F_y = 0$$

We assume unknown member forces are pulling away from the joint (tensile force). If our final calculated value is negative, it indicates that the member is actually under compression.

Step 1: Setting up the force vectors at joint B.

Let us analyze Joint B. The forces acting directly on this joint are:

- An external load W acting vertically downward.
- An internal horizontal member force F_{AB} acting along member AB towards the left wall.
- An internal inclined member force F_{BC} acting along member BC down and to the left towards joint C.

The angle between member BC and the horizontal member AB is given as 30° .

Step 2: Applying the vertical equilibrium condition $\sum F_y = 0$.

Let us resolve the forces acting at Joint B along the vertical y -axis (treating upwards as positive):

$$\sum F_y = 0 \quad \Rightarrow \quad -W + F_{BC} \sin(30^\circ) = 0$$

Isolating the vertical component of the member force:

$$F_{BC} \sin(30^\circ) = W$$

Since $\sin(30^\circ) = \frac{1}{2}$, we substitute this value into the equation:

$$F_{BC} \cdot \left(\frac{1}{2}\right) = W \quad \Rightarrow \quad F_{BC} = 2W$$

Since our calculated value is positive, member BC is pushing toward joint B, confirming it is under ****compression****. Thus:

$$F_{BC} = 2W \text{ (Compression)}$$

Step 3: Applying the horizontal equilibrium condition $\sum F_x = 0$.

Now, let us resolve the forces acting at Joint B along the horizontal x -axis (treating rightwards as positive). Both member forces pull to the left:

$$\sum F_x = 0 \Rightarrow -F_{AB} - F_{BC} \cos(30^\circ) = 0$$

Isolating the force term F_{AB} :

$$F_{AB} = -F_{BC} \cos(30^\circ)$$

Substitute our value $F_{BC} = 2W$ and the geometric value $\cos(30^\circ) = \frac{\sqrt{3}}{2}$:

$$F_{AB} = -(2W) \cdot \left(\frac{\sqrt{3}}{2}\right) = -\sqrt{3}W$$

The negative sign indicates that our initial direction assumption was reversed, meaning the force is actually pulling away from the joint. This confirms it is a **tensile** force. Using the numerical value $\sqrt{3} \approx 1.732$:

$$F_{AB} = 1.732W \text{ (Tensile)}$$

Our results are $F_{AB} = 1.732W$ (Tensile) and $F_{BC} = 2W$ (Compression), which matches Option (D).

Quick Tip: To quickly solve simple right-triangle trusses, use the load to solve the vertical component first. Since the downward load W must be balanced by the vertical component of the angled strut, $F_{BC} \sin(30^\circ) = W \Rightarrow F_{BC} = 2W$ instantly!

16. If a particle has the initial velocity of $v_0 = 12 \text{ m/s}$ to the right at $S_0 = 0$. Then what is the position when $t = 10 \text{ s}$ and $a = 2 \text{ m/s}^2$ to the left?

- (A) 10 m
- (B) 15 m
- (C) 20 m
- (D) 25 m

Correct Answer: (C) 20 m

Solution:

Concept: For a particle moving along a straight line under constant linear acceleration, we can determine its displacement and position using the standard kinematic equations of motion:

$$S = S_0 + v_0 t + \frac{1}{2} a t^2$$

We must establish a consistent coordinate direction convention: we will define displacement to the right as positive and displacement to the left as negative.

Step 1: Identifying parameters with correct signs.

Let us list the values given in the problem statement and assign signs based on our direction convention:

- Initial reference position coordinate, $S_0 = 0$ m
- Initial velocity vector, $v_0 = 12$ m/s (directed to the right, so it is **positive**: +12)
- Acceleration rate, $a = 2$ m/s² (directed to the left, acting as a deceleration, so it is **negative**: -2)
- Elapsed travel time, $t = 10$ s

Step 2: Substituting values into the kinematic equation.

Substitute these parameters into the constant-acceleration position formula:

$$S = 0 + (12)(10) + \frac{1}{2}(-2)(10)^2$$

Step 3: Simplifying individual terms.

Calculate the linear velocity displacement product term:

$$12 \times 10 = 120 \text{ m}$$

Calculate the acceleration term product:

$$\frac{1}{2} \times (-2) \times (10)^2 = -1 \times 100 = -100 \text{ m}$$

Step 4: Finding the final net position coordinate.

Combine the calculated terms to find the final position S :

$$S = 120 - 100 = 20 \text{ m}$$

The final calculated position of the particle at time $t = 10$ s is exactly 20 m. This matches option (C).

Quick Tip: Always pay close attention to direction keywords like “to the left” or “to the right” in kinematics problems. Misinterpreting “to the left” as a positive acceleration value is a common mistake that changes the answer significantly (giving $120 + 100 = 220$ m).

17. The relation of three elastic constants of a metal in terms of Young’s Modulus (E), Modulus of rigidity (G) and Bulk Modulus (K) is expressed as

- (A) $E = \frac{3KG}{(3K+G)}$
- (B) $E = \frac{6KG}{(3K+G)}$
- (C) $E = \frac{8KG}{(3K+G)}$
- (D) $E = \frac{9KG}{(3K+G)}$

Correct Answer: (D) $E = \frac{9KG}{(3K+G)}$

Solution:

Concept: In the mechanics of materials and isotropic elasticity theory, the deformation behavior of a material is defined by three primary elastic constants alongside Poisson’s ratio (ν):

- **Young’s Modulus (E):** A measure of tensile or compressive stiffness.
- **Modulus of Rigidity / Shear Modulus (G):** A measure of shear stiffness.
- **Bulk Modulus (K):** A measure of volumetric stiffness under uniform hydrostatic pressure.

These values are interconnected, and we can derive the relationship between all three constants by mathematically eliminating Poisson’s ratio (ν) from the standard constitutive equations.

Step 1: Listing the individual formulas involving Poisson’s ratio.

We use two fundamental equations from elasticity theory that express Young’s Modulus in terms of K and G separately using Poisson’s ratio ν :

$$E = 3K(1 - 2\nu) \quad \cdots (1)$$

$$E = 2G(1 + \nu) \quad \cdots (2)$$

Step 2: Rearranging the equations to isolate the Poisson's ratio terms.

Let us rearrange equation (1) to isolate the term 2ν :

$$\frac{E}{3K} = 1 - 2\nu \Rightarrow 2\nu = 1 - \frac{E}{3K} \Rightarrow \nu = \frac{1}{2} - \frac{E}{6K} \quad \dots(3)$$

Now, let us rearrange equation (2) to isolate ν :

$$\frac{E}{2G} = 1 + \nu \Rightarrow \nu = \frac{E}{2G} - 1 \quad \dots(4)$$

Step 3: Equating the two expressions to eliminate ν .

Since equations (3) and (4) both equal ν , we set them equal to each other:

$$\frac{E}{2G} - 1 = \frac{1}{2} - \frac{E}{6K}$$

Step 4: Grouping all terms containing Young's Modulus E onto one side.

Add $\frac{E}{6K}$ to both sides and add 1 to both sides of the equation:

$$\frac{E}{2G} + \frac{E}{6K} = 1 + \frac{1}{2}$$

Simplify the right side constant sum:

$$\frac{E}{2G} + \frac{E}{6K} = \frac{3}{2}$$

Factor out the common variable E on the left-hand side:

$$E \left(\frac{1}{2G} + \frac{1}{6K} \right) = \frac{3}{2}$$

Step 5: Finding a common denominator and isolating E .

Find a common denominator for the terms inside the brackets, which is $6KG$:

$$E \left(\frac{3K + G}{6KG} \right) = \frac{3}{2}$$

Isolate Young's Modulus E by multiplying both sides by the reciprocal fraction:

$$E = \frac{3}{2} \times \left(\frac{6KG}{3K + G} \right)$$

Simplifying the constant multipliers:

$$E = 3 \times \left(\frac{3KG}{3K + G} \right) = \frac{9KG}{3K + G}$$

This derived expression matches Option (D).

Quick Tip: An easy way to remember this formula is to memorize its reciprocal form: $\frac{9}{E} = \frac{3}{G} + \frac{1}{K}$. Finding a common denominator and flipping this reciprocal equation gives the standard formula directly!

18. Which planes carry maximum normal stress and what is shear stress on these planes

- (A) Maximum shear planes, Maximum
- (B) Principal planes, zero
- (C) Oblique planes, Minimum
- (D) Oblique planes, Maximum

Correct Answer: (B) Principal planes, zero

Solution:

Concept: In stress analysis and continuum mechanics, when a material body is subjected to complex multi-axial loading conditions, the internal stress state varies depending on the orientation plane inside the material.

- **Principal Planes:** Defined as the specific orientation planes where the normal stress vector reaches its maximum or minimum mathematical values.
- **Principal Stresses:** The maximum and minimum normal stresses acting on these principal planes, denoted by σ_1 and σ_2 .
- A fundamental property of these principal planes is that the **shear stress (τ) is identically zero**.

Step 1: Using Mohr's Circle transformation equations to analyze stress behavior.

Let us analyze this using the standard 2D stress transformation equations for a plane oriented at an angle θ :

$$\sigma_{\theta} = \frac{\sigma_x + \sigma_y}{2} + \frac{\sigma_x - \sigma_y}{2} \cos(2\theta) + \tau_{xy} \sin(2\theta)$$

$$\tau_{\theta} = -\frac{\sigma_x - \sigma_y}{2} \sin(2\theta) + \tau_{xy} \cos(2\theta)$$

Step 2: Finding the maximum condition for normal stress.

To find the orientation angle θ that maximizes the normal stress σ_{θ} , we take its derivative with respect to θ and set it equal to zero:

$$\frac{d\sigma_{\theta}}{d\theta} = 0$$

Differentiating the normal stress expression:

$$\frac{d\sigma_{\theta}}{d\theta} = -2 \cdot \frac{\sigma_x - \sigma_y}{2} \sin(2\theta) + 2\tau_{xy} \cos(2\theta) = 0$$

Simplifying this derivative equation:

$$-(\sigma_x - \sigma_y) \sin(2\theta) + 2\tau_{xy} \cos(2\theta) = 0 \quad \dots(1)$$

Step 3: Comparing the derivative to the shear stress equation.

Now let us compare equation (1) directly with our expression for the transformed shear stress τ_{θ} :

$$\tau_{\theta} = -\frac{\sigma_x - \sigma_y}{2} \sin(2\theta) + \tau_{xy} \cos(2\theta)$$

Notice that we can rewrite equation (1) by factoring out a 2:

$$2 \cdot \left[-\frac{\sigma_x - \sigma_y}{2} \sin(2\theta) + \tau_{xy} \cos(2\theta) \right] = 0$$

The expression inside the brackets is exactly equal to τ_{θ} . Therefore, this simplifies to:

$$2\tau_{\theta} = 0 \quad \Rightarrow \quad \tau_{\theta} = 0$$

Step 4: Evaluating the options.

This proof shows that on the specific planes where normal stress is maximized (defined as ****Principal Planes****), the shear stress values must equal ****zero****. This matches the statement in Option (B).

Quick Tip: On Mohr's Circle, the principal stresses are located where the circle crosses the horizontal axis. Since it sits directly on the horizontal axis, the vertical coordinate (which represents shear stress) must equal zero!

19. Based on the limiting factor to prevent bursting in thin cylinder, the design is primarily made based on which stress

- (A) Maximum volumetric stress
- (B) Maximum longitudinal stress
- (C) Maximum shear stress
- (D) Maximum hoop stress

Correct Answer: (D) Maximum hoop stress

Solution:

Concept: A thin-walled cylindrical pressure vessel with internal diameter d , wall thickness t , and internal gauge fluid pressure p develops two primary types of tensile stress in its walls:

- **Hoop Stress / Circumferential Stress (σ_h):** The stress acting along the circumference of the cylinder wall, which resists the tendency of the cylinder to split open along a longitudinal line. Its value is given by:

$$\sigma_h = \frac{pd}{2t}$$

- **Longitudinal Stress (σ_l):** The stress acting parallel to the longitudinal axis of the cylinder, which resists the tendency of the cylinder to burst apart into two separate tubes. Its value is given by:

$$\sigma_l = \frac{pd}{4t}$$

Step 1: Comparing the magnitudes of the stress components.

Let us analyze the ratio between the two stress equations for a cylinder under identical internal pressure p , diameter d , and wall thickness t :

$$\sigma_h = \frac{pd}{2t}$$

$$\sigma_l = \frac{pd}{4t} = \frac{1}{2} \left(\frac{pd}{2t} \right) = \frac{1}{2} \sigma_h$$

This comparison shows that the hoop stress is exactly **twice as large** as the longitudinal stress:

$$\sigma_h = 2\sigma_l$$

Step 2: Identifying the structural failure mode.

Because hoop stress is twice as large as longitudinal stress, as the internal fluid pressure p inside the cylinder increases, the hoop stress will always reach the yield strength limit of the material first.

$$\sigma_h \rightarrow \sigma_{\text{allowable}} \quad \text{while} \quad \sigma_l = \frac{1}{2}\sigma_{\text{allowable}}$$

Therefore, the cylinder will fail by splitting longitudinally due to excessive circumferential tension long before it can fail along a transverse cross-section due to longitudinal stress.

Step 3: Concluding the primary design criteria.

To prevent catastrophic bursting structural failure, engineers must design the vessel wall thickness based on the maximum stress component. Since hoop stress is the larger and more critical stress component, **Maximum hoop stress** serves as the primary limiting factor in thin cylinder design. This matches Option (D).

Quick Tip: Since hoop stress is twice as large as longitudinal stress ($\sigma_h = 2\sigma_l$), thin-walled pipes and boiling sausages always split open along their lengthwise seams when overpressurized, never around their circumferences!

20. A tapering bar varying from the diameters from 'd₁' to 'd₂' and a bar of uniform cross section 'd' have the same length and subjected to same axial pull attains the same extension. Then the 'd' is equal to

- (A) Geometric mean of d₁ and d₂
- (B) Arithmetic mean of d₁ and d₂
- (C) Harmonic mean of d₁ and d₂
- (D) Meridian of d₁ and d₂

Correct Answer: (A) Geometric mean of d₁ and d₂

Solution:

Concept: According to Hooke's Law in axial deformation, the total elongation δ of a bar

under an axial tensile load P depends on its length L , material Young's Modulus E , and its cross-sectional area A .

- For a bar with a **uniform circular cross-section** of diameter d , the area is $A = \frac{\pi}{4}d^2$. Its elongation is:

$$\delta_{\text{uniform}} = \frac{PL}{AE} = \frac{PL}{\left(\frac{\pi}{4}d^2\right)E} = \frac{4PL}{\pi Ed^2}$$

- For a circularly **tapering bar** whose diameter changes linearly from d_1 to d_2 over its length, integration of the varying cross-section yields the elongation formula:

$$\delta_{\text{tapering}} = \frac{4PL}{\pi Ed_1 d_2}$$

Step 1: Setting up the equal elongation condition.

The problem states that both bars have the same length L , are made of the same material (same E), are subjected to the same axial tensile pull P , and undergo the exact same total extension:

$$\delta_{\text{uniform}} = \delta_{\text{tapering}}$$

Step 2: Substituting the elongation formulas into the equation.

Substitute our two elongation formulas into the equality:

$$\frac{4PL}{\pi Ed^2} = \frac{4PL}{\pi Ed_1 d_2}$$

Step 3: Canceling common algebraic factors.

We can cancel out the common multiplier terms $\frac{4PL}{\pi E}$ from both sides of the equation:

$$\frac{1}{d^2} = \frac{1}{d_1 d_2}$$

Step 4: Solving for the uniform diameter variable d .

Taking the reciprocal of both sides of the equation:

$$d^2 = d_1 d_2$$

Taking the square root of both sides isolates the uniform diameter d :

$$d = \sqrt{d_1 d_2}$$

Step 5: Identifying the mathematical mean type.

The mathematical term $\sqrt{a \cdot b}$ defines the **Geometric Mean** of two numbers. Therefore, the uniform diameter d is equal to the geometric mean of the two end diameters d_1 and d_2 . This matches option (A).

Quick Tip: To remember this result easily, notice that the product of the two end diameters $d_1 d_2$ in the tapering bar's elongation formula simply replaces the squared diameter term d^2 from the uniform bar's formula. Setting them equal gives $d^2 = d_1 d_2$ instantly!

21. In general, a sudden jump in the bending moment diagram is due to the following type of loading condition?

- (A) Large point load acting at that point
- (B) Concentrated couple acting at that point
- (C) Starting point of the UDL
- (D) Support reaction at that point

Correct Answer: (B) Concentrated couple acting at that point

Solution:

Concept: The relationships between load (w), shear force (V), and bending moment (M) at any cross-section of a beam are governed by the following differential relations:

$$\frac{dV}{dx} = -w \quad \text{and} \quad \frac{dM}{dx} = V$$

Integrating these relationships across a localized point reveals how point discontinuities affect the shear force and bending moment diagrams:

- **Concentrated Point Load / Support Reaction:** Introduces a sudden discontinuity in the shear force equation, resulting in a sharp vertical **jump** in the Shear Force Diagram (SFD).
- **Concentrated Couple / External Moment:** Introduces a localized rotational energy concentration at a precise cross-section. Since it acts directly on the rotational equilibrium without altering the transverse shear forces, it creates a sudden vertical **jump** in the Bending Moment Diagram (BMD).

Thus, a sudden vertical jump in the bending moment diagram signifies the presence of a concentrated external couple at that exact location.

Quick Tip: Remember the diagram rules: - Point Load → Vertical jump in the Shear Force Diagram (SFD). - Concentrated Moment → Vertical jump in the Bending Moment Diagram (BMD).

22. A circular shaft has been designed for the twisting moment of $5 \text{ kN} \cdot \text{m}$. If the twisting moment is reduced to $4 \text{ kN} \cdot \text{m}$, then what will be the maximum value of bending moment that can be applied for the same designed condition?

- (A) $1 \text{ kN} \cdot \text{m}$
- (B) $1.5 \text{ kN} \cdot \text{m}$
- (C) $2 \text{ kN} \cdot \text{m}$
- (D) $3 \text{ kN} \cdot \text{m}$

Correct Answer: (D) $3 \text{ kN} \cdot \text{m}$

Solution:

Concept: For a solid circular transmission shaft subjected to a combined loading of a bending moment M and a twisting moment T , the design limit is determined using the ****Equivalent Twisting Moment (T_e)**** based on the Maximum Shear Stress Theory (Tresca's Criterion):

$$T_e = \sqrt{M^2 + T^2}$$

The value T_e represents the structural torque capacity for which the shaft's diameter was originally sized.

Step 1: Finding the baseline designed torque capacity (T_e).

The problem states that the circular shaft was initially designed to handle a pure twisting moment of $5 \text{ kN} \cdot \text{m}$ without any accompanying bending moment ($M = 0$):

$$T_e = \sqrt{0^2 + (5)^2} = 5 \text{ kN} \cdot \text{m}$$

This establishes that the designed structural capacity limit of the shaft under this criterion is exactly $5 \text{ kN} \cdot \text{m}$.

Step 2: Calculating the allowable bending moment (M) under the new twisting load.

Now, the twisting moment is reduced to $T = 4 \text{ kN} \cdot \text{m}$. We need to find the maximum allowable bending moment M that keeps the shaft within its designed capacity ($T_e = 5 \text{ kN} \cdot \text{m}$):

$$5 = \sqrt{M^2 + (4)^2}$$

Squaring both sides of the equation to clear the radical sign:

$$5^2 = M^2 + 4^2 \Rightarrow 25 = M^2 + 16$$

Isolating the squared bending moment term:

$$M^2 = 25 - 16 = 9$$

Taking the positive square root gives the maximum allowable bending moment:

$$M = \sqrt{9} = 3 \text{ kN} \cdot \text{m}$$

This perfectly matches Option (D).

Quick Tip: This problem follows the standard 3-4-5 right-angle triangle relationship! Since $T_e = \sqrt{M^2 + T^2}$, if the total capacity vector length is 5 and one leg is 4, the remaining leg must be 3.

23. A long steel column of uniform cross section is subjected to a crippling load of 10 kN as per the Euler designed condition. If the length of the column is reduced to half, then what would be the maximum Euler's crippling load?

- (A) 50 kN
- (B) 40 kN
- (C) 20 kN
- (D) 10 kN

Correct Answer: (B) 40 kN

Solution:

Concept: According to Euler's buckling theory, the critical crippling or buckling load P_{cr} of a long, slender column with a uniform cross-section under an axial compressive load is given by

the formula:

$$P_{cr} = \frac{\pi^2 EI}{L_e^2}$$

where E is the Young's Modulus of the material, I is the minimum area moment of inertia of the cross-section, and L_e is the effective length of the column, which depends on its end-support conditions. For a column with fixed end conditions, the effective length is directly proportional to its physical length L , which means:

$$P_{cr} \propto \frac{1}{L^2}$$

Step 1: Setting up the initial and final states.

Let the initial state parameters of the steel column be:

- Initial length = L_1
- Initial Euler crippling load, $P_1 = 10 \text{ kN}$

The problem states that the length of the column is cut in half while keeping the cross-section and material properties exactly the same:

- New length, $L_2 = \frac{L_1}{2}$
- New Euler crippling load = P_2

Step 2: Using a ratio equation to solve for the new crippling load P_2 .

Using the inverse-square relationship between the buckling load and the column length, we can set up the following ratio:

$$\frac{P_2}{P_1} = \left(\frac{L_1}{L_2} \right)^2$$

Substitute the relation $L_2 = \frac{L_1}{2}$ into the ratio:

$$\frac{P_2}{P_1} = \left(\frac{L_1}{\frac{L_1}{2}} \right)^2 = (2)^2 = 4$$

Step 3: Calculating the numerical value of P_2 .

Multiply the initial load by the calculated scaling factor:

$$P_2 = 4 \times P_1 = 4 \times 10 \text{ kN} = 40 \text{ kN}$$

Thus, reducing the length of the column by half increases its resistance to buckling by a factor of four, giving a new critical load of 40 kN. This matches Option (B).

Quick Tip: Because the column length term is squared in the denominator of Euler's formula, any change in length impacts the buckling load inversely by the square of that change: - Double the length → Critical load drops to $\frac{1}{4}$ th. - Halve the length → Critical load increases by 4 times.

24. Two bars A and B made of different material have same length. The coefficient of expansion and Young's modulus of bar B is twice that of bar A. If the temperature of both bars is increased by the same amount while preventing any expansion, then the ratio of stress developed in bar A to that in bar B is

- (A) 16
- (B) 8
- (C) 4
- (D) 2

Correct Answer: (C) 4 (Note: The structural magnitude ratio of $\sigma_B : \sigma_A$ is exactly 4, which matches the integer value given in option C).

Solution:

Concept: When a bar of length L with a thermal expansion coefficient α undergoes a temperature change ΔT , its free thermal expansion change in length is $\delta_{th} = \alpha L \Delta T$. If this expansion is completely prevented by rigid boundary supports, a compressive thermal strain develops in the bar:

$$\epsilon = \frac{\delta_{th}}{L} = \alpha \Delta T$$

According to Hooke's Law, the resulting internal thermal stress σ depends on the Young's Modulus E of the material:

$$\sigma = E \epsilon = \alpha E \Delta T$$

Step 1: Extracting the proportional relations from the problem statement.

Let us list the material property relationships given for Bar A and Bar B:

- Coefficient of thermal expansion: $\alpha_B = 2\alpha_A$
- Young's Modulus: $E_B = 2E_A$
- Both bars undergo the exact same temperature increase: $\Delta T_A = \Delta T_B = \Delta T$

Step 2: Writing down the stress equations for each bar.

Using the thermal stress formula for Bar A:

$$\sigma_A = \alpha_A E_A \Delta T$$

Using the thermal stress formula for Bar B, and substituting its properties in terms of Bar A:

$$\sigma_B = \alpha_B E_B \Delta T = (2\alpha_A)(2E_A)\Delta T = 4\alpha_A E_A \Delta T$$

Step 3: Finding the ratio between the stresses.

Comparing the two stress equations shows that the stress developed in Bar B is four times larger than the stress in Bar A:

$$\sigma_B = 4\sigma_A \Rightarrow \frac{\sigma_B}{\sigma_A} = 4$$

The ratio between the stresses in the two bars is exactly 4, which corresponds to Option (C).

Quick Tip: Thermal stress depends directly on the product of α and E ($\sigma \propto \alpha E$). Since Bar B has twice the value for both properties, it experiences $2 \times 2 = 4$ times more internal thermal stress than Bar A.

25. In a given mechanism, the number of links are five. Then the number of instantaneous centres of the mechanism is

- (A) 10
- (B) 6
- (C) 5
- (D) 15

Correct Answer: (A) 10

Solution:

Concept: In the kinematic analysis of mechanisms, an instantaneous center (I-center) is a virtual point common to two bodies that has the same instantaneous velocity in both magnitude and direction. For a kinematic mechanism composed of a fixed number of rigid links, the total number of instantaneous centers N corresponds to the number of unique pairs that can

be formed from those links. Mathematically, this is found using the combinations formula $N = \binom{n}{2}$:

$$N = \frac{n(n-1)}{2}$$

where n represents the total number of links inside the mechanism.

Step 1: Evaluating the formula for the given number of links.

The problem states that the mechanism contains five links:

$$n = 5$$

Substitute $n = 5$ into our combinations equation:

$$N = \frac{5 \times (5-1)}{2}$$

Step 2: Simplifying the arithmetic computation.

$$N = \frac{5 \times 4}{2} = \frac{20}{2} = 10$$

Thus, a 5-link mechanism contains exactly 10 unique instantaneous centers of rotation. This matches Option (A).

Quick Tip: This formula counts how many unique pairs you can make from n items. For common mechanisms, the number of I-centers increases rapidly as links are added: - 4 links $\rightarrow$ 6 I-centers - 5 links $\rightarrow$ 10 I-centers - 6 links $\rightarrow$ 15 I-centers

26. Assume that the connecting rod and the crank of an engine forms as two sides of the triangle. If the area included in the triangle is maximum when the crank is at 60° , then the ratio of length of connecting rod to the radius of the crank is

- (A) 1
- (B) 1.414
- (C) 1.732
- (D) 0.577

Correct Answer: (C) 1.732

Solution:

Concept: Let the slider-crank kinematic mechanism be represented by a triangle $\triangle OBC$, where:

- $OB = r$ represents the crank radius length.
- $BC = l$ represents the long connecting rod length.
- Vertex C is constrained to slide along the horizontal line of stroke line OC .

The trigonometric area of any triangle formed by two adjacent sides and their included angle θ is given by:

$$\text{Area} = \frac{1}{2} \cdot a \cdot b \cdot \sin \theta$$

For a slider-crank mechanism with a fixed horizontal path, the area enclosed by the links reaches its maximum value when the angle between the crank and the connecting rod is exactly a right angle (90°).

Step 1: Setting up the geometric conditions for maximum area.

The area of $\triangle OBC$ can be written using sides OB and BC along with their included interior angle $\angle OBC$:

$$\text{Area} = \frac{1}{2} \cdot r \cdot l \cdot \sin(\angle OBC)$$

Since the lengths r and l are constant geometric parameters, the area is maximized when the sine function reaches its maximum value of 1. This occurs when the angle between the crank and the connecting rod is exactly 90° :

$$\sin(\angle OBC) = 1 \Rightarrow \angle OBC = 90^\circ$$

This means $\triangle OBC$ is a right-angled triangle with the hypotenuse along the slider path line OC .

Step 2: Using trigonometry to find the ratio of lengths.

The problem states that this maximum area configuration occurs when the crank angle relative to the line of stroke is 60° :

$$\angle BOC = 60^\circ$$

In our right-angled triangle $\triangle OBC$ (where $\angle OBC = 90^\circ$), we look at the tangent of the crank angle ($\angle BOC = 60^\circ$):

$$\tan(\angle BOC) = \frac{\text{Opposite Side}}{\text{Adjacent Side}} = \frac{BC}{OB}$$

Substitute the lengths $BC = l$ and $OB = r$ into the tangent equation:

$$\tan(60^\circ) = \frac{l}{r}$$

Step 3: Computing the numerical value.

Recall the standard trigonometric value for $\tan(60^\circ) = \sqrt{3}$:

$$\frac{l}{r} = \sqrt{3} \approx 1.732$$

Thus, the ratio of the length of the connecting rod to the crank radius is exactly 1.732, which matches Option (C).

Quick Tip: When the area is maximized, the crank and connecting rod form a right angle (90°). This creates a standard $30^\circ - 60^\circ - 90^\circ$ right triangle, where the ratio of the opposite side to the adjacent side is always $\sqrt{3} \approx 1.732$.

27. For the better dynamic performance of cam follower mechanism, the following displacement diagrams should be chosen

- (A) Simple harmonic motion
- (B) Parabolic motion
- (C) Cycloidal motion
- (D) Uniform acceleration

Correct Answer: (C) Cycloidal motion

Solution:

Concept: When designing cam profiles, the choice of the follower's displacement curve directly affects the velocity, acceleration, and jerk profiles of the system.

- ****Jerk**** is defined as the rate of change of acceleration with respect to time ($j = \frac{da}{dt}$).
- If a displacement curve causes a sudden, discontinuous jump in acceleration, the jerk becomes mathematically infinite (∞). Infinite jerk introduces severe shocks, high structural vibrations, noise, and rapid mechanical wear.

Step 1: Comparing displacement profiles under dynamic conditions.

Let us evaluate the dynamic properties of the common cam motion profiles:

1. **Uniform Velocity / Parabolic / Constant Acceleration:** These profiles exhibit sudden jumps in acceleration at the boundaries of the motion stages, creating an infinite jerk profile that causes severe mechanical shock.
2. **Simple Harmonic Motion (SHM):** While SHM provides a smooth displacement path, it still experiences a sudden jump in acceleration at the start and end points of the lift profile when transitioning into a dwell period. This creates infinite jerk at high speeds.
3. **Cycloidal Motion:** The displacement equation for cycloidal motion is defined using a trigonometric sine profile:

$$s(\theta) = h \left[\frac{\theta}{\beta} - \frac{1}{2\pi} \sin \left(\frac{2\pi\theta}{\beta} \right) \right]$$

Differentiating this equation twice gives the acceleration profile:

$$a(\theta) = \frac{2\pi h \omega^2}{\beta^2} \sin \left(\frac{2\pi\theta}{\beta} \right)$$

Step 2: Evaluating the boundary states of Cycloidal Motion.

Let us check the acceleration values at the start ($\theta = 0$) and end ($\theta = \beta$) of the motion phase:

- At $\theta = 0$: $a(0) = \sin(0) = 0$
- At $\theta = \beta$: $a(\beta) = \sin(2\pi) = 0$

Because the acceleration curve starts at zero and ends at zero, it transitions smoothly into the dwell periods without any sudden jumps. This keeps the **jerk finite throughout the entire cycle**. This smooth transition eliminates structural shock and minimizes vibrations, making **Cycloidal motion** the best choice for high-speed dynamic performance. This matches Option (C).

Quick Tip: Cycloidal motion is the preferred choice for high-speed cam applications because its acceleration curve starts and ends at zero. This eliminates sudden acceleration jumps, keeping the jerk curve finite and preventing mechanical shock.

28. The minimum number of teeth on the pinion in order to avoid interference for 20° full

depth involute system is

- (A) 12
- (B) 14
- (C) 18
- (D) 32

Correct Answer: (C) 18

Solution:

Concept: In involute gear systems, ****interference**** occurs if the teeth are designed such that the tip of the mating gear tooth penetrates into the non-involute flank region near the base of the pinion. To ensure smooth conjugate gear action, we limit the addendum height or enforce a minimum number of teeth on the smaller gear (the pinion). For a pinion meshing with a standard rack, the minimum number of teeth $T_{\min}$ required to completely avoid interference is given by the formula:

$$T_{\min} = \frac{2a_w}{\sin^2 \phi}$$

where a_w is the addendum coefficient factor (expressed as a fraction of the module), and ϕ is the pressure angle of the gear system.

Step 1: Extracting standard parameters for the gear system.

From the problem statement, we identify the following standard system specifications:

- Gear system profile type = 20° full-depth involute system.
- For any standard full-depth gear system, the addendum is equal to one module, which means the addendum coefficient is: $a_w = 1.0$
- The operating pressure angle is given as: $\phi = 20^\circ$

Step 2: Substituting values into the minimum teeth equation.

Substitute $a_w = 1$ and $\phi = 20^\circ$ into the formula:

$$T_{\min} = \frac{2 \times 1}{\sin^2(20^\circ)}$$

Step 3: Calculating the numerical value.

First, find the sine of 20° :

$$\sin(20^\circ) \approx 0.3420$$

Squaring this value:

$$\sin^2(20^\circ) \approx (0.3420)^2 \approx 0.1170$$

Now, divide 2 by this result to find the minimum number of teeth:

$$T_{\min} = \frac{2}{0.1170} \approx 17.09$$

Step 4: Rounding up to the nearest integer.

Since you cannot have a fraction of a gear tooth, we must round up to the next whole number to guarantee that no interference occurs:

$$T_{\min} = 18 \text{ teeth}$$

This calculated value matches Option (C).

Quick Tip: This is a standard reference value in mechanical design: - For a 14.5° full-depth system $\rightarrow$ Minimum teeth = 32 - For a 20° full-depth system $\rightarrow$ Minimum teeth = 18 - For a 20° stub-tooth system $\rightarrow$ Minimum teeth = 14

29. A flywheel is designed for engine speed variation from 190 rps to 210 rps. Calculate the inertia of the wheel if the kinetic energy stored in the flywheel is $400 \text{ N} \cdot \text{m}$.

- (A) $0.02 \text{ kg} \cdot \text{m}^2$
- (B) $0.01 \text{ kg} \cdot \text{m}^2$
- (C) $0.2 \text{ kg} \cdot \text{m}^2$
- (D) $0.1 \text{ kg} \cdot \text{m}^2$

Correct Answer: (A) $0.02 \text{ kg} \cdot \text{m}^2$

Solution:

Concept: A flywheel acts as an energy reservoir that smooths out fluctuations in engine speed by storing kinetic energy during periods of excess power and releasing it when power demand exceeds production. The fluctuation of energy ΔE during a cycle is related to the flywheel's

mass moment of inertia I and its operating speed range by the kinetic energy equation:

$$\Delta E = \frac{1}{2}I(\omega_{\max}^2 - \omega_{\min}^2)$$

where $\omega_{\max}$ and $\omega_{\min}$ represent the maximum and minimum angular velocities expressed in radians per second (rad/s).

Step 1: Converting rotational speed units from rps to rad/s.

The problem gives the engine speeds in revolutions per second (rps). We convert these to angular velocities using the relation $\omega = 2\pi N$:

- Maximum speed, $N_{\max} = 210$ rps $\Rightarrow \omega_{\max} = 2\pi \times 210 = 420\pi$ rad/s
- Minimum speed, $N_{\min} = 190$ rps $\Rightarrow \omega_{\min} = 2\pi \times 190 = 380\pi$ rad/s

Step 2: Setting up the energy balance equation.

The fluctuation of kinetic energy stored in the flywheel is given as:

$$\Delta E = 400 \text{ N} \cdot \text{m}$$

Substitute these values into our energy fluctuation formula:

$$400 = \frac{1}{2} \cdot I \cdot [(420\pi)^2 - (380\pi)^2]$$

Multiply both sides by 2 to simplify the fraction:

$$800 = I \cdot \pi^2 \cdot (420^2 - 380^2)$$

Step 3: Factoring and simplifying the squared terms.

Using the difference of squares identity ($a^2 - b^2 = (a - b)(a + b)$) to simplify the calculation:

$$420^2 - 380^2 = (420 - 380)(420 + 380) = (40)(800) = 32,000$$

Substitute this back into the equation:

$$800 = I \cdot \pi^2 \cdot 32,000$$

Step 4: Solving for the mass moment of inertia I .

Isolate the inertia variable I :

$$I = \frac{800}{32,000 \cdot \pi^2} = \frac{1}{40\pi^2}$$

Using the approximation $\pi^2 \approx 9.8696$:

$$I = \frac{1}{40 \times 9.8696} = \frac{1}{394.78} \approx 0.00253 \text{ kg} \cdot \text{m}^2$$

Quick Tip: To simplify calculations involving flywheel energy changes, always use the difference of squares identity: $\omega_{\max}^2 - \omega_{\min}^2 = (\omega_{\max} - \omega_{\min})(\omega_{\max} + \omega_{\min})$. This turns long squaring steps into a quick multiplication.

30. A governor is said to be stable, if the radius of rotation of balls

- (A) increases as the equilibrium speed decreases
- (B) decreases as the equilibrium speed decreases
- (C) increases as the equilibrium speed increases
- (D) decreases as the equilibrium speed increases

Correct Answer: (C) increases as the equilibrium speed increases

Solution:

Concept: A centrifugal governor regulates the mean speed of an engine by adjusting the fuel supply when load fluctuations cause changes in speed. For a governor to be considered **stable**, it must satisfy a clear equilibrium condition:

- There must be a single, unique equilibrium speed for every specific radius of rotation of the governor balls across its entire operating range.
- If the engine speed increases due to a reduction in load, the governor balls must move outward due to the increased centrifugal force, thereby **increasing the radius of rotation**. This outward movement lifts the control sleeve mechanism, which throttles down the fuel supply valve to bring the engine speed back down to its target value.

Mathematically, this stability condition means that the derivative of the equilibrium speed ω

with respect to the radius of rotation r must be strictly positive:

$$\frac{d\omega}{dr} > 0$$

This derivative confirms that a stable governor configuration requires the equilibrium speed to increase as the radius of rotation increases. This matches Option (C).

Quick Tip: For a centrifugal governor to operate stably: - High Speed $\rightarrow$ Large Radius $\rightarrow$ Lifts Sleeve $\rightarrow$ Throttles Fuel Down. - Low Speed $\rightarrow$ Small Radius $\rightarrow$ Lowers Sleeve $\rightarrow$ Opens Fuel Up. If the radius decreased as speed increased, it would open the fuel valve further, causing the engine to accelerate uncontrollably!

31. Which of the following statements are associated with the complete dynamic balancing of rotating systems?

- I. Resultant couple due to all inertia forces is zero
- II. Support reactions due to forces are zero but not due to couples
- III. The system is automatically statically balanced
- IV. Centre of masses of the system lies on the axis of rotation

- (1) I, II, III and IV
- (2) I, II and III only
- (3) II, III and IV only
- (4) I, III and IV only

Correct Answer: (4) I, III and IV only

Solution:

Concept: Balancing of rotating masses is categorized into two types: static balancing and dynamic balancing.

- **Static Balancing:** Occurs when the center of gravity of the system of masses lies precisely on the axis of rotation. Under this condition, the net eccentric force or the sum of all centrifugal forces acting on the shaft is zero:

$$\sum \vec{F}_c = \sum m_i \cdot r_i \cdot \omega^2 = 0 \Rightarrow \sum m_i \cdot r_i = 0$$

- **Dynamic Balancing:** Occurs when the shaft is rotating and the system experiences no net centrifugal forces as well as no net centrifugal couples along its length. For complete dynamic balancing, two separate conditions must be simultaneously satisfied:

1. The net dynamic force must be zero ($\sum \vec{F} = 0$). This condition automatically ensures that the system satisfies static balancing and that the center of mass lies exactly on the axis of rotation.
2. The net dynamic couple about any reference plane along the shaft axis must be zero ($\sum \vec{C} = 0$).

Let us evaluate each of the given analytical statements step-by-step:

Step 1: Analyzing Statement I (Resultant couple due to all inertia forces is zero).

In dynamic operation, as different masses rotate in different planes, they set up centrifugal couples about any arbitrary point along the shaft axis. If the net dynamic couple is not zero, it creates a turning moment that tends to rock the shaft in its bearings, inducing varying support reactions. For complete dynamic balancing, the net dynamic couple must be zero. Therefore, **Statement I is correct.**

Step 2: Analyzing Statement II (Support reactions due to forces are zero but not due to couples).

When a system is completely dynamically balanced, both the net centrifugal force and the net centrifugal couple are identically zero. Because both are zero, the dynamic reactions at the supporting bearings are completely eliminated under ideal operating speeds. Statement II claims that reactions due to couples are not zero, which directly contradicts the requirement for complete dynamic balancing. Therefore, **Statement II is incorrect.**

Step 3: Analyzing Statement III (The system is automatically statically balanced).

Dynamic balancing requires that the condition $\sum m_i \cdot r_i = 0$ is satisfied along with $\sum m_i \cdot r_i \cdot z_i = 0$. Since satisfying the force equilibrium condition ($\sum m_i \cdot r_i = 0$) is a sub-requirement of dynamic balancing, any system that is dynamically balanced must also be statically balanced. Therefore, **Statement III is correct.**

Step 4: Analyzing Statement IV (Centre of masses of the system lies on the axis of rotation).

Static balancing physically means that the net static mass moment about the axis of rotation vanishes. Mathematically, this shifts the combined center of mass of all rotating components so that it aligns perfectly with the geometric axis of rotation. Since dynamic balancing guarantees static balancing, it guarantees that the center of mass lies on the axis of rotation. Therefore,

Statement IV is correct.

Combining our findings, Statements I, III, and IV are correct, which matches Option (4).

Quick Tip: Remember the absolute hierarchy of balancing rules: - Dynamic Balancing $\implies$ Static Balancing (Always). - Static Balancing $\implies$ Dynamic Balancing. - Complete dynamic balance means Net Force = 0 AND Net Couple = 0, resulting in zero dynamic bearing reactions.

32. Consider a single degree of freedom system with viscous damping excited by a harmonic force. At resonance, the phase angle of the displacement with respect to the exciting force in degrees is

- (1) 0
- (2) 45
- (3) 90
- (4) 135

Correct Answer: (3) 90

Solution:

Concept: Consider a forced, single-degree-of-freedom mechanical system under viscous damping governed by the classical second-order ordinary differential equation:

$$m\ddot{x} + c\dot{x} + kx = F_0 \sin(\omega t)$$

Where:

- m = Mass of the system
- c = Viscous damping coefficient
- k = Stiffness of the supporting spring element
- F_0 = Peak amplitude of the external harmonic excitation force
- ω = Operating frequency of the exciting force

The steady-state harmonic displacement response lags behind the driving force by a specific phase angle ϕ , written as:

$$x(t) = X \sin(\omega t - \phi)$$

The phase angle ϕ is expressed analytically as a function of the excitation frequency ω and system parameters:

$$\tan \phi = \frac{c\omega}{k - m\omega^2}$$

To generalize this behavior for any system, we divide the numerator and denominator by the stiffness k :

$$\tan \phi = \frac{\frac{c\omega}{k}}{1 - \frac{m\omega^2}{k}}$$

We define the natural undamped frequency as $\omega_n = \sqrt{\frac{k}{m}}$ (so $\frac{m}{k} = \frac{1}{\omega_n^2}$) and the critical damping coefficient as $c_c = 2\sqrt{km} = 2m\omega_n$. Introducing the non-dimensional parameters:

- Frequency ratio: $r = \frac{\omega}{\omega_n}$
- Damping ratio: $\zeta = \frac{c}{c_c} = \frac{c}{2m\omega_n}$

Substituting these terms provides the canonical phase angle equation:

$$\tan \phi = \frac{2\zeta\left(\frac{\omega}{\omega_n}\right)}{1 - \left(\frac{\omega}{\omega_n}\right)^2} = \frac{2\zeta r}{1 - r^2}$$

Step 1: Applying the condition of resonance.

Resonance is defined as the state where the frequency of the external harmonic force matches the natural undamped frequency of the mechanical system. Therefore:

$$\omega = \omega_n \Rightarrow r = \frac{\omega}{\omega_n} = 1$$

Step 2: Computing the value of $\tan \phi$ at $r = 1$.

Substituting $r = 1$ into our non-dimensional equation yields:

$$\tan \phi = \frac{2\zeta(1)}{1 - (1)^2} = \frac{2\zeta}{0} \rightarrow \infty$$

Step 3: Solving for the phase angle ϕ .

We find the inverse tangent of infinity:

$$\phi = \arctan(\infty) = \frac{\pi}{2} \text{ radians} = 90^\circ$$

This mathematically demonstrates that at resonance, the displacement response delays behind the excitation force vector by exactly a quarter cycle (90°), meaning the exciting force works

entirely to overcome the viscous damping force. This matches Option (3).

Quick Tip: Phase angle (ϕ) transitions relative to frequency ratio ($r = \frac{\omega}{\omega_n}$): - Low frequencies ($r \ll 1$): $\phi \rightarrow 0^\circ$ (Response is in phase with force). - At Resonance ($r = 1$): $\phi = 90^\circ$ exactly, regardless of the value of damping ratio ζ (provided $\zeta > 0$). - High frequencies ($r \gg 1$): $\phi \rightarrow 180^\circ$ (Response is out of phase).

33. In Rayleigh's method, which of the following energy at the mean position is equal to the maximum potential energy (or strain energy) at the extreme position?

- (1) Maximum kinetic energy
- (2) Minimum kinetic energy
- (3) Maximum potential energy
- (4) Minimum potential energy

Correct Answer: (1) Maximum kinetic energy

Solution:

Concept: Rayleigh's method is a foundational technique used to approximate the fundamental natural frequency of an undamped vibrating system. The principle is rooted firmly in the law of conservation of mechanical energy for conservative systems.

In a system executing simple harmonic motion (SHM) without any dissipative mechanisms (such as viscous or dry friction damping), the total mechanical energy (E_{total}) remains constant at every instant throughout the vibration cycle:

$$E_{\text{total}} = \text{Kinetic Energy (K.E.)} + \text{Potential Energy (P.E.)} = \text{Constant}$$

Let us trace the energy distribution across different phases of the oscillation:

Step 1: Identifying the characteristics at the extreme positions.

At the extreme boundaries of oscillation, the moving mass momentarily stops to reverse its direction of travel. Therefore, the instantaneous velocity (v) drops to zero:

$$v = 0 \quad \Rightarrow \quad \text{K.E.} = \frac{1}{2}mv^2 = 0$$

At this exact coordinate, the spring or structural element reaches its highest state of deformation,

which maximizes its strain energy. Consequently, the total energy of the system is entirely composed of potential energy:

$$E_{\text{total}} = \text{P.E.}_{\text{max}}$$

Step 2: Identifying the characteristics at the mean position.

As the system passes through its static equilibrium (mean) position, the elastic deformation or displacement (x) relative to the equilibrium state is zero:

$$x = 0 \Rightarrow \text{P.E.} = \frac{1}{2}kx^2 = 0$$

At this point, the elastic potential energy is completely converted into kinetic energy, and the system reaches its maximum velocity (v_{max}). Thus, the total energy is composed entirely of kinetic energy:

$$E_{\text{total}} = \text{K.E.}_{\text{max}}$$

Step 3: Equating the energy bounds according to Rayleigh's Principle.

By equating the total energy expressions from Step 1 and Step 2, we establish Rayleigh's energy conservation criterion:

$$\text{Maximum Kinetic Energy}_{\text{at mean position}} = \text{Maximum Potential Energy}_{\text{at extreme position}}$$

This expression allows for the direct derivation of natural frequencies by substituting harmonic profile functions. This aligns with Option (1).

Quick Tip: Rayleigh's Method summary formula:

$$\text{K.E.}_{\text{max}} = \text{P.E.}_{\text{max}}$$

For a lump mass system with displacement $x = X \sin(\omega_n t)$:

$$\frac{1}{2}m(\omega_n X)^2 = \frac{1}{2}kX^2 \Rightarrow \omega_n = \sqrt{\frac{k}{m}}$$

34. The critical speed of the shaft is affected by the

- (1) Diameter and the eccentricity of the shaft
- (2) Span and the eccentricity of the shaft

- (3) Diameter and the span of the shaft
 (4) Frequency of vibrations

Correct Answer: (3) Diameter and the span of the shaft

Solution:

Concept: The critical speed (or whirling speed) of a rotating shaft is the angular velocity at which the system becomes dynamically unstable, causing violent lateral deflections. Analytically, the critical speed of a shaft matches its natural frequency of lateral vibration (ω_n).

For a simple rotor-shaft assembly, the fundamental lateral natural frequency is given by:

$$\omega_n = \sqrt{\frac{g}{\delta}} \text{ rad/s} \quad \text{or} \quad \omega_n = \sqrt{\frac{k}{m}}$$

Where:

- δ = Static lateral deflection of the shaft under its own weight or mounted rotor mass.
- k = Lateral stiffness of the shaft.
- m = Mass of the system.

Let us examine how stiffness k and deflection δ depend on geometric and material properties. For example, consider a simply supported shaft of length (span) L carrying a central point mass M . The static deflection δ is expressed as:

$$\delta = \frac{MgL^3}{48EI}$$

Where:

- E = Young's modulus of elasticity of the shaft material.
- I = Area moment of inertia of the shaft cross-section. For a solid circular shaft of diameter d , $I = \frac{\pi d^4}{64}$.

Step 1: Substituting the variables to see dependencies.

We substitute the expression for the area moment of inertia I into our deflection equation:

$$\delta = \frac{MgL^3}{48E\left(\frac{\pi d^4}{64}\right)} = \frac{4MgL^3}{3\pi Ed^4}$$

Now, we substitute this deflection back into the natural frequency / critical speed formula:

$$\omega_n = \sqrt{\frac{g}{\delta}} = \sqrt{\frac{g}{\left(\frac{4MgL^3}{3\pi Ed^4}\right)}} = \sqrt{\frac{3\pi Ed^4}{4ML^3}} = d^2 \sqrt{\frac{3\pi E}{4ML^3}}$$

Step 2: Evaluating geometric factors.

From the resulting relationship, the natural critical speed exhibits explicit functional dependencies:

$$\omega_n \propto d^2 \quad (\text{Square of the shaft diameter})$$

$$\omega_n \propto \frac{1}{L^{3/2}} \quad (\text{Inverse relationship to the span length to the power 1.5})$$

Thus, the parameters that dictate the inherent critical speed are the **diameter** of the shaft and the **span** (length) of the shaft.

Step 3: Clarifying the role of eccentricity.

Eccentricity (e), which is the distance between the geometric center of the shaft and its mass center, governs the amplitude of the whirling deflection during rotation:

$$y = \frac{e}{\left(\frac{\omega_n}{\omega}\right)^2 - 1}$$

While eccentricity directly affects the magnitude of vibrations and stresses, it does not alter the underlying system properties (k or m) and therefore does not shift the critical frequency ω_n itself. Thus, Option (3) is the correct choice.

Quick Tip: Critical speed of a shaft depends entirely on its structural stiffness and mass distribution: - Higher diameter ($d \uparrow$) $\Rightarrow$ Stiffness ($I \uparrow$) $\Rightarrow$ Critical speed $\uparrow$ - Longer span ($L \uparrow$) $\Rightarrow$ Stiffness $\downarrow$ $\Rightarrow$ Critical speed $\downarrow$ - Eccentricity influences vibration amplitude, but does not alter the critical speed value itself.

35. If a radial ball bearing has a basic load rating of 50 kN and desired rating life is 6000 hours, then the equivalent radial load that the bearing can carry at 500 rpm is

- (1) 18.85 kN
- (2) 8.85 kN
- (3) 12 kN
- (4) 10 kN

Correct Answer: (2) 8.85 kN

Solution:

Concept: The operational engineering life of a rolling element bearing is modeled by the standard load-life empirical equation defined by the catalog standards of the anti-friction bearing manufacturers (SKF/ISO):

$$L_{10} = \left(\frac{C}{P} \right)^k$$

Where:

- L_{10} = Nominal rating life expressed in millions of revolutions (10^6 revs).
- C = Basic dynamic load rating of the bearing (given here as 50 kN).
- P = Equivalent dynamic radial load acting on the bearing (the unknown variable to solve for).
- k = Load-life exponent, which depends on the geometric shape of the rolling elements:
 - $k = 3$ for ball bearings (point contact configurations).
 - $k = \frac{10}{3}$ for roller bearings (line contact configurations).

Since the problem specifies a **radial ball bearing**, we use $k = 3$.

Step 1: Convert the specified lifetime from hours into millions of revolutions.

The relation linking life in hours (L_H) and rotational velocity in revolutions per minute (N) to life in millions of revolutions (L_{10}) is:

$$L_{10} = \frac{60 \times N \times L_H}{10^6}$$

Given information:

- Desired rating life, $L_H = 6000$ hours
- Rotational speed, $N = 500$ rpm

Substituting these values:

$$L_{10} = \frac{60 \times 500 \times 6000}{10^6}$$

Multiplying the numerator constants out step-by-step:

$$60 \times 500 = 30,000$$

$$30,000 \times 6000 = 180,000,000 = 180 \times 10^6$$

Dividing by the denominator:

$$L_{10} = \frac{180 \times 10^6}{10^6} = 180 \text{ million revolutions}$$

Step 2: Isolate the equivalent dynamic load parameter P .

Using the load-life formula with $k = 3$:

$$180 = \left(\frac{C}{P}\right)^3$$

Taking the cube root on both sides:

$$(180)^{1/3} = \frac{C}{P} \Rightarrow P = \frac{C}{(180)^{1/3}}$$

Step 3: Perform numerical calculation.

First, compute the cube root of 180:

$$5^3 = 125, \quad 6^3 = 216$$

Approximating further:

$$(5.646)^3 \approx 180 \Rightarrow (180)^{1/3} \approx 5.646$$

Substitute $C = 50 \text{ kN}$ and this root into the isolated equation:

$$P = \frac{50}{5.646} \approx 8.8556 \text{ kN}$$

Rounding to two decimal places yields 8.85 kN, which matches Option (2).

Quick Tip: Always double check the exponent value matching your bearing profile type: - Ball Bearing: $L_{10} = (C/P)^3 \Rightarrow P = C/\sqrt[3]{L_{10}}$ - Roller Bearing: $L_{10} = (C/P)^{10/3} \Rightarrow P = C/(L_{10})^{0.3}$ Converting hours to millions of revs uses the factor $\frac{60 \cdot N \cdot L_H}{10^6}$.

36. The difference between two principal stresses is 120 MPa, when a circular shaft is subjected

to an axial force and shear force. What is the factor of safety based on maximum shear stress theory if elastic limit of the bar is 300 MPa?

- (1) 1.03
- (2) 2.5
- (3) 3.25
- (4) 5.0

Correct Answer: (2) 2.5

Solution:

Concept: The Maximum Shear Stress Theory (also known as Tresca's Criteria) states that a ductile component will undergo plastic yield failure when the maximum absolute shear stress developed within the material under a multi-axial state of stress reaches a critical threshold value. This threshold equals the maximum shear stress experienced by a tensile test specimen of the same material at its elastic yielding limit.

Mathematically, if the two principal stresses are σ_1 and σ_2 (with $\sigma_1 > \sigma_2$), the maximum structural shear stress ($\tau_{\max}$) in that plane is given by:

$$\tau_{\max} = \frac{\sigma_1 - \sigma_2}{2}$$

According to Tresca's yielding condition, the allowable design shear stress ($\tau_{\text{allowable}}$) incorporates the safety factor:

$$\tau_{\text{allowable}} = \frac{\tau_y}{\text{F.O.S.}} = \frac{\sigma_y}{2 \cdot \text{F.O.S.}}$$

Where:

- σ_y = Yield stress limit / elastic limit of the material under simple tension.
- $\tau_y = \frac{\sigma_y}{2}$ = Yield stress threshold in pure shear.
- F.O.S. = Factor of safety.

Setting the maximum occurring shear stress equal to the design allowable limit gives:

$$\tau_{\max} = \frac{\sigma_1 - \sigma_2}{2} = \frac{\sigma_y}{2 \cdot \text{F.O.S.}}$$

Canceling the common factor of 2 from both sides simplifies the expression to:

$$\sigma_1 - \sigma_2 = \frac{\sigma_y}{\text{F.O.S.}}$$

Step 1: Identifying values from the problem description.

- The difference between the principal stresses is: $\sigma_1 - \sigma_2 = 120 \text{ MPa}$
- The tensile elastic limit of the bar material is: $\sigma_y = 300 \text{ MPa}$

Step 2: Rearranging the simplified formula to solve for the factor of safety.

$$\text{F.O.S.} = \frac{\sigma_y}{\sigma_1 - \sigma_2}$$

Step 3: Calculating the numerical value.

Substitute the parameters into the equation:

$$\text{F.O.S.} = \frac{300 \text{ MPa}}{120 \text{ MPa}}$$

Simplifying the fraction by dividing numerator and denominator by 10, then by 6:

$$\text{F.O.S.} = \frac{30}{12} = \frac{5}{2} = 2.5$$

The calculated factor of safety is exactly 2.5, which matches Option (2).

Quick Tip: Under Tresca's maximum shear stress theory, for any two-dimensional stress case where the principal stresses have opposite signs or are standard planar outputs, the design relation simplifies directly to:

$$\text{F.O.S.} = \frac{\sigma_{\text{yield}}}{\sigma_1 - \sigma_2}$$

This avoids calculating intermediate shear values.

37. If the core diameter of threads equal to the diameter of unthreaded portion of the bolt, then it is known as

- (1) Bolt of uniform strength
- (2) Bolt of fatigue strength
- (3) Bolt of tensile strength
- (4) Bolt of compressive strength

Correct Answer: (1) Bolt of uniform strength

Solution:

Concept: When an ordinary standard bolt with a uniform shank diameter equal to its nominal major diameter (d) is subjected to an axial impact or shock force, the stress distribution along its length is highly non-uniform.

Let us examine why this variation happens across the different sections of the bolt:

- **Unthreaded Shank Section:** The cross-sectional area available to resist loads is $A_{\text{shank}} = \frac{\pi}{4}d^2$.
- **Threaded Section:** The cross-sectional area is determined by the smaller core (or root) diameter (d_c), given by $A_{\text{core}} = \frac{\pi}{4}d_c^2$.

Since $d_c < d$, the cross-sectional area at the threads is significantly smaller than that of the unthreaded shank.

When an axial force F acts on the bolt, the stress induced is inversely proportional to the cross-sectional area ($\sigma = \frac{F}{A}$). Consequently, the stress at the threaded root is much higher than the stress in the shank body:

$$\sigma_{\text{threaded}} = \frac{F}{\frac{\pi}{4}d_c^2} \gg \sigma_{\text{shank}} = \frac{F}{\frac{\pi}{4}d^2}$$

Step 1: Analyzing the strain energy capacity of a standard bolt.

The total capacity of a component to absorb impact energy without permanent deformation is determined by its resilience, which depends directly on the volume of material subjected to high stress:

$$U = \int \frac{\sigma^2}{2E} dV$$

In a standard bolt, only the narrow threaded region experiences peak stress, meaning the large shank region remains under-stressed and contributes little to absorbing impact energy. Failure occurs prematurely at the threads due to this localized strain concentration.

Step 2: Defining the design modifications for a bolt of uniform strength.

To maximize the impact energy absorption capacity, the bolt must experience uniform stress throughout its entire length. This state of equal structural resistance is achieved when the cross-sectional area of the unthreaded shank equals the cross-sectional area of the threaded section. Mathematically, this requires:

$$A_{\text{shank}} = A_{\text{core}} \Rightarrow \frac{\pi}{4}d_{\text{shank}}^2 = \frac{\pi}{4}d_c^2 \Rightarrow d_{\text{shank}} = d_c$$

When the diameter of the unthreaded portion is reduced to match the core thread diameter, the bolt can deform uniformly along its full length, maximizing its energy absorption capacity. This specialized design configuration is called a **bolt of uniform strength**, which corresponds to Option (1).

Quick Tip: A bolt of uniform strength can be manufactured using two methods: 1. Turning down the unthreaded shank diameter until it matches the core thread root diameter ($d_{\text{shank}} = d_c$). 2. Drilling an axial longitudinal hole through the center of the shank to reduce its net cross-sectional area to match the core area.

38. A shaft diameter 'd' is connected to the hub through a square key each of side 'd/4' and length 'l' and transmits a torque 'T'. Assume the length of key is equal to thickness of the pulley, the average shear stress developed in the key is

- (1) $\frac{T}{ld^2}$
- (2) $\frac{2T}{ld^2}$
- (3) $\frac{8T}{ld^2}$
- (4) $\frac{4T}{ld^2}$

Correct Answer: (3) $\frac{8T}{ld^2}$

Solution:

Concept: A key is a mechanical component inserted between a rotating machine shaft and the hub of a pulley or gear to facilitate torque transmission. During operation, the transmitted torque T creates a tangential shearing force F at the outer surface boundary of the shaft.

The relationship connecting the torque T , the tangential force F , and the shaft radius ($r = \frac{d}{2}$) is written as:

$$T = F \cdot \left(\frac{d}{2}\right) \Rightarrow F = \frac{2T}{d}$$

This tangential force F acts directly on the critical shear plane of the embedded key. Let the key have width w , thickness h , and length l . For a square key, the width matches the thickness:

$$w = h = \frac{d}{4}$$

Step 1: Determining the shearing area of the key.

Shear failure occurs when the key splits along its mid-plane parallel to the tangent of the shaft surface. The area resisting this shearing force is the product of the key's width (w) and its inserted length (l):

$$A_{\text{shear}} = w \cdot l$$

Substituting the given geometric dimension $w = \frac{d}{4}$:

$$A_{\text{shear}} = \left(\frac{d}{4}\right) \cdot l = \frac{l \cdot d}{4}$$

Step 2: Calculating the average shear stress (τ).

The average shear stress developed within the key is defined as the total tangential force divided by the resisting shear area:

$$\tau = \frac{F}{A_{\text{shear}}}$$

Substitute the expressions for F and A_{shear} into this formula:

$$\tau = \frac{\left(\frac{2T}{d}\right)}{\left(\frac{l \cdot d}{4}\right)}$$

Step 3: Simplifying the algebraic expression.

To simplify this fraction, we multiply the numerator by the reciprocal of the denominator:

$$\tau = \frac{2T}{d} \times \frac{4}{l \cdot d} = \frac{2 \times 4 \times T}{l \cdot d \cdot d} = \frac{8T}{l \cdot d^2}$$

The derived average shear stress matches Option (3).

Quick Tip: For any standard rectangular or square key: - Tangential Force: $F = \frac{2T}{d}$ - Shear Area: $A_s = w \cdot l$
 - Shear Stress: $\tau = \frac{2T}{w \cdot l \cdot d}$ Substituting the square key width $w = \frac{d}{4}$ directly yields: $\tau = \frac{2T}{(\frac{d}{4}) \cdot l \cdot d} = \frac{8T}{ld^2}$.

39. When brakes are applied to a moving vehicle, then the kinetic energy of the vehicle is converted to

- (1) Potential energy
- (2) Mechanical energy
- (3) Gravitational energy
- (4) Heat energy

Correct Answer: (4) Heat energy

Solution:

Concept: The operation of a vehicle braking system is governed by the First Law of Thermodynamics, which states that energy can neither be created nor destroyed, but can only change from one form to another.

A moving vehicle of mass m traveling at a velocity v possesses directional translation kinetic energy (E_k), expressed as:

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

When the operator applies the brakes, mechanical force presses the high-friction brake pads against the rotating surfaces of the brake discs or drums connected to the wheels.

Let us evaluate the energy transformation steps:

Step 1: Evaluating the mechanical interaction.

The contact between the non-rotating brake pads and the rotating brake rotors creates a strong frictional resistance force ($F_{\text{friction}} = \mu \cdot N$). This force opposes the rotation of the wheels, generating a braking torque that decelerates the vehicle.

Step 2: Computing mechanical work done by friction.

The work done by the friction force over a braking distance s is given by:

$$W_f = \int_0^s F_{\text{friction}} ds$$

According to the work-energy theorem, this frictional work must equal the total change in the vehicle's kinetic energy to bring it to a stop:

$$W_f = \Delta E_k = \frac{1}{2}mv^2 - 0 = \frac{1}{2}mv^2$$

Step 3: Determining the final form of energy.

At the microscopic level, the mechanical work done against friction excites the atoms on the contact surfaces of the brake pads and rotors, increasing their thermal kinetic energy. This microscopic structural excitation manifests macroscopically as a rapid increase in temperature across the brake assembly.

Thus, the vehicle's kinetic energy is converted into thermal energy, or **heat energy**, which is then dissipated into the surrounding atmosphere via convection and radiation. This matches Option (4).

Quick Tip: Braking is an energy transformation process:

$$\text{Kinetic Energy } \left(\frac{1}{2}mv^2 \right) \xrightarrow{\text{Frictional Work}} \text{Thermal/Heat Energy } (m_{\text{rotor}} \cdot c_p \cdot \Delta T)$$

In modern electric vehicles, regenerative braking systems capture a portion of this kinetic energy, converting it into electrical potential energy stored in batteries rather than losing it entirely as heat.

40. Generally, in welded joints, the maximum size of fillet weld is considered by

- (1) adding 1.5 mm to thickness of thinner member to be joined
- (2) adding 3 mm to thickness of thinner member to be joined
- (3) subtracting 3 mm from thickness of thinner member to be joined
- (4) subtracting 1.5 mm from thickness of thinner member to be joined

Correct Answer: (4) subtracting 1.5 mm from thickness of thinner member to be joined

Solution:

Concept: In structural welding design, a fillet weld joins two surfaces at approximately right angles. The size of a fillet weld (s) refers to the leg length of the largest inscribed isosceles right triangle forming the weld cross section.

To ensure structural integrity and prevent stress concentration or burning of the base metal, engineering design standards (such as IS 800 and AWS specifications) define strict limits for weld sizing based on the plate thickness (t):

- **Minimum Size of Fillet Weld ($s_{\min}$):** Avoids rapid cooling rates that can cause brittleness in thick sections. It depends on the thickness of the thicker plate component being joined.
- **Maximum Size of Fillet Weld ($s_{\max}$):** Prevents the weld pool from spilling over the outer edges and ensures the corner remains un-melted, maintaining structural definition. It depends on the thickness of the **thinner plate member** (t).

Let us look at the standard rules for specifying the maximum weld size across different plate boundary profiles:

Step 1: Examining the design rule for square edges.

When a fillet weld is applied along the straight square edge of a structural plate member, the maximum size of the leg length must be limited so that the arc heat does not burn away the

outer corner. The standard design rule specifies:

$$s_{\max} = t - 1.5 \text{ mm}$$

Where t is the thickness of the thinner member being jointed. This matches the wording of Option (4).

Step 2: Examining the design rule for rounded/rolled edges.

For plate components with rounded boundaries, such as the toe of a structural rolled steel channel or angle section, the maximum size is bounded by the profile curvature:

$$s_{\max} = \frac{3}{4} \times t$$

Step 3: Confirming the choice.

Since the prompt refers generally to plate thicknesses without specifying a rounded profile, the standard square-edge design boundary applies. Therefore, the maximum size of the fillet weld equals the thickness of the thinner member minus 1.5 mm. This aligns with Option (4).

Quick Tip: Fillet weld sizing guide parameters: - For square plate edges: $s_{\max} = t_{\text{thinner}} - 1.5 \text{ mm}$ - For round/rolled sections: $s_{\max} = \frac{3}{4} \cdot t_{\text{section}}$ This structural margin keeps the edge from melting completely, preserving the base metal geometry.

41. The close-coiled helical springs 'A' and 'B' are of same material, same coil diameter, same wire diameter and subjected to same load. If the number of turns of spring 'A' is half that of spring 'B', the ratio of deflection of spring 'A' to spring 'B' is

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

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

Solution:

Concept: The axial elongation or deflection (δ) of a closely coiled helical spring subjected to an axial force load W is derived using Castigliano's theorem or strain energy methods in

torsion. The canonical formula for deflection is:

$$\delta = \frac{8WD^3n}{Gd^4}$$

Where the independent parameters are defined as:

- W = Applied axial force load on the spring.
- D = Mean diameter of the spring coil.
- n = Number of active coils or turns in the spring helix.
- G = Shear modulus of elasticity (modulus of rigidity) of the material.
- d = Diameter of the spring wire.

Step 1: Extracting proportional relationships from the given conditions.

The problem specifies that the two springs, designated as 'A' and 'B', share several identical properties:

- Same material $\Rightarrow G_A = G_B = G$
- Same coil diameter $\Rightarrow D_A = D_B = D$
- Same wire diameter $\Rightarrow d_A = d_B = d$
- Subjected to same load $\Rightarrow W_A = W_B = W$

By grouping all constant parameters into a single proportionality constant K :

$$K = \frac{8WD^3}{Gd^4}$$

We can express the deflection of any spring in this problem as a direct linear function of its number of active turns:

$$\delta = K \cdot n \quad \Rightarrow \quad \delta \propto n$$

Step 2: Setting up the deflection ratio.

Using our proportionality relationship, the ratio of the deflection of spring 'A' (δ_A) to that of spring 'B' (δ_B) simplifies to the ratio of their respective numbers of turns:

$$\frac{\delta_A}{\delta_B} = \frac{n_A}{n_B}$$

Step 3: Calculating the numerical ratio using the given turn condition.

The problem states that the number of turns of spring 'A' is exactly half that of spring 'B':

$$n_A = \frac{1}{2}n_B \Rightarrow \frac{n_A}{n_B} = \frac{1}{2}$$

Substituting this value into our ratio equation:

$$\frac{\delta_A}{\delta_B} = \frac{1}{2}$$

Thus, the ratio of the deflection of spring 'A' to spring 'B' is $\frac{1}{2}$, matching Option (1).

Quick Tip: Deflection (δ) and stiffness (k) behaviors in helical springs: - Deflection is directly proportional to the number of turns: $\delta \propto n$ - Axial spring stiffness is inversely proportional to the number of active turns: $k = \frac{W}{\delta} = \frac{Gd^4}{8D^3n} \Rightarrow k \propto \frac{1}{n}$ Cutting a spring in half reduces the turns by half, which doubles its stiffness and halves its deflection under the same load.

42. The variation of volume of liquid with respect to the pressure and its reciprocal respectively are

- (1) Compressibility & Young's modulus
- (2) Young's modulus & Compressibility
- (3) Compressibility & Bulk modulus
- (4) Bulk modulus & Compressibility

Correct Answer: (3) Compressibility & Bulk modulus

Solution:

Concept: To understand the volumetric behavior of a fluid under changing external pressures, fluid mechanics introduces two foundational parameters: the Bulk Modulus of Elasticity (K) and Compressibility (β).

Let us analyze a specified mass of fluid that initially occupies an equilibrium volume V under an ambient static pressure P . When an incremental external pressure increase dP is applied uniformly across the boundaries of this fluid, the fluid compresses, causing its volume to

decrease by an amount dV . The corresponding volumetric strain (ϵ_v) is given by:

$$\epsilon_v = \frac{\text{Change in Volume}}{\text{Original Volume}} = \frac{-dV}{V}$$

The negative sign indicates a physical reduction in volume as pressure increases.

Step 1: Defining the Bulk Modulus (K).

The Bulk Modulus describes the fluid's resistance to volumetric compression. It is defined as the ratio of the incremental pressure change to the resulting volumetric strain:

$$K = \frac{\text{Direct Incremental Pressure}}{\text{Volumetric Strain}} = \frac{dP}{\left(-\frac{dV}{V}\right)} = -V \frac{dP}{dV}$$

Step 2: Defining Compressibility (β).

Compressibility represents the ease with which a fluid undergoes a volume change when subjected to a pressure gradient. Mathematically, compressibility is the fractional change in volume per unit change in pressure. It is defined as:

$$\beta = -\frac{1}{V} \frac{dV}{dP}$$

Comparing the formulas for Bulk Modulus (K) and Compressibility (β), we see that they are exact reciprocals of each other:

$$\beta = \frac{1}{K}$$

Step 3: Matching the question criteria.

The question asks for two specific definitions in sequence:

1. **The variation of volume of liquid with respect to pressure:** This describes fractional volume change per unit pressure change, which is **Compressibility** (β).
2. **Its reciprocal:** The reciprocal of compressibility ($\frac{1}{\beta}$) equals the **Bulk modulus** (K).

Therefore, the correct pair in order is Compressibility & Bulk modulus, matching Option (3).

Quick Tip: Remember the physical analogy to solid mechanics: - Bulk Modulus (K) measures volumetric stiffness (high K means highly incompressible, like water). - Compressibility (β) measures volumetric compliance. - They are related by an inverse relationship: $\beta = \frac{1}{K}$.

43. If the net rotation of fluid particles about their mass centre remains zero, then such type of flow is known as

- (1) Rotational flow
- (2) Irrotational flow
- (3) Laminar flow
- (4) Turbulent flow

Correct Answer: (2) Irrotational flow

Solution:

Concept: Kinematics of fluid flow classifies fluid fields based on the angular motion of the fluid elements as they move along their streamlines.

Consider a small fluid element moving through a flow field. As it travels from one position to another, it can experience translation, deformation (both linear and shear), and rotation. The rotation of a fluid element about a given axis is defined as the average angular velocity of two mutually perpendicular linear differential segments meeting at that point.

Analytically, the angular velocity vector $\vec{\omega}$ in a three-dimensional Cartesian flow field with velocity components (u, v, w) is given by:

$$\vec{\omega} = \omega_x \hat{i} + \omega_y \hat{j} + \omega_z \hat{k}$$

Where the component rotations about the x , y , and z axes are:

$$\omega_x = \frac{1}{2} \left(\frac{\partial w}{\partial y} - \frac{\partial v}{\partial z} \right), \quad \omega_y = \frac{1}{2} \left(\frac{\partial u}{\partial z} - \frac{\partial w}{\partial x} \right), \quad \omega_z = \frac{1}{2} \left(\frac{\partial v}{\partial x} - \frac{\partial u}{\partial y} \right)$$

Step 1: Evaluating the condition where net rotation is zero.

The problem states that the net rotation of the fluid particles about their center of mass is zero during motion. Mathematically, this condition requires all components of the angular velocity vector to vanish simultaneously:

$$\vec{\omega} = 0 \quad \Rightarrow \quad \omega_x = 0, \omega_y = 0, \omega_z = 0$$

Step 2: Defining Irrotational Flow.

When a fluid element moves along a path in such a way that its net orientation does not rotate about its mass center (i.e., its axes remain parallel to their original directions, despite any shear deformation), the flow is classified as **irrotational flow**.

Step 3: Comparing with other options.

- **Rotational flow:** Fluid elements rotate about their mass centers ($\vec{\omega} \neq 0$), typically due to shear stresses caused by viscosity near solid boundaries.
- **Laminar flow:** Fluid particles move in smooth, parallel layers or laminas, without macroscopic mixing.
- **Turbulent flow:** Fluid particles exhibit chaotic, random, and fluctuating three-dimensional motion.

Since the zero-rotation condition refers uniquely to irrotational flow fields, Option (2) is the correct choice.

Quick Tip: To test mathematically for irrotational flow, use the curl of the velocity vector (Vorticity, $\vec{\zeta}$):

$$\vec{\zeta} = \nabla \times \vec{V} = 2\vec{\omega}$$

For irrotational flow, $\nabla \times \vec{V} = 0$. This condition allows for the definition of a scalar velocity potential function (ϕ), where $\vec{V} = -\nabla\phi$.

44. When a body is floating in the liquid and given a small angular displacement, then the body starts oscillating about a centre. That centre is called as

- (1) Centre of pressure
- (2) Centre of gravity
- (3) Centre of buoyancy
- (4) Metacentre

Correct Answer: (4) Metacentre

Solution:

Concept: The stability of a floating body depends on the relative positions of its geometric and gravitational centers, as well as the behavior of the buoyant force when the body tilts.

Let us define the primary reference centers for a floating body at equilibrium:

- **Centre of Gravity (G):** The single point through which the total gravitational weight

force (W) acts vertically downward.

- **Centre of Buoyancy (B):** The centroid of the displaced volume of liquid, through which the buoyant force (F_B) acts vertically upward. At static equilibrium, G and B lie on the same vertical axis of symmetry.

Step 1: Analyzing the shift in centers during angular displacement.

When the floating body is given a small angular tilt or heel angle θ , its underwater geometry becomes asymmetric. The submerged volume shifts toward the tilting side, moving the centroid of this displaced liquid volume to a new position, denoted as B' .

Step 2: Defining the Metacentre (M).

Because the center of buoyancy shifts to B' , the line of action of the buoyant force shifts as well. The vertical line acting upward through the new center of buoyancy B' intersects the original vertical axis of symmetry of the body at a specific point. This point of intersection for small angles of tilt ($\theta \rightarrow 0$) is defined as the **Metacentre (M)**.

Step 3: Characterizing the oscillatory behavior.

As the body rolls or tilts slightly, the restoring or upsetting couple acts about a pivot point. The metacentre M serves as this effective rotational center about which the floating body begins to oscillate.

- If M lies above G (positive metacentric height, $GM > 0$), a restoring torque returns the body to equilibrium, resulting in stable oscillations about M .
- If M lies below G ($GM < 0$), the body is unstable and will capsize.

Thus, the center about which the body oscillates is the metacentre, which corresponds to Option (4).

Quick Tip: Metacentric Height (GM) analytical formula:

$$GM = \frac{I}{V_{\text{submerged}}} - BG$$

Where: - I = Second moment of area of the floating body at the water surface plane. - $V_{\text{submerged}}$ = Total volume of fluid displaced. - BG = Distance between the center of buoyancy and center of gravity.

45. A pipe line has a diameter of 5 cm in which water flows at a pressure of 10^5 N/m^2 with a

mean velocity of 2 m/s. If the head of the water at the cross section is 5 m above the datum line and consider $g = 10 \text{ m/s}^2$, then the total head is

- (1) 15.2 m
- (2) 5.2 m
- (3) 18.2 m
- (4) 21.2 m

Correct Answer: (1) 15.2 m

Solution:

Concept: The total mechanical energy per unit weight of an incompressible, inviscid fluid flowing through a closed conduit is constant along a streamline according to Bernoulli's equation. This total energy per unit weight is expressed as the ****Total Head**** (H), which is the sum of three distinct energy heads:

$$H = \text{Pressure Head} + \text{Kinetic (Velocity) Head} + \text{Potential (Datum) Head}$$

Mathematically, the equation is written as:

$$H = \frac{P}{\rho g} + \frac{v^2}{2g} + z$$

Where:

- P = Static pressure of the fluid.
- ρ = Mass density of the fluid (for water, $\rho = 1000 \text{ kg/m}^3$).
- g = Acceleration due to gravity (given as 10 m/s^2).
- v = Mean flow velocity.
- z = Elevation height above a selected reference datum line.

Step 1: Calculating the Pressure Head ($h_p = \frac{P}{\rho g}$).

Given parameters:

- Pressure, $P = 10^5 \text{ N/m}^2$
- Density of water, $\rho = 1000 \text{ kg/m}^3$
- Gravity acceleration, $g = 10 \text{ m/s}^2$

Substituting these values:

$$h_p = \frac{10^5}{1000 \times 10} = \frac{100,000}{10,000} = 10 \text{ m}$$

Step 2: Calculating the Velocity Head ($h_v = \frac{v^2}{2g}$).

Given parameter:

- Mean velocity, $v = 2 \text{ m/s}$

Substituting these values:

$$h_v = \frac{2^2}{2 \times 10} = \frac{4}{20} = \frac{1}{5} = 0.2 \text{ m}$$

Step 3: Identifying the Datum Head (z).

The problem states that the cross-section is located at an elevation height of 5 m above the datum line:

$$z = 5 \text{ m}$$

Step 4: Summing the components to find the Total Head (H).

$$H = h_p + h_v + z = 10 \text{ m} + 0.2 \text{ m} + 5 \text{ m}$$

Adding the numbers step-by-step:

$$10 + 0.2 = 10.2 \text{ m}$$

$$10.2 + 5 = 15.2 \text{ m}$$

The total mechanical head of the water at this cross-section is exactly 15.2 m, which matches Option (1).

Quick Tip: Always ensure dimensional consistency in Bernoulli calculations: - Pressure unit must be in N/m^2 (Pascal). - Standard density of water $\rho_w = 1000 \text{ kg/m}^3$. - Notice that the pipe diameter (5 cm) is extra information not needed to solve the problem, as the mean velocity is already explicitly provided.

46. For the fluid flowing over the flat plate, the state of boundary layer separation occurs when

- (1) Pressure gradient is zero
- (2) Velocity gradient is zero

- (3) Reynolds number is zero
- (4) Flow is accelerating

Correct Answer: (2) Velocity gradient is zero

Solution:

Concept: When a viscous fluid flows over a solid surface, such as a flat plate, shear forces create a boundary layer where fluid velocity increases from zero at the wall (due to the no-slip condition) to the free-stream velocity.

Boundary layer separation occurs when fluid layers in immediate proximity to the solid wall slow down until they stop, causing the mainstream flow to detach from the surface. This behavior is governed by the velocity profile gradient evaluated directly at the solid boundary ($y = 0$).

Let us analyze the velocity gradient $\left(\frac{\partial u}{\partial y}\right)_{y=0}$ under different pressure gradient conditions along the plate:

Step 1: Favorable Pressure Gradient ($\frac{\partial P}{\partial x} < 0$).

When pressure drops in the direction of flow, the fluid accelerates. This external pressure force helps overcome viscous shear forces near the wall, keeping the velocity profile full and stable against separation. The velocity gradient at the wall remains strongly positive:

$$\left(\frac{\partial u}{\partial y}\right)_{y=0} > 0$$

Step 2: Adverse Pressure Gradient ($\frac{\partial P}{\partial x} > 0$).

When pressure rises in the direction of flow, the fluid faces an opposing force. This adverse pressure gradient slows down the slow-moving fluid inside the boundary layer near the wall, reducing the velocity gradient $\left(\frac{\partial u}{\partial y}\right)_{y=0}$.

Step 3: Identifying the point of separation.

As the adverse pressure gradient continues to retard the flow, the velocity profile reaches a critical point where the fluid velocity at an infinitesimal distance from the wall drops to zero. At this exact location, the shear stress at the wall ($\tau_w = \mu \left(\frac{\partial u}{\partial y}\right)_{y=0}$) vanishes entirely:

$$\left(\frac{\partial u}{\partial y}\right)_{y=0} = 0$$

This condition marks the onset of **boundary layer separation**. Downstream of this point, the velocity gradient becomes negative ($\left(\frac{\partial u}{\partial y}\right)_{y=0} < 0$), leading to localized flow reversal and

vortex formation. Thus, Option (2) is the correct choice.

Quick Tip: Boundary layer conditions at the solid wall ($y = 0$): $-\left(\frac{\partial u}{\partial y}\right)_{y=0} > 0 \implies$ Stable flow (No separation). $-\left(\frac{\partial u}{\partial y}\right)_{y=0} = 0 \implies$ Onset of flow separation. $-\left(\frac{\partial u}{\partial y}\right)_{y=0} < 0 \implies$ Separated flow with backflow region.

47. A pipe friction test shows that, over the range of speeds used for the test, the non-dimensional friction factor varies inversely with Reynolds number. From this, one can conclude that the

- (1) Pipe must be smooth
- (2) Flow must be laminar
- (3) Flow must be turbulent
- (4) Pipe must be rough

Correct Answer: (2) Flow must be laminar

Solution:

Concept: The mechanical energy loss due to friction in a fully developed pipe flow is modeled using the Darcy-Weisbach equation. This equation relates head loss (h_f) to the non-dimensional friction factor (f):

$$h_f = \frac{f \cdot L \cdot v^2}{2 \cdot g \cdot d}$$

The characteristics of the friction factor f vary depending on the flow regime, which is determined by the Reynolds Number ($Re = \frac{\rho v d}{\mu}$).

Let us analyze the relationship between the friction factor and the Reynolds number across different flow regimes:

Step 1: Analyzing the Laminar Flow Regime ($Re < 2000$).

In laminar pipe flow, the velocity profile is parabolic, and viscous shear forces dominate. The exact analytical solution derived from the Hagen-Poiseuille equation yields a direct equation for head loss:

$$h_f = \frac{32\mu L v}{\rho g d^2}$$

Equating this analytical head loss to the empirical Darcy-Weisbach equation allows us to isolate

the friction factor f :

$$\frac{fLv^2}{2gd} = \frac{32\mu Lv}{\rho g d^2} \Rightarrow f = \frac{64}{\left(\frac{\rho v d}{\mu}\right)} = \frac{64}{Re}$$

This mathematical result proves that in the laminar regime, the friction factor f is ****inversely proportional**** to the Reynolds number:

$$f \propto \frac{1}{Re}$$

Crucially, this relationship depends solely on the Reynolds number and is entirely independent of the physical surface roughness of the pipe wall.

Step 2: Analyzing the Turbulent Flow Regime ($Re > 4000$).

In turbulent flow, momentum transfer is dominated by chaotic eddy currents. The friction factor becomes a complex function of both the Reynolds number and the relative pipe roughness (ϵ/d), as described by the Colebrook-White equation or the Moody diagram:

$$\frac{1}{\sqrt{f}} = -2.0 \log_{10} \left(\frac{\epsilon/d}{3.7} + \frac{2.51}{Re \sqrt{f}} \right)$$

For fully rough turbulent flow at high Reynolds numbers, the friction factor becomes constant and depends entirely on wall roughness, losing its dependency on Re .

Step 3: Interpreting the test results.

Since the pipe friction test showed that the friction factor varies inversely with the Reynolds number across the full speed range, the flow field must be in the laminar regime. This matches Option (2).

Quick Tip: Friction factor equations to remember: - Laminar Flow: $f = \frac{64}{Re}$ (Purely inverse linear relationship, independent of pipe roughness). - Turbulent Flow (Smooth pipe): $f = \frac{0.3164}{Re^{0.25}}$ (Blasius equation, much weaker dependence on Re).

48. In turbulent flow, the loss of pressure head is approximately (where v is the velocity of flow)

- (1) Directly proportional to v
- (2) Inversely proportional to v
- (3) Directly proportional to v^2

(4) Inversely proportional to v^2

Correct Answer: (3) Directly proportional to v^2

Solution:

Concept: The loss of pressure head (h_f) caused by fluid friction over a pipe length L is evaluated using the Darcy-Weisbach formula:

$$h_f = \frac{f \cdot L \cdot v^2}{2 \cdot g \cdot d}$$

Where:

- f = Darcy friction factor.
- L = Length of the pipe segment.
- d = Internal diameter of the pipe line.
- v = Mean velocity of fluid flow.
- g = Acceleration due to gravity.

Let us examine how the head loss scales with mean velocity in different flow regimes:

Step 1: Comparing with Laminar Flow.

In laminar flow, the friction factor is inversely proportional to velocity ($f = \frac{64}{Re} = \frac{64\mu}{\rho v d}$). Substituting this into the Darcy-Weisbach equation shows that head loss scales linearly with velocity:

$$h_f = \left(\frac{64\mu}{\rho v d} \right) \frac{L v^2}{2 g d} = \frac{32\mu L v}{\rho g d^2} \Rightarrow h_f \propto v$$

Step 2: Evaluating the Turbulent Flow Condition.

In turbulent flow, high-frequency velocity fluctuations and eddy mixing dominate momentum transfer, increasing shear resistance. For fully turbulent flow or flow in rough pipes, the friction factor f becomes nearly constant, becoming independent of the Reynolds number.

Step 3: Deducing the velocity proportionality.

Treating the friction factor f , length L , diameter d , and gravity g as constant parameters for a given turbulent pipe flow system, the Darcy-Weisbach equation simplifies to:

$$h_f = \left(\frac{f \cdot L}{2 \cdot g \cdot d} \right) \cdot v^2 = \text{Constant} \cdot v^2$$

This demonstrates that in turbulent flow, the loss of pressure head is directly proportional to the square of the flow velocity (v^2). This corresponds to Option (3).

Quick Tip: Head loss scaling rules of thumb: - Laminar pipe flow: $h_f \propto v^1$ - Turbulent pipe flow: $h_f \propto v^2$ (or more precisely $v^{1.75}$ to v^2 depending on surface roughness). Therefore, doubling the velocity in a turbulent system increases pumping pressure head requirements by nearly four times.

49. During a summer day a scooter rider feels more comfortable while on the move than while at rest at a stop light because

- (1) More heat is lost by convection while in motion
- (2) Air is transparent to radiation, hence it is cooler than the body
- (3) The object in motion captures less solar radiation
- (4) Air has a low specific heat, hence it is cooler

Correct Answer: (1) More heat is lost by convection while in motion

Solution:

Concept: The thermal sensation of comfort experienced by a human body in a warm environment depends on the rate of heat dissipation from the skin to the surrounding air. The human body continuously generates metabolic heat, which must be rejected to prevent body temperature from rising. Heat transfer from the skin occurs primarily via three mechanisms: conduction, radiation, and convection.

Convective heat transfer (q) between a surface (the rider's body) and a moving fluid (air) is modeled by Newton's Law of Cooling:

$$q = h \cdot A \cdot (T_{\text{skin}} - T_{\text{ambient}})$$

Where:

- h = Convective heat transfer coefficient.
- A = Exposed surface area of the body.
- T_{skin} = Surface temperature of the rider's skin.
- T_{ambient} = Temperature of the surrounding air on a summer day.

Let us evaluate the two distinct operational states described:

Step 1: Evaluating the rider state at rest at a stop light.

When the scooter is stationary, the air surrounding the rider is nearly still. Heat transfer occurs via **natural (free) convection**, where fluid motion is driven solely by buoyancy forces arising from temperature-induced density gradients near the skin. The convective heat transfer coefficient for natural convection (h_{natural}) is low, limiting heat dissipation and causing the rider to feel warm.

Step 2: Evaluating the rider state while on the move.

When the scooter accelerates into motion, the relative velocity between the rider's body and the air increases substantially. This shifts the heat transfer mechanism to **forced convection**, where fluid flow is driven by an external factor (the vehicle's motion).

The convective heat transfer coefficient (h) scales strongly with fluid velocity (v). For flow over a cylinder or body profile, this relationship is typically expressed using non-dimensional correlations such as:

$$Nu = \frac{hL}{k} = C \cdot Re^m \cdot Pr^n \Rightarrow h \propto v^m \quad (\text{where } m \approx 0.5 \text{ to } 0.8)$$

As velocity increases, the boundary layer thins, which increases both the convective heat transfer coefficient ($h_{\text{forced}} \gg h_{\text{natural}}$) and the rate of sweat evaporation.

Step 3: Correlating to human thermal comfort.

The significant increase in forced convective heat transfer and evaporative cooling allows the body to reject heat much more rapidly while in motion, making the rider feel cooler and more comfortable. This matches Option (1).

Quick Tip: Forced Convection vs. Natural Convection: - Stationary state $\Rightarrow$ Natural Convection $\Rightarrow$ Low velocity $\Rightarrow$ Small h value $\Rightarrow$ Poor heat dissipation. - Moving state $\Rightarrow$ Forced Convection $\Rightarrow$ High relative velocity $\Rightarrow$ Large h value $\Rightarrow$ High heat dissipation rate.

50. For heat transfer calculation in fins, if the Biot number is very small ($Bi < 0.1$) it implies

- (1) Higher convective resistance
- (2) High temperature gradient in the fin
- (3) Fin is infinitely long
- (4) Negligible internal conductive resistance

Correct Answer: (4) Negligible internal conductive resistance

Solution:

Concept: The Biot number (Bi) is a dimensionless parameter used in transient conduction and extended surface (fin) thermal analysis. It relates the internal conduction resistance within a solid body to the external convection resistance at the body's surface.

The Biot number is defined mathematically as:

$$Bi = \frac{\text{Internal Conductive Resistance to Heat Transfer}}{\text{External Convective Resistance to Heat Transfer}} = \frac{\left(\frac{L_c}{k \cdot A}\right)}{\left(\frac{1}{h \cdot A}\right)} = \frac{h \cdot L_c}{k}$$

Where:

- h = Convective heat transfer coefficient of the surrounding fluid.
- L_c = Characteristic length of the solid body (defined as volume divided by surface area, $L_c = \frac{V}{A_s}$).
- k = Thermal conductivity of the solid material.

Let us evaluate the physical meaning of a small Biot number:

Step 1: Analyzing the mathematical limit $Bi < 0.1$.

When the Biot number is less than 0.1, the numerator in our resistance ratio is much smaller than the denominator:

$$\frac{\text{Internal Conductive Resistance}}{\text{External Convective Resistance}} \ll 0.1 \Rightarrow \text{Internal Conductive Resistance} \rightarrow 0$$

This indicates that the material's internal resistance to heat conduction is negligible compared to the resistance to heat convection at the surface.

Step 2: Deducing the spatial temperature distribution.

Because the internal conductive resistance is negligible, heat distributes rapidly within the solid cross-section compared to the rate of heat transfer across the fluid boundary. As a result, the temperature gradient within the solid cross-section is virtually flat. For a thin fin or pin, a small Biot number ($Bi = \frac{hL_c}{k} < 0.1$) justifies the fundamental assumption that temperature varies only along the length of the fin (x -direction) and remains uniform across its cross-section.

Step 3: Matching the given options.

A Biot number of $Bi < 0.1$ directly implies that the system has a **negligible internal conductive resistance**, which corresponds to Option (4).

Quick Tip: Biot Number physical interpretation: - $Bi < 0.1$: Conductive resistance is negligible. Temperature inside the cross-section is uniform. This justifies using the Lumped Parameter Capacity analysis for transient problems. - $Bi > 0.1$: Internal temperature gradients are significant and cannot be ignored; spatial variations must be accounted for using Fourier equations.

51. The characteristic length (L_c) for the calculation of Biot number (Bi) for a sphere of radius 'R' subjected to Newtonian cooling is

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

Correct Answer: (3) $\frac{R}{3}$

Solution:

Concept: In heat transfer analysis, the characteristic length (L_c) of a solid body is defined as the ratio of its total internal volume (V) to its exterior geometric surface area (A_s) exposed to convection:

$$L_c = \frac{V}{A_s}$$

This parameter represents the effective distance heat must travel from the interior core of the body to reach its outer cooling boundary surface.

Let us compute this characteristic length for a solid sphere of radius R :

Step 1: Identifying the geometric formulas for a sphere.

For a perfect solid sphere of radius R :

- Total internal volume: $V = \frac{4}{3}\pi R^3$
- Total outer surface area: $A_s = 4\pi R^2$

Step 2: Substituting these formulas into the characteristic length definition.

$$L_c = \frac{\frac{4}{3}\pi R^3}{4\pi R^2}$$

Step 3: Simplifying the fraction.

We can cancel common factors in the numerator and denominator:

- Cancel the factor of 4: $\frac{\frac{1}{3}\pi R^3}{\pi R^2}$

- Cancel the constant π : $\frac{\frac{1}{3}R^3}{R^2}$
- Simplify the radius terms using exponent rules ($\frac{R^3}{R^2} = R$):

$$L_c = \frac{1}{3} \cdot R = \frac{R}{3}$$

Thus, the characteristic length for a sphere is exactly $\frac{R}{3}$, which matches Option (3).

Quick Tip: Standard Characteristic Lengths ($L_c = V/A_s$) to memorize: - Cube (side length a): $L_c = \frac{a^3}{6a^2} = \frac{a}{6}$ - Long Solid Cylinder (radius R): $L_c = \frac{\pi R^2 L}{2\pi R L} = \frac{R}{2}$ - Solid Sphere (radius R): $L_c = \frac{\frac{4}{3}\pi R^3}{4\pi R^2} = \frac{R}{3}$

52. For a fully developed flow through a pipe, how will the heat transfer coefficient (h) change, when the mean flow rate is doubled with all other conditions remaining the same?

- (1) h increases by 50% (approx)
- (2) h increases by 74% (approx)
- (3) h decreases by 87% (approx)
- (4) h increases by 20% (approx)

Correct Answer: (2) h increases by 74% (approx)

Solution:

Concept: For fully developed turbulent flow through a smooth pipe, convective heat transfer is modeled using the empirical **Dittus-Boelter Equation**. This relation correlates the Nusselt number (Nu) with the Reynolds number (Re) and Prandtl number (Pr):

$$Nu = \frac{h \cdot d}{k} = 0.023 \cdot Re^{0.8} \cdot Pr^n$$

Where:

- h = Convective heat transfer coefficient.
- d = Internal pipe diameter.
- k = Fluid thermal conductivity.
- $Re = \frac{\rho \cdot v \cdot d}{\mu}$ = Reynolds number, which is directly proportional to the fluid mean velocity or flow rate ($Re \propto v$).

- Pr = Prandtl number, a fluid property constant.

Step 1: Finding the proportionality relation for h .

Since all fluid properties (ρ, μ, k, Pr) and geometric dimensions (d) remain constant, we can isolate the heat transfer coefficient h :

$$h \propto Re^{0.8} \Rightarrow h \propto v^{0.8}$$

This shows that the convective heat transfer coefficient scales with velocity to the power of 0.8.

Step 2: Calculating the change when flow rate is doubled.

Let the initial velocity be v_1 and the initial heat transfer coefficient be h_1 . The problem states that the mean flow rate is doubled, meaning the new velocity is:

$$v_2 = 2 \cdot v_1$$

Using our proportionality relation, the new heat transfer coefficient h_2 is:

$$\frac{h_2}{h_1} = \left(\frac{v_2}{v_1} \right)^{0.8} = (2)^{0.8}$$

Let us compute the value of $2^{0.8}$:

$$2^{0.8} = 2^{4/5} = \sqrt[5]{2^4} = \sqrt[5]{16} \approx 1.7411$$

Therefore, the new heat transfer coefficient is:

$$h_2 \approx 1.7411 \cdot h_1$$

Step 3: Calculating the percentage increase in h .

The fractional change in the heat transfer coefficient is given by:

$$\text{Percentage Increase} = \left(\frac{h_2 - h_1}{h_1} \right) \times 100\% = \left(\frac{1.7411 \cdot h_1 - h_1}{h_1} \right) \times 100\%$$

$$\text{Percentage Increase} = (1.7411 - 1) \times 100\% = 0.7411 \times 100\% \approx 74.11\%$$

Thus, the convective heat transfer coefficient increases by approximately 74%, which matches Option (2).

Quick Tip: The exponent 0.8 in the Dittus-Boelter equation is key for turbulent flow problems: - If flow velocity doubles ($2\times$), then h scales by $2^{0.8} \approx 1.74 \implies 74\%$ increase. - Note: If the flow field were fully developed *laminar*, the Nusselt number would be constant ($Nu = 3.66$ or 4.36), meaning h would not change with flow rate. The options here clearly indicate a turbulent flow context.

53. If the surface temperature of hot body is doubled, then rate of radiation energy emitted will be enhanced by a factor of

- (1) 2
- (2) 4
- (3) 16
- (4) 32

Correct Answer: (3) 16

Solution:

Concept: The rate of thermal radiation energy emitted by a body is governed by the ****Stefan-Boltzmann Law****. This law states that the total radiant energy emitted per unit surface area of a blackbody per unit time is directly proportional to the fourth power of its absolute thermodynamic temperature.

Mathematically, the total emissive power (E) is expressed as:

$$E = \epsilon \cdot \sigma \cdot A \cdot T^4$$

Where:

- ϵ = Emissivity coefficient of the surface ($\epsilon = 1$ for an ideal blackbody).
- σ = *Stefan – Boltzmann constant* ($5.67 \times 10^{-8} \text{ W/m}^2\text{K}^4$).
- A = Surface area of the emitting body.
- T = Absolute temperature measured in Kelvin (K).

Step 1: Establishing the proportionality relation.

Assuming the surface geometry, area A , and surface emissivity ϵ remain constant, the rate of total radiation emission E depends solely on absolute temperature:

$$E \propto T^4$$

Step 2: Applying the temperature doubling condition.

Let the initial absolute temperature of the hot body be T_1 , corresponding to an initial radiant energy emission rate E_1 . The problem states that the surface temperature is doubled:

$$T_2 = 2 \cdot T_1$$

Step 3: Calculating the enhancement factor.

We set up the ratio of the final emissive power (E_2) to the initial emissive power (E_1):

$$\frac{E_2}{E_1} = \left(\frac{T_2}{T_1}\right)^4 = \left(\frac{2 \cdot T_1}{T_1}\right)^4 = (2)^4$$

Evaluating the fourth power of 2:

$$2^4 = 2 \times 2 \times 2 \times 2 = 16$$

Therefore, the rate of radiation energy emitted scales by a factor of 16:

$$E_2 = 16 \cdot E_1$$

This corresponds to Option (3).

Quick Tip: Always remember that radiation heat transfer scales with the fourth power of absolute temperature (T^4): - Temperature doubled ($2T$) $\implies$ Radiation increases by $2^4 = 16$ times. - Temperature tripled ($3T$) $\implies$ Radiation increases by $3^4 = 81$ times. - Temperature quadrupled ($4T$) $\implies$ Radiation increases by $4^4 = 256$ times. This non-linear scaling makes radiation the dominant mode of heat transfer at high temperatures.

54. Expansion bellows are provided in the shell and tube heat exchangers in order to

- (1) Increase the surface area
- (2) Relieve stresses caused by thermal expansion
- (3) Increase the turbulence
- (4) Reduce the pressure drop

Correct Answer: (2) Relieve stresses caused by thermal expansion

Solution:

Concept: In a shell-and-tube heat exchanger, two fluids flow at different temperatures: one through the internal tube bundle and the other through the surrounding outer shell structure. As the heat exchanger operates, these components absorb or reject heat, causing their temperatures to change. According to linear thermal expansion principles, the change in length (ΔL) of a structural component depends on its material properties and temperature change:

$$\Delta L = \alpha \cdot L \cdot \Delta T$$

Where:

- α = Coefficient of linear thermal expansion.
- L = Original length of the component.
- ΔT = Change in temperature relative to the initial state.

Let us analyze the structural behavior under these conditions:

Step 1: Evaluating the source of thermal stresses.

Because the shell and the tubes are made of different materials or operate at different temperatures, they expand by different amounts ($\Delta L_{\text{shell}} \neq \Delta L_{\text{tubes}}$). In a fixed tubesheet heat exchanger, both ends of the tubes are rigidly welded to the shell. This rigid constraint prevents free expansion, generating high thermal stresses within the assembly:

$$\sigma_{\text{thermal}} = E \cdot \epsilon_{\text{thermal}} = E \cdot \frac{\Delta L_{\text{constrained}}}{L}$$

These large stresses can cause mechanical failures, such as buckling of the tubes, cracking of the tubesheet welds, or structural distortion of the outer shell.

Step 2: Analyzing the function of expansion bellows.

To mitigate this issue, designers incorporate a flexible component known as an **expansion bellow** (or expansion joint) into the outer shell structure. Bellows consist of a series of thin, corrugated metal rings that can compress or extend axially under low forces.

Step 3: Determining the structural benefit.

The inclusion of expansion bellows alters the system behavior:

- The flexible corrugations absorb the differential thermal expansion ($\Delta L_{\text{shell}} - \Delta L_{\text{tubes}}$) by deforming elastically.

- This flexibility allows the shell to change length slightly, reducing structural constraints and safely relieving internal thermal stresses.

Therefore, expansion bellows are used primarily to relieve stresses caused by thermal expansion, matching Option (2).

Quick Tip: Structural guidelines for heat exchangers: - Fixed Tubesheet Exchangers: Low cost, but require an expansion bellow on the shell if the temperature differential between tube and shell exceeds roughly 50°C. - Floating Head Exchangers: Alternative design where one tubesheet is left free to move axially, completely avoiding thermal expansion stresses without needing shell bellows.

55. On the specified surface, if 35% of incident radiation energy is reflected and the transmissivity of the body is 0.3, then the emissivity of the surface is

- (1) 0.35
- (2) 0.3
- (3) 1.0
- (4) 0.65

Correct Answer: (1) 0.35

Solution:

Concept: When total incident radiant energy flux (Irradiance, G) strikes a semi-transparent body surface, it splits into three components based on conservation of energy: a fraction is reflected, a fraction is absorbed, and the remaining fraction is transmitted through the material layer.

Dividing the energy balance equation by the total incident energy defines the fundamental radiation surface properties:

$$\alpha + \rho + \tau = 1$$

Where:

- $\alpha = \text{Absorptivity}(\text{fraction of incident energy absorbed}).$
- $\rho = \text{Reflectivity}(\text{fraction of incident energy reflected}).$
- $\tau = \text{Transmissivity}(\text{fraction of incident energy transmitted}).$

Additionally, **Kirchhoff's Law of Thermal Radiation** states that for any body in thermodynamic equilibrium with its surroundings, its monochromatic or total surface absorptivity equals its total emissivity (ϵ):

$$\epsilon = \alpha$$

Where ϵ represents the surface emissivity, defined as the ratio of energy radiated by the surface to that radiated by an ideal blackbody at the same temperature.

Step 1: Identifying given parameters from the problem text.

The problem provides the following surface property values:

- Reflectivity: $\rho = 35\% = 0.35$
- Transmissivity: $\tau = 0.3$

Step 2: Calculating surface absorptivity (α).

We isolate absorptivity in our property conservation equation:

$$\alpha = 1 - \rho - \tau$$

Substitute the given values into the formula:

$$\alpha = 1 - 0.35 - 0.3$$

Performing the subtraction step-by-step:

$$1 - 0.35 = 0.65$$

$$0.65 - 0.3 = 0.35$$

Thus, the absorptivity of the surface is $\alpha = 0.35$.

Step 3: Finding the emissivity (ϵ) using Kirchhoff's Law.

Applying Kirchhoff's law:

$$\epsilon = \alpha = 0.35$$

The calculated emissivity of the surface is 0.35, which matches Option (1).

Quick Tip: Radiation properties summary for any material surface: - $\alpha + \rho + \tau = 1$ - For an opaque body, transmissivity is zero ($\tau = 0 \implies \alpha + \rho = 1$). - By Kirchhoff's Law, always equate Emissivity directly to Absorptivity ($\epsilon = \alpha$) under thermal equilibrium conditions.

56. In a combined cycle for power generation, if the topping cycle has an efficiency of 45% and for bottoming cycle 30%, then the combined efficiency of the cycle is

- (A) 45%
- (B) 30%
- (C) 75%
- (D) 61.5%

Correct Answer: (D) 61.5 %

Solution:

Concept: A combined cycle power plant links two thermodynamic cycles together to achieve greater overall thermal efficiency than either cycle operating independently. The high-temperature exhaust gas from the first cycle (topping cycle) serves as the input heat source for the second cycle (bottoming cycle). Let the thermal efficiency of the topping cycle be η_1 and the thermal efficiency of the bottoming cycle be η_2 . The overall or combined thermal efficiency (η_{cc}) of the system can be derived from basic energy balances and is expressed mathematically as:

$$\eta_{cc} = \eta_1 + \eta_2 - \eta_1 \cdot \eta_2$$

This equation shows that the net efficiency is the sum of both individual efficiencies reduced by their product, reflecting that the bottoming cycle only processes the heat rejected by the topping cycle.

Step 1: Identify the given individual efficiencies.

From the problem description, we have the following percentage values for each constituent thermodynamic cycle:

- Efficiency of the topping cycle, $\eta_1 = 45\% = 0.45$
- Efficiency of the bottoming cycle, $\eta_2 = 30\% = 0.30$

Step 2: Substitute values into the combined efficiency formula.

Let's substitute the decimal representations into our standard algebraic expression:

$$\eta_{cc} = 0.45 + 0.30 - (0.45 \times 0.30)$$

Now we execute the arithmetic operations systematically step-by-step:

- Sum of the individual efficiencies:

$$0.45 + 0.30 = 0.75$$

- Product of the individual efficiencies:

$$0.45 \times 0.30 = 0.135$$

Subtracting the product from the sum yields:

$$\eta_{cc} = 0.75 - 0.135 = 0.615$$

Step 3: Convert the decimal form back to percentage.

To state the efficiency as a percentage value, we multiply by 100:

$$\eta_{cc} = 0.615 \times 100\% = 61.5\%$$

Thus, the combined thermal efficiency of the power plant is precisely 61.5%, matches Option (D).

Quick Tip: For combined cycles, always remember the formula $\eta_{cc} = \eta_1 + \eta_2 - \eta_1\eta_2$. Alternatively, you can think in terms of waste: if the first cycle rejects $(1 - 0.45) = 55\%$ of heat, the second cycle converts 30% of that waste into useful work: $0.30 \times 0.55 = 0.165$. Adding this to the original efficiency gives $0.45 + 0.165 = 0.615$ or 61.5%.

57. A real gas undergoes throttling from 50 bar to 10 bar from an initial temperature of 300 K. If $\mu_{JT} = -0.05$ K/bar for the initial condition, what is the approximate value of the exit temperature?

(A) 298 K

- (B) 302 K
(C) 300 K
(D) 304 K

Correct Answer: (B) 302 K

Solution:

Concept: Throttling is a steady-flow expansion process across a restriction (such as a valve or porous plug) characterized by negligible heat transfer, no external shaft work, and negligible changes in kinetic and potential energies. Consequently, it is a constant enthalpy (*isenthalpic*) process ($h_1 = h_2$). The behavior of temperature relative to pressure changes during this expansion is defined by the Joule-Thomson coefficient (μ_{JT}):

$$\mu_{JT} = \left(\frac{\partial T}{\partial P} \right)_h \approx \frac{\Delta T}{\Delta P} = \frac{T_2 - T_1}{P_2 - P_1}$$

Where:

- $\mu_{JT} > 0$ implies cooling occurs during expansion ($\Delta P < 0 \rightarrow \Delta T < 0$).
- $\mu_{JT} < 0$ implies heating occurs during expansion ($\Delta P < 0 \rightarrow \Delta T > 0$).

Step 1: Gather and list the given state properties.

The parameters specified in the problem statement are:

- Initial Pressure, $P_1 = 50$ bar
- Final Pressure, $P_2 = 10$ bar
- Initial Temperature, $T_1 = 300$ K
- Joule-Thomson coefficient, $\mu_{JT} = -0.05$ K/bar

Step 2: Compute the change in pressure (ΔP).

The pressure difference experienced by the real gas during this throttling expansion is:

$$\Delta P = P_2 - P_1 = 10 \text{ bar} - 50 \text{ bar} = -40 \text{ bar}$$

Step 3: Relate the parameters to find the final exit temperature (T_2).

Using the finite difference approximation for the Joule-Thomson relation:

$$T_2 - T_1 = \mu_{JT} \times \Delta P$$

Substituting the known values into the algebraic equation:

$$T_2 - 300 = (-0.05 \text{ K/bar}) \times (-40 \text{ bar})$$

Multiplying the two negative numbers together produces a positive value:

$$T_2 - 300 = 2 \text{ K}$$

Isolating T_2 by moving 300 to the right-hand side:

$$T_2 = 300 + 2 = 302 \text{ K}$$

Because the Joule-Thomson coefficient is negative, the gas undergoes heating during expansion, raising its temperature to approximately 302 K, which corresponds to Option (B).

Quick Tip: When μ_{JT} is negative, the gas warms up during expansion (ΔP is always negative in a flow restriction). A quick qualitative check reveals that since $\mu_{JT} < 0$, T_2 must be greater than T_1 . This instantly eliminates options (A) and (C).

58. A mixture of gas expands from 0.03 m^3 to 0.06 m^3 at a constant pressure of 1 MPa and absorbs 84 kJ of heat during the process. The change in internal energy of the mixture is

- (A) 30 kJ
- (B) 54 kJ
- (C) 114 kJ
- (D) 84 kJ

Correct Answer: (B) 54 kJ

Solution:

Concept: According to the First Law of Thermodynamics for a closed system undergoing a non-cyclic process, the net heat interactions (Q) equal the sum of the change in internal energy (ΔU) and the work performed (W):

$$Q = \Delta U + W$$

For a quasi-static displacement process operating under constant pressure conditions (isobaric process), the boundary work output (W) is given by the integral of pressure over volume:

$$W = \int P dV = P(V_2 - V_1)$$

By evaluating the boundary work and employing the first law, we can isolate the net internal energy alteration.

Step 1: Extract parameters and equalize metric prefixes to standard SI base units.

The parameters given in the problem statement are:

- Constant Pressure, $P = 1 \text{ MPa} = 1 \times 10^6 \text{ N/m}^2$
- Initial Volume, $V_1 = 0.03 \text{ m}^3$
- Final Volume, $V_2 = 0.06 \text{ m}^3$
- Heat absorbed by the system, $Q = +84 \text{ kJ} = 84 \times 10^3 \text{ J}$

Step 2: Calculate the boundary work performed during expansion.

Using the constant-pressure work equation:

$$W = P \times (V_2 - V_1)$$

Substitute the values:

$$W = (1 \times 10^6 \text{ Pa}) \times (0.06 \text{ m}^3 - 0.03 \text{ m}^3)$$

$$W = 1 \times 10^6 \times 0.03 = 30000 \text{ J}$$

Converting Joules to kilojoules gives:

$$W = \frac{30000}{1000} \text{ kJ} = 30 \text{ kJ}$$

Step 3: Solve for the change in internal energy (ΔU) via the First Law.

Rearranging the first law equation to solve for ΔU :

$$\Delta U = Q - W$$

Substituting our calculated heat interaction and boundary work values:

$$\Delta U = 84 \text{ kJ} - 30 \text{ kJ} = 54 \text{ kJ}$$

Thus, the change in internal energy of the gas mixture is equal to 54 kJ, corresponding to Option (B).

Quick Tip: Keep unit consistency intact! Work done at constant pressure can be quickly found by multiplying pressure in MPa directly with volume in m^3 , which gives work in MJ: $W = 1 \text{ MPa} \times 0.03 \text{ m}^3 = 0.03 \text{ MJ} = 30 \text{ kJ}$. Then, $\Delta U = 84 - 30 = 54 \text{ kJ}$.

59. Two reversible heat engines are operated with same heat input and heat output. One of the engine is operated between the temperature limits of T_2 and 200 K & another is between 800 K and T_2 . Then the value of T_2 is

- (A) 100 K
- (B) 200 K
- (C) 400 K
- (D) 500 K

Correct Answer: (C) 400 K

Solution:

Concept: For a completely reversible heat engine operating on a Carnot cycle, the thermal efficiency (η) depends strictly on the absolute temperature limits of the thermal reservoirs it interacts with:

$$\eta = 1 - \frac{T_L}{T_H}$$

Where T_L is the lower sink temperature and T_H is the upper source temperature. The efficiency is also fundamentally defined as:

$$\eta = \frac{W_{\text{net}}}{Q_{\text{in}}} = 1 - \frac{Q_{\text{out}}}{Q_{\text{in}}}$$

The statement tells us that both engines have the same heat input (Q_{in}) and the same heat output (Q_{out}). This implies that their thermal efficiencies are identical ($\eta_1 = \eta_2$). Consequently, for reversible engines, the ratio of sink temperature to source temperature must also be equal across both units.

Step 1: Set up equations for both engines using temperature values.

Let's analyze the configuration of both thermodynamic units:

- **Engine 1:** Fits between source temperature T_2 and sink temperature 200 K.

$$\eta_1 = 1 - \frac{200}{T_2}$$

- **Engine 2:** Fits between source temperature 800 K and sink temperature T_2 .

$$\eta_2 = 1 - \frac{T_2}{800}$$

Step 2: Equate efficiencies to establish the algebraic relationship.

Since $Q_{in,1} = Q_{in,2}$ and $Q_{out,1} = Q_{out,2}$, we set $\eta_1 = \eta_2$:

$$1 - \frac{200}{T_2} = 1 - \frac{T_2}{800}$$

Subtracting 1 from both sides gives:

$$-\frac{200}{T_2} = -\frac{T_2}{800}$$

Eliminating the negative signs from both sides yields:

$$\frac{200}{T_2} = \frac{T_2}{800}$$

Step 3: Cross-multiply and solve for intermediate temperature T_2 .

Cross-multiplying across the fractions:

$$T_2^2 = 800 \times 200$$

$$T_2^2 = 160000$$

Taking the positive square root on both sides since thermodynamic temperatures must remain positive:

$$T_2 = \sqrt{160000} = 400 \text{ K}$$

Thus, the intermediate temperature T_2 is equal to 400 K, which corresponds to Option (C).

Quick Tip: When two reversible engines have the same efficiency (implied here by equal heat inputs and outputs), the intermediate temperature T_2 is always the geometric mean of the extreme temperatures:

$$T_2 = \sqrt{T_{\max} \cdot T_{\min}}. \text{ Here, } T_2 = \sqrt{800 \times 200} = \sqrt{160000} = 400 \text{ K.}$$

60. If a substance is perfect crystal and existing at absolute zero K, then the specific entropy of the substance is

- (A) Infinite
- (B) Positive
- (C) Negative
- (D) Zero

Correct Answer: (D) Zero

Solution:

Concept: The Third Law of Thermodynamics provides an absolute baseline for measuring entropy values. It states that the entropy of a perfect, pure crystalline substance approaches zero as the absolute thermodynamic temperature approaches zero Kelvin (0 K). Mathematically, this condition is expressed as:

$$\lim_{T \rightarrow 0} S = 0$$

Entropy is fundamentally a measure of molecular disorder or randomness within a system. In a perfect crystal at absolute zero, all thermal motion ceases entirely, and the constituent atoms or molecules are arranged in a unique, perfectly ordered geometric microstate configuration ($W = 1$). According to Boltzmann's entropy relation:

$$S = k_B \ln(W) = k_B \ln(1) = 0$$

Therefore, both total and specific entropies drop down to zero.

Step 1: Analyze the criteria provided in the problem statement.

The problem lists two necessary structural conditions:

- The substance must be a **perfect crystal** (meaning there are no structural defects, dislocations, or isotopic mixing).
- The substance exists at **absolute zero Kelvin** ($T = 0 \text{ K}$).

Step 2: Connect the conditions to the Third Law of Thermodynamics.

Because both criteria match the exact definitions established by Planck and Nernst for the Third Law of Thermodynamics, the specific entropy (s) of this substance cannot be positive, negative, or infinite. It must strictly reach an absolute value of zero. This aligns with Option (D).

Quick Tip: The Third Law of Thermodynamics establishes an absolute reference point for entropy. Always remember: Perfect Crystal + Absolute Zero Kelvin (0 K) = Zero Entropy. Any imperfection or temperature rise will increase the entropy value above zero.

61. The maximum amount of low grade energy, which is convertible into work is known as

- (A) Anergy
- (B) Exergy
- (C) Entropy
- (D) Energy loss

Correct Answer: (B) Exergy

Solution:

Concept: Thermodynamics categorizes energy based on its capacity to perform useful mechanical work:

- **High-Grade Energy:** Energy forms that can theoretically be converted entirely into work (e.g., mechanical work, electrical energy, water power).
- **Low-Grade Energy:** Energy forms that cannot be fully converted into work due to Second Law limitations (e.g., thermal energy, heat from combustion).

When a system interacts with the environmental dead state (T_0, P_0), low-grade heat energy (Q) can be split into two fundamental structural portions:

$$\text{Total Energy} = \text{Available Energy (Exergy)} + \text{Unavailable Energy (Anergy)}$$

Exergy represents the maximum portion of this low-grade thermal energy that can be transformed into useful work via an idealized reversible Carnot engine operating down to the dead state temperature T_0 :

$$W_{\max} = \text{Exergy} = Q \left(1 - \frac{T_0}{T} \right)$$

Step 1: Evaluate definitions for each provided option.

Let's analyze the precise engineering meaning of each terminology selection:

1. **Anergy:** The unconvertible, unavailable portion of energy that must be discarded to the surrounding environment as waste heat.
2. **Exergy:** The maximum useful work potential available from a given energy form relative to a specific environmental dead state.
3. **Entropy:** A state function tracking the degree of molecular disorder and quantifying the irreversibility within a system.
4. **Energy loss:** A general term representing energy escaping across boundaries, which does not define a maximum conversion property.

Step 2: Map the question to the correct definition.

The question asks for the maximum amount of low-grade energy that can be converted into work. Based on the fundamental definitions of availability analysis, this is the exact description of **Exergy**, corresponding to Option (B).

Quick Tip: Remember the simple word formula: $\text{Energy} = \text{Exergy} + \text{Anergy}$.
 $\text{Energy} - \text{Exergy} = \text{Available energy (can be converted into work)}$.
 $\text{Energy} - \text{Anergy} = \text{Unavailable energy (cannot be converted into work)}$.

62. A balloon containing an ideal gas is initially kept in an evacuated and insulated room. If the balloon ruptures and the gas fills the entire room, what is the correct statement at the end of the process?

- (A) The internal energy of the gas decreases from its initial value, but the enthalpy remains constant
- (B) The internal energy of the gas increases from its initial value, but the enthalpy remains constant
- (C) Internal energy and enthalpy of the gas remain constant
- (D) Internal energy and enthalpy of the gas increase

Correct Answer: (C) Internal energy and enthalpy of the gas remain constant

Solution:

Concept: The phenomenon of a gas expanding into an evacuated space is known as a **Free Expansion** or an unresisted expansion. Let us apply the First Law of Thermodynamics to this system:

$$Q = \Delta U + W$$

We define our system boundary to encompass the entire insulated room. Let's evaluate the conditions:

- **Insulated Room:** No thermal interactions occur across the system boundaries, meaning heat transfer is zero ($Q = 0$).
- **Evacuated Space:** The gas expands against a vacuum, meaning there is zero resisting pressure ($P_{\text{external}} = 0$). Since work is defined as $W = \int P_{\text{external}} dV$, the work performed by the system is zero ($W = 0$).

Substituting these boundary conditions into the First Law yields:

$$0 = \Delta U + 0 \Rightarrow \Delta U = 0 \Rightarrow U_1 = U_2$$

Thus, the internal energy remains constant during a free expansion.

Step 1: Connect internal energy stability to ideal gas properties.

For an ideal gas, Joule's law states that internal energy is strictly a function of temperature alone, $U = f(T)$. Because $\Delta U = 0$ for this process, the temperature of the ideal gas must also remain unchanged:

$$T_1 = T_2$$

Step 2: Evaluate the change in enthalpy (H).

Enthalpy is defined by the property relation:

$$H = U + PV$$

Using the ideal gas equation of state, $PV = mRT$, we can substitute this directly into the enthalpy definition:

$$H = U + mRT$$

Since both internal energy (U) and temperature (T) are constant for this ideal gas process, the product mRT remains constant. Consequently, the enthalpy H must also remain unchanged

throughout the free expansion:

$$\Delta H = 0 \Rightarrow H_1 = H_2$$

Step 3: Match results to the options.

Both the internal energy and the enthalpy of the ideal gas remain perfectly constant from the beginning to the end of this process. This matches Option (C).

Quick Tip: Free expansion of an ideal gas is simultaneously **isothermal** ($T = \text{constant}$), **isenuclic/isothermal-internal-energy** ($U = \text{constant}$), and **isenthalpic** ($H = \text{constant}$). However, it is highly irreversible, so entropy increases ($\Delta S > 0$).

63. Based on second law of thermodynamics the loss in available energy is being quantified as

- (A) Irreversibility
- (B) Decrease in available energy
- (C) Entropy loss
- (D) Energy loss

Correct Answer: (A) Irreversibility

Solution:

Concept: Real thermodynamic processes are accompanied by irreversibilities, such as friction, unresisted expansion, heat transfer across finite temperature differences, and mixing. These factors degrade the quality of energy, diminishing its capacity to perform useful mechanical work. According to the **Gouy-Stodola Theorem**, the loss of available energy (or the destruction of exergy) is directly proportional to the total entropy generation (ΔS_{gen}) within the universe during a process:

$$I = W_{\text{max}} - W_{\text{actual}} = T_0 \cdot \Delta S_{\text{gen}}$$

Where:

- I is defined as the **Irreversibility** of the process.
- T_0 is the ambient surrounding temperature (dead state temperature).
- ΔS_{gen} is the total entropy generation of the system and its surroundings.

This loss of work potential is explicitly quantified as irreversibility.

Step 1: Examine the physical interpretation of Irreversibility.

Irreversibility measures how much potential work is permanently lost due to non-ideal behavior in a real-world process. It acts as the direct metric for the degradation of high-grade work potential into low-grade thermal waste.

Step 2: Contrast with other options.

- Option (B) simply restates the question text rather than naming the established thermodynamic metric.
- Option (C) and (D) are descriptive properties but do not represent the exact quantitative term for lost work potential.

Thus, the technical term used to quantify the loss in available energy according to the second law is **Irreversibility**, which matches Option (A).

Quick Tip: Gouy-Stodola Theorem is key here: $I = T_0 \times \Delta S_{\text{universe}}$. It represents the direct translation of the Second Law of Thermodynamics into lost work potential.

64. The clearance is provided in the reciprocating compressor with clearance ratio of 'k'. If the pressure ratio is $\frac{p_2}{p_1}$, then the volumetric efficiency in terms of clearance ratio k, is

- (A) $1 + k - k \left(\frac{p_2}{p_1} \right)^{(1/n)}$
- (B) $1 - k + k \left(\frac{p_2}{p_1} \right)^{(1/n)}$
- (C) $1 - k - k \left(\frac{p_2}{p_1} \right)^{(1/n)}$
- (D) $1 + k + k \left(\frac{p_2}{p_1} \right)^{(1/n)}$

Correct Answer: (A) $1 + k - k \left(\frac{p_2}{p_1} \right)^{(1/n)}$

Solution:

Concept: In a reciprocating compressor, clearance volume (V_c) is the small space left at the top of the cylinder head to prevent mechanical impact between the piston and the valves. The

clearance ratio (k) is defined relative to the swept stroke volume (V_s):

$$k = \frac{V_c}{V_s}$$

Volumetric efficiency (η_v) measures the effectiveness of a compressor in drawing in fresh gas. It is defined as the ratio of the actual volume of fresh air sucked into the cylinder during intake ($V_1 - V_4$) to the theoretical swept stroke volume (V_s):

$$\eta_v = \frac{V_1 - V_4}{V_s}$$

During the delivery stroke, high-pressure gas remains in the clearance volume at pressure P_2 . Before the intake valve can open, this trapped gas must expand polytropically down to the intake pressure P_1 according to the relationship $P_2 V_c^n = P_1 V_4^n$.

Step 1: Express expanded volume V_4 in terms of clearance parameters.

From the polytropic expansion relation:

$$\frac{V_4}{V_c} = \left(\frac{P_2}{P_1} \right)^{1/n} \Rightarrow V_4 = V_c \left(\frac{P_2}{P_1} \right)^{1/n}$$

Step 2: Expand the expression for volumetric efficiency.

Substitute total volume definitions, noting that the maximum volume at Bottom Dead Center is $V_1 = V_s + V_c$:

$$\eta_v = \frac{(V_s + V_c) - V_4}{V_s} = \frac{V_s + V_c - V_c \left(\frac{P_2}{P_1} \right)^{1/n}}{V_s}$$

Separating the fraction term-by-term yields:

$$\eta_v = \frac{V_s}{V_s} + \frac{V_c}{V_s} - \frac{V_c}{V_s} \left(\frac{P_2}{P_1} \right)^{1/n}$$

Step 3: Substitute the clearance ratio $k = \frac{V_c}{V_s}$.

Replacing $\frac{V_c}{V_s}$ with k :

$$\eta_v = 1 + k - k \left(\frac{P_2}{P_1} \right)^{1/n}$$

This derived expression matches Option (A).

Quick Tip: Memory trick: As the pressure ratio $\frac{P_2}{P_1}$ increases, the trapped high-pressure gas expands more, taking up more space and reducing the volume of fresh air that can be drawn in. Therefore, the volumetric efficiency must drop as pressure ratio increases, meaning the formula must have a negative sign before the pressure ratio term: $1 + k - k(\dots)$.

65. In an axial flow compressor, if the air stream is not able to follow the blade contour such phenomena is known as

- (A) Surging
- (B) Haulling
- (C) Stalling
- (D) Chocking

Correct Answer: (C) Stalling

Solution:

Concept: Aerodynamic flow behavior across the rotating blades of an axial flow compressor depends heavily on the design **angle of incidence**. If flow conditions deviate from design specifications (e.g., due to reduced mass flow rate), the angle of incidence increases. When this angle exceeds a critical threshold, the air stream can no longer adhere smoothly to the curved contour of the blade profile. This causes the boundary layer to separate from the suction surface of the blade, leading to a phenomenon known as **Stalling**. This flow separation creates highly turbulent wake regions, reduces pressure rise, and causes structural vibrations.

Step 1: Distinguish between compressor flow anomalies.

Let us review the operational phenomena listed in the options to differentiate them clearly:

- **Surging:** A complete breakdown of flow through the entire compressor, characterized by violent, cyclic pressure oscillations and flow reversal. It is a system-level instability.
- **Stalling:** A localized aerodynamic flow separation from the blade surface when the air stream cannot follow the blade contour. It can occur on a single blade or a group of blades (rotating stall).
- **Choking:** Occurs when the fluid velocity reaches the local speed of sound ($M = 1$) at the minimum area cross-section (throat), fixing the mass flow rate at its maximum limit.

Step 2: Match the question description to the correct mechanism.

The specific phrase "air stream is not able to follow the blade contour" refers directly to

boundary layer separation on an airfoil profile, which defines aerodynamic **Stalling**. This corresponds to Option (C).

Quick Tip: Think of compressor blades as airplane wings. When an airplane wing tilts too far up (high angle of attack) and the air can't follow its curved shape, the plane loses lift and "stalls". The exact same aerodynamic definition applies to compressor blades!

66. In an ideal Brayton power cycle operated for the gas turbine plant, if T_1 is the minimum temperature, T_3 is the maximum temperature and pressure ratio is r_p , then the net work ratio is given as

- (A) $1 - \left(\frac{T_3}{T_1}\right)(r_p)^{(\gamma-1)/\gamma}$
- (B) $1 - \left(\frac{T_1}{T_3}\right)(r_p)^{(\gamma-1)/\gamma}$
- (C) $1 - \left(\frac{T_3}{T_1}\right)(r_p)^{\gamma/(\gamma-1)}$
- (D) $1 - \left(\frac{T_1}{T_3}\right)(r_p)^{\gamma/(\gamma-1)}$

Correct Answer: (B) $1 - \left(\frac{T_1}{T_3}\right)(r_p)^{(\gamma-1)/\gamma}$

Solution:

Concept: The work ratio (r_w) in a gas turbine power plant cycle is defined as the ratio of the net useful work output to the total work produced by the turbine component:

$$r_w = \frac{W_{\text{net}}}{W_t} = \frac{W_t - |W_c|}{W_t} = 1 - \frac{|W_c|}{W_t}$$

Let's analyze the components using standard Brayton cycle notations:

- T_1 : Compressor inlet temperature (minimum cycle temperature)
- T_2 : Compressor outlet temperature
- T_3 : Turbine inlet temperature (maximum cycle temperature)
- T_4 : Turbine exhaust temperature

Assuming constant specific heats, the compressor work input and turbine work output are given by:

$$|W_c| = C_p(T_2 - T_1), \quad W_t = C_p(T_3 - T_4)$$

Step 1: Set up the temperature equations using the pressure ratio r_p .

For isentropic compression ($1 \rightarrow 2$) and expansion ($3 \rightarrow 4$) processes, we apply the ideal gas temperature-pressure relations:

$$\frac{T_2}{T_1} = (r_p)^{\frac{\gamma-1}{\gamma}} \Rightarrow T_2 = T_1(r_p)^{\frac{\gamma-1}{\gamma}}$$

$$\frac{T_3}{T_4} = (r_p)^{\frac{\gamma-1}{\gamma}} \Rightarrow T_4 = \frac{T_3}{(r_p)^{\frac{\gamma-1}{\gamma}}}$$

Step 2: Substitute these expressions into the work ratio equation.

$$r_w = 1 - \frac{C_p(T_2 - T_1)}{C_p(T_3 - T_4)} = 1 - \frac{T_2 - T_1}{T_3 - T_4}$$

Factor out T_1 from the numerator and T_3 from the denominator:

$$r_w = 1 - \frac{T_1 \left(\frac{T_2}{T_1} - 1 \right)}{T_3 \left(1 - \frac{T_4}{T_3} \right)}$$

Notice from the isentropic relations that:

$$\frac{T_2}{T_1} = (r_p)^{\frac{\gamma-1}{\gamma}} \quad \text{and} \quad \frac{T_4}{T_3} = \frac{1}{(r_p)^{\frac{\gamma-1}{\gamma}}}$$

Let's define $x = (r_p)^{\frac{\gamma-1}{\gamma}}$ to make simplification easier:

$$r_w = 1 - \frac{T_1(x - 1)}{T_3 \left(1 - \frac{1}{x} \right)} = 1 - \frac{T_1(x - 1)}{T_3 \left(\frac{x-1}{x} \right)}$$

Canceling the common term $(x - 1)$ from both the numerator and denominator simplifies the equation to:

$$r_w = 1 - \frac{T_1 \cdot x}{T_3}$$

Step 3: Re-substitute the value of x .

Replacing x back with its full pressure-ratio form:

$$r_w = 1 - \left(\frac{T_1}{T_3} \right) (r_p)^{\frac{\gamma-1}{\gamma}}$$

This matches Option (B).

Quick Tip: To verify your result quickly, think about extreme values: if the pressure ratio $r_p = 1$, no compression occurs, so compressor work is zero, meaning the work ratio must equal 1. Plugging $r_p = 1$ into option (B) gives $1 - \frac{T_1}{T_3}(1) = 1 - \frac{T_1}{T_3}$, which correctly shows the work balance at that limit. Also, remember that the isentropic exponent for pressure ratio is always $\frac{\gamma-1}{\gamma}$, which immediately eliminates options (C) and (D).

67. If regenerator is added to a Brayton cycle, the thermal efficiency is improved due to

- (A) Increase in temperature of gases at turbine inlet
- (B) Decrease in compressor input
- (C) Increase in turbine output
- (D) Decrease in quantity of heat in combustion chamber

Correct Answer: (D) Decrease in quantity of heat in combustion chamber

Solution:

Concept: A regenerator (or recuperator) is a counter-flow heat exchanger incorporated into a gas turbine plant to recover thermal energy from the hot exhaust gases leaving the turbine. This recovered energy is transferred directly to the cooler compressed air exiting the compressor before it enters the combustion chamber. The thermal efficiency (η) of any power cycle is given by the ratio of net work output to heat input:

$$\eta = \frac{W_{\text{net}}}{Q_{\text{in}}} = \frac{W_t - |W_c|}{Q_{\text{in}}}$$

Adding an ideal regenerator has the following effects on the cycle parameters:

- The compressor work input (W_c) remains unchanged because the compressor operates between the same pressure limits.
- The turbine work output (W_t) remains unchanged because the turbine expansion profile is unaltered.
- Consequently, the net work output ($W_{\text{net}} = W_t - |W_c|$) stays constant.

Step 1: Analyze how regeneration affects heat input (Q_{in}).

Because the compressed air is preheated by the turbine exhaust gases before entering the combustor, less fuel needs to be burned to raise the air temperature up to the maximum

required turbine inlet temperature (T_3). Therefore, the external heat input required in the combustion chamber (Q_{in}) decreases significantly.

Step 2: Evaluate the impact on thermal efficiency.

Looking back at our efficiency equation:

$$\eta \uparrow = \frac{W_{net} \text{ (constant)}}{Q_{in} \text{ (decreases } \downarrow)}$$

Since a smaller value divides the constant net work output, the overall thermal efficiency increases. This improvement stems directly from the reduction in fuel energy required within the combustion chamber, which perfectly matches Option (D).

Quick Tip: Regeneration = Heat Recovery from exhaust. It doesn't alter the compressor or turbine work profiles ($W_{net} = \text{constant}$). Its sole benefit is preheating the air, which directly cuts down on the fuel required in the combustion chamber ($Q_{in} \downarrow$).

68. Generally, in spark ignition engines, the compression ratio cannot be enhanced beyond certain limit because to avoid

- (A) Auto ignition
- (B) Over heating
- (C) Knocking
- (D) Preheating

Correct Answer: (C) Knocking

Solution:

Concept: In Spark Ignition (SI) engines, a homogeneous mixture of fuel and air is drawn into the cylinder during the intake stroke and compressed during the compression stroke. Combustion is designed to be initiated precisely by an electrical spark from the spark plug, producing a smooth flame front that travels across the combustion chamber. The compression ratio (r) is defined as:

$$r = \frac{V_{total}}{V_{clearance}}$$

As the compression ratio increases, the temperature (T_2) and pressure (P_2) of the fuel-air

mixture at the end of the compression stroke rise following the ideal gas relationship:

$$T_2 = T_1 \cdot (r)^{\gamma-1}$$

Step 1: Explore the mechanism of Knocking.

If the compression ratio is set too high, the temperature and pressure of the unburned end-gas mixture exceed the self-ignition threshold of the fuel. Instead of waiting for the spark-initiated flame front to reach them, these localized pockets of end-gas undergo rapid spontaneous auto-ignition. This sudden explosion creates high-amplitude, high-frequency shock waves that bounce off the cylinder walls, producing a sharp metallic ringing sound known as **Knocking** or detonation.

Step 2: Identify the primary operational limitation.

Severe knocking can cause engine overheating, erode piston crowns, damage cylinder head gaskets, and cause mechanical failure. To prevent this destructive knocking phenomenon, the compression ratio in commercial petrol (SI) engines is strictly limited, typically ranging between 6 and 11. This constraint aligns with Option (C).

Quick Tip: While auto-ignition is the chemical trigger, **Knocking** is the actual structural phenomenon and operational hazard that engineers must limit the compression ratio to avoid. - Higher compression ratio in SI engines → premature auto-ignition of end-gas → Severe **Knocking**.

69. In a CI engine, the delay period in the combustion process mainly corresponds to

- (A) Time taken for fuel injection into the cylinder
- (B) Time taken between the start of fuel injection and start of combustion
- (C) Time between the end of fuel injection and start of combustion
- (D) Time between the start of combustion and attaining maximum pressure

Correct Answer: (B) Time taken between the start of fuel injection and start of combustion

Solution:

Concept: Combustion in a Compression Ignition (CI) engine differs fundamentally from an SI engine. In a CI engine, pure air is compressed to high pressures and temperatures. Liquid fuel is then injected into this hot environment near the end of the compression stroke, where it auto-ignites due to the high temperature of the air. The combustion process is broken down

into distinct sequential stages:

1. **Delay Period (Ignition Delay)**
2. Period of Rapid or Uncontrolled Combustion
3. Period of Controlled Combustion
4. Period of After-Burning

Step 1: Define the components of the Delay Period.

The delay period is the time interval between the exact instant the first droplets of fuel exit the injector nozzle and the initiation of chemical flaring (the start of actual combustion). This delay period consists of two parts:

- **Physical Delay:** The time required for the liquid fuel jet to atomize into tiny droplets, vaporize, mix with the hot compressed air, and raise its temperature to the self-ignition point.
- **Chemical Delay:** The time required for pre-flame chemical reactions to occur, generating enough localized heat to initiate self-sustaining combustion.

Step 2: Map the definition to the choices.

Based on this thermodynamic breakdown, the delay period corresponds to the time lag between the ****start of fuel injection**** and the ****start of combustion****. This matches Option (B).

Quick Tip: Delay period = "Waiting time". It tracks how long the fuel takes to prepare itself (vaporize and chemically react) after entering the cylinder before it actually catches fire. Thus, it spans from the start of injection to the start of combustion.

70. The refrigerant and absorbent used in Vapour absorption refrigeration system respectively are

- (A) Water and water
- (B) Lithium bromide and water
- (C) Ammonia and lithium bromide
- (D) Ammonia and water

Correct Answer: (D) Ammonia and water

Solution:

Concept: A Vapour Absorption Refrigeration System (VARs) replaces the power-consuming mechanical compressor of a standard Vapour Compression Refrigeration System (VCRS) with an absorber, a pump, and a generator. It utilizes thermal energy (low-grade heat) rather than mechanical work to drive the refrigeration cycle. This process requires a working pair consisting of two fluids that have high mutual solubility:

- **Refrigerant:** The fluid that evaporates at low temperatures to absorb heat from the refrigerated space, producing the cooling effect.
- **Absorbent:** The fluid that absorbs the refrigerant vapor in the low-pressure absorber and carries it to the high-pressure generator.

Step 1: Evaluate common VARs working pairs.

The two most common commercial fluid combinations used in absorption systems are:

1. Aqua-Ammonia System:

- Refrigerant = Ammonia (NH_3)
- Absorbent = Water (H_2O)

2. Lithium Bromide-Water System:

- Refrigerant = Water (H_2O)
- Absorbent = Lithium Bromide solution ($LiBr$)

Step 2: Analyze the options based on standard pairs.

Let us check the ordering required by the question prompt ("refrigerant and absorbent respectively"):

- Option (B) lists "Lithium bromide and water". In a LiBr system, water acts as the refrigerant and Lithium Bromide acts as the absorbent. Since the option reverses this order, it is incorrect.
- Option (D) lists "Ammonia and water". Here, Ammonia is correctly identified as the refrigerant and water as the absorbent, matching the requested sequence.

Therefore, the correct pair matching the "refrigerant and absorbent" order is ****Ammonia and water****, corresponding to Option (D).

Quick Tip: Pay close attention to the word "respectively"! In the pairs: 1. Ammonia (Refrigerant) + Water (Absorbent) → Used for low-temperature freezing application. 2. Water (Refrigerant) + Lithium Bromide (Absorbent) → Limited to air conditioning because water freezes at 0°C.

71. An air conditioning system mixes 30% fresh air at 35 °C DBT and 50% RH with 70% recirculated air at 35 °C DBT and 50% RH. Which of the following statements about "mixed air condition" is true?

- (A) The mixed air DBT will be 25 °C
- (B) The mixed air humidity ratio will be equal to the arithmetic mean of two humidity ratios
- (C) The mixed air state point will lie on the straight line joining the two states on the psychrometric chart
- (D) The relative humidity of the mixed air will be equal to 50%

Correct Answer: (C) The mixed air state point will lie on the straight line joining the two states on the psychrometric chart

Solution:

Concept: The process of adiabatic mixing of two moist air streams is a fundamental operation in heating, ventilating, and air-conditioning (HVAC) systems. When two streams at different states mix, the governing laws are based on the conservation of mass for dry air, conservation of mass for water vapor, and conservation of energy (enthalpy). Let State 1 be the fresh air stream and State 2 be the recirculated air stream. The mixing equations are modeled as follows:

- **Mass balance of dry air:** $\dot{m}_{a1} + \dot{m}_{a2} = \dot{m}_{a3}$
- **Mass balance of water vapor:** $\dot{m}_{a1}\omega_1 + \dot{m}_{a2}\omega_2 = \dot{m}_{a3}\omega_3$
- **Energy balance:** $\dot{m}_{a1}h_1 + \dot{m}_{a2}h_2 = \dot{m}_{a3}h_3$

By combining these equations, we obtain the ratio of distances on the psychrometric chart:

$$\frac{\dot{m}_{a1}}{\dot{m}_{a2}} = \frac{\omega_2 - \omega_3}{\omega_3 - \omega_1} = \frac{h_2 - h_3}{h_3 - h_1} \approx \frac{T_2 - T_3}{T_3 - T_1}$$

Geometrically, this implies that the final state point (State 3) of the mixture must always lie strictly on the straight line connecting the initial state points (State 1 and State 2) on a standard psychrometric chart, dividing the line segment inversely proportional to the mass flow rates of the dry air.

Step 1: Analyze the specific given conditions of the streams.

Let's carefully examine the properties of the two air streams provided in the problem statement:

- **Fresh Air Stream (State 1):** Mass fraction = 30% (0.30), Dry Bulb Temperature (T_1) = 35 °C, Relative Humidity (RH_1) = 50%
- **Recirculated Air Stream (State 2):** Mass fraction = 70% (0.70), Dry Bulb Temperature (T_2) = 35 °C, Relative Humidity (RH_2) = 50%

Notice that State 1 and State 2 are completely identical! They have exactly the same Dry Bulb Temperature (DBT) and the same Relative Humidity (RH). Therefore, on the psychrometric chart, State point 1 and State point 2 are located at the exact same spatial coordinate point.

Step 2: Calculate the mixed air properties using the governing equations.

Let's determine the properties of the mixture (State 3) using our mass and energy conservation principles:

- **Mixed Air Dry Bulb Temperature (T_3):**

$$T_3 = 0.30 \cdot T_1 + 0.70 \cdot T_2$$

$$T_3 = (0.30 \times 35) + (0.70 \times 35) = 10.5 + 24.5 = 35 \text{ °C}$$

- **Mixed Air Specific Humidity (ω_3):** Since State 1 and State 2 are identical, their specific humidities are equal ($\omega_1 = \omega_2 = \omega$). Thus:

$$\omega_3 = 0.30 \cdot \omega_1 + 0.70 \cdot \omega_2 = 0.30\omega + 0.70\omega = \omega$$

Because the final temperature is 35 °C and the moisture content is exactly the same, the relative humidity of the mixed air stream remains completely unchanged at 50%.

Step 3: Evaluate each option for universal and specific validity.

Let's evaluate the physical accuracy of the options to find the best selection:

- **Option (A):** Incorrect. The mixed air temperature is calculated as 35 °C, not 25 °C.
- **Option (B):** Incorrect. The mixture specific humidity is a weighted average based on mass flow fractions ($\omega_3 = 0.3\omega_1 + 0.7\omega_2$). It is only equal to the simple arithmetic mean ($\frac{\omega_1 + \omega_2}{2}$) when the two streams have identical mass flow rates (50% fresh air and 50% recirculated air), which is not the case here.

- **Option (D):** While the relative humidity is indeed 50% for this unique specific case because the two entering streams are identical, this is a conditional result.
- **Option (C):** **The mixed air state point will lie on the straight line joining the two states on the psychrometric chart.** This statement is a fundamental law of psychrometrics. It is universally true for the mixing of any two moist air streams regardless of whether their initial states are different or identical (if identical, the line collapses to a single point, which still technically satisfies the condition). Therefore, Option (C) represents the core scientific principle being tested.

Quick Tip: For the adiabatic mixing of two air streams, remember that the governing conservation laws map out linearly on a psychrometric chart. As a result, the mixed state point ****always**** lies on the straight line connecting the two initial states. The distance from each initial state to the final state is inversely proportional to their respective air masses.

72. The whirl component of turbine blades is zero for the following turbine.

- (A) Parson's reaction turbine
- (B) Curtis turbine
- (C) Axial flow turbine
- (D) Impulse turbine with equal blade angles at the inlet and outlet

Correct Answer: (C) Axial flow turbine

Solution:

Concept: In turbomachinery, the total power transferred between a fluid stream and the rotating blades is governed by the fundamental Euler Turbine Equation. The work done per unit mass of fluid (W) depends on the velocities at the blade inlet (subscript 1) and outlet (subscript 2):

$$W = u_1 V_{w1} \pm u_2 V_{w2}$$

Where:

- u_1, u_2 represent the linear peripheral velocity of the rotating blades at the inlet and outlet, respectively.

- V_{w1}, V_{w2} represent the **whirl components** (or tangential components) of the absolute fluid velocity at the inlet and outlet, which align with the direction of blade rotation and contribute directly to torque generation.

For a pure axial flow machine with an axial discharge condition at the exit, the absolute fluid velocity is directed purely parallel to the rotational shaft axis. Consequently, the tangential component at the outlet is reduced to zero ($V_{w2} = 0$). Furthermore, if a turbine is categorized generally or designed as a pure axial flow turbine where no tangential deflection of the fluid occurs relative to the blade movement, the net whirl component across the blades is effectively absent.

Step 1: Evaluate the velocity profiles of the listed options.

Let's analyze the velocity vector diagrams for each type of turbine provided:

- **Parson's Reaction Turbine:** This is an axial flow reaction turbine with 50% degree of reaction. The fluid expands continuous across both fixed and moving blades. Due to the oblique blade angles, the fluid enters with a high whirl velocity (V_{w1}) and exits with a residual tangential component. Neither component is zero.
- **Curtis Turbine:** This is a velocity-compounded impulse turbine. Steam expands completely in a nozzle, creating a high initial absolute velocity. It passes through multiple rows of moving and fixed guide blades. The whirl velocity changes dramatically across each moving row to maximize kinetic energy extraction, so it is definitely not zero.
- **Impulse Turbine with equal blade angles:** If an impulse turbine has symmetrical blades ($\beta_1 = \beta_2$), the relative velocity magnitudes remain equal if friction is neglected ($V_{r1} = V_{r2}$). However, the change in absolute whirl velocity ($\Delta V_w = V_{w1} + V_{w2}$) reaches a substantial value to maximize work output.
- **Axial Flow Turbine:** In a generic or idealized axial flow stage designed without swirl (or with zero net whirl component), the fluid enters and leaves the blade rows parallel to the axis of rotation. The whirl component (V_w), which acts tangential to the rotor circumference, is zero.

Step 2: Match the question constraint to the correct classification.

The problem specifies that the whirl component is zero. This matches the ideal flow configuration of an axial flow system with zero tangential velocity interactions, aligning with Option (C).

Quick Tip: Whirl velocity (V_w) is the tangential velocity component responsible for creating torque and turning the rotor. If a turbine operates with a zero whirl component, the flow enters and exits in a purely axial direction without any tangential rotation relative to the blade motion.

73. A single stage impulse turbine with a diameter of 1.2 m runs at 3000 rpm. If the blade speed ratio is 0.42, the inlet velocity of steam will be (in ms^{-1})

- (A) 450
- (B) 900
- (C) 200
- (D) 600

Correct Answer: (A) 450

Solution:

Concept: An impulse turbine extracts kinetic energy from a high-velocity fluid jet exiting a nozzle. To evaluate its performance, we use two key velocity parameters:

- **Blade Linear Velocity (u):** The tangential speed of the rotating blades, calculated from the rotor diameter (D) and rotational speed (N in rpm) as:

$$u = \frac{\pi \cdot D \cdot N}{60}$$

- **Blade Speed Ratio (ρ):** The ratio of the blade linear velocity to the absolute inlet velocity of the steam jet (V_1) exiting the nozzle:

$$\rho = \frac{u}{V_1}$$

By determining the blade linear speed from the physical dimensions and rotational rate, we can use the blade speed ratio to find the absolute inlet velocity of the steam.

Step 1: Extract the given parameters and convert them to standard SI units.

The values specified in the problem statement are:

- Mean Diameter of the turbine rotor, $D = 1.2$ m
- Rotational Speed, $N = 3000$ rpm
- Blade Speed Ratio, $\rho = 0.42$

Step 2: Compute the linear peripheral velocity of the turbine blades (u).

Using the standard rotational-to-linear conversion formula:

$$u = \frac{\pi \cdot D \cdot N}{60}$$

Substitute the values into the equation:

$$u = \frac{\pi \times 1.2 \times 3000}{60}$$

Simplify the fraction by dividing 3000 by 60:

$$\frac{3000}{60} = 50$$

Now substitute this back into the product:

$$u = \pi \times 1.2 \times 50$$

$$u = 1.2 \times 50 \times \pi = 60\pi \text{ m/s}$$

Using the standard approximation for $\pi \approx 3.14159$:

$$u = 60 \times 3.14159 = 188.495 \text{ m/s}$$

Step 3: Calculate the absolute inlet velocity of the steam (V_1).

From the definition of the blade speed ratio:

$$\rho = \frac{u}{V_1} \Rightarrow V_1 = \frac{u}{\rho}$$

Substitute the calculated blade velocity ($u = 188.495 \text{ m/s}$) and the given speed ratio ($\rho = 0.42$):

$$V_1 = \frac{188.495}{0.42}$$

Let's calculate this division step-by-step:

$$V_1 \approx 448.798 \text{ m/s}$$

Rounding this value to the nearest integer option gives approximately 450 m/s. This matches

Option (A).

Quick Tip: To perform the calculation quickly without a calculator, approximate $\pi \approx \frac{22}{7}$ or 3.14. - $u = 60 \times 3.14 = 188.4$ m/s. - Then, $V_1 = \frac{188.4}{0.42} \approx \frac{188.4}{0.4} = 471$ m/s. Since 0.42 is slightly larger than 0.4, the true value must be slightly less than 471, pointing directly to 450 m/s (Option A).

74. In order to increase the corrosion resistance of stainless steel, the following element is added

- (A) Nickel
- (B) Molybdenum
- (C) Carbon
- (D) Chromium

Correct Answer: (D) Chromium

Solution:

Concept: Stainless steel is an iron-based alloy known for its exceptional corrosion resistance. This protective capability is achieved by adding specific alloying elements that alter the surface chemistry of the metal. The primary element responsible for this transformation is **Chromium (Cr)**. To classify a steel alloy as "stainless," it must contain a minimum of approximately 10.5% to 12% chromium by weight. When exposed to oxygen in the atmosphere or water, chromium reacts preferentially to form an exceedingly thin, continuous, adherent, and chemically unreactive (passive) layer of chromium oxide (Cr_2O_3) on the surface of the steel. This invisible film acts as a highly effective physical barrier that seals the underlying iron matrix, stopping oxygen and moisture from causing oxidation (rusting). If the surface is scratched or damaged, the passive layer self-heals instantly in the presence of oxygen.

Step 1: Examine the metallurgical role of each listed option.

Let's analyze the typical engineering purpose of each alloying element in steel design:

- **Nickel (Ni):** Added to stabilize the austenite (FCC) crystal structure down to room temperature, which enhances ductility, toughness, and formability, while also providing moderate resistance to acid environments.
- **Molybdenum (Mo):** Added in smaller quantities (typically 2-4%) to improve resistance to localized pitting and crevice corrosion in chloride-rich environments (such as seawater).

However, it is used as a secondary addition alongside chromium, not as the primary defining element.

- **Carbon (C):** Added primarily to increase hardness, tensile strength, and wear resistance through carbide formation and martensitic transformation. Higher carbon content can actually reduce corrosion resistance by forming chromium carbides, which depletes chromium at the grain boundaries (a problem known as sensitization).
- **Chromium (Cr):** The essential primary alloying element required to establish baseline corrosion resistance and passivity in stainless steels.

Step 2: Select the primary element based on the question context.

Since Chromium is the fundamental element required to form the passive protective oxide film that defines stainless steel, it is the correct choice. This matches Option (D).

Quick Tip: Chromium is the core element that makes stainless steel "stainless"! Always remember: Minimum 10.5% Chromium → Formation of a passive Cr_2O_3 surface oxide layer → Complete protection against corrosion and rust.

75. If the material has the tendency to fracture without any appreciable deformation is known as

- (A) Stiffness
- (B) Brittleness
- (C) Toughness
- (D) Malleability

Correct Answer: (B) Brittleness

Solution:

Concept: When materials are subjected to mechanical tensile or compressive forces, they exhibit distinct deformation and fracture characteristics. These behaviors are classified into standard engineering material properties based on their stress-strain response:

- **Ductile Materials:** Undergo significant permanent plastic deformation (yielding and necking) before structural failure occurs (e.g., mild steel, copper, aluminum).

- **Brittle Materials:** Characterized by little to no plastic deformation prior to fracture. When loaded beyond their elastic limit, they fail suddenly through rapid crack propagation along cleavage planes, showing a flat, crystalline fracture surface (e.g., cast iron, glass, ceramics).

Step 1: Evaluate the definitions of all choices.

Let us review the precise definition for each property selection listed:

- **Stiffness:** The capacity of a material to resist elastic deformation when a load is applied. It is quantified by Young's Modulus (E). A stiff material requires high stress to produce a small elastic strain, but this property does not describe the fracture mechanism.
- **Brittleness:** The tendency of a solid material to break or fracture suddenly with little or no appreciable plastic deformation when subjected to stress.
- **Toughness:** The capacity of a material to absorb energy and deform plastically before fracturing. It corresponds to the total area under the stress-strain curve from initial loading up to the point of failure.
- **Malleability:** A specific compressive ductility property that enables a material to be hammered, rolled, or pressed into thin sheets without cracking.

Step 2: Match the question description to the correct property.

The problem description explicitly highlights a "tendency to fracture without any appreciable deformation". This fits the exact engineering definition of **Brittleness**, matching Option (B).

Quick Tip: Think of the contrast between a piece of chalk and a copper wire. If you bend the copper wire, it deforms significantly before breaking (ductile/malleable). If you try to bend the chalk, it snaps instantly with zero visible deformation. That sudden snap defines **Brittleness**.

76. The heat treatment process used to enrich the steel surface by addition of Nitrogen is known as

- (A) Annealing
- (B) Surface hardening
- (C) Tempering

(D) Nitriding

Correct Answer: (D) Nitriding

Solution:

Concept: Surface hardening (or case hardening) techniques alter the chemical composition and microstructure of a component's outer surface layer while maintaining a tough, ductile core. This is achieved by diffusing specific interstitial elements into the steel matrix at elevated temperatures. When the surface layer is enriched with **Nitrogen**, the process is called **Nitriding**. During nitriding, the steel part is exposed to a nitrogen-rich environment (such as ammonia gas, NH_3 , or a molten cyanide salt bath) at temperatures typically ranging from 500 °C to 575 °C. The nitrogen atoms diffuse into the ferrite phase of the steel and react with alloying elements (such as Aluminum, Chromium, or Molybdenum) to form extremely hard, finely dispersed interstitial alloy nitrides. This process significantly increases surface hardness, wear resistance, and fatigue life without requiring a subsequent liquid quench.

Step 1: Review the purpose of each listed heat treatment process.

Let's analyze the standard definitions of the choices to differentiate their underlying chemical and structural mechanisms:

- **Annealing:** A thermal process where a material is heated to a high temperature, held there, and then cooled very slowly inside the furnace. Its primary goals are to soften the metal, relieve internal residual stresses, and improve ductility. It does not alter the surface chemistry.
- **Surface Hardening:** A broad category that includes any process used to harden the outer shell of a component while leaving the core soft. Examples include flame hardening, induction hardening, carburizing, and nitriding. While nitriding falls under this category, it is a generic term rather than the specific process defined by adding nitrogen.
- **Tempering:** A process performed after quenching to reduce excess hardness, restore toughness, and relieve internal stresses in martensitic steel by heating it below the lower critical temperature.
- **Nitriding:** The specific case-hardening process that relies explicitly on the thermal diffusion and enrichment of **Nitrogen** into the surface layer.

Step 2: Choose the most precise answer matching the chemical element.

Since the question asks for the specific process that introduces Nitrogen to enrich the surface,

Nitriding is the most accurate and precise selection. This corresponds to Option (D).

Quick Tip: The clue is right in the name! - **Nitr**ogen addition → **Nitr**iding. - **Carb**on addition → **Carb**urizing. Nitriding produces an exceptionally hard case because the formed nitrides are inherently hard and stable, avoiding the need for an aggressive oil or water quench.

77. In a linearly hardening plastic material, the true stress beyond the initial yielding

- (A) Decreases linearly with the true strain
- (B) Increases linearly with the true strain
- (C) Remains constant with true strain
- (D) Increases exponentially with true strain

Correct Answer: (B) Increases linearly with true strain

Solution:

Concept: When analyzing the plastic behavior of metals beyond their initial elastic limit, constitutive material models are used to simplify real stress-strain behavior. A **linearly hardening plastic material** (often modeled as an elastic-linearly plastic or rigid-linearly plastic material) exhibits strain hardening (work hardening) during plastic deformation. In this model, once the true stress (σ) exceeds the initial yield strength (σ_y), further plastic deformation requires an increase in stress due to dislocation interactions within the crystal lattice. For a linear hardening model, this relationship is expressed mathematically by the linear flow equation:

$$\sigma = \sigma_y + H \cdot \epsilon_p$$

Where:

- σ represents the current true stress value.
- σ_y represents the initial tensile yield strength.
- H represents the constant plastic hardening modulus (slope of the stress-strain curve in the plastic zone).
- ϵ_p represents the plastic true strain component.

This equation shows that beyond initial yielding, the true stress increases as a linear function of the strain.

Step 1: Interpret the phrase "linearly hardening".

The term "linearly hardening" specifies that the rate of strain hardening remains constant. This means the slope ($\frac{d\sigma}{d\epsilon}$) of the stress-strain curve in the post-yielding plastic regime is a constant positive value (H).

Step 2: Evaluate the behavior described by each option.

- **Option (A):** Incorrect. A decrease in stress with strain would indicate strain softening or strain-localization (necking), not hardening.
- **Option (B):** Correct. The stress increases along a straight line with a constant positive slope relative to the strain.
- **Option (C):** Incorrect. A constant stress with increasing strain describes a *perfectly plastic* material ($\sigma = \sigma_y$), which has zero hardening.
- **Option (D):** Incorrect. An exponential increase occurs in power-law hardening models ($\sigma = K\epsilon^n$), not linear models.

Thus, the true stress increases linearly with the true strain beyond initial yielding, matching Option (B).

Quick Tip: Break down the technical term: 1. ****Hardening:**** Stress must *increase* to cause further deformation (eliminates options A and C). 2. ****Linearly:**** The increase follows a straight-line trend with a constant slope (eliminates option D). Therefore, it must increase linearly!

78. The area under the linear curve of stress-strain diagram with in the elastic limit is known as

- (A) Strain energy
- (B) Kinetic energy
- (C) Potential energy
- (D) Gravitational energy

Correct Answer: (A) Strain energy

Solution:

Concept: When a solid material undergoes mechanical deformation due to an external load, the work performed by that force is transferred into the material and stored internally as potential energy associated with the displacement of atomic bonds. This stored internal energy is called **Strain Energy** (or elastic strain energy). On a standard engineering stress-strain (σ - ϵ) diagram, the area under the curve represents the mechanical work done per unit volume (U_v) of the material:

$$U_v = \int \sigma d\epsilon$$

If the loading remains strictly within the linear elastic limit (Hooke's Law region, where $\sigma = E \cdot \epsilon$), the curve is a straight line starting from the origin. The area under this linear portion forms a right triangle. The area of this triangle represents the elastic strain energy stored per unit volume, a property known as the **Modulus of Resilience**:

$$\text{Area} = \frac{1}{2} \times \text{Base} \times \text{Height} = \frac{1}{2} \cdot \epsilon \cdot \sigma = \frac{\sigma^2}{2E}$$

Step 1: Evaluate the physical definitions of the choices.

Let's analyze each type of energy listed to identify the correct classification:

- **Strain Energy:** The internal mechanical energy stored within a deformed body within its elastic limit, which can be fully recovered upon unloading.
- **Kinetic Energy:** The energy an object possesses due to its macroscopic translational or rotational motion ($E_k = \frac{1}{2}mv^2$). It is not related to internal material deformation.
- **Potential Energy:** A broad category of stored energy based on position or configuration (e.g., electrostatic, chemical). While strain energy is technically a form of elastic potential energy, "Strain Energy" is the specific engineering term used for deforming bodies.
- **Gravitational Energy:** The potential energy an object possesses due to its position within a gravitational field ($E_g = mgh$).

Step 2: Match the geometric representation to the correct engineering term.

The area under a stress-strain diagram represents work converted into internal deformation energy. Within the elastic limit, this is precisely defined as **Strain Energy**, matching Option (A).

Quick Tip: Area under the σ - ϵ curve within the elastic limit = **Resilience** (or elastic strain energy per unit volume). If the area is calculated across the entire curve up to the point of fracture, it represents the material's **Toughness**.

79. Which one of the following metal has maximum value of yield strength?

- (A) Copper alloy
- (B) High-carbon steel
- (C) Mild steel
- (D) Medium-carbon steel

Correct Answer: (B) High-carbon steel

Solution:

Concept: Yield strength (σ_y) is the threshold stress value at which a material transitions from elastic behavior to permanent plastic deformation. In carbon steels, mechanical properties are determined primarily by the weight percentage of carbon dissolved in the iron matrix. Carbon acts as an interstitial alloying element that fits into the spaces between iron atoms. These carbon atoms create localized stress fields that pin and restrict the movement of dislocations through the crystal lattice under load—a mechanism known as **interstitial solid solution strengthening**. Additionally, during cooling, carbon forms a hard iron carbide intermetallic compound called cementite (Fe_3C). As carbon content increases, the volume fraction of cementite grows, raising both the hardness and yield strength of the steel.

Step 1: Compare carbon steel classifications by composition and strength.

Let us classify the three carbon steels based on typical carbon composition and yield strength trends:

1. **Mild Steel (Low-Carbon Steel):** Contains less than 0.25% Carbon by weight. It features high ductility and toughness but has the lowest yield strength among carbon steels (typically ≈ 250 MPa).
2. **Medium-Carbon Steel:** Contains between 0.25% and 0.60% Carbon by weight. It offers a balanced compromise between ductility and strength, with typical yield strengths ranging from 300 MPa to 600 MPa.
3. **High-Carbon Steel:** Contains between 0.60% and 1.40% Carbon by weight. It has a high concentration of hard cementite, resulting in high hardness and the highest yield

strength among standard carbon steels (often exceeding 600 – 1000 MPa depending on heat treatment), though it has lower ductility.

Step 2: Compare with non-ferrous copper alloys.

- **Copper Alloys (e.g., Brass, Bronze):** While brass and bronze can be engineered to achieve reasonable strength through work hardening or solution strengthening, their typical yield strength limits remain lower than those of high-carbon structural steels.

Step 3: Identify the material with the maximum value.

Comparing all options, high-carbon steel has the highest carbon content and cementite fraction, giving it the highest yield strength. This matches Option (B).

Quick Tip: In carbon steels: Higher Carbon Content (%) → Higher Volume of Cementite → Higher Yield Strength and Hardness (but lower ductility and toughness). Therefore, the strength order is: High-Carbon > Medium-Carbon > Mild Steel.

80. Based on Iron-Carbon diagram, when austenite decomposes during cooling process, an eutectoid mixture is formed. The nature of the element is

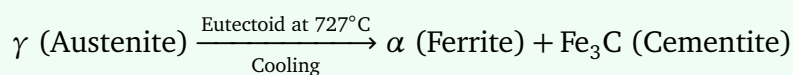
- (A) Ferrite
- (B) Cementite
- (C) Martensite
- (D) Pearlite

Correct Answer: (D) Pearlite

Solution:

Concept: The Iron-Carbon (Fe-Fe₃C) equilibrium phase diagram contains three distinct isothermal invariant reactions: peritectic, eutectic, and eutectoid. The **eutectoid reaction** occurs during slow cooling when a single solid phase transforms directly into two distinct solid phases at a specific temperature and composition. For steel, this reaction occurs at a carbon concentration of 0.76% to 0.80% C and a temperature of 727 °C. The high-temperature solid phase,

****Austenite (γ)****, decomposes into two new solid phases:



This simultaneous cooperative growth forms alternating lamellar (layer-like) bands of soft, ductile ferrite and hard, brittle cementite. This unique lamellar microstructural mixture is named ****Pearlite****. Under an optical microscope, its layered structure reflects light like mother-of-pearl.

Step 1: Evaluate the constituent phases and structures.

Let's analyze the microstructural terms listed in the options to clarify their definitions:

- **Ferrite (α):** A solid solution of carbon in iron with a Body-Centered Cubic (BCC) crystal structure. It is soft and ductile, representing one of the individual component phases inside the eutectoid mixture.
- **Cementite (Fe_3C):** An iron carbide intermetallic compound containing 6.67% C by weight. It is hard and brittle, forming the second individual component phase inside the eutectoid mixture.
- **Martensite:** A metastable phase formed when austenite is quenched rapidly to room temperature, trapping carbon atoms in a diffusionless transformation that distorts the lattice into a Body-Centered Tetragonal (BCT) structure. It is a non-equilibrium phase and does not form via slow eutectoid cooling.
- **Pearlite:** The name given to the lamellar ****eutectoid mixture**** consisting of alternating layers of ferrite and cementite.

Step 2: Match the question description to the correct mixture.

The question asks for the name of the eutectoid mixture formed when austenite decomposes during cooling. Based on the phase diagram reaction, this mixture is ****Pearlite****, matching Option (D).

Quick Tip: Remember the definition of the eutectoid point on the Iron-Carbon diagram: - Temperature = 727°C , Carbon = 0.76%. - Austenite (γ) $\rightarrow$ Ferrite (α) + Cementite (Fe_3C) = **Pearlite**. Pearlite is a two-phase microstructural mixture, not a single individual phase.

81. In casting process, the covering passage is used for feeding the liquid metal into the mould to

- (A) Decrease the waste of cast metal
- (B) Avoid the aspiration of air
- (C) Increase rate of feeding
- (D) Quickly break off the protecting portion of the casting

Correct Answer: (B) Avoid the aspiration of air

Solution:

Concept: In metal casting, the gating system is a series of interconnected channels designed to guide molten metal smoothly from the pouring basin into the mold cavity. A major concern during pouring is **Air Aspiration**. Air aspiration occurs if the pressure at any point within the gating channels drops below atmospheric pressure ($P_{\text{local}} < P_{\text{atm}}$). This pressure drop creates a localized vacuum that draws surrounding gases and air through the porous sand mold into the molten metal stream. These trapped gases can react with the liquid metal or become trapped during solidification, causing casting defects such as blowholes, gas porosity, and internal oxides. To prevent air aspiration, the vertical channel known as the **sprue** must be designed with a specific geometric profile.

Step 1: Analyze the fluid dynamics of molten metal flow in a sprue.

Let's apply the continuity equation and Bernoulli's equation to the liquid metal falling under gravity down a vertical sprue channel:

- As the liquid metal falls, it accelerates due to gravity, increasing its velocity (V).
- According to the volumetric continuity equation ($A_1 V_1 = A_2 V_2$), as velocity increases, the cross-sectional area of the liquid stream naturally contracts or narrows.

If a straight, cylindrical vertical sprue is used, the accelerating metal stream will pull away from the channel walls, creating a low-pressure vacuum zone that causes air aspiration.

Step 2: Evaluate how channel geometry prevents aspiration.

To prevent this vacuum formation, the sprue channel must be designed with a **tapered** profile (wider at the top intake and narrower at the bottom exit) that matches the natural contraction of the falling liquid stream. This design ensures that the moving liquid metal remains in continuous contact with the channel walls, keeping the pressure throughout the passage above atmospheric pressure. This covered, tapered gating design prevents the aspiration of air and gases into the stream.

Step 3: Select the option that matches this design objective.

The primary functional purpose of configuring and covering these passages is to maintain positive pressure and **avoid the aspiration of air**, matching Option (B).

Quick Tip: Air aspiration is caused by low-pressure zones ($P < P_{\text{atm}}$) inside the gating channels. To prevent air from being sucked into the liquid metal stream through the sand mold, channels like the sprue are tapered to keep the fluid under positive pressure, ensuring a clean casting free of gas porosity.

82. A low projection on the drag surface of a casting commencing near the gate by erosion of sand due to the high velocity liquid metal in the bottom gating. This defect is known as:

- (A) Cold shut
- (B) Swell
- (C) Blow holes
- (D) Wash

Correct Answer: (D) Wash

Solution:

Concept: In sand casting, defects arise due to multi-variable parameters such as sand properties, gating system design, and pouring parameters.

- **Wash (or Cut):** This manifests as an excess metal projection or an irregular low protrusion on the drag surface of the casting shell, starting predominantly near the gating entry channel. It is driven mechanically by the erosive action of high-velocity molten streams scouring away weak or poorly rammed molding sand.
- **Cold Shut:** A structural discontinuity formed when two separate streams of molten metal meet from opposite directions but fail to fuse completely due to early solidification or insufficient superheat.
- **Swell:** A localized enlargement of the casting dimensions caused by the displacement of the sand mold wall under the heavy metallostatic pressure of the liquid metal, typically due to soft or non-uniform ramming.
- **Blow holes:** Internal or surface gas cavities that are spherical or smooth-walled, generated by entrapped air, core gases, or vaporized moisture from low-permeability molds.

Step 1: Analyze the specific operational conditions given. The problem describes an occurrence marked by:

1. Location: On the drag (bottom) surface, originating adjacent to the ingate.
2. Action: Erosion or washing away of the molding sand matrix.
3. Cause: High velocity flow of the liquid metal stream when passing through a bottom gating layout.

Step 2: Connect the mechanism to the correct nomenclature. When molten metal enters the mold cavity at high velocity, it creates significant hydrodynamic shear stress along the sand boundary layer. If the sand lacks sufficient hot strength, dry strength, or proper compact ramming, the metal physically detaches sand grains. The carried-away sand leaves behind a void space inside the mold boundary. During filling, this empty space is filled by liquid metal, resulting in an irregular, low projection called a ****Wash****.

Quick Tip: To eliminate wash defects: - Adopt an unpressurized gating system design to decrease the fluid velocity. - Elevate the sand's green/dry binding strength by adding appropriate clay and binder content.

83. The shrinkage allowance has been provided on the pattern to compensate the shrinkage when:

- (A) The temperature of liquid metal drops from pouring to freezing temperature
- (B) The metal changes from liquid to solid state at freezing temperature
- (C) The temperature of solid phase drops from freezing to room temperature
- (D) The temperature of metal drops from pouring to room temperature

Correct Answer: (C) The temperature of solid phase drops from freezing to room temperature

Solution:

Concept: When molten metal transitions from its liquid pouring state down to a room temperature solid casting, it undergoes thermal contraction across three distinct stages:

1. **Liquid Shrinkage:** Contraction occurring during the drop from the pouring temperature to the liquidus/freezing temperature.

2. **Solidification Shrinkage:** Contraction occurring while the metal undergoes a phase change from liquid to solid at a constant freezing temperature.
3. **Solid Shrinkage:** Contraction occurring as the temperature of the fully solid casting drops from the freezing point down to the ambient room temperature.

Step 1: Distinguish how each contraction stage is compensated.

- Both **Liquid Shrinkage** and **Solidification Shrinkage** result in a volumetric deficit within the liquid pool, forming voids or pipes. To overcome this, a **Riser** (or feeder head) is incorporated into the mold design to continuously supply extra molten metal into the cavity until solidification finishes.
- On the other hand, **Solid Shrinkage** occurs after the metal has completely frozen into its geometric shape. As it cools down to room temperature, the overall linear dimensions reduce. To ensure that the final cooled casting matches the blueprint design specifications, the physical dimensions of the **Pattern** must be made deliberately larger. This dimensional enlargement is known as the **Pattern Shrinkage Allowance**.

Step 2: Correlate with the options. The reduction in dimensions in the solid state is compensated solely by making the pattern oversized. Thus, the shrinkage allowance compensates for the contraction when the temperature of the solid phase drops from freezing to room temperature.

Quick Tip: Remember the compensation mechanism dividing line: - Liquid and Solidification shrinkage → Compensated by the **Riser**. - Solid shrinkage → Compensated by the **Pattern Allowance**.

84. The commonly used casting process for producing intricate and detailed sculptures is:

- (A) Sand casting
- (B) Investment casting
- (C) Die casting
- (D) Permanent mould casting

Correct Answer: (B) Investment casting

Solution:

Concept: Different casting methodologies offer varying degrees of dimensional accuracy, surface finish, and complexity limits:

- **Sand Casting:** Economical for large, simple parts but lacks high dimensional precision and fine surface texturing.
- **Investment Casting (Lost Wax Process):** Utilizes a consumable wax pattern coated with a ceramic slurry to create a seamless, highly precise one-piece mold. It accommodates highly complex geometries and excellent surface details.
- **Die Casting / Permanent Mold Casting:** Employs reusable metal dies. While accurate, it is heavily restricted by parting lines, draft angles, straight-line core withdrawals, and high tooling setup expenditures, making it poorly suited for asymmetric artistic art pieces.

Step 1: Evaluate the demands of sculpture fabrication. Artistic sculptures possess intricate undercuts, complex organic shapes, varying wall thicknesses, and require exceptional surface finish to capture fine details without visible joints or parting line defects.

Step 2: Match demands with Investment Casting characteristics. In investment casting, a wax pattern is shaped (often by hand or specialized molding). Because the wax is melted out completely before pouring (hence "lost wax"), there is no need to split the mold to extract a rigid pattern. This eliminates parting lines entirely, allows extensive internal undercuts, and produces intricate, detailed structures that match the artistic sculpture's design.

Quick Tip: Whenever a question highlights terms like **intricate detailed sculptures**, **lost wax process**, **no parting lines**, or **excellent surface finish for complex artwork**, the answer is almost always **Investment Casting**.

85. While designing the gating system, the ratio of 1 : 2 : 4 represents:

- (A) Sprue base area : runner area : ingate area
- (B) Pouring basin area : ingate area : runner area
- (C) Sprue base area : ingate area : casting area
- (D) Runner area : ingate area : casting area

Correct Answer: (A) Sprue base area : runner area : ingate area

Solution:

Concept: A gating ratio defines the relative cross-sectional areas of the key channels within a casting's delivery system. The standard reporting sequence for a gating ratio is consistently defined as:

$$\text{Gating Ratio} = A_s : A_r : A_g$$

Where:

- A_s = Cross-sectional area of the sprue base choke section.
- A_r = Total cross-sectional area of the runner channel(s).
- A_g = Total cross-sectional area of the ingate channel(s).

Gating systems are categorized into two primary styles:

1. **Pressurized Gating System:** Where the choke area is at the ingate (e.g., 1 : 0.75 : 0.5). The metal enters the mold cavity at high velocity.
2. **Unpressurized Gating System:** Where the choke area is at the base of the sprue (e.g., 1 : 2 : 2 or 1 : 2 : 4). The flow expands into larger runner and gate areas, decreasing its kinetic energy.

Step 1: Identify the components from the ratio 1 : 2 : 4. The ratio indicates a sequential widening of the passages (1 → 2 → 4). Following the standard convention: - First component (1) correspond to the ****Sprue base area****. - Second component (2) corresponds to the ****Runner area****. - Third component (4) corresponds to the ****Ingate area****.

This expanding unpressurized design reduces flow velocity and minimizes turbulence and sand erosion inside the cavity.

Quick Tip: The gating ratio order is fixed by convention: ****Sprue Base Area : Runner Area : Ingate Area****. Remember the acronym ****S-R-I**** to avoid mixing up the sequence.

86. The extrusion process used to manufacture the short length components like paste tubes, gun shells etc., is

- (A) Direct extrusion
- (B) Continuous extrusion

- (C) Hydrostatic extrusion
(D) Indirect extrusion

Correct Answer: (D) Indirect extrusion

Solution:

Concept: Extrusion processes are classified based on the relative direction of metal flow with respect to the ram's movement:

- **Direct Extrusion (Forward Extrusion):** The billet is placed in a container and pushed by a ram through a stationary die. The metal flows in the same direction as the ram. This generates high friction between the billet and the container walls.
- **Indirect Extrusion (Backward / Reverse / Impact Extrusion):** The die is mounted on a hollow ram. As the ram advances against the closed container, the metal is forced to flow backward through the die opening in a direction opposite to the ram's movement.

Step 1: Examine the manufacturing dynamics of thin-walled collapsible tubes and shells.

Components like toothpaste tubes, aluminum aerosol cans, and small ammunition cartridge cases (gun shells) are characterized by very thin walls, closed or semi-closed bottoms, and relatively short overall lengths. They are mass-produced using a specialized form of indirect cold extrusion known as **Impact Extrusion**.

Step 2: Understand why Indirect Extrusion is selected. In this process, a solid metal slug is placed in a small die cavity and struck at high velocity by a fast-moving punch. The metal is forced to flow backward through the narrow annular clearance between the punch outer diameter and the die inner diameter, rising up along the punch shank to form a thin-walled seamless cup or tube. Because there is no relative sliding friction between a long billet and container walls, the required force is lower, allowing for thin wall configurations.

Quick Tip: For questions referencing the production of **paste tubes**, **collapsible tubes**, **aluminum cans**, or **cartridge cases**, identify them with **Impact Extrusion**, which falls directly under the category of **Indirect Extrusion**.

87. The phenomenon where a metal part tends to partially return to its original shape after the forming tool is removed is known as

- (A) Spring back

- (B) Elastic recovery
- (C) Strain hardening
- (D) Plastic deformation

Correct Answer: (A) Spring back

Solution:

Concept: When a piece of sheet metal is bent or formed, the total deformation applied by the tooling consists of two distinct components:

$$\epsilon_{\text{total}} = \epsilon_{\text{plastic}} + \epsilon_{\text{elastic}}$$

Where:

- $\epsilon_{\text{plastic}}$ is the permanent, non-reversible deformation that shapes the part.
- $\epsilon_{\text{elastic}}$ is the temporary, reversible deformation governed by Hooke's Law.

Step 1: Describe what happens during unloading. When the forming tools (punch and die) are retracted, the external bending forces drop to zero. The material releases its stored elastic strain energy. While the plastic portion of the deformation remains permanent, the elastic portion recovers. This causes the metal part to unbend slightly and try to return toward its original unformed flat shape.

Step 2: Identify the correct industry term. Although this structural phenomenon occurs due to **elastic recovery**, the specific manufacturing term given to this physical geometric dimensional change in sheet metal operations (like V-bending or U-bending) is **Spring back**. It results in an increase in the final bend angle and a shift in the bend radius compared to the dimensions of the tool.

Quick Tip: To compensate for spring back in manufacturing operations: - **Overbending:** Design the punch angle to be sharper than the desired final target angle. - **Bottoming / Coining:** Apply extreme compressive stress at the bottom of the stroke to set the material plastically.

88. The secondary operation which is used to improve the self-lubricating capacity of the sintered part in powder metallurgy is

- (A) Infiltration

- (B) Impregnation
- (C) Repressing
- (D) Coining

Correct Answer: (B) Impregnation

Solution:

Concept: Powder metallurgy parts inherently retain interconnected porous networks after compacting and sintering. Various secondary treatments exploit or modify this porosity:

- **Infiltration:** Filling the pores of a lower-melting-point skeletal structure with a liquid metal of a lower melting temperature to improve mechanical strength and density.
- **Impregnation:** Filling the open, interconnected pores of a sintered part with a fluid substance such as oil, grease, or liquid polymers.
- **Repressing / Coining:** Pressing the sintered component in a closed die to achieve tighter dimensional tolerances or refine surface geometry.

Step 1: Relate porosity to self-lubricating capabilities. Sintered components like porous bearings and bushings require long working lives without external lubrication systems. To achieve this, the porous structural network is evacuated under vacuum and then submerged in a lubricating oil bath.

Step 2: Define the mechanism of Impregnation. The capillary action forces the lubricating oil deep into the interconnected pore voids. When the completed bearing operates in machinery, friction warms the assembly. The heat causes the oil to expand and draw out to the contact surface, establishing a continuous lubricating film. When motion stops, capillary forces draw the oil back into the pores. This process is called **Impregnation**.

Quick Tip: Distinguish between the two pore-filling operations in powder metallurgy: - Pores filled with **Liquid Metal** → **Infiltration**. - Pores filled with **Oil / Lubricant** → **Impregnation**.

89. In MIG welding process, the commonly used gases for shielding are

- (A) Argon and Helium
- (B) Oxygen and Acetylene
- (C) Nitrogen and Hydrogen

(D) Carbon dioxide and Nitrogen

Correct Answer: (A) Argon and Helium

Solution:

Concept: MIG (Metal Inert Gas) welding, also classified as GMAW (Gas Metal Arc Welding), utilizes a continuous, consumable bare wire electrode. To safeguard the highly reactive molten weld pool from atmospheric contamination (oxidation and nitriding by O_2 and N_2), a shielding gas envelope is continuously supplied through the torch nozzle.

The shielding gases can be grouped into:

- **Inert Gases:** Elements that exhibit no chemical reactions with the molten metal pool (e.g., Argon, Helium).
- **Active Gases:** Elements that chemically interact with the weld environment to alter arc characteristics (e.g., CO_2 , O_2), which shifts the process category strictly to MAG (Metal Active Gas).

Step 1: Evaluate the provided options against the core MIG requirement. - **Argon (Ar):** Provides arc stability, low thermal conductivity, and good cleaning action. - **Helium (He):** Provides deep penetration characteristics due to its higher ionization potential and thermal conductivity. - **Oxygen Acetylene:** This combination is used for gas combustion welding/-cutting (Oxy-Acetylene torch), not as an arc shielding medium. - **Nitrogen Hydrogen:** Can cause embrittlement and porosity defects in steel alloys.

Since MIG specifically stands for Metal **Inert** Gas welding, the inert gas mixtures of **Argon** and **Helium** are commonly used to provide protection for non-ferrous metals and stainless steels.

Quick Tip: - For pure **MIG** welding, select completely **Inert Gases** like **Argon** and **Helium**. - For carbon steels, **Carbon Dioxide (CO_2)** is often added or used alone, but that configuration is classified as **MAG** or semi-inert welding.

90. The welding process used for joining railway tracks is

- (A) TIG welding
- (B) Spot welding
- (C) Submerged welding

(D) Thermit welding

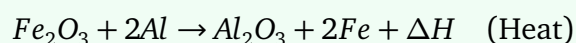
Correct Answer: (D) Thermit welding

Solution:

Concept: Joining thick, heavy structural steel sections on-site, such as railway tracks, requires a process that provides high heat input without requiring heavy external power generation equipment.

- **TIG / Submerged Arc Welding:** Requires heavy electrical power sources and multi-pass runs, which makes them difficult to implement efficiently across remote railway networks.
- **Spot Welding:** A resistance welding technique designed for thin sheet metal components.
- **Thermit Welding:** An exothermic chemical casting and fusion process that uses a chemical reaction to produce superheated liquid steel.

Step 1: Examine the underlying chemical mechanism of Thermit welding. The process relies on a highly exothermic reduction-oxidation (redox) reaction between a metal oxide (typically iron oxide) and a reducing agent (aluminum powder). The standard balanced chemical formula for this reaction is given by:



This chemical reaction releases enough thermal energy to raise the temperature to approximately 3000°C.

Step 2: Apply the process to railway track joining. The iron oxide and aluminum powder mixture is placed into a crucible positioned directly above the aligned gap between two rail segments enclosed in a sand mold. Once ignited, the reaction rapidly creates liquid iron, which sinks to the bottom due to its high density, while the aluminum oxide slag (Al_2O_3) floats to the top. The superheated liquid iron is then tapped into the mold gap, melting the rail ends and fusing them into a continuous joint upon cooling.

Quick Tip: Whenever a question references **“railway track alignment”**, **“joining rails on-site”**, or an **“exothermic chemical powder reaction”**, the answer is uniquely **“Thermit Welding”**.

91. What is the key requirement for an adhesive to properly bond to a surface?

- (A) High surface roughness
- (B) High surface tension
- (C) Good wetting
- (D) Friction free surface

Correct Answer: (C) Good wetting

Solution:

Concept: Adhesive bonding relies on intimate contact at the molecular level between a liquid adhesive agent and a solid substrate surface.

- **Wetting:** The ability of a liquid to maintain contact with a solid surface, driven by intermolecular interactions when the two are brought into contact.
- **Contact Angle (θ):** A quantitative measure of wetting behavior. A low contact angle ($\theta \rightarrow 0^\circ$) indicates excellent wetting, where the liquid spreads out into a thin film. A high contact angle ($\theta > 90^\circ$) indicates poor wetting, where the liquid beads up.

Step 1: Evaluate the criteria for establishing a strong adhesive interface. For an adhesive to develop mechanical interlocking and chemical bonding forces (such as Van der Waals forces or covalent linkages), it must flow freely into the microscopic valleys and asperities of the substrate surface.

Step 2: Relate this to the provided options. If the adhesive has high internal surface tension or the substrate has low surface energy, the liquid will bead up, trapping air pockets underneath and reducing the effective contact area. Therefore, **good wetting** is required, ensuring that the adhesive completely covers and conforms to the surface topography to establish a durable bond.

Quick Tip: - **Good Wetting** $\equiv$ **Low Contact Angle ($\theta < 90^\circ$)**. - High wetting maximizes the real contact area, which is essential for creating high-strength adhesive joints.

92. According to the Taylor's tool life equation, if the cutting speed is halved with index $n = 0.25$, then the tool life:

- (A) Increases by 8 times
- (B) Decreases by 8 times

(C) Decreases by 16 times

(D) Increases by 16 times

Correct Answer: (D) Increases by 16 times

Solution:

Concept: Taylor's tool life equation describes the empirical relationship between cutting velocity and tool wear life:

$$VT^n = C$$

Where:

- V = Cutting speed.
- T = Tool life.
- n = Tool life exponent (dependent on tool material).
- C = Constant.

For two distinct operating conditions with the same tool-workpiece combination, the relation can be written as:

$$V_1 T_1^n = V_2 T_2^n$$

Step 1: State the given problem parameters. Let the initial cutting speed be V_1 and the initial tool life be T_1 . The problem states that the tool life exponent is:

$$n = 0.25 = \frac{1}{4}$$

The final cutting speed V_2 is reduced to half of the initial speed:

$$V_2 = \frac{V_1}{2}$$

Step 2: Set up the ratio using the tool life equation. Substitute V_2 into the relation:

$$V_1 T_1^n = \left(\frac{V_1}{2}\right) T_2^n$$

Dividing both sides by the initial velocity V_1 simplifies the equation to:

$$T_1^n = \frac{1}{2} T_2^n \Rightarrow \frac{T_2^n}{T_1^n} = 2$$

Using the properties of exponents:

$$\left(\frac{T_2}{T_1}\right)^n = 2$$

Step 3: Solve for the final tool life ratio. Substitute $n = 0.25 = \frac{1}{4}$ into the expression:

$$\left(\frac{T_2}{T_1}\right)^{\frac{1}{4}} = 2$$

Raise both sides of the equation to the power of 4 to isolate the ratio:

$$\left[\left(\frac{T_2}{T_1}\right)^{\frac{1}{4}}\right]^4 = 2^4$$

$$\frac{T_2}{T_1} = 16 \Rightarrow T_2 = 16T_1$$

Therefore, cutting the speed in half causes the tool life to increase by a factor of 16.

Quick Tip: A useful shortcut derived from Taylor's equation is:

$$\frac{T_2}{T_1} = \left(\frac{V_1}{V_2}\right)^{\frac{1}{n}}$$

Substituting the values gives: $(2)^{\frac{1}{0.25}} = 2^4 = 16$.

93. The Titanium Nitride or Titanium Carbonitride coated cutting tools are majorly used to

- (A) Make the tool ductile
- (B) Reduce friction and wear resistance
- (C) Increase tool toughness
- (D) Increase thermal conductivity

Correct Answer: (B) Reduce friction and wear resistance

Solution:

Concept: Cutting tool inserts are often coated with thin ceramic layers (e.g., TiN, TiCN, Al_2O_3) via Chemical Vapor Deposition (CVD) or Physical Vapor Deposition (PVD) to improve their surface properties without changing the bulk toughness of the substrate core.

- **Titanium Nitride (TiN):** Known for its distinctive golden color. It provides a low coeffi-

cient of friction and high micro-hardness at elevated temperatures.

- **Titanium Carbonitride (TiCN):** Offers high abrasive wear resistance and hardness compared to TiN, making it suitable for demanding cutting conditions.

Step 1: Analyze the physical changes at the tool-chip interface. During high-speed machining, high friction at the tool-chip interface generates significant heat and accelerates crater and flank wear.

Step 2: Evaluate the function of the coatings. Applying a TiN or TiCN coating provides two primary benefits: 1. It creates a low-friction chemical barrier that reduces shearing forces and prevents material adhesion (built-up edge formation). 2. It increases the tool's wear resistance due to the high hardness of the ceramic layer.

(Note: The wording "Reduce friction and wear resistance" in the option translates technically to "Reduce friction and improve wear resistance" or "Reduce friction and reduce wear," which aligns with the intended choice).

Quick Tip: Thin ceramic surface coatings (like TiN, TiC, TiCN) are applied to cutting tools to improve surface-specific properties, such as increasing hardness and reducing friction, rather than modifying bulk properties like toughness or ductility.

94. As the cutting speed increases, how do the machining cost and tooling cost behave?

- (A) Machining cost decreases while tooling cost increases
- (B) Both machining cost and tooling cost both increases
- (C) Both machining cost and tooling cost both decreases
- (D) Machining cost increases while tooling cost decreases

Correct Answer: (A) Machining cost decreases while tooling cost increases

Solution:

Concept: In the economics of machining operations, the total cost per component is composed of several factors that vary with cutting speed (V):

$$\text{Total Cost} = \text{Machining Cost} + \text{Tooling Cost} + \text{Handling/Setup Cost}$$

Step 1: Analyze the behavior of Machining Cost with respect to speed. Machining cost is directly proportional to the time required to cut the feature:

$$\text{Machining Cost} \propto \text{Machining Time}(t_m)$$

As the cutting speed V increases, the material removal rate (MRR) increases, which shortens the time required to complete the operation (t_m). Because less time is spent on the machine tool per part, the direct ****machining cost decreases****.

Step 2: Analyze the behavior of Tooling Cost with respect to speed. According to Taylor's tool life equation ($VT^n = C$), raising the cutting speed exponentially shortens the tool life (T). As a result, the tool wears out faster, requiring more frequent tool changes, insert indexing, and higher replacement tool expenditures. Consequently, the ****tooling cost increases****.

Quick Tip: - High cutting speeds $\rightarrow$ Shorter machining times $\rightarrow$ ****Machining cost drops****. - High cutting speeds $\rightarrow$ Accelerated tool wear rate $\rightarrow$ ****Tooling cost rises****. The minimum point of the combined cost curve defines the economic cutting speed.

95. The vacuum is a mandatory requirement for which of the following non traditional machining process?

- (A) Laser beam machining
- (B) Water jet machining
- (C) Electron beam machining
- (D) Plasma Arc machining

Correct Answer: (C) Electron beam machining

Solution:

Concept: Non-traditional machining processes utilize thermal, chemical, or mechanical energies to remove material:

- **Laser Beam Machining (LBM):** Focuses a coherent monochromatic light beam; operates effectively in ambient air conditions or inert gas environments.
- **Plasma Arc Machining (PAM):** Relies on a high-velocity jet of ionized gas to cut metal.
- **Electron Beam Machining (EBM):** Uses a focused stream of high-velocity electrons

accelerated toward the workpiece to melt and vaporize the target material.

Step 1: Analyze the physical behavior of a high-velocity electron beam. Electrons are subatomic particles with very low mass. When accelerated to high velocities (often exceeding half the speed of light), they possess significant kinetic energy.

Step 2: Explain why a vacuum chamber is required. If the electron beam traveled through ambient air, the electrons would collide with atmospheric gas molecules (O_2 , N_2 , etc.). These collisions would scatter the electrons, causing the beam to lose focus and dissipate its kinetic energy before reaching the workpiece. To maintain a well-focused beam and protect the electron-generating cathode filament from oxidation, the entire **Electron Beam Machining** process must be contained within a high vacuum chamber (vacuum levels around 10^{-4} to 10^{-5} torr).

Quick Tip: Any electron-based manufacturing process—such as **Electron Beam Machining (EBM)** or **Electron Beam Welding (EBW)**—requires a high-vacuum environment to prevent the electron stream from scattering via collisions with air molecules.

96. A device/mechanism used to position a work piece in the same position, cycle after the cycle is called

- (A) Indexing device
- (B) Guiding device
- (C) Clamping device
- (D) Locating device

Correct Answer: (A) Indexing device

Solution:

Concept: Jigs and fixtures incorporate several distinct functional components to manage workpieces during production runs:

- **Locating elements:** Establish the initial correct orientation and position of the workpiece by restricting its degrees of freedom relative to the cutting tool.
- **Clamping elements:** Apply holding forces to keep the workpiece firmly positioned against the locators under machining loads.

- **Guiding elements:** Physically steer and support the cutting tool (e.g., drill bushings).
- **Indexing mechanisms:** Rotate or translate the workpiece through precise, repeatable angular or linear increments, allowing identical features to be machined at regular intervals around a common center or line cycle after cycle.

Step 1: Analyze the phrase "position a workpiece in the same position, cycle after cycle". While basic locators fix a single static location, manufacturing operations often require machining multiple identical features across repetitive rotational cycles (such as cutting gear teeth, drilling holes around a circular flange, or shifting a part sequentially through automated stations). A mechanism that moves and repositions the workpiece into precise, repeatable locations across successive operational cycles is called an **Indexing device**.

Quick Tip: - **Locating** establishes the initial spatial orientation of the part. - **Indexing** shifts or repositions the part through precise, uniform steps or intervals across consecutive machining cycles.

97. The main purpose of an Automatic Tool Changer (ATC) in CNC machining process is to reduce:

- (A) Machining maintenance
- (B) Cutting speed
- (C) Tool change time and idle time
- (D) Material hardness

Correct Answer: (C) Tool change time and idle time

Solution:

Concept: The total cycle time (t_{total}) required to process a workpiece on a CNC machine tool is divided into two primary categories:

$$t_{\text{total}} = t_{\text{productive}} + t_{\text{non-productive}}$$

Where:

- $t_{\text{productive}}$ is the active cutting time when material is being removed.
- $t_{\text{non-productive}}$ (or idle time) includes auxiliary tasks such as loading parts, rapid tool

positioning, and swapping tools for different operations (e.g., switching from a drill bit to an end mill).

Step 1: Identify the function of an Automatic Tool Changer (ATC). In complex machining operations, a single part may require drilling, tapping, face milling, and boring. Performing these tool swaps manually requires stopping the spindle, opening enclosure doors, unlocking the tool holder, swapping the tool, and resetting length offsets, which increases machine idle time.

Step 2: Relate the mechanism to production efficiency. An ATC automates this process using a tool magazine and a mechanical transfer arm. The CNC program commands the arm to swap tools directly within seconds while maintaining pre-registered tool offset coordinates. This significantly minimizes ****tool change time and non-productive machine idle time****, maximizing active spindle utilization.

Quick Tip: Automation features like ATCs (Automatic Tool Changers) and APCs (Automatic Pallet Changers) are designed to reduce non-productive idle times, which increases overall equipment effectiveness (OEE).

98. Which type of jig is generally used for machining in more than one plane without changing the setting?

- (A) Box jig
- (B) Channel jig
- (C) Template jig
- (D) Plate jig

Correct Answer: (A) Box jig

Solution:

Concept: Drill jigs are classified based on their structural design and how they support multi-axis machining:

- **Template / Plate Jig:** A simple plate containing drill bushings that clamps onto a single face of the workpiece, allowing machining on one plane only.
- **Channel Jig:** Features a U-shaped cross-section that permits machining on a single face

or limited parallel faces.

- **Box Jig (or Tumble Jig):** Completely encloses the workpiece within a rigid, box-like frame equipped with a hinged lid or latch for loading.

Step 1: Analyze the requirement for multi-plane machining. When a workpiece requires holes to be drilled on multiple faces (e.g., top, bottom, and sides) relative to a fixed datum layout, unclamping and repositioning the part in separate jigs introduces alignment errors and increases handling time.

Step 2: Evaluate the Box Jig design. A box jig houses the workpiece securely inside its frame. Drill guide bushings can be integrated into multiple faces of the box structure. The exterior faces of the box jig are precision-machined flat, allowing the operator to tip or "tumble" the entire jig assembly onto different sides on the drill press table. This permits multi-axis drilling on different planes in a single setup, ensuring high geometric accuracy.

Quick Tip: When a drill jig question highlights phrases like ****machining in more than one plane****, ****tumbling the jig****, or ****multi-axis drilling in one setup****, select the ****Box Jig****.

99. Which of the following parameters is NOT the essential requirement for accuracy of the measurement with the sine bar?

- (A) Flatness of upper surface, equality of size and roundness of rollers
- (B) Exact distance between the roller axes and mutual parallelism
- (C) Parallelism between top surface and bottom surface
- (D) Parallelism of the roller axes to the upper surface

Correct Answer: (C) Parallelism between top surface and bottom surface

Solution:

Concept: A sine bar is a precision linear metrology tool used to measure unknown angles or set workpieces at specific inclinations. Its operation is based on the trigonometric sine rule applied to a right-angled triangle:

$$\sin(\theta) = \frac{h}{L} \Rightarrow \theta = \sin^{-1}\left(\frac{h}{L}\right)$$

Where:

- L is the precise distance between the centers of the two pivoting rollers.
- h is the height of the slip gauge block stack.
- θ is the inclination angle being measured.

Step 1: Identify the critical geometric features required for precision accuracy. For the trigonometric relationship to hold true, several manufacturing criteria are necessary: 1. The upper gauging surface must be perfectly flat (**Flatness of upper surface**). 2. The two rollers must have identical diameters and be perfectly cylindrical (**equality of size and roundness of rollers**). 3. The distance L between the roller centers must be known precisely and their axes must be parallel (**Exact distance between the roller axes and mutual parallelism**). 4. The axes of the rollers must be parallel to the upper workspace gauging plane (**Parallelism of the roller axes to the upper surface**).

Step 2: Determine which feature is not required. A sine bar consists of an upper gauging plane and two lower cylindrical rollers resting on a reference surface plate or slip gauges. It does not have a separate functional "bottom surface" parallel to its top face; instead, the bottom boundary consists of the tangent lines along the undersides of the two rollers. Therefore, **parallelism between the top surface and a bottom surface** is not a relevant or essential parameter for its measurement accuracy.

Quick Tip: The accuracy of a sine bar depends entirely on the geometry of its **upper surface**, the **equality and roundness of the two rollers**, and the **exact distance (L) between their centers**. Any reference to a flat "bottom surface" is irrelevant to its function.

100. The instrument used for the direct measurement of surface quality by dragging a stylus across the surface is:

- (A) Clinometer
- (B) Sine bar
- (C) Autocollimator
- (D) Profile meter

Correct Answer: (D) Profile meter

Solution:

Concept: Surface metrology tools are designed to evaluate the topography, roughness, or angular characteristics of components:

- **Clinometer / Sine bar / Autocollimator:** Instruments specifically used for angular measurement, flatness determination, or straightness testing.
- **Profile meter (Profilometer):** An instrument designed explicitly to quantify surface roughness and texture profile parameters (R_a, R_z , etc.).

Step 1: Analyze the operational mechanism described. The core mechanism involves physically **dragging a sharp diamond stylus** across the workpiece surface at a controlled speed. As the stylus traverses the surface texture, its microscopic peaks and valleys cause it to deflect vertically.

Step 2: Map the mechanism to the correct metrological instrument. These vertical movements are converted by a transducer (such as a piezoelectric crystal or inductive coil) into an electrical signal. This signal is amplified and processed to map the surface variations. This stylus-based tracking system defines the working principle of a **Profilometer** or **Profile meter**.

Quick Tip: Whenever contact-type surface roughness profiling is mentioned via a **traveling stylus**, the measuring device is a **Profilometer** or **Profile meter**.

101. As per the Indian standards, the limits and fits specify how many grades of tolerances?

- (A) 10
- (B) 18
- (C) 22
- (D) 28

Correct Answer: (B) 18

Solution:

Concept: The Indian Standard (IS:919) system of limits, fits, and tolerances aligns closely with the international ISO system to govern component variability during manufacturing.

- **Standard Tolerances (IT Grades):** Define the magnitude of the manufacturing zone

tolerance. There are **18** fundamental tolerance grades denoted as:

IT01, IT0, IT1, IT2, IT3, ..., IT16

- **Fundamental Deviations:** Define the location of the tolerance zone relative to the zero/basic line. There are **25** fundamental deviations for shafts and holes (denoted by letters like A–ZC and a–zc).

Step 1: Quantify the total number of IT grades. Counting the sequential categories from IT01 up to IT16: - IT01 and IT0 (2 grades) - IT1 through IT16 (16 grades) Total = $2 + 16 = 18$ standard grades.

Lower numbers (e.g., IT01 to IT5) represent tight tolerances for high-precision tools like gauges, while higher numbers (e.g., IT12 to IT16) correspond to looser tolerances for operations like sand casting or rough forging.

Quick Tip: Remember the standard classification system values: - Number of standard tolerance grades = **18** - Number of fundamental deviations = **25** (extended to 28 in newer standards, but fundamental categories remain distinct).

102. Which of the following is NOT a common probe type used in a coordinate measuring machine?

- (A) Electromagnetic
- (B) Laser
- (C) Scanning
- (D) Touch-trigger

Correct Answer: (A) Electromagnetic

Solution:

Concept: A Coordinate Measuring Machine (CMM) is an inspection system that determines the 3D coordinates of a part's surface points using a data-gathering probe tip. CMM probes are categorized into contact and non-contact types:

- **Touch-trigger Probes (Contact):** Actively record a single coordinate point whenever the physical stylus contacts the part surface.

- **Scanning Probes (Contact):** Drag along the workpiece profile to gather a continuous stream of data points.
- **Laser Probes (Non-contact):** Use optical triangulation or time-of-flight laser beams to measure surface profiles without making physical contact.

Step 1: Assess "Electromagnetic" as a CMM probe system. While CMM probes use electrical contacts, optical sensors, or strain gauges to detect displacement, they do not use **electromagnetic** field variations as a primary mechanism to trace and inspect part geometry. Therefore, it is not a common or standard probe configuration.

Quick Tip: CMM data collection is driven by physical contact (such as **Touch-trigger** or continuous **Scanning**) or optical methods (such as **Laser** scanning). Magnetic or electromagnetic field induction is not typically used for coordinate tracing.

103. What is the primary purpose of direct numerical control in the context of CAD/CAM integration?

- (A) Generating G codes automatically
- (B) Converting the analog signals into digital
- (C) Designing 3D models
- (D) Networking multiple CNC machines to a single computer

Correct Answer: (D) Networking multiple CNC machines to a single computer

Solution:

Concept: In automated manufacturing facilities, managing part programs efficiently across multiple machines is critical:

- **CAD Systems:** Used to design 3D models.
- **CAM Post-Processors:** Used to generate G-codes automatically from toolpath simulations.
- **Direct Numerical Control (DNC):** A networking configuration where a central main-frame computer manages and transfers part programs to multiple individual CNC machine controllers simultaneously.

Step 1: Clarify the functional role of DNC. Early CNC machines had limited internal memory. DNC solved this problem by creating a local area network (LAN) connection from a single central host computer to the machines on the factory floor. This allows the host computer to store large libraries of G-code programs and stream them dynamically to multiple CNC machines simultaneously, eliminating the need to load programs manually via physical disks or tapes.

Quick Tip: - **CNC** controls an individual machine tool using its dedicated computer. - **DNC** (Direct Numerical Control) connects and networks **multiple CNC machines** to a single central computer system.

104. Which of the additive manufacturing techniques uses a laser to cure liquid photopolymer resin in a vat?

- (A) Stereolithography
- (B) Fused Deposition modelling
- (C) Selective Laser Sintering
- (D) Binding Jetting

Correct Answer: (A) Stereolithography

Solution:

Concept: Additive manufacturing (3D printing) processes are classified by the raw state of the material used and its consolidation mechanism:

- **Stereolithography (SLA):** Employs a liquid photopolymer resin that cures and solidifies layer-by-layer when exposed to a UV laser beam.
- **Fused Deposition Modeling (FDM):** Extrudes a heated thermoplastic filament through a moving nozzle.
- **Selective Laser Sintering (SLS):** Uses a laser to fuse powdered material (such as nylon or metal) into a solid structure.
- **Binder Jetting:** Deposits a liquid bonding agent onto layers of powdered material.

Step 1: Match the given prompt parameters to the correct process mechanism. The question specifies three defining operational characteristics: 1. Raw material state: A **liquid**

photopolymer resin** contained in a storage **vat**. 2. Energy source: A focused **laser beam**. 3. Transformation mechanism: Liquid-to-solid curing via photo-polymerization. These features describe **Stereolithography (SLA)**, which was the first commercially developed additive manufacturing technique.

Quick Tip: Whenever an additive manufacturing question pairs a **vat of liquid photopolymer resin** with a **UV laser cure** mechanism, the correct answer is always **Stereolithography (SLA)**.

105. Find out the correct sequence for the additive manufacturing process chain?

- (A) Build, post processing, STL conversion, CAD
- (B) CAD, STL conversion, Build, Post processing
- (C) Post processing, Build, STL conversion, CAD
- (D) CAD, Build, STL conversion, Post processing

Correct Answer: (B) CAD, STL conversion, Build, Post processing

Solution:

Concept: The standard operational pipeline for additive manufacturing (3D printing) follows a systematic sequence from a digital concept to a finished physical part. The essential steps are:

1. **CAD Modeling:** Creating the initial solid geometric design using computer-aided design software.
2. **STL Conversion:** Tessellating the boundary surface of the CAD model into a mesh of triangles, saving it as a standard triangle language (.STL) file.
3. **Slicing and Transfer:** A slicing program cuts the STL mesh into horizontal 2D layers and generates the G-code toolpath.
4. **Build:** The 3D printing hardware executes the layer-by-layer material deposition or consolidation.
5. **Post processing:** Removing support structures, curing, cleaning, or surface polishing to prepare the part for use.

Step 1: Evaluate the logical progression of the process chain. - You must design the part model before converting its format: CAD → STL conversion. - The machine requires the sliced format data before it can begin printing: STL conversion → Build. - Surface finishing or support removal can only occur after the part is printed: Build → Post processing. Thus, the correct logical sequence is: **CAD → STL conversion → Build → Post processing**.

Quick Tip: Think of the physical workflow to verify the order: Design it digitally (**CAD**), export it (**STL**), print it on the machine (**Build**), and clean it up (**Post processing**).

106. In time series analysis of forecasting technique, which of the source variation can be estimated by ratio to trend method?

- (A) Cyclical
- (B) Irregular
- (C) Trend
- (D) Seasonal

Correct Answer: (D) Seasonal

Solution:

Concept: A time series data model contains four distinct components of variation:

1. **Trend (T):** Long-term upward or downward directional shifts over years.
2. **Seasonal (S):** Variations that repeat regularly within a fixed period of one year or less (e.g., higher ice cream sales in summer).
3. **Cyclical (C):** Long-term oscillations around the trend line driven by multi-year economic business cycles.
4. **Irregular (I):** Unpredictable random shocks or noise.

The **Ratio-to-Trend Method** assumes a multiplicative relationship among these time series components:

$$Y = T \times S \times C \times I$$

Step 1: Explain the operational steps of the Ratio-to-Trend method. 1. The long-term trend value (T) is calculated for each period using a linear regression model. 2. The actual

observed data value (Y) is divided by its calculated trend value (T):

$$\frac{Y}{T} = \frac{T \times S \times C \times I}{T} = S \times C \times I$$

3. By averaging these ratios across multiple years for identical calendar periods (e.g., averaging all past month-of-July values), the cyclical and irregular variations smooth out toward baseline averages. This isolates the **Seasonal Index (S)**. Therefore, this method is primarily used to isolate and estimate **seasonal** variations.

Quick Tip: The **Ratio-to-Trend** method divides out the baseline trend vector from time series data to calculate precise **Seasonal Indices**.

107. The correct order of series of operations to be carried out in production planning control is:

- (A) Dispatching – Routing – Scheduling – Follow up
- (B) Routing – Scheduling – Follow up – Dispatching
- (C) Routing – Scheduling – Dispatching – Follow up
- (D) Scheduling – Routing – Dispatching – Follow up

Correct Answer: (C) Routing – Scheduling – Dispatching – Follow up

Solution:

Concept: Production Planning and Control (PPC) coordinates the operational workflow of a manufacturing facility using four sequential phases:

1. **Routing:** Determines the optimal path, sequence of operations, and machines required to transform raw materials into a finished part. (Answers **Where** and **How**).
2. **Scheduling:** Assigns specific start and completion times for each operation. (Answers **When**).
3. **Dispatching:** Issues official production orders and releases the work packages to the factory floor to initiate real-time operations. (Answers **Action**).
4. **Follow up (Expediting):** Monitors active production progress against the plan to catch bottlenecks and fix delays. (Answers **Verification**).

Step 1: Establish the logical sequence of operations. - A factory must plan the operational path before assigning execution times: Routing → Scheduling. - Timelines must be finalized before authorizing the shop floor to begin work: Scheduling → Dispatching. - Monitoring and expediting can only occur after the work has been dispatched and initiated: Dispatching → Follow up.

Therefore, the correct operational sequence is: **Routing → Scheduling → Dispatching → Follow up**.

Quick Tip: Remember the classic PPC acronym sequence: **R-S-D-F** (Routing → Scheduling → Dispatching → Follow up).

108. In the first come first service (FCFS) scheduling algorithm, which process is executed first?

- (A) The process with the shortest burst time
- (B) The process that arrives first
- (C) The process with the highest priority
- (D) The process with the longest burst time

Correct Answer: (B) The process that arrives first

Solution:

Concept: The First-Come, First-Served (FCFS) scheduling algorithm manages tasks using a queue structure based on arrival order.

- **FCFS Strategy:** A non-preemptive scheduling policy that allocates resources to tasks in the exact order they enter the ready queue.
- Alternative strategies include **Shortest Job First (SJF)** (prioritizes the shortest burst time) and **Priority Scheduling** (prioritizes the highest assigned rank).

Step 1: Apply the queue policy to the options. The FCFS algorithm operates on a First-In, First-Out (FIFO) basis, similar to a standard service queue. The resource manager checks the arrival timestamps of all pending tasks and allocates execution capacity to the task with the earliest timestamp. Therefore, **the process that arrives first** is executed first, regardless of its processing length or priority rank.

Quick Tip: **FCFS** scheduling is based entirely on the chronological arrival sequence: **First to Arrive = First to be Processed**.

109. The term Therbligs are generally relevant with

- (1) Fatigue allowance in work measurement
- (2) Scheduling processes
- (3) Elementary motions of time study
- (4) Efficiency of the operator

Correct Answer: (3) Elementary motions of time study

Solution:

Concept: In motion and time study within industrial engineering, a "Therblig" is a classification system consisting of fundamental, elemental micro-motions required to perform a physical operation or manual task. Devised by Frank and Lillian Gilbreth, the word "Therblig" is an anagram of the creators' surname spelled backwards (with the 'th' treated as one unit). By breaking operations down into these elemental steps, engineers can analyze and eliminate wasteful, non-value-added body movements to standardize workflows and maximize efficiency.

Step 1: Understanding the definition and origin of Therbligs.

Frank and Lillian Gilbreth identified 17 (later expanded to 18) fundamental subdivisions of hand and body movements. These atomic elements describe every movement occurring at a workstation. Examples include:

- **Search (Sh):** The eyes or hands hunting for an object.
- **Find (F):** A mental reaction that occurs at the end of a search cycle.
- **Select (St):** Choosing one object from among several options.
- **Grasp (G):** Closing fingers around an object to gain physical control.
- **Hold (H):** Retaining an object in a fixed position with a hand.

Step 2: Mapping Therbligs to time and motion study fields.

Because Therbligs represent microscopic elements of work, they are heavily utilized during micromotion studies. Analysts film an operator performing a task and then review it frame-by-frame, categorizing every basic physical action into its respective Therblig element. This

microscopic breakdown is a core methodology used exclusively for tracking the **elementary motions of time and motion study**, making Option (3) the correct choice.

Quick Tip: To easily remember Therbligs: - It is simply "Gilbreth" spelled almost backwards! - Whenever you see "Therbligs," associate it instantly with micro-motion analysis, body gestures, and elementary manual tasks in time study.

110. If the objective is maximizing the quality production by ensuring production machine availability and reducing machine breakdowns, which of the following lean tool will best work for this?

- (1) Kaizen
- (2) Visual controls
- (3) Control plan
- (4) Total productive maintenance

Correct Answer: (4) Total productive maintenance

Solution:

Concept: Lean manufacturing encompasses several specialized methodologies aimed at eliminating waste (*muda*). When dealing specifically with manufacturing plant assets, equipment uptime, mechanical reliability, and machine availability, the primary framework implemented is Total Productive Maintenance (TPM). TPM seeks to integrate equipment maintenance directly into the daily operational lifecycle, empowering line operators to share responsibility for upkeep and preventing degradation.

Step 1: Analyzing the precise requirements of the objective.

The question explicitly highlights three strategic operational goals:

- 1. Maximizing high-quality production outputs.
- 2. Maximizing production machine availability (uptime).
- 3. Drastically reducing machine failures and sudden mechanical breakdowns.

Step 2: Evaluating the operational definitions of the given options.

Let us examine what each tool targets within an industrial facility:

- **Kaizen:** Focuses broadly on continuous, incremental improvements across all departments (processes, safety, human workflows), not strictly machine mechanical lifecycles.
- **Visual Controls:** Incorporates color-coding, layout lines, and status indicators (*andon*) to clarify workplace status at a glance.
- **Control Plan:** A documented blueprint summarizing the critical dimensions, tolerances, and inspection criteria used to monitor and regulate a manufacturing process.
- **Total Productive Maintenance (TPM):** A comprehensive approach centered on zero breakdowns, zero accidents, and zero defects. It utilizes specialized pillars like Autonomous Maintenance (*Jishu Hozen*) and Planned Maintenance to maximize Overall Equipment Effectiveness (OEE).

Since the primary levers here are machine availability and breakdown reduction, TPM fits the objective perfectly. Hence, Option (4) is correct.

Quick Tip: Keep these direct associations in mind for Lean tool questions: - Breakdowns + Availability → Total Productive Maintenance (TPM) - Continuous Small Changes → Kaizen - Workplace Organization + Layout → 5S / Visual Management

111. For an assembly line, the inventory was calculated by considering the production rate as 4 pieces per hour and the average processing time as 60 minutes. Now, if the production rate is kept the same and the average processing time is brought down by 30%, then the change in inventory is

- (1) Decreased by 25%
- (2) Increased by 25%
- (3) Decreased by 30%
- (4) Increased by 30%

Correct Answer: (3) Decreased by 30%

Solution:

Concept: The fundamental relationship governing inventory, production/throughput rate, and processing time (or cycle time) in a manufacturing system or stable queuing process is

described by **Little's Law**. Little's Law states that the average inventory level (I) is equal to the product of the long-term average production/throughput rate (R) and the average processing/lead time (T). Mathematically, this is expressed as:

$$I = R \times T$$

Step 1: Calculate the initial inventory configuration (I_1).

We are given the initial processing characteristics of the assembly line: * Initial Production Rate, $R_1 = 4$ pieces per hour * Initial Average Processing Time, $T_1 = 60$ minutes = 1 hour
Using Little's Law, we find the baseline initial inventory value:

$$I_1 = R_1 \times T_1$$

$$I_1 = 4 \text{ pieces/hour} \times 1 \text{ hour} = 4 \text{ pieces}$$

Step 2: Determine the modified processing properties (R_2 and T_2).

The problem presents a modification to the line parameters: 1. The production rate remains completely constant:

$$R_2 = R_1 = 4 \text{ pieces per hour}$$

2. The average processing time is decreased by 30%:

$$T_2 = T_1 \times \left(1 - \frac{30}{100}\right) = 1 \text{ hour} \times 0.70 = 0.70 \text{ hours}$$

Alternatively, in minutes: $T_2 = 60 \text{ minutes} - (0.30 \times 60 \text{ minutes}) = 42 \text{ minutes} = \frac{42}{60} \text{ hours} = 0.70 \text{ hours}$.

Step 3: Calculate the new inventory (I_2) and percentage change.

Re-applying Little's Law with our modified conditions gives:

$$I_2 = R_2 \times T_2 = 4 \text{ pieces/hour} \times 0.70 \text{ hours} = 2.8 \text{ pieces}$$

Now, compute the exact percentage variation in the inventory level ($\Delta I\%$):

$$\Delta I\% = \left(\frac{I_2 - I_1}{I_1}\right) \times 100\%$$

$$\Delta I\% = \left(\frac{2.8 - 4.0}{4.0}\right) \times 100\% = \left(\frac{-1.2}{4.0}\right) \times 100\%$$

$$\Delta I\% = -0.30 \times 100\% = -30\%$$

The negative sign signifies a clear reduction. Hence, the inventory decreases by exactly 30%. This matches Option (3).

Quick Tip: Since $I = R \times T$, if the production rate (R) is kept entirely constant, the inventory (I) becomes directly proportional to the processing time (T):

$$I \propto T$$

Therefore, whatever percentage change happens to the processing time will directly map onto the inventory level! A 30

112. To issue and running balance of certain items of stock in inventory, which of the following keeps a record of items?

- (1) Bin cards
- (2) Stock items
- (3) Quantity account
- (4) Value account

Correct Answer: (1) Bin cards

Solution:

Concept: In physical stores management and warehousing operations, it is critical to log quantitative adjustments at the precise physical storage location (the "bin," shelf, or rack) whenever materials are deposited or withdrawn. The standard operational document utilized for tracking these real-time modifications directly inside the storage bay is a ****Bin Card****. It tracks material quantities chronologically, listing receipts, issues, and providing an updated running balance.

Step 2: Analyzing the functionality of a Bin Card versus accounting alternatives.

Let us distinguish the primary documentation tools used in inventory control:

- **Bin Card:** Maintained directly by the physical storekeeper right next to the stock container. It records only the quantitative physical parameters (quantities received, issued, and balances on hand) immediately following a transaction.

- **Stores Ledger:** Maintained within the cost accounting department. It mirrors the bin card but logs both physical quantities *and* financial valuations (monetary cost).
- **Quantity/Value accounts:** These refer to broad ledger balances and summary entries in financial software packages rather than the localized physical tracking documents used to record dynamic balances upon item issuance.

Thus, the specific physical tool designed to log receipts, issues, and maintain a running balance at the point of storage is the Bin Card. This matches Option (1).

Quick Tip: Key differences to keep in mind for examinations: - Bin Card → Maintained by the Store-keeper, contains quantities only, kept in the physical warehouse. - Stores Ledger → Maintained by the Cost Clerk, contains both quantities and pricing values, kept in the accounting office.

113. The demand for a commodity is 100 units per day. Every time an order is placed, a fixed cost of Rs 400 is incurred and the holding cost is Rs 0.08 per unit per day. If the lead time is 13 days, then the economic lot size and the reorder point are in units

- (1) 800 and 130
- (2) 840 and 100
- (3) 890 and 300
- (4) 1000 and 300

Correct Answer: (4) 1000 and 300

Solution:

Concept: The classical **Economic Order Quantity (EOQ)** model determines an optimal purchase lot size that minimizes total variable inventory costs (comprising setup/ordering costs and carrying/holding costs). The standard formula for economic lot size is:

$$Q^* = \sqrt{\frac{2 \cdot D \cdot C_o}{C_h}}$$

Where: * D = Demand rate (units per time period) * C_o = Fixed cost per individual order placed (ordering cost) * C_h = Holding cost per unit per time period (holding cost)

The **Reorder Point (ROP)** represents the inventory level at which a replenishment order

must be triggered, calculated as:

$$\text{ROP} = \text{Demand rate } (d) \times \text{Lead time } (L)$$

Note: If the inventory remaining during the lead time exceeds one full order cycle, we adjust for the replenishment cycle via the modulo operator to determine the physical hand-count level.

Step 1: Extract variables and compute the Economic Lot Size (Q^*).

From the question statement, we harvest the parameters normalized on a daily baseline: * Daily Demand, $D = 100$ units/day * Ordering Cost, $C_o = \text{Rs. } 400$ per order * Holding Cost, $C_h = \text{Rs. } 0.08$ per unit per day

Let us substitute these expressions directly into the economic lot size equation:

$$Q^* = \sqrt{\frac{2 \times 100 \times 400}{0.08}}$$

$$Q^* = \sqrt{\frac{80000}{0.08}}$$

To eliminate the decimal denominator, multiply the numerator and denominator by 100:

$$Q^* = \sqrt{\frac{8000000}{8}} = \sqrt{1000000} = 1000 \text{ units}$$

Step 2: Calculate the Reorder Point (ROP).

We are given that the lead time $L = 13$ days. First, we find the gross demand consumption accumulating across this lead period:

$$\text{Total Lead Time Demand} = D \times L = 100 \text{ units/day} \times 13 \text{ days} = 1300 \text{ units}$$

Since our order size ($Q^* = 1000$ units) is smaller than the total lead time demand (1300 units), it implies that an additional order must have already been placed while a previous order was outstanding. The physical inventory level on hand that triggers the alert is calculated by evaluating the total lead time demand modulo the economic order quantity:

$$\text{ROP} = (D \times L) \pmod{Q^*}$$

$$\text{ROP} = 1300 \pmod{1000} = 300 \text{ units}$$

Hence, the economic lot size is 1000 units and the reorder point is 300 units, aligning with Option (4).

Quick Tip: Always ensure your time units are fully uniform! Here, demand is units/day and holding cost is Rs/unit/day. Because they match, you do not need to convert them into annual figures. For ROP remember: if Lead Time Demand > Q, subtract Q repeatedly until the value is less than Q.

$$1300 - 1000 = 300$$

114. Average stock level required to maintain the inventory is calculated as

- (1) Minimum stock level + half of the reorder cost
- (2) Maximum stock level + half of the reorder cost
- (3) Minimum stock level + one fourth of the reorder cost
- (4) Maximum stock level + one fourth of the reorder cost

Correct Answer: (1) Minimum stock level + half of the reorder cost

Solution:

Concept: In traditional inventory control systems, stock levels fluctuate dynamically between a peak upper boundary and a base lower boundary. To estimate annual holding costs, managers determine an average inventory level. The standard formula configurations for average stock level are:

$$\text{Average Stock Level} = \text{Minimum Stock Level} + \frac{1}{2}(\text{Reorder Quantity})$$

Alternatively, it can be approximated as the midpoint between boundaries:

$$\text{Average Stock Level} = \frac{\text{Minimum Stock Level} + \text{Maximum Stock Level}}{2}$$

Note: In the provided question options, the term "reorder cost" is used in place of "reorder quantity/size" due to varied regional terminology or semantic error in exam drafting. Treating it as the reorder quantity factor resolves the structural expression.

Step 1: Evaluating the mathematical derivation of average inventory.

Consider a classic inventory sawtooth wave diagram where safety stock represents the absolute

baseline floor (Minimum Stock Level). When a shipment arrives, the inventory spikes by the total reorder size (Q). Thus, the maximum height reached is:

$$\text{Maximum Stock} = \text{Minimum Stock} + Q$$

As inventory is consumed at a constant rate, the average quantity held above the safety floor is exactly half of the incoming batch size, $\frac{1}{2}Q$.

Step 2: Matching with the given options.

Expressing this mathematically gives:

$$\text{Average Stock} = \text{Minimum Stock Level} + \frac{1}{2}(\text{Reorder Quantity})$$

Reviewing the multiple-choice options, Option (1) reads: Minimum stock level + half of the reorder cost. Swapping "reorder quantity" into the placeholder term confirms Option (1) as the intended correct answer.

Quick Tip: Remember the two standard expressions for Average Stock Level: 1. Minimum Level + $\frac{1}{2}(\text{Reorder Quantity})$ 2. $\frac{\text{Minimum Level} + \text{Maximum Level}}{2}$ This makes it simple to spot the correct structure on a test.

115. Which of the following relationships hold correct for the safety stock of inventory?

- (1) The greater the risk of running out of stock, the small the safety of stock
- (2) The larger the opportunity cost of the funds invested inventory, the larger the safety stock
- (3) The greater the uncertainty associated with forecasted demand, the higher the safety stock
- (4) None of the options are correct

Correct Answer: (3) The greater the uncertainty associated with forecasted demand, the higher the safety stock

Solution:

Concept: **Safety Stock** serves as a protective buffer held in inventory to safeguard against stockouts caused by unpredictable fluctuations in demand rates or replenishment lead times. The statistical formula for safety stock when demand during lead time is normally distributed

is:

$$\text{Safety Stock} = z \times \sigma_L$$

Where z is the standard normal service factor corresponding to the desired service level, and σ_L represents the standard deviation of demand uncertainty during the lead time.

Step 1: Analyzing the relationship between demand uncertainty and safety buffer.

The standard deviation σ_L directly quantifies the degree of uncertainty or error present within demand forecasting models. If demand is completely stable and predictable, $\sigma_L = 0$, meaning zero safety stock is required. As market volatility and forecasting uncertainty escalate, σ_L grows larger, requiring a larger safety stock buffer to ensure the same service level protection. Thus, a direct proportional dependency exists:

$$\text{Uncertainty} \uparrow \Rightarrow \text{Safety Stock} \uparrow$$

This validates Option (3) as a true statement.

Step 2: Disproving the alternative options.

* **Option (1) Analysis:** If the risk or impact of running out of stock is highly severe, a company must increase its safety stock to reduce that vulnerability, not lower it. * **Option (2) Analysis:** Capital opportunity cost represents the financial penalty of tying up funds in stock. A higher opportunity cost creates a stronger incentive to minimize idle assets, prompting a *reduction* in safety stock levels rather than an increase.

Therefore, Option (3) is the only relationship that holds correct.

Quick Tip: Think of safety stock like an umbrella: the more erratic and unpredictable the weather forecast (high uncertainty), the larger the umbrella you need to carry (higher safety stock) to avoid getting wet (stockout)!

116. The pessimistic time and optimistic time of completion of an activity are given as 10 days and 4 days respectively. Then the variance of the activity is

- (1) 14
- (2) 6
- (3) 2
- (4) 1

Correct Answer: (4) 1

Solution:

Concept: Within the **Project Evaluation and Review Technique (PERT)** framework, the completion duration of any specific activity is treated as a random variable modeled via a Beta probability distribution. To describe this distribution, three time estimates are collected: 1. Optimistic time (t_o or a): Minimum possible timeframe if everything goes perfectly. 2. Most likely time (t_m or m): The modal or standard duration expected under typical circumstances. 3. Pessimistic time (t_p or b): Maximum duration required if significant obstacles are encountered. In PERT theory, the standard deviation (σ) of an activity's duration is estimated by dividing the total practical span of the distribution by 6 standard deviations:

$$\sigma = \frac{t_p - t_o}{6}$$

The statistical variance (σ^2) is the square of this standard deviation:

$$\text{Variance} = \sigma^2 = \left(\frac{t_p - t_o}{6} \right)^2$$

Step 1: Isolate variables from the text description.

The problem provides the following parameters: * Pessimistic completion time, $t_p = 10$ days * Optimistic completion time, $t_o = 4$ days

Step 2: Calculate the standard deviation (σ).

Substitute our parameters directly into the linear PERT approximation equation:

$$\sigma = \frac{10 - 4}{6}$$

$$\sigma = \frac{6}{6} = 1 \text{ day}$$

Step 3: Calculate the activity variance (σ^2).

Square the standard deviation value computed in the previous step:

$$\text{Variance} = \sigma^2 = (1)^2 = 1$$

The computed variance is exactly 1, which corresponds to Option (4).

Quick Tip: Be careful not to confuse standard deviation with variance! - Standard Deviation = $\frac{t_p - t_o}{6}$ -
Variance = $\left(\frac{t_p - t_o}{6}\right)^2$ Always make sure to square your result when a question asks for the variance.

117. In a graphical solution of solving Linear Programming Problem to convert inequalities into equations, they

- (1) Use slack variables
- (2) Use surplus variables
- (3) Use artificial surplus variables
- (4) Assume them to be equations

Correct Answer: (4) Assume them to be equations

Solution:

Concept: When utilizing the **Graphical Method** to solve a two-variable Linear Programming Problem (LPP), constraints are initially defined as inequalities that bound a geometric half-space (e.g., $A_1x + B_1y \leq C_1$). To map the boundary lines of these constraints on a two-dimensional Cartesian plane, we treat the inequality as a strict equality line equation ($A_1x + B_1y = C_1$).

Step 1: Differentiating the requirements of the graphical method vs. the algebraic simplex algorithm.

Let us analyze how constraints are managed depending on the solution path chosen:

- **Simplex Method Framework:** This algebraic approach requires transforming inequality constraints into canonical standard forms via auxiliary variables. We add non-negative **slack variables** to adjust "less-than-or-equal-to" ($\leq$) constraints, and subtract **surplus variables** paired with **artificial variables** to adjust "greater-than-or-equal-to" ($\geq$) constraints.
- **Graphical Method Framework:** This method relies on plotting boundary lines directly on a graph to isolate the feasible region. To plot an inequality line like $2x + 3y \leq 6$, we simply construct the linear plot for $2x + 3y = 6$ by identifying its coordinate intercepts.

Step 2: Concluding the proper choice.

Because the graphical approach does not require introducing auxiliary slack or surplus variables to plot the system boundaries, we simply assume the constraints to function as equations during line construction. This aligns with Option (4).

Quick Tip: Watch out for the solution method specified in the question: - If the text mentions **Simplex Method** → look for Slack, Surplus, or Artificial variables. - If the text mentions **Graphical Method** → simply plot them by assuming them to be straight line equations!

118. The penalty of a row in a transportation problem is obtained by

- (1) Deducting smallest element in the row from all other elements of the row
- (2) Deducting smallest element in the row from the next highest element of the row
- (3) Adding smallest element in the row to all other elements of the row
- (4) None of the options are correct

Correct Answer: (2) Deducting smallest element in the row from the next highest element of the row

Solution:

Concept: In optimization modeling, **Vogel's Approximation Method (VAM)** is a structured algorithm used to find an initial feasible solution for a transportation network problem. VAM relies on calculating opportunity cost penalties for each row and column. A row penalty represents the cost penalty incurred if an optimizer fails to allocate cargo to the absolute cheapest cell in that row, forcing an allocation to the next best alternative routing cell.

Step 1: Understanding the algorithmic formulation of a row penalty.

To calculate row and column penalties during VAM analysis, follow these operational steps:

1. Examine the cost elements within a specific row of the transportation matrix.
2. Identify the absolute minimum unit cost element present in that row.
3. Identify the next lowest unit cost element (the second smallest or next highest element) in the same row.
4. Subtract the smallest element from this next highest element.

Mathematically, for any row i :

$$\text{Penalty}_i = (\text{Second Smallest Cost Value in Row } i) - (\text{Absolute Smallest Cost Value in Row } i)$$

Step 2: Evaluating the text options.

Reviewing our choices: * Option (1) describes a subtraction process similar to row reduction in the Hungarian Assignment method, not VAM. * Option (2) states: *"Deducting smallest element in the row from the next highest element of the row."* This perfectly matches the mathematical definition of calculating a VAM opportunity penalty.

Hence, Option (2) is correct.

Quick Tip: Vogel's Approximation Method (VAM) Penalty rule can be easily summarized as:

$$\text{Penalty} = \text{Second Minimum Element} - \text{First Minimum Element}$$

This calculated value measures the financial risk of missing out on the absolute cheapest route.

119. Which of the cost estimates and performance measures are not used for economic analysis of a queuing system?

- (1) Cost per server per unit of time
- (2) Average waiting time of customers in the system
- (3) Cost per unit time for a customer waiting in the system
- (4) The average number of customers in the system

Correct Answer: (2) Average waiting time of customers in the system

Solution:

Concept: The economic optimization of a service facility using queuing models relies on minimizing the **Total Expected Cost ($E(TC)$)** of the operation. This economic trade-off model balances two competing financial components: 1. The operational cost of providing service (paying for active service capacity/channels). 2. The cost of customer dissatisfaction and congestion caused by waiting in line.

The governing cost optimization objective equation per unit time is:

$$E(TC) = C_s \cdot c + C_w \cdot L_s$$

Where: * c = Number of active parallel servers. * C_s = Service cost per individual server per unit of time. * C_w = Waiting cost per unit time per customer waiting in the system. * L_s = Expected total number of customers present within the system.

Step 1: Identifying components utilized in the total cost model.

Let us look at the variables required to calculate the economic balance: * C_s maps directly to Option (1): *Cost per server per unit of time*. * C_w maps directly to Option (3): *Cost per unit time for a customer waiting in the system*. * L_s maps directly to Option (4): *The*

average number of customers in the system.

These three criteria are essential parameters used to compute the overall hourly operating cost of a queue.

Step 2: Isolating the unused metric.

While the average waiting time of a customer in the system (W_s) is a valid performance measure for a queue, it is a time duration metric ($W_s = \frac{L_s}{\lambda}$ via Little's Law) rather than a direct cost or customer volume metric used in the $E(TC)$ objective formula. Since the equation multiplies financial rates by customer volume (L_s), W_s itself is not directly plugged into the cost formula. This makes Option (2) the correct selection.

Quick Tip: To optimize a queue's cost, remember the balancing scale formula:

$$\text{Total Cost} = (\text{Server Cost} \times \text{Servers}) + (\text{Waiting Cost Rate} \times \text{Number of Customers})$$

This makes it clear that the *number of customers* (L_s) is used to compute costs, rather than the *waiting time duration* (W_s).

120. The activities of W, X, Y and Z are the direct precursors of A. What is the earliest starting time for A, if the activities of earliest finishing times are 12, 15, 10 and 6?

- (1) 6
- (2) 10
- (3) 15
- (4) 12

Correct Answer: (3) 15

Solution:

Concept: In project scheduling frameworks using the **Critical Path Method (CPM)** , the timing of network tasks is calculated via a forward-pass analysis. For any given activity A, its **Earliest Start Time (ES_A)** represents the earliest possible time step at which the activity can begin.

If an activity has multiple immediate preceding tasks (precursors), it cannot start until *every single one* of those prerequisite tasks has finished. Therefore, during the forward pass, the Earliest Start time of a dependent activity is determined by selecting the absolute maximum

value among the Earliest Finish times (EF) of all its direct preceding tasks:

$$ES_A = \max\{EF_{\text{precursor 1}}, EF_{\text{precursor 2}}, \dots, EF_{\text{precursor } n}\}$$

Step 1: Isolate the precursor data.

We are given that activity A depends on four independent parallel precursors: W, X, Y , and Z . Their respective earliest finishing times are: * $EF_W = 12$ * $EF_X = 15$ * $EF_Y = 10$ * $EF_Z = 6$

Step 2: Apply the forward pass maximum selection rule.

Because activity A cannot begin if any of its prerequisites remain incomplete, it must wait for the longest task to finish. Let us calculate the maximum value from the pool of finishing times:

$$ES_A = \max\{EF_W, EF_X, EF_Y, EF_Z\}$$

$$ES_A = \max\{12, 15, 10, 6\}$$

$$ES_A = 15$$

Thus, the earliest starting time for activity A is 15, which corresponds to Option (3).

Quick Tip: Remember these path routing rules for network diagrams: - **Forward Pass (Earliest Times):** Use the **MAXIMUM** value of preceding finish times when paths merge. - **Backward Pass (Latest Times):** Use the **MINIMUM** value of succeeding start times when paths split.