



RBI DEPR PHASE 2 PAPER 1

DISCLAIMER — PLEASE READ

This is an UNOFFICIAL, memory-based practice paper. It has been reconstructed from topics and questions recalled by several candidates who appeared for the RBI DEPR Phase II 2026 examination (across both shifts), and is meant purely as an exam-preparation aid.

PAPER 1 — ECONOMICS**Instructions**

Total marks: 100. Time: 3 hours.

This paper has TWO sections — Section A (Micro-Economics) and Section B (Macro-Economics), each with 4 questions.

Attempt any FIVE questions in all, with a MINIMUM of TWO questions from each section.

Each question carries 20 marks. Support your answers with diagrams, derivations and economic reasoning wherever applicable.

SECTION A — MICRO-ECONOMICS

Q1. [Total: 20 marks]

(a) State and prove the three conditions necessary for the attainment of Pareto optimality (efficiency in production, efficiency in consumption, and efficiency in product-mix) in a competitive economy. [10]

(b) “Pareto efficiency is, at best, a necessary and not a sufficient criterion for social welfare.” Discuss this statement with reference to Amartya Sen’s capability approach as an alternative welfare criterion. Also explain briefly Sen’s Liberal Paradox. [10]

Answer:

(a) Pareto optimality (or Pareto efficiency) describes a state of resource allocation in which it is impossible to reallocate resources so as to make at least one economic agent better off without making at least one other agent worse off. In a simple economy with two goods (X and Y), two factors (Labour L and Capital K), and two consumers (A and B), attaining a Pareto-optimal allocation requires three separate marginal conditions to hold simultaneously — one governing production, one governing consumption, and one linking the two. Each is proved below using constrained optimisation, followed by its economic intuition.

Condition 1: Efficiency in Production

Statement: The Marginal Rate of Technical Substitution between labour and capital must be equal in the production of both goods:

MRTS of Labour for Capital in producing X = MRTS of Labour for Capital in producing Y

Proof: With fixed total factor endowments ($\bar{L}$ and $\bar{K}$), efficiency requires maximising output of X for any given output level of Y. Setting up the problem: maximise $X = f(L_x, K_x)$ subject to $Y = g(\bar{L} - L_x, \bar{K} - K_x) = \bar{Y}$ (a fixed target level of Y).

Using a Lagrange multiplier (λ) and taking first-order conditions with respect to L_x and K_x gives two equations:

- Marginal product of labour in X = $\lambda \times$ Marginal product of labour in Y
- Marginal product of capital in X = $\lambda \times$ Marginal product of capital in Y

Dividing the first equation by the second cancels out λ , leaving:

(Marginal product of labour in X) / (Marginal product of capital in X) = (Marginal product of labour in Y) / (Marginal product of capital in Y)

This ratio, by definition, is the MRTS of labour for capital in each sector — proving $MRTS^x = MRTS^y$.

Intuition and Example: Suppose the textile sector (X) has an MRTS of 3 (can give up 3 units of capital for 1 extra unit of labour and keep output constant), while the steel sector (Y) has an MRTS of 1. Then shifting one unit of capital from steel to textiles, and simultaneously shifting labour from textiles to steel, allows textiles to give up less labour than steel needs — raising total combined output at no extra cost. Only when both sectors' MRTS are equal (as with tangent isoquants in the Edgeworth production box) is this kind of costless reallocation exhausted. The locus of all such tangency points is called the Production Contract Curve, and it maps directly into the economy's Production Possibility Frontier (PPF).

Condition 2: Efficiency in Consumption

Statement: The Marginal Rate of Substitution between the two goods must be equal across all consumers:

MRS of X for Y for Consumer A = MRS of X for Y for Consumer B

Proof: With a fixed total endowment of goods ($\bar{X}$ and $\bar{Y}$), consumption efficiency requires maximising Consumer A's utility for any given level of Consumer B's utility. Maximise $U^A(X_A, Y_A)$ subject to $U^B(\bar{X}-X_A, \bar{Y}-Y_A) = a$ fixed target level of B's utility.

The first-order conditions (using multiplier μ) give:

- Marginal utility of X for A = $\mu \times$ Marginal utility of X for B
- Marginal utility of Y for A = $\mu \times$ Marginal utility of Y for B

Dividing one by the other cancels μ :

$$\frac{(\text{Marginal utility of X for A})}{(\text{Marginal utility of Y for A})} = \frac{(\text{Marginal utility of X for B})}{(\text{Marginal utility of Y for B})}$$

This ratio is, by definition, each consumer's MRS — proving $MRS^A = MRS^B$.

Intuition and Example: Suppose Consumer A (a vegetarian household) values an extra unit of rice at 4 units of vegetables given up, while Consumer B values it at only 1 unit of vegetables. A mutually beneficial trade exists — B can sell rice to A for, say, 2 units of vegetables, making both better off. Trade continues until both consumers' MRS converge to the same rate — the point of tangency between their indifference curves in the Edgeworth consumption box, tracing the Consumption Contract Curve. Once MRS is equalised, all gains from voluntary exchange are exhausted.

Condition 3: Efficiency in Product-Mix

Statement: The Marginal Rate of Transformation between the two goods (the slope of the PPF) must equal the common Marginal Rate of Substitution of consumers:

MRT of X for Y (in production) = MRS of X for Y (in consumption)

Proof: Given the PPF derived from Condition 1 (representing all production-efficient combinations of X and Y), and the common MRS from Condition 2, overall efficiency requires choosing the specific point on the PPF that maximises consumer welfare. Maximising Consumer A's utility subject to the PPF and holding B's utility fixed yields, at the optimum:

$$\text{MRT of X for Y} = \text{MRS of X for Y}$$

where MRT is the slope of the PPF (the opportunity cost, in units of Y foregone, of producing one more unit of X), and MRS is consumers' willingness to trade Y for X.

Intuition and Example: Suppose the economy can convert 1 unit of X into 2 units of Y through reallocating production ($\text{MRT} = 2$), but consumers are willing to give up only 1 unit of Y for 1 extra unit of X ($\text{MRS} = 1$). Since consumers value the resource cost of X less than what production actually sacrifices to make it, society is over-producing X relative to what consumers want — shifting production toward Y raises total welfare. Only when MRT equals MRS — say, an economy producing wheat and cars where the opportunity cost of one more car (in wheat foregone) exactly matches what consumers are willing to pay for that car in terms of wheat — is the product-mix itself optimal.

Conclusion

The three conditions — equal MRTS across sectors (production efficiency), equal MRS across consumers (consumption/exchange efficiency), and $\text{MRT} = \text{MRS}$ (product-mix efficiency) — must hold simultaneously for an allocation to be Pareto optimal. This is the analytical content of the First Fundamental Theorem of Welfare Economics: under perfect competition, with no externalities and standard convexity assumptions, profit-maximising firms (equating MRTS to the factor-price ratio w/r) and utility-maximising consumers (equating MRS to the goods-price ratio P_x/P_y) automatically satisfy all three conditions through the price mechanism, since firms also produce where price equals marginal cost (ensuring $\text{MRT} = P_x/P_y = \text{MRS}$). This is the formal justification for the classical claim that competitive markets, left undistorted, achieve productive and allocative efficiency.

(b) Pareto optimality tells us only that no further mutually beneficial reallocation is possible — it says nothing about whether the resulting distribution is fair, just, or good for society as a whole. This is why Pareto efficiency is widely regarded as a necessary but not sufficient condition for social welfare: any acceptable social outcome should be efficient, but not every efficient outcome is acceptable. Amartya Sen's work — both his capability approach and his Liberal Paradox — are among the most influential critiques of relying on Pareto efficiency (or its underlying utilitarian/welfarist framework) as a complete guide to social welfare.

Why Pareto Efficiency Is Not Sufficient

1. **Indifference to distribution:** For any economy, there exists an entire Contract Curve — infinitely many Pareto-optimal allocations, ranging from highly equal to highly unequal. An allocation in which one person owns virtually everything and others live in destitution can still be Pareto optimal, since no reallocation can help the deprived without taking something from the wealthy person. Pareto's criterion is silent on which of these many efficient points society should choose.
2. **The Second Welfare Theorem's silence on equity:** While the Second Fundamental Theorem shows that any Pareto-optimal allocation can, in principle, be reached as a competitive equilibrium through an appropriate lump-sum redistribution of initial endowments, the theorem does not tell us which redistribution is desirable — that judgment must come from outside the Pareto framework, via some social welfare function or ethical criterion.
3. **Focus on outcomes, not processes or rights:** Pareto efficiency evaluates only the final allocation of goods, ignoring whether the process that produced it was fair, and ignoring individual freedoms and entitlements — a gap Sen's work addresses directly.

Amartya Sen's Capability Approach

Sen argued that both welfarism (judging social states purely by utility, as Pareto-based welfare economics does) and resourcism (judging well-being purely by income or the bundle of goods a person holds, as in Rawls's "primary goods") are inadequate measures of genuine well-being. He proposed instead to evaluate welfare in terms of:

- **Functionings:** the things a person actually values being or doing — being well-nourished, being literate, participating in community life, avoiding preventable illness.
- **Capabilities:** the real freedom or opportunity a person has to achieve these valuable functionings — i.e., the full set of functioning combinations a person could feasibly choose from, not merely the one they end up with.

Why this matters — example: Two individuals may have identical income, say ₹20,000 per month, yet have very different real freedoms. A person with a physical disability may need a much larger share of that income just to achieve basic mobility, so identical income does not mean identical capability. Similarly, a person who has lived in chronic deprivation may report high subjective satisfaction simply because they have adjusted their expectations downward — Sen calls this the problem of adaptive preferences. A purely utility-based (welfarist/Pareto) measure would wrongly read this adjusted contentment as evidence that well-being is adequate, missing the underlying deprivation entirely.

Practical influence: This approach directly shaped the UNDP's Human Development Index (HDI), which Sen helped conceptualise alongside Mahbub ul Haq. The HDI deliberately moves away from ranking countries by per-capita income (GDP) alone, instead combining life expectancy, education, and standard of living — a direct, real-world application of evaluating welfare through capabilities rather than resources or income, and a practical rejection of relying on efficiency-based or purely monetary measures of social progress.

Sen's Liberal Paradox (The Impossibility of a Paretian Liberal, 1970)

Sen showed that the Pareto principle — often treated as an uncontroversial minimal requirement for any social choice rule — can actually conflict with even a minimal notion of individual liberty.

He set out three conditions a social choice rule might reasonably be expected to satisfy together:

- **Unrestricted domain:** the rule must work for any possible set of individual preferences.
- **Minimal liberalism:** each individual should be decisive over at least one purely personal choice — for example, what they themselves read, or how they decorate their own room — regardless of what others prefer.
- **The Pareto principle:** if everyone prefers outcome X to outcome Y, society should rank X above Y.

Illustration: Consider a Prude and a Lascivious reader deciding who should read a racy novel. The Prude most prefers nobody reads it, but if someone must, prefers reading it himself rather than have the Lascivious reader do so (to "protect" him from further corruption). The Lascivious reader most enjoys the idea of the Prude being made to read it, and next prefers reading it himself, least preferring that nobody reads it. Minimal liberalism would assign each person decisiveness over whether they themselves read the book — this points toward the outcome where nobody reads it. But both individuals actually prefer a different outcome (the Prude reading it) to that one — meaning the Pareto principle recommends overriding what minimal liberalism would otherwise dictate. Sen proved that no social choice rule can simultaneously satisfy unrestricted domain, minimal liberalism, and the Pareto principle without generating such a contradiction.

Significance: The Liberal Paradox demonstrates that even the seemingly harmless Pareto principle can conflict with individual rights once social choice extends beyond simple resource allocation — reinforcing Sen's broader argument that efficiency (Pareto optimality) cannot be treated as the sole or overriding criterion for social welfare, and that liberty, rights, and capability must be built into the welfare criterion independently, not derived from aggregate efficiency alone.

Conclusion

Pareto optimality remains analytically important as a baseline requirement — an allocation that leaves unrealised gains on the table is unambiguously inferior — but it cannot, by itself, resolve questions of distribution, rights, or genuine well-being. Sen's capability approach supplies a richer, freedom-based alternative that has reshaped how institutions like the UNDP measure development, while his Liberal Paradox exposes a deeper tension between efficiency and liberty even within the logical structure of social choice theory. Together, these contributions justify treating Pareto efficiency as a necessary, but decisively not sufficient, criterion for evaluating social welfare — a conclusion of considerable relevance for how a central bank or regulator (such as SEBI or the RBI) should weigh efficiency against equity and rights-based considerations in designing economic policy.

Q2. [Total: 20 marks]

(a) Explain why convexity of indifference curves is an important assumption in consumer theory. How does this assumption help explain observed consumer behaviour (e.g., the diminishing marginal rate of substitution)? [10]

(b) (i) Why is diminishing marginal productivity of a factor not, by itself, a sufficient condition to ensure convexity of an isoquant? (ii) Can a production function display constant returns to scale (CRS) even when each individual factor exhibits diminishing marginal productivity? Explain with a suitable example. [10]

Answer:

(a) In consumer theory, indifference curves depict combinations of two goods that yield the same level of utility to a consumer. One of the central assumptions imposed on consumer preferences is convexity — that indifference curves bow inward toward the origin. This assumption is not a mere mathematical convenience; it is what allows consumer theory to generate economically sensible predictions about how people actually behave, particularly the well-observed tendency for consumers to prefer balanced combinations of goods over extreme concentrations of just one.

An indifference curve is convex to the origin if, for any two consumption bundles that yield equal utility, a bundle formed by averaging the two (a mix, rather than either extreme) is preferred to — or at least as good as — either of the original extreme bundles. Formally, this reflects the assumption of a diminishing Marginal Rate of Substitution (MRS): as a consumer moves along an indifference curve consuming more of good X and less of good Y, the amount of Y they are willing to give up for one additional unit of X steadily falls.

Importance of Convexity:

1. It reflects diminishing marginal utility interacting across goods: As a consumer holds progressively less of good Y (having traded it away for more X), each remaining unit of Y becomes relatively more precious to them, while each additional unit of X becomes relatively less urgently needed (since they already have more of it). This makes the consumer increasingly reluctant to give up further units of Y — precisely the diminishing MRS that convexity captures.
2. It ensures a unique, well-behaved consumer optimum: With convex indifference curves and a linear budget constraint, the point of tangency between the indifference curve and the budget line represents a genuine utility-maximum, and this optimum is unique. If indifference curves were instead concave (bowed outward), the tangency point would represent a utility minimum, and the true optimum would lie at a corner (spending the entire budget on just one good) — a far less realistic prediction of typical consumer behaviour.
3. It guarantees the existence of interior (diversified) solutions: Convexity is what generates the standard prediction that consumers typically buy a mix of goods rather than specialising entirely in one — consistent with everyday observation (people buy both food and clothing, rather than spending their entire income on just one).
4. It underlies comparative statics results: Convexity ensures demand curves slope downward in a well-behaved way and supports standard results such as the Slutsky decomposition of price effects into substitution and income effects — the substitution effect is unambiguously negative precisely because of convex preferences.

How Convexity Explains Diminishing MRS — with an Example

Consider a consumer choosing between rice and vegetables. Suppose they currently consume a bundle heavily weighted toward rice — say, 20 units of rice and 2 units of vegetables. At this point, since vegetables are scarce in their diet, they would be willing to give up a large quantity of rice (say, 8 units) for just 1 additional unit of vegetables — their MRS of rice for vegetables is high. Now suppose, through successive trades, they move to a more balanced bundle — 10 units of rice and 8 units of vegetables. At this new point, rice is no longer abundant relative to vegetables, so they are willing to give up far less rice (say, only 1.5 units) for one more unit of vegetables — the MRS has fallen. This declining willingness to trade one good for another as its relative abundance changes is exactly what produces the characteristic convex, bowed-in shape of the indifference curve, and it matches everyday intuition: the value we place on an additional unit of something depends on how much of it — and its substitutes — we already have.

Conclusion:

Convexity of indifference curves is not an arbitrary technical restriction; it formalises the observed regularity that consumers value variety and balance in consumption, and that their willingness to substitute one good for another diminishes as they acquire progressively more of it. This assumption is what allows consumer theory to yield a unique, interior, and economically realistic optimum, and it forms the foundation for standard demand theory results used throughout applied microeconomics and policy analysis (e.g., in modelling consumer response to relative price changes, such as GST rate adjustments across goods categories).

(b-i) It is a common misconception that diminishing marginal productivity of each factor automatically implies that isoquants are convex to the origin. In fact, these are two conceptually distinct properties of a production function, and the former does not, by itself, guarantee the latter.

The Key Distinction

- Diminishing marginal productivity is a one-dimensional property: it describes what happens to output as one input (say, labour) is increased while the other input (capital) is held fixed. It requires only that the marginal product of labour falls as labour rises (holding capital constant), and separately, that the marginal product of capital falls as capital rises (holding labour constant).
- Convexity of the isoquant is a fundamentally two-dimensional (joint) property: it requires that the Marginal Rate of Technical Substitution ($MRTS = \text{Marginal Product of Labour} \div \text{Marginal Product of Capital}$) diminishes as we move along the isoquant itself — i.e., as both inputs are adjusted simultaneously to keep output constant, substituting one for the other.

Why the Two Are Not Equivalent

Whether the MRTS actually diminishes along an isoquant depends not only on how quickly each factor's own marginal product falls (the "own" effects), but critically also on the cross-effect — how the marginal product of labour changes when capital changes (and vice versa). Formally, the condition for a convex isoquant (technically, quasi-concavity of the production function) involves a combination of the two own second-derivative terms and this cross-partial term together — not the own terms alone. Diminishing marginal productivity guarantees only that the own terms are negative; it says nothing about the magnitude or sign of the cross-effect, which can still cause the MRTS to fail to diminish smoothly, or even become non-monotonic, in some regions.

Illustrative Example

Consider the extreme (limiting) case of a Leontief (fixed-proportions) production function — e.g., producing a fixed quantity of cement always requires exactly 2 units of labour combined with exactly 1 unit of capital, in strict fixed ratio, with no substitution possible between them. In this case, increasing labour alone (beyond the point where capital becomes the binding constraint) adds zero extra output — an extreme, degenerate form of "diminishing" (in fact, exhausted) marginal productivity. Yet the resulting isoquant is L-shaped, not smoothly convex — the MRTS jumps discontinuously from infinity to zero at the corner rather than diminishing gradually. This demonstrates that a production relationship can exhibit strongly diminishing (indeed vanishing) marginal returns to each factor individually, while the isoquant itself does not display the smooth, continuously diminishing MRTS that "convexity" in the standard interior sense requires. More generally, whenever the two factors are strongly complementary (a large positive cross-effect between them), the ease of substitution along the isoquant can behave quite differently from what the own diminishing-returns behaviour alone would suggest.

Conclusion

Diminishing marginal productivity is therefore a necessary building block but not, by itself, a sufficient condition for isoquant convexity — economists must separately assume (or verify) that the production function is quasi-concave, a condition that additionally restricts how the marginal products of the two factors interact with each other, not merely how each behaves in isolation.

(b-ii) At first glance, it may seem paradoxical that a production function could simultaneously display diminishing returns to each individual factor and constant returns to scale (CRS) overall. In fact, this combination is not only possible but is the standard, empirically common case in economics — the two concepts operate along entirely different dimensions and are fully compatible.

Resolving the Apparent Paradox

- Diminishing marginal productivity describes what happens when one factor is increased while the other is held fixed.
- Constant returns to scale describes what happens when all factors are increased together, in the same proportion — e.g., doubling both labour and capital simultaneously exactly doubles output.

Since these describe different experiments (varying one input alone, versus scaling all inputs together), there is no logical contradiction in a production function exhibiting both properties at once.

Illustrative Example: The Cobb-Douglas Production Function

Consider the widely used Cobb-Douglas production function: $\text{Output} = A \times (\text{Labour})^{0.5} \times (\text{Capital})^{0.5}$, where the exponents sum to 1 ($0.5 + 0.5 = 1$).

Checking CRS: If both labour and capital are scaled up by a factor t (e.g., doubled, $t = 2$):
 new output $= A \times (tL)^{0.5} \times (tK)^{0.5} = A \times t^{(0.5+0.5)} \times L^{0.5} \times K^{0.5} = t \times (\text{original output})$. Since output scales exactly proportionately with the scaling factor t , this function exhibits constant returns to scale.

Checking diminishing marginal product of labour (holding capital fixed): Suppose capital is fixed at 100 units. The marginal product of labour at successive levels of labour is:

- At Labour = 1: Marginal Product ≈ 5.0
- At Labour = 4: Marginal Product ≈ 2.5
- At Labour = 9: Marginal Product ≈ 1.67

The marginal product of labour clearly falls as labour increases, holding capital constant — confirming diminishing marginal productivity of labour. By the symmetric structure of the function, diminishing marginal productivity of capital holds equally when labour is held fixed.

Why This Makes Economic Sense:

This is intuitive: if a factory doubles both its workers and its machines, output doubles too (CRS) — nothing has changed in the proportion of capital available per worker, so there is no reason for output to more or less than double. But if the factory adds only more workers to the same fixed stock of machines, each additional worker has progressively less capital to work with, so each successive worker adds less to output than the last (diminishing marginal product of labour) — a classic case of workers "crowding" a fixed set of machines. This is also consistent with Euler's theorem for linearly homogeneous (degree-1, i.e., CRS) production functions, which guarantees that total output is exactly exhausted by paying each factor its marginal product times its quantity used ($\text{Output} = \text{Marginal Product of Labour} \times \text{Labour} + \text{Marginal Product of Capital} \times \text{Capital}$) — a result that holds regardless of whether individual factors show diminishing marginal returns.

Conclusion

Yes, constant returns to scale and diminishing marginal productivity of each individual factor can — and typically do — coexist, as demonstrated by the standard Cobb-Douglas production function whenever its exponents sum to one. This combination is, in fact, the empirically standard assumption used in most macroeconomic growth models (e.g., the Solow-Swan growth model), where CRS in capital and labour jointly ensures well-defined per-worker production functions, while diminishing returns to capital alone (holding labour/effective labour fixed) is precisely what drives the model's convergence predictions.

Q3. [Total: 20 marks]

(a) A monopolist faces the market demand curve $P = 100 - 4Q$ and total cost function $C = 50 + 20Q$. (i) Find the profit-maximising price, output and level of profit. (ii) Suppose the government now imposes a specific (per-unit) tax of Rs 10 on every unit sold. Find the new profit-maximising price, output and profit. [10]

(b) Compare the welfare and resource-allocation implications of a specific (per-unit) tax with those of a lump-sum tax of equal revenue yield, both imposed on the monopolist in part (a). Why is a lump-sum tax able to raise the same revenue while allocating resources more efficiently than a specific tax? [10]

Answer:

(a) A profit-maximising monopolist sets output where Marginal Revenue (MR) equals Marginal Cost (MC), and then charges the price the market demand curve dictates for that output level. This differs fundamentally from a competitive firm, which is a price-taker and produces where Price = MC. We solve the monopolist's problem first without a tax, then examine how a specific (per-unit) tax alters the outcome by effectively raising the monopolist's marginal cost.

(i) Profit-Maximising Price, Output, and Profit — No Tax

Given: Demand: $P = 100 - 4Q$; Total Cost: $C = 50 + 20Q$

Step 1: Derive Total Revenue and Marginal Revenue

Total Revenue (TR) = $P \times Q = (100 - 4Q) \times Q = 100Q - 4Q^2$

Marginal Revenue (MR) = derivative of TR with respect to $Q = 100 - 8Q$

Step 2: Derive Marginal Cost

Marginal Cost (MC) = derivative of Total Cost with respect to $Q = 20$ (constant, since cost is linear in Q)

Step 3: Set $MR = MC$ and Solve for Q

$$100 - 8Q = 20$$

$$8Q = 80$$

$$Q = 10 \text{ units}$$

Step 4: Find Price from the Demand Curve

$$P = 100 - 4(10) = 100 - 40$$

$$P = ₹60$$

Step 5: Compute Profit

$$\text{Total Revenue} = P \times Q = 60 \times 10 = ₹600$$

$$\text{Total Cost} = 50 + 20(10) = 50 + 200 = ₹250$$

$$\text{Profit} = 600 - 250 = ₹350$$

Result: Without any tax, the monopolist produces 10 units, charges ₹60 per unit, and earns a profit of ₹350.

(ii) Profit-Maximising Price, Output, and Profit — With a Specific Tax of ₹10 per Unit

A specific (per-unit) tax of ₹10 on every unit sold adds ₹10 to the cost of producing each unit — it is economically equivalent to a ₹10 increase in marginal cost.

Step 1: New Total Cost and Marginal Cost

$$\text{New Total Cost} = 50 + 20Q + 10Q = 50 + 30Q$$

$$\text{New Marginal Cost} = 30$$

Step 2: Set MR = New MC and Solve for Q

$$100 - 8Q = 30$$

$$8Q = 70$$

$$Q = 8.75 \text{ units}$$

Step 3: Find the New Price

$$P = 100 - 4(8.75) = 100 - 35$$

$$P = ₹65$$

Step 4: Compute the New Profit

$$\text{Total Revenue} = 65 \times 8.75 = ₹568.75$$

$$\text{Total Cost (including tax)} = 50 + 30(8.75) = 50 + 262.5 = ₹312.5$$

$$\text{Profit} = 568.75 - 312.5 = ₹256.25$$

$$\text{Government's Tax Revenue} = ₹10 \times 8.75 = ₹87.50$$

Summary Table

	No Tax	With ₹10 Specific Tax
Output (Q)	10 units	8.75 units
Price (P)	₹60	₹65
Monopolist's Profit	₹350	₹256.25
Government Revenue	—	₹87.50

Conclusion

The specific tax raises the monopolist's effective marginal cost from ₹20 to ₹30, shifting the $MR = MC$ intersection leftward. Consequently, output falls (from 10 to 8.75 units), price rises (from ₹60 to ₹65) — but by less than the full ₹10 tax, since the monopolist absorbs part of the tax burden through the fall in profit — and profit falls sharply (from ₹350 to ₹256.25), reflecting both the direct tax payment and the loss from reduced output.

(b) Having established that a specific tax of ₹10 per unit raises government revenue of ₹87.50 while reducing monopoly output to 8.75 units, we now compare this outcome to a lump-sum tax designed to raise the identical ₹87.50 in revenue, and examine why the two taxes — despite collecting the same amount — have very different implications for resource allocation and social welfare.

Effect of a Lump-Sum Tax on the Monopolist's Decision

A lump-sum tax is a fixed amount (here, ₹87.50) that the monopolist must pay regardless of how much it produces — it does not depend on output at all. Critically, this means a lump-sum tax behaves exactly like an increase in fixed cost, not marginal cost.

Step 1: New Total Cost Under the Lump-Sum Tax

$$\text{New Total Cost} = 50 + 87.50 + 20Q = 137.50 + 20Q$$

Step 2: Marginal Cost Is Unchanged

Marginal Cost = derivative of Total Cost with respect to $Q = 20$ (unchanged, since the lump-sum tax does not vary with Q)

Step 3: Set $MR = MC$ — Same Condition as Before

$$100 - 8Q = 20$$

$Q = 10$ units (unchanged from the no-tax case)

$P = ₹60$ (unchanged from the no-tax case)

Step 4: Compute Profit

$$\text{Total Revenue} = 60 \times 10 = ₹600$$

$$\text{Total Cost} = 137.50 + 20(10) = 137.50 + 200 = ₹337.50$$

$$\text{Profit} = 600 - 337.50 = ₹262.50$$

Comparison Table — Equal Government Revenue of ₹87.50

	Specific Tax (₹10/unit)	Lump-Sum Tax (₹87.50)
Output (Q)	8.75 units	10 units (unchanged)
Price (P)	₹65	₹60 (unchanged)
Monopolist's Profit	₹256.25	₹262.50
Government Revenue	₹87.50	₹87.50

Welfare (Surplus) Comparison

Using the underlying demand and cost structure, we can compute Consumer Surplus (CS) and total surplus (CS + Producer Surplus + Government Revenue, treating tax as a transfer rather than a resource loss) in each scenario, excluding the sunk fixed cost of ₹50 for comparability:

- No tax: CS = ₹200; Producer surplus (excl. fixed cost) = ₹400; Total surplus = ₹600
- Specific tax: CS = ₹153.13; Producer surplus (excl. tax and fixed cost) = ₹306.25 (before the ₹87.50 tax deduction, i.e., ₹393.75 gross); Total surplus = ₹546.88
- Lump-sum tax: CS = ₹200 (unchanged); Producer surplus (excl. tax and fixed cost) = ₹400 (unchanged); Total surplus = ₹600 (unchanged)

Key Result: The specific tax causes total surplus to fall from ₹600 to ₹546.88 — an additional deadweight loss of about ₹53.13, over and above the deadweight loss the monopoly already creates relative to a competitive outcome. The lump-sum tax, by contrast, creates no additional deadweight loss at all — total surplus remains exactly ₹600; the ₹87.50 is a pure transfer from the monopolist's profit to the government, with consumers entirely unaffected (price and output stay exactly where they were before the tax).

Why the Lump-Sum Tax Is More Efficient — Explanation

1. A specific tax distorts the marginal decision: By raising marginal cost, a specific tax changes the incentive at the margin — the monopolist's calculation of "should I produce one more unit?" is altered, so it rationally produces less. This shrinks the market below its already-restricted monopoly output, compounding the existing monopoly distortion (monopoly output is already below the competitive/efficient level of 20 units, where price would equal marginal cost of ₹20) with a fresh distortion from the tax.

2. A lump-sum tax does not touch the marginal decision: Since it is a fixed payment independent of output, it does not change MR or MC at all — the profit-maximising condition ($MR = MC$) is exactly the same with or without the lump-sum tax. The monopolist's optimal output and price are therefore completely unaffected; only its residual profit shrinks.
3. Economic principle — taxing the margin vs taxing the "prize": A specific tax effectively taxes each unit of activity (production/sale), discouraging that activity at the margin. A lump-sum tax instead taxes the fixed prize of being a monopolist (its economic profit) without penalising any additional unit produced — it is analogous to a tax on pure economic rent, which by definition does not distort behaviour since the rent-earner's optimal quantity decision is unaffected by how much of the rent is subsequently taxed away.
4. Practical caveat: Despite its efficiency advantage, lump-sum taxation is rarely used in practice for real-world monopoly regulation, because it requires the tax authority to know the firm's profit function accurately, and because a sufficiently large lump-sum tax could in principle drive the firm out of business altogether (if it exceeds even short-run viable profit) — whereas a specific/per-unit tax adjusts automatically with the scale of production. Additionally, lump-sum taxes on profit can raise concerns of practical administrability and are more easily evaded through profit-shifting, which specific/ad-valorem sales taxes are comparatively harder to avoid.

Conclusion

For the same government revenue of ₹87.50, the specific tax reduces output, raises price, and destroys an additional ₹53.13 of social surplus through the deadweight loss triangle created by distorting the monopolist's marginal production decision. The lump-sum tax, by contrast, raises the identical revenue purely by transferring existing monopoly profit to the government, without altering output, price, or the allocation of resources at all — making it the more allocatively efficient instrument. This illustrates a foundational principle in public economics: taxes on economic rent/pure profit (lump-sum in nature) are non-distortionary, while taxes on marginal economic activity (specific/unit taxes) create deadweight loss — a principle directly relevant to real-world tax design debates, such as the case for taxing supernormal/monopoly profits directly (e.g., windfall taxes) rather than through per-unit levies that further restrict output.

Q4. [Total: 20 marks]

(a) Distinguish between the Cournot and Bertrand models of duopoly. What mechanism allows firms to earn positive economic profit in Cournot competition, while price (Bertrand) competition drives economic profit down to zero? Explain the underlying economic intuition. [10]

(b) Define Nash equilibrium and dominant strategy, illustrating each with a suitable payoff matrix. Briefly distinguish a Nash equilibrium from a subgame-perfect equilibrium. [10]

Answer:

(a) Both the Cournot and Bertrand models describe markets with a small number of firms (duopoly, in the simplest case) that recognise their strategic interdependence — each firm's profit depends not only on its own choices but on its rival's. The two models differ fundamentally in the strategic variable firms are assumed to choose, and this single difference produces starkly different equilibrium outcomes: Cournot competition typically sustains positive economic profit above the competitive level, while Bertrand competition (even with just two firms) can drive economic profit all the way down to zero, matching the perfectly competitive outcome. Understanding why requires examining the different mechanisms of rivalry each model assumes.

The Cournot Model

In the Cournot model, firms simultaneously and independently choose the quantity they will produce, and the market price is then determined by the demand curve given total industry output. Each firm treats its rival's output as fixed when deciding its own best response, and the equilibrium is the point where both firms' quantity choices are mutual best responses (a Nash equilibrium in quantities).

The Bertrand Model

In the Bertrand model, firms simultaneously and independently choose the price they will charge for an identical (homogeneous) product, and consumers buy entirely from whichever firm charges the lower price (splitting demand equally if prices are tied). Each firm treats its rival's price as fixed when deciding its own best response.

Why Cournot Competition Sustains Positive Economic Profit

- 1. Quantity decisions are strategic substitutes with "soft" competition: When a firm increases its output in the Cournot model, this puts some downward pressure on price (since total market supply rises), but the firm still retains a captured share of demand that its rival cannot instantly seize — output cannot be costlessly and infinitely undercut the way price can. Each firm effectively controls only its own segment of quantity, softening the intensity of rivalry.**

2. No discontinuous "all-or-nothing" competition: In Cournot, a firm producing slightly more than its rival does not capture the entire market — it merely marginally increases its own market share and depresses price for both firms modestly. This gradual, continuous trade-off allows both firms to retain positive markup over marginal cost in equilibrium.
3. Illustrative numbers: Using the earlier demand curve $P = 100 - 4Q$ (industry) with each firm having marginal cost of 20, standard Cournot analysis shows each firm produces a positive quantity and earns positive profit — industry output under Cournot duopoly typically lies between the monopoly output (10 units, from Q_3) and the competitive output (20 units), and price similarly lies between the monopoly price (₹60) and the competitive price (₹20) — reflecting more intense rivalry than pure monopoly, but still short of full competitive dissipation of profit.

Why Bertrand Competition Drives Profit to Zero — The "Bertrand Paradox"

1. Price is a discontinuous, all-or-nothing weapon: With homogeneous products, if Firm 1 charges even one paisa less than Firm 2, it captures the entire market demand instantly — there is no partial gain from a small price cut; the reward is total. This creates an overwhelming incentive to always undercut a rival's price as long as price still exceeds marginal cost.
2. The undercutting logic drives price to marginal cost: Starting from any price above marginal cost, each firm has an incentive to slightly undercut its rival to capture the whole market, and this process of mutual undercutting continues until price equals marginal cost — at which point neither firm has any further incentive to cut price (since selling below marginal cost would mean losses), nor any incentive to raise price (since doing so would lose all customers to the rival).
3. Result — the competitive outcome, with just two firms: In equilibrium, both firms charge Price = Marginal Cost (₹20 in our earlier numerical example), output rises to the competitive level (20 units, from $P = 100 - 4Q = 20$), and economic profit falls to exactly zero — an outcome identical to what would emerge under perfect competition with many firms, despite there being only two competitors. This surprising result is known as the Bertrand Paradox, since it seems counterintuitive that just two firms can fully replicate the competitive outcome.

Underlying Economic Intuition — Summary

The key distinction lies in how aggressively a marginal move by one firm affects its rival: in Cournot, expanding output only gradually erodes the rival's position (a "soft" strategic interaction, since quantities are imperfect substitutes for capturing the whole market), whereas in Bertrand, cutting price is a winner-takes-all move (a "hard" strategic interaction) that makes any price above marginal cost unsustainable. This is why real-world markets with capacity constraints, product differentiation, or quantity-setting behaviour (e.g., where firms build production capacity in advance and cannot instantly

flood the market) tend to exhibit Cournot-like outcomes with sustained margins, while markets with easily adjustable prices and homogeneous, easily comparable products (e.g., certain undifferentiated online retail goods) can exhibit intense, near-zero-margin Bertrand-like price competition.

Conclusion

The Cournot and Bertrand models illustrate how the nature of the strategic variable firms compete over — quantity versus price — fundamentally shapes market outcomes, even holding the number of competing firms constant. This has important real-world regulatory relevance: competition authorities assessing likely post-merger pricing behaviour, or analysing an industry's competitive intensity, must consider not just the number of firms but how those firms actually compete (capacity/quantity commitments versus flexible pricing) to correctly predict whether meaningful market power and above-competitive profit are likely to persist.

(b) Game theory provides the formal tools to analyse strategic interaction where each player's optimal choice depends on what others are expected to do. Two of the most fundamental solution concepts are the dominant strategy and the Nash equilibrium, and a further refinement — the subgame-perfect equilibrium — becomes essential once we move from simultaneous, one-shot games to sequential (multi-stage) games.

Dominant Strategy — Definition and Example

A dominant strategy is a strategy that yields a player the highest payoff regardless of what strategy the other player chooses — it is optimal against every possible action by the rival, not just a specific one.

Example: The Classic Prisoner's Dilemma

Consider two arrested suspects (Firm A and Firm B, in a cartel-pricing context) who must simultaneously decide whether to "Cooperate" (maintain high cartel prices) or "Defect" (secretly undercut and cheat on the cartel agreement). Payoffs (profits, in ₹ crore) are shown as (Firm A's payoff, Firm B's payoff):

	Firm B: Cooperate	Firm B: Defect
Firm A: Cooperate	(10, 10)	(2, 15)
Firm A: Defect	(15, 2)	(5, 5)

Analysis: Consider Firm A's choice. If Firm B cooperates, Firm A earns 15 by defecting versus 10 by cooperating — defecting is better. If Firm B defects, Firm A earns 5 by defecting versus 2 by cooperating — defecting is again better. Since "Defect" gives Firm A a higher payoff no matter what Firm B does, Defect is Firm A's dominant strategy. By

symmetry, Defect is also Firm B's dominant strategy. Hence both firms defect in equilibrium, each earning only 5 — even though both would have earned more (10 each) had they cooperated. This captures the classic tension in cartel behaviour: even when collusion would be jointly profitable, each firm has an individual incentive to cheat, making cartels inherently unstable without external enforcement.

Nash Equilibrium — Definition and Example

A Nash equilibrium is a combination of strategies, one for each player, such that no player can improve their own payoff by unilaterally changing their own strategy, given the strategy the other player has chosen. Unlike a dominant strategy, a Nash equilibrium strategy need only be optimal given the rival's specific equilibrium choice — not against every conceivable choice the rival might make.

Example: Coordination Game — Choice of Technology Standard

Consider two banks (Bank A and Bank B) deciding which digital payment standard to adopt — UPI-compatible or a proprietary system — where compatibility with each other yields network benefits:

	Bank B: UPI-standard	Bank B: Proprietary
Bank A: UPI-standard	(8, 8)	(2, 3)
Bank A: Proprietary	(3, 2)	(5, 5)

Analysis: Here, neither bank has a dominant strategy — Bank A's best choice depends on what Bank B does (if B chooses UPI, A's best response is UPI, earning 8 vs 3; if B chooses Proprietary, A's best response is Proprietary, earning 5 vs 2). Checking each cell: (UPI, UPI) is a Nash equilibrium, since neither bank gains by unilaterally switching away (Bank A switching to Proprietary alone would fall from 8 to 2; Bank B similarly falls from 8 to 2). Similarly, (Proprietary, Proprietary) is also a Nash equilibrium (5,5), since unilateral deviation again reduces payoff for whichever bank deviates. This game has two Nash equilibria — illustrating that, unlike a dominant strategy equilibrium (which if it exists is unique and does not depend on coordination), Nash equilibria can be multiple, and here reflect a genuine coordination problem between the two competing standards.

Relationship Between the Two Concepts: Every dominant-strategy equilibrium is automatically a Nash equilibrium (since a strategy optimal against every rival action is certainly optimal against the rival's actual equilibrium action) — as seen in the Prisoner's Dilemma, where (Defect, Defect) is both the dominant-strategy outcome and a Nash equilibrium. However, the reverse is not true: a Nash equilibrium need not involve dominant strategies at all, as the coordination-game example shows.

Distinguishing Nash Equilibrium from Subgame-Perfect Equilibrium

The Nash equilibrium concept, as defined above, works well for simultaneous-move games, but can produce implausible predictions in sequential (multi-stage) games, because it does not require that a player's strategy be optimal at every decision point (node) of the game — including points that are never actually reached along the equilibrium path.

- A Subgame-Perfect Nash Equilibrium (SPNE) is a refinement that requires the players' strategies to constitute a Nash equilibrium in every subgame of the larger game — that is, the strategy must be credible and optimal not just along the anticipated equilibrium path, but at every possible point the game could reach, including off-equilibrium branches.
- Why this matters — example: Consider an incumbent firm that threatens to fight a price war (an unprofitable strategy for itself) if a new entrant enters the market, hoping this threat deters entry. If, upon actually reaching the "entry has occurred" node, fighting is genuinely not in the incumbent's own interest (because fighting is more costly than simply accommodating the entrant), then this threat is not credible — a rational entrant, anticipating this, would call the bluff and enter anyway. A pure Nash equilibrium analysis of the overall game might still support the (threat, don't-enter) outcome as an equilibrium (since neither player benefits from unilaterally deviating given the other's strategy), but this equilibrium relies on an incredible, non-credible threat that the incumbent would not actually carry out if truly called upon to do so. The subgame-perfect equilibrium concept eliminates such implausible equilibria by requiring optimality at every subgame, typically solved using backward induction — working from the last decision node of the game backward to the first — which correctly identifies that the incumbent would actually accommodate entry, so the credible equilibrium involves entry occurring.

Conclusion

Dominant strategy and Nash equilibrium are foundational solution concepts for analysing strategic interaction — a dominant strategy equilibrium is the stronger, more restrictive concept (optimal against any rival action), while Nash equilibrium is more general and can admit multiple equilibria, including coordination problems. When games unfold over multiple stages, however, plain Nash equilibrium can sustain outcomes built on non-credible threats or promises; the subgame-perfect equilibrium refinement, solved via backward induction, restores economic plausibility by requiring that every player's strategy be genuinely optimal at every point in the game, not merely along the path that is expected to be played — a distinction with direct relevance to real-world strategic settings such as regulatory commitment problems, entry-deterrence in antitrust analysis, and central bank credibility in monetary policy.

SECTION B — MACRO-ECONOMICS

Q5. [Total: 20 marks]

(a) Explain Kaldor's theory of income distribution/growth. How does the functional distribution of income between wages and profits act as the equilibrating mechanism that ensures ex-ante savings equal ex-ante investment ($S = I$) in the economy? [10]

(b) Using Marx's notation, define constant capital (c), variable capital (v) and surplus value (s). Explain how the capitalist class's continuous drive to raise the rate of surplus value (the rate of exploitation) leads, in Marx's analysis, to a secular tendency for the rate of profit to fall. [10]

Answer:

(a) Nicholas Kaldor developed his theory of income distribution in the 1950s as part of the post-Keynesian tradition, seeking to answer a question left open by Keynes's General Theory: in a growing economy, what mechanism ensures that ex-ante (planned) savings equal ex-ante (planned) investment, given that savings and investment decisions are typically made by different economic agents for different reasons? Kaldor's distinctive answer was that the functional distribution of income between wages and profits — rather than adjustments in interest rates (as in the loanable funds tradition) or output/employment adjustments (as in the simple Keynesian multiplier model) — is the equilibrating variable that reconciles savings with investment, particularly in a full-employment growth context.

The Building Blocks of Kaldor's Model

Kaldor's model rests on two key behavioural assumptions about the propensity to save out of different types of income:

- 1. Differential savings propensities by income class/type:** Kaldor assumed that the propensity to save out of profit income (S_p) is significantly higher than the propensity to save out of wage income (S_w) — formally, $S_p > S_w$. This reflects the classical/Keynesian observation that capitalists (recipients of profit) save a much larger fraction of their income than workers (recipients of wages), who tend to consume most of their earnings.
- 2. Investment as exogenously/independently determined:** In line with the Keynesian tradition, Kaldor treated the investment decision (I) as determined independently by entrepreneurs' expectations, animal spirits, and technological opportunities — not directly by the savings decision. This creates the central problem: how does the economy ensure that whatever level of investment is autonomously decided upon is exactly matched by an equal volume of savings?

The Equilibrating Mechanism — How Income Distribution Adjusts S to Equal I

National income (Y) is split between wages (W) and profits (P): $Y = W + P$. Total savings in the economy is the sum of savings out of each type of income:

$$\text{Total Savings (S)} = S_w \times W + S_p \times P$$

Since $S_p > S_w$, the aggregate/average savings ratio of the economy depends on how income is distributed between wages and profits — an economy with a higher profit share will, other things equal, generate a higher aggregate savings ratio, and vice versa.

The Adjustment Process:

1. Suppose entrepreneurs autonomously decide to raise investment (I). At the initial income distribution, planned savings would fall short of this higher planned investment (excess demand for investible resources).
2. This excess demand for goods (relative to the initial output/distribution) causes prices to rise relative to money wages (since money wages are typically more sticky/slow to adjust than prices in the short run) — this is the crucial Kaldorian mechanism.
3. As prices rise faster than wages, real wages fall and the share of profits in national income rises — income is effectively redistributed away from wage-earners (low savers) toward profit-earners (high savers).
4. Since profit-earners save a larger fraction of their (now larger) income share, aggregate savings automatically rises as the profit share increases — even though total output/income may not have changed much.
5. This process of redistribution continues until the new, higher aggregate savings ratio, given the new profit share, generates exactly enough savings to match the higher level of investment — restoring the equality $S = I$, but now at a higher profit share and a lower real wage share than before.

Formal Expression

At equilibrium, $S = I$ implies: $S_w \times W + S_p \times P = I$

Since $Y = W + P$, this can be rewritten in terms of the profit share (P/Y):

$$I/Y = S_w + (S_p - S_w) \times (P/Y)$$

Rearranging for the profit share:

$$P/Y = [(I/Y) - S_w] / (S_p - S_w)$$

This equation shows explicitly that a higher desired investment-to-income ratio (I/Y) requires a correspondingly higher equilibrium profit share (P/Y) — the mechanism by which the economy's income distribution adjusts to validate whatever level of investment entrepreneurs have autonomously chosen. Conversely, if entrepreneurs invest less, the required profit share falls, real wages rise relative to prices, and the wage share increases correspondingly.

Illustrative Example

Suppose $S_w = 0.10$ (workers save 10% of wage income) and $S_p = 0.40$ (capitalists save 40% of profit income). If the investment-to-income ratio (I/Y) is 0.20 (investment is 20% of national income), then applying the formula: Profit share = $(0.20 - 0.10) / (0.40 - 0.10) = 0.10/0.30 \approx 0.33$, i.e., profits would need to constitute roughly 33% of national income to generate enough aggregate savings to match this investment rate. If entrepreneurs became more optimistic and raised I/Y to 0.25, the required profit share would rise further, to $(0.25 - 0.10)/0.30 \approx 0.50$ — illustrating how more ambitious investment plans are validated, in Kaldor's framework, by a redistribution of income toward profits (and away from wages) via the price-wage adjustment mechanism.

Conditions and Limitations

- This mechanism requires the economy to be operating at, or reasonably close to, full employment, so that excess demand manifests primarily as price and distributional changes rather than being absorbed through output/employment expansion (as in the simple Keynesian model with unemployed resources).
- It also requires $S_p > S_w$ to hold meaningfully, and the profit share must remain within economically plausible bounds (between the pure wage-share limit and pure profit-share limit) for the mechanism to operate as described — Kaldor himself noted the model applies within a defined range of feasible profit shares.

Conclusion

Kaldor's theory offers a distinctively Post-Keynesian, distribution-based resolution to the savings-investment equality problem, standing apart from both the classical loanable-funds mechanism (interest-rate adjustment) and the simple Keynesian multiplier mechanism (output/employment adjustment). By making the functional distribution of income between wages and profits the equilibrating variable — operating through the relative stickiness of money wages compared to prices — Kaldor's model links macroeconomic growth theory directly to questions of income distribution and inequality, a connection of considerable relevance to policy debates on how growth-oriented investment strategies interact with wage share and inequality trends in developing economies.

(b) Karl Marx's analysis of capitalism, developed in Capital (particularly Volume I and Volume III), rests on the labour theory of value and a distinctive set of categories describing how capital is deployed in production and how surplus is extracted from labour. From these building blocks, Marx derived one of his most famous — and most debated — propositions: the tendency for the rate of profit to fall as capitalism develops, driven paradoxically by capitalists' own rational, profit-seeking behaviour.

Marx's Basic Categories

1. **Constant Capital (c):** The value of capital invested in means of production — raw materials, machinery, plant, and equipment. Marx termed this "constant" because, in his framework, these inputs merely transfer their own existing value to the final product during production; they do not themselves create new value in the production process.
2. **Variable Capital (v):** The value of capital invested in labour power — i.e., the wages paid to workers. Marx termed this "variable" because, unlike constant capital, labour power is the unique source of new value creation: workers, once employed, can be made to work for longer than the value embodied in their own wages (their subsistence requirements), and it is this capacity that generates the surplus.
3. **Surplus Value (s):** The new value created by labour over and above the value of variable capital (wages) paid to it — i.e., the value workers produce during the portion of the working day beyond what is needed to reproduce their own wage (the "necessary labour time"). This surplus value, appropriated by the capitalist rather than returned to the worker, is Marx's technical term for the source of profit — reflecting his central claim that profit originates from the exploitation of labour (workers being paid the value of their labour power, but made to work longer than that value requires).

Key Ratios

- **Rate of Surplus Value (Rate of Exploitation) = s / v** — measures the ratio of surplus value extracted to the wages paid, i.e., how intensively labour is being exploited.
- **Organic Composition of Capital = c / v** — measures the ratio of capital invested in machinery/materials relative to capital invested in labour, reflecting the "technical" or capital intensity of production.
- **Rate of Profit = $s / (c + v)$** — the surplus value expressed as a return on the total capital advanced (both constant and variable), which is the relevant profitability measure from the capitalist's perspective.

Deriving the Tendency of the Rate of Profit to Fall

The rate of profit can be rewritten algebraically by dividing both numerator and denominator by v :

$$\text{Rate of Profit} = s/(c+v) = (s/v) / [(c/v) + 1]$$

This expression shows the rate of profit depends on two competing forces: the rate of surplus value (s/v) in the numerator, and the organic composition of capital (c/v) in the denominator.

The Mechanism Marx Describes

1. **Competitive pressure drives capitalists to mechanise:** In Marx's analysis, individual capitalists, competing against one another, are continuously driven to adopt more machinery and labour-saving technology (raising c relative to v) in order to increase labour productivity, cut unit costs, and gain a temporary competitive edge over rivals — this is the process driving up the organic composition of capital (c/v) over time as capitalism develops.
2. **Capitalists also try to raise the rate of exploitation:** Simultaneously, capitalists attempt to increase s/v — by lengthening the working day, intensifying labour, or suppressing wages — to extract more surplus value per unit of variable capital invested.
3. **Why the tendency toward falling profit rate dominates:** Marx argued that, over the long run, the rise in the organic composition of capital (c/v) tends to outpace the rise in the rate of surplus value (s/v), because there are practical/social limits to how far the rate of exploitation can be pushed (a limit on how much the working day can be lengthened, or wages compressed below subsistence, without provoking resistance or undermining the labour force's reproduction), whereas capital-deepening through mechanisation faces no such inherent limit and is continuously reinforced by inter-capitalist competition. As c/v rises persistently while s/v rises more slowly (or plateaus), the denominator $[(c/v) + 1]$ in the profit-rate formula grows faster than the numerator (s/v), causing the overall rate of profit to exhibit a secular tendency to decline — even as the absolute mass of surplus value (and profit) may continue to grow with the scale of production.

Illustrative Numerical Example

Suppose initially: $c = 100$, $v = 100$, $s = 100$, giving rate of surplus value $s/v = 1.0$ (100%), and rate of profit $= 100/(100+100) = 0.50$ (50%).

Now suppose competitive mechanisation raises the organic composition of capital sharply: $c = 400$, $v = 100$.

Suppose capitalists also succeed in raising the rate of exploitation to $s/v = 1.5$ (150%), meaning $s = 150$.

The new rate of profit = $150/(400+100) = 150/500 = 0.30$ (30%).

Despite capitalists successfully raising the rate of exploitation (from 100% to 150%), the much larger increase in the organic composition of capital (from 1.0 to 4.0) causes the overall rate of profit to fall from 50% to 30% — illustrating the paradox at the heart of Marx's argument: capitalists' own individually rational strategies (mechanisation to gain competitive advantage, and pushing up exploitation) collectively produce an outcome — a falling profit rate — that undermines the system's own long-run vitality.

Counteracting Tendencies Marx Acknowledged

Marx himself noted several "counteracting factors" that could offset or delay this tendency, including cheapening of constant capital elements (machinery becoming cheaper due to productivity gains, partially restraining the rise in c/v), increased intensity/duration of labour exploitation, the depression of wages below the value of labour power, foreign trade opening cheaper input sources, and the expansion into new markets — meaning the "law" operates as a tendency, not a mechanical certainty, and Marx explicitly framed it as "the law of the tendency of the rate of profit to fall," subject to real-world counter-pressures.

Conclusion

Marx's analysis links the microeconomic behaviour of individual capitalists (competitive mechanisation and the drive to raise exploitation) to a macro-level, systemic tendency — a declining rate of profit — that he saw as an inherent contradiction of the capitalist mode of production: the very competitive process by which capitalists seek to gain individual advantage collectively erodes the profitability of the system as a whole over time, a conclusion Marx viewed as central to understanding capitalism's long-run crises and periodic tendencies toward economic disruption, even as the theory remains a subject of extensive debate among economists regarding its empirical validity and internal consistency (e.g., the Okishio theorem's later challenge to the necessity of this tendency under viable technical change).

Q6. [Total: 20 marks]

(a) An economy is described by the following IS and LM equations: IS: $Y = 300 - 200r$
LM: $Y = 100 + 400r$, where Y is income and r is the rate of interest (in %). Express this system in matrix form and solve for the equilibrium values of income (Y) and the interest rate (r) using matrix inversion / Cramer's Rule. [8]

(Candidates recalled being asked to find Y and r 'through the Hessian matrix'. For a two-equation linear system such as this, the technique that actually applies is matrix inversion/Cramer's Rule — a Hessian, strictly, is the matrix of second-order partial derivatives used to check concavity/convexity in optimisation problems, not to solve a linear system. Practise the matrix-inversion method here; it is what the question requires.)

(b) In the same economy, $MPC = 0.8$ and there are no taxes or other leakages besides saving. If the government wishes to raise equilibrium income by Rs 500 crore, use the appropriate multiplier to calculate the required change in government expenditure. [7]

(c) Distinguish between the Cambridge cash-balance approach and Fisher's quantity theory of money as explanations of the demand for money. [5]

Answer:

(a) The IS-LM model describes simultaneous equilibrium in the goods market (IS curve) and the money market (LM curve). Given two linear equations in two unknowns (Y and r), the equilibrium values can be found by expressing the system in matrix form and solving via Cramer's Rule (or equivalently, matrix inversion). Note, as flagged in the question, that a Hessian matrix is a distinct concept — it comprises second-order partial derivatives used to check concavity/convexity in optimisation problems, and has no role in solving a linear simultaneous-equations system like this one; the correct technique here is matrix inversion/Cramer's Rule.

Step 1: Rearranging the Equations into Standard Linear Form

Given:

$$IS: Y = 300 - 200r$$

$$LM: Y = 100 + 400r$$

Rearranging both equations so that Y and r appear on the left-hand side, with constants on the right:

$$Y + 200r = 300 \dots (IS, \text{rearranged})$$

$$Y - 400r = 100 \dots (LM, \text{rearranged})$$

Step 2: Expressing in Matrix Form

This system can be written as $A \cdot X = B$, where:

$$A = \begin{bmatrix} 1 & 200 \\ 1 & -400 \end{bmatrix}$$

$$X = \begin{bmatrix} Y \\ r \end{bmatrix}$$

$$B = \begin{bmatrix} 300 \\ 100 \end{bmatrix}$$

$$[r]$$

$$B = \begin{bmatrix} 300 \\ 100 \end{bmatrix}$$

$$[100]$$

Step 3: Compute the Determinant of A

$$\text{Determinant } (D) = (1 \times -400) - (200 \times 1) = -400 - 200 = -600$$

Since $D \neq 0$, the system has a unique solution — confirming a well-defined equilibrium exists.

Step 4: Solve for Y Using Cramer's Rule

Replace the first column of A with B to form matrix A_Y:

$$A_Y = \begin{bmatrix} 300 & 200 \\ 100 & -400 \end{bmatrix}$$

$$\text{Determinant of } A_Y = (300 \times -400) - (200 \times 100) = -120,000 - 20,000 = -140,000$$

$$Y^* = \text{Determinant of } A_Y \div D = (-140,000) \div (-600) = 233.33$$

Step 5: Solve for r Using Cramer's Rule

Replace the second column of A with B to form matrix A_r:

$$A_r = \begin{bmatrix} 1 & 300 \\ 1 & 100 \end{bmatrix}$$

$$\text{Determinant of } A_r = (1 \times 100) - (300 \times 1) = 100 - 300 = -200$$

$$r^* = \text{Determinant of } A_r \div D = (-200) \div (-600) = 0.33$$

Step 6: Verification

$$\text{Substituting back into the IS equation: } Y = 300 - 200(0.33) = 300 - 66.67 = 233.33$$

$$\text{Substituting into the LM equation: } Y = 100 + 400(0.33) = 100 + 133.33 = 233.33$$

Both equations are satisfied, confirming the solution is correct.

Conclusion

The equilibrium level of income is $Y = ₹233.33$ crore*, and the equilibrium interest rate is $r = 0.33$ (i.e., approximately 1/3, in the percentage units defined by the model)**. This is the unique combination of income and the interest rate at which both the goods market (savings = investment, implicit in the IS curve) and the money market (money demand = money supply, implicit in the LM curve) clear simultaneously.

(b) The simple Keynesian expenditure multiplier measures how much equilibrium income changes for a given change in autonomous spending (such as government expenditure), when the only leakage from the circular flow is savings (i.e., no taxes, no imports). This part asks us to use that multiplier to determine how large a change in government expenditure (ΔG) is needed to raise equilibrium income by a specified amount.

Step 1: Compute the Multiplier

Since there are no taxes and no other leakages besides saving, the simple expenditure multiplier (k) is:

$$k = 1 \div (1 - MPC)$$

Given MPC = 0.8:

$$k = 1 \div (1 - 0.8) = 1 \div 0.2 = 5$$

This means every ₹1 crore of new autonomous spending (such as government expenditure) ultimately raises equilibrium income by ₹5 crore, once all successive rounds of induced consumption spending are accounted for.

Step 2: Compute the Required Change in Government Expenditure

We are told the desired change in equilibrium income, $\Delta Y = ₹500$ crore. Using the multiplier relationship:

$$\Delta Y = k \times \Delta G$$

$$500 = 5 \times \Delta G$$

$$\Delta G = 500 \div 5 = ₹100 \text{ crore}$$

Conclusion

The government would need to raise its expenditure by ₹100 crore to achieve the desired ₹500 crore increase in equilibrium income, given the multiplier of 5. It is worth noting, as a conceptual caveat, that this simple multiplier calculation implicitly holds the interest rate fixed (as in the basic Keynesian cross model). In the fuller IS-LM framework of part (a), however, an increase in government spending would shift the IS curve rightward, raising both income and the interest rate — and this induced rise in the interest rate would crowd out some private investment, meaning the actual increase in equilibrium income from a given ΔG in a full IS-LM setting would typically be somewhat smaller than the simple multiplier of 5 suggests. The ₹100 crore figure above is therefore the correct answer under the simple (fixed-interest-rate) multiplier framework as the question specifies, but would understate the required ΔG if crowding-out effects from the LM curve were to be factored in.

(c) Both the Cambridge cash-balance approach and Fisher's quantity theory of money belong to the classical quantity-theory tradition, and both ultimately arrive at broadly similar conclusions about the relationship between the money supply and the price level. However, they differ significantly in their underlying analytical approach — one emphasising the mechanical transactions role of money, and the other emphasising individuals' deliberate portfolio/asset-holding decision to hold money as a form of wealth.

Fisher's Quantity Theory (Transactions Approach)

Fisher's equation of exchange is: $M \times V = P \times T$

where M is the money supply, V is the transactions velocity of money (the number of times a unit of currency changes hands in transactions during the period), P is the general price level, and T is the total volume of transactions (or real output, Y , in a common simplified variant: $M \times V = P \times Y$).

- **Focus:** Fisher's approach is essentially mechanical and transactions-based — money is viewed purely as a medium of exchange, and V is treated as institutionally determined (by payment habits, banking practices, frequency of income payments) and roughly constant in the short run.
- **Implication:** With V and T (or Y) assumed constant, changes in M translate proportionately into changes in P — the classical neutrality of money result.

The Cambridge Cash-Balance Approach:

The Cambridge equation, associated with Marshall and Pigou, is: $M = k \times P \times Y$

where k represents the fraction of nominal income ($P \times Y$) that individuals choose to hold as money balances (cash and readily spendable deposits), rather than in other forms of wealth (bonds, physical assets).

- **Focus:** The Cambridge approach shifts attention to the demand for money as an asset — it asks why individuals want to hold money, emphasising money's role in providing convenience, liquidity, and a buffer against uncertainty, not merely its mechanical role in facilitating transactions. The parameter k (sometimes called the Cambridge k , or Marshallian k) is essentially the inverse of Fisher's velocity ($k \approx 1/V$), but is derived from individual behavioural choice rather than institutional/mechanical necessity.
- **Implication:** Since k reflects a genuine economic decision about how much wealth to hold in liquid form, the Cambridge approach opens the door to k (and hence money demand) responding to economic variables such as the interest rate (the opportunity cost of holding non-interest-bearing money) — a crucial conceptual bridge later built upon explicitly by Keynes's liquidity preference theory, which formalised the interest-sensitivity of money demand that the mechanical Fisherian approach could not easily accommodate.

Key Points of Distinction

Basis	Fisher's Quantity Theory	Cambridge Cash-Balance Approach
Core equation	$M \times V = P \times T$	$M = k \times P \times Y$
Perspective	Money as medium of exchange (transactions/flow view)	Money as a store of value/asset (stock, portfolio-choice view)

Basis	Fisher's Quantity Theory	Cambridge Cash-Balance Approach
Key behavioural parameter	Velocity (V) — largely institutional/mechanical	Cambridge k — reflects individual choice, can vary with interest rate/expectations
Analytical emphasis	Supply-side, mechanical relationship between money and transactions	Demand-side, microeconomic decision of how much wealth to hold as money
Link to later theory	Foundation for the classical dichotomy/money neutrality	Direct conceptual precursor to Keynes's liquidity preference theory

Conclusion

While both approaches share the classical conclusion that money supply changes are ultimately linked to the price level, the Cambridge cash-balance approach represents a conceptual advance over Fisher's more mechanical transactions-based formulation by explicitly framing money demand as a deliberate portfolio decision by individuals — a shift in emphasis from the velocity of circulation to the desired proportion of income held as money balances that laid essential groundwork for the more sophisticated, interest-rate-sensitive theory of money demand developed later by Keynes and subsequently formalised within the LM curve of the very IS-LM framework used in part (a) of this question.

Q7. [Total: 20 marks]

(a) From the following data (Rs crore) for an economy, calculate (i) Net Domestic Product at Factor Cost (NDP_{fc}), (ii) Gross Domestic Product at Market Price (GDP_{mp}), and (iii) National Income (NNP_{fc}), using the Income Method: Compensation of employees: 4,500; Operating surplus: 3,200; Mixed income of self-employed: 1,800; Net indirect taxes: 900; Consumption of fixed capital: 600; Net factor income from abroad: (-)50. [10]

(b) The actual output of an economy is \$140 million while its potential output is \$182 million. (i) Calculate the output gap as a percentage of potential output. (ii) Discuss the monetary and fiscal policy measures that could be undertaken to close this gap. [10]

Answer:

(a) The Income Method estimates national income by summing all factor incomes generated within the domestic economy — compensation of employees, operating surplus (rent, interest, profit), and mixed income of the self-employed. From this base figure, we can derive various related aggregates by adding or subtracting depreciation, net indirect taxes, and net factor income from abroad, depending on which specific aggregate (domestic vs. national, factor cost vs. market price, net vs. gross) is required.

Given Data (₹ crore)

- Compensation of employees = 4,500
- Operating surplus = 3,200
- Mixed income of self-employed = 1,800
- Net indirect taxes = 900
- Consumption of fixed capital (depreciation) = 600
- Net factor income from abroad = (-) 50

(i) Net Domestic Product at Factor Cost (NDPfc)

NDPfc is the sum of all factor incomes earned by residents and non-residents within the domestic territory (before adjusting for depreciation, indirect taxes, or foreign factor income):

NDPfc = Compensation of employees + Operating surplus + Mixed income of self-employed

$$\text{NDPfc} = 4,500 + 3,200 + 1,800$$

$$\text{NDPfc} = ₹9,500 \text{ crore}$$

(ii) Gross Domestic Product at Market Price (GDPmp)

To move from NDPfc to GDPmp, two adjustments are required:

- Add back Consumption of fixed capital (to convert Net → Gross)
- Add Net indirect taxes (to convert Factor Cost → Market Price, since market prices include indirect taxes net of subsidies)

Step 1 — Gross Domestic Product at Factor Cost (GDPfc):

$$\text{GDPfc} = \text{NDPfc} + \text{Consumption of fixed capital} = 9,500 + 600 = 10,100$$

Step 2 — GDPmp:

$$\text{GDPmp} = \text{GDPfc} + \text{Net indirect taxes} = 10,100 + 900$$

$$\text{GDPmp} = ₹11,000 \text{ crore}$$

(iii) National Income (NNP_{fc})

National Income, conventionally defined as Net National Product at Factor Cost (NNP_{fc}), adjusts the domestic aggregate (NDP_{fc}) to a national aggregate by adding Net Factor Income from Abroad (NFIA) — which accounts for income earned by residents from assets/employment abroad, minus income earned by non-residents from assets/employment within the domestic economy:

$$\text{NNP}_{fc} = \text{NDP}_{fc} + \text{Net factor income from abroad}$$

$$\text{NNP}_{fc} = 9,500 + (-50)$$

$$\text{NNP}_{fc} = ₹9,450 \text{ crore}$$

Summary Table

Aggregate	Formula	Value (₹ crore)
NDP _{fc}	Compensation + Operating surplus + Mixed income	9,500
GDP _{fc}	NDP _{fc} + Consumption of fixed capital	10,100
GDP _{mp}	GDP _{fc} + Net indirect taxes	11,000
NNP _{fc} (National Income)	NDP _{fc} + Net factor income from abroad	9,450

Conclusion

Starting from the domestic factor-income base of ₹9,500 crore, the two adjustments move in different directions depending on the aggregate sought: accounting for depreciation and indirect taxes raises the domestic aggregate to a market-price, gross basis (₹11,000 crore), while incorporating the (negative) net factor income from abroad — reflecting that this economy pays out more factor income to the rest of the world than it receives — lowers the domestic aggregate to arrive at National Income of ₹9,450 crore, which is less than NDP_{fc} precisely because NFIA is negative.

(b) The output gap measures the difference between an economy's actual output and its potential output (the maximum sustainable level of output the economy can produce using its resources — labour, capital, technology — without generating accelerating inflation, i.e., output at full employment of resources). A negative output gap (actual output below potential) indicates the economy is operating below capacity, with associated unemployment and unutilised productive resources, calling for expansionary policy; a positive output gap indicates overheating.

(i) Calculating the Output Gap

Given: Actual Output = \$140 million; Potential Output = \$182 million

Output Gap (as % of Potential Output) = $[(\text{Potential Output} - \text{Actual Output}) / \text{Potential Output}] \times 100$

$$\text{Output Gap} = [(182 - 140) / 182] \times 100$$

$$\text{Output Gap} = (42 / 182) \times 100$$

$$\text{Output Gap} \approx -23.08\%$$

The negative sign indicates a recessionary/deflationary gap — actual output falls short of potential output by approximately 23.08%, indicating substantial slack (unemployed labour and idle capacity) in the economy.

(ii) Policy Measures to Close the Output Gap

Since this is a large recessionary gap, both monetary and fiscal policy should adopt an expansionary stance to stimulate aggregate demand and bring actual output back toward potential.

Monetary Policy Measures

1. **Reduce the policy repo rate:** Lowering the central bank's benchmark interest rate reduces the cost of borrowing for both firms and households, encouraging higher investment (e.g., in plant, machinery, housing) and consumption (e.g., durable goods purchases on credit), thereby boosting aggregate demand.
2. **Reduce the Cash Reserve Ratio (CRR) / Statutory Liquidity Ratio (SLR):** Lowering these reserve requirements increases the loanable funds available with commercial banks, expanding credit supply to the economy.
3. **Open Market Operations (OMOs) — purchase of government securities:** The central bank buying government bonds injects liquidity directly into the banking system, lowering interest rates further and supporting credit expansion.
4. **Quantitative easing (in more severe/prolonged output gaps):** If conventional rate cuts are insufficient (e.g., when policy rates approach the zero lower bound), the central bank may undertake large-scale asset purchases to directly lower long-term interest rates and support financial conditions.
5. **Forward guidance:** Central bank communication committing to maintain accommodative policy for an extended period can lower expected future interest rates, encouraging current investment and consumption decisions based on that expectation.

Fiscal Policy Measures

1. **Increase government expenditure:** Direct increases in public spending — on infrastructure, capital projects, or public employment schemes — raise aggregate demand directly and, through the multiplier effect (as computed in Q6(b)), generate a magnified increase in equilibrium income, helping close the output gap.
2. **Tax cuts:** Reducing personal income tax or corporate tax rates raises disposable income for households and post-tax profits for firms, stimulating consumption and investment respectively — though the impact tends to be smaller than direct government spending increases, since some of the tax cut may be saved rather than spent (depending on the marginal propensity to consume).
3. **Targeted transfer payments/social spending:** Increasing transfers to lower-income households (who typically have a higher marginal propensity to consume) can be a particularly effective way to boost aggregate demand quickly.
4. **Automatic stabilisers:** Allowing built-in fiscal mechanisms (unemployment benefits, progressive taxation) to operate without offsetting fiscal tightening helps cushion the downturn without requiring discrete new policy action.
5. **Public investment in productive capacity:** Beyond simply boosting short-run demand, well-targeted public infrastructure investment can also expand the economy's potential output over time (narrowing the gap from both sides — raising actual output toward potential, and potentially also affecting the trajectory of potential output itself).

Coordination and Caveats

- Given the scale of the gap (over 23% of potential output — a very large shortfall by real-world standards), a combination of both monetary and fiscal expansion would likely be necessary rather than relying on either instrument alone, since monetary policy can face limits (e.g., near-zero interest rate floors, or a liquidity trap where further rate cuts fail to stimulate demand), and fiscal expansion alone can raise concerns about rising public debt-to-GDP ratios if not calibrated carefully.
- Policymakers must also monitor for potential inflationary pressure as the gap closes — since expansionary policy that overshoots and pushes actual output above potential would generate a positive output gap and demand-pull inflationary pressures, requiring the stimulus to be calibrated and gradually withdrawn as the economy approaches full employment.

Conclusion

With a substantial negative output gap of approximately 23.08% of potential output, this economy requires a coordinated expansionary policy response — monetary easing (rate cuts, liquidity infusion) working alongside fiscal stimulus (higher government spending and/or tax relief) — to raise aggregate demand and bring actual output back toward its potential level, while remaining alert to the need to moderate this stimulus as the gap narrows to avoid triggering demand-pull inflation once the economy approaches full capacity.

Q8. [Total: 20 marks]

(a) Using the AD-AS framework, distinguish between the effects of an anticipated versus an unanticipated increase in the money supply on output and the price level, in the New Classical (rational expectations) framework. [7]

(b) “Creative destruction driven by innovation is the fundamental source of long-run economic growth” (Philippe Aghion). Critically examine this statement with reference to the Schumpeterian/Aghion–Howitt model of growth. [7]

(c) Explain the effectiveness of monetary policy under conditions of perfect capital mobility (the Mundell–Fleming framework), contrasting outcomes under fixed versus flexible exchange rate regimes. [6]

Answer:

(a) The New Classical school, building on the rational expectations hypothesis (Lucas, Sargent, Barro), reached a striking conclusion that departs sharply from both Keynesian and even adaptive-expectations monetarist analysis: whether a change in the money supply affects real output depends critically on whether that change was anticipated (correctly foreseen by economic agents) or unanticipated (a genuine surprise). This is the essence of the Policy Ineffectiveness Proposition associated with Sargent and Wallace, and it rests on Lucas's "surprise" or "signal-extraction" supply function.

The Lucas Aggregate Supply Function

The New Classical aggregate supply curve can be represented as:

Actual Output = Potential (natural-rate) Output + a positive coefficient × (Actual Price Level – Expected Price Level)

This says output deviates from its natural/potential level only when the actual price level differs from what agents had expected — because individual suppliers, observing a rise in the price of their own product, cannot immediately tell whether this reflects a genuine relative-price change (justifying more production) or simply general inflation (in which case real output should not change). Agents with rational expectations, however, use all

available information — including their understanding of the central bank's systematic policy rule — to form their price expectations.

Case 1: Anticipated Increase in Money Supply

If an increase in the money supply is fully anticipated — say, the central bank has a known, credible policy rule, or has clearly pre-announced the expansion — rational agents incorporate this information into their price expectations before the change actually occurs.

- Both aggregate demand (AD) and the expected price level shift simultaneously.
- Because the expected price level rises in exact step with the actual price level, the term (Actual Price – Expected Price) in the Lucas supply function remains zero.
- Result: The AD curve shifts rightward, but the (vertical, in the fully rational-expectations, no-surprise case) Aggregate Supply curve does not generate any output response — the entire effect is absorbed by a rise in the price level, with output remaining at its natural/potential level. Money is fully neutral, even in the short run.

Case 2: Unanticipated Increase in Money Supply

If the money supply increase is a genuine surprise — unexpected by agents, perhaps due to a discretionary, non-systematic policy action — then expected prices do not adjust in advance.

- The AD curve shifts rightward, but expected prices remain unchanged (based on old information).
- Actual prices now exceed expected prices, so the term (Actual Price – Expected Price) is positive.
- Result: Firms and workers, observing higher prices/wages for their own output without immediately realising this reflects general inflation, respond by increasing output above the natural rate — both output and the price level rise in the short run, exactly as in a conventional upward-sloping short-run AS curve.
- However, this effect is temporary: once agents observe the true, higher price level and revise their expectations accordingly, the short-run AS curve shifts, and output returns to its potential level, with the price level settling at its new, fully-adjusted higher value.

Illustrative Example

Suppose the RBI has a well-known, credible policy of never surprising markets and always pre-announcing changes to the repo rate/money supply growth well in advance. If it announces a planned monetary expansion for the coming quarter, rational agents adjust wage and price-setting behaviour in anticipation — inflation rises roughly in line with the money supply increase, but real GDP growth and employment remain largely unaffected. By contrast, if the RBI unexpectedly and secretly expanded the money supply through an unannounced, discretionary action, firms observing a sudden rise in demand for their products (without immediately attributing it to general inflation) would temporarily expand output and hiring — a genuine short-run stimulus — until the surprise is "discovered" and expectations catch up.

Conclusion — The Policy Ineffectiveness Proposition

The central policy implication is that only unanticipated (surprise) monetary policy can affect real output, even in the short run — a systematic, rule-based, and correctly anticipated monetary policy is powerless to move output away from its natural rate, since rational agents will have already priced in its effects. This is a much stronger neutrality result than the traditional monetarist view (which allowed monetary policy real effects in the short run even under adaptive expectations), and it provided the theoretical foundation for later arguments favouring rule-based over discretionary monetary policy and for the importance of central bank credibility and transparency — since only surprises generate real effects, and central banks cannot systematically and repeatedly "fool" rational agents.

(b) Philippe Aghion and Peter Howitt (1992), building on Joseph Schumpeter's concept of "creative destruction," developed a formal endogenous growth model in which long-run economic growth is driven not by exogenous technological progress (as in the Solow model) or simple capital accumulation, but by a continuous, competitive process of innovation that destroys the profits of incumbent producers even as it creates new value for society.

The Core Mechanism of the Aghion–Howitt Model

1. **Quality-ladder innovation:** Firms undertake costly R&D in the hope of discovering a new, higher-quality product or production process that will render the previous "state-of-the-art" technology obsolete. Successful innovators temporarily earn monopoly profits (as the sole producer of the improved product/process), which is the incentive that motivates R&D investment in the first place.
2. **Creative destruction:** When a new innovation succeeds, it destroys the monopoly rents of the previous innovator, whose technology is now obsolete — hence "creative destruction": the creation of new value simultaneously destroys old value/incumbent positions. This is fundamentally different from the Solow model, where technological progress is a smooth, exogenous, non-competitive process.

3. **Growth as a sequence of temporary monopolies:** Aggregate long-run growth in this framework emerges as the expected outcome of a stochastic sequence of such innovations occurring across many sectors of the economy — each individual innovation is a discrete "step" up the quality ladder, but the aggregate economy-wide growth rate reflects the average frequency and size of these innovations.
4. **Endogenous determination of the growth rate:** Because the rate of innovation depends on economic variables — the size of the market/expected monopoly profit, the cost of R&D, the interest rate (discounting future profits), and the degree of competition — the long-run growth rate is genuinely endogenous, determined by economic incentives, market structure, and policy, rather than being an unexplained exogenous parameter as in the Solow model.

Strengths of the Framework — Supporting Aghion's Statement

1. **Explains the empirical persistence of firm turnover and "business dynamism":** The model matches real-world patterns of firms rising, becoming dominant, and then being displaced by new entrants with superior technology — e.g., how firms in digital photography were displaced by smartphone camera technology, or how legacy taxi services faced disruption from ride-hailing platforms — a pattern the smooth, non-competitive Solow-style technology treatment cannot capture.
2. **Provides a genuine microfoundation for R&D-driven growth:** Unlike the Solow model's residual ("Solow residual" / Total Factor Productivity as an unexplained catch-all), the Aghion–Howitt framework explicitly models the profit-motivated decision to invest in R&D, connecting growth theory to industrial organisation and market structure.
3. **Generates rich policy implications:** The model can analyse how competition policy, intellectual property protection (patent length/breadth), and R&D subsidies affect the long-run growth rate — issues of direct relevance to innovation and competition policy, which the Solow model is largely silent on.

Critical Qualifications and Limitations

1. The "business-stealing" externality can lead to excessive innovation incentives in some specifications, alongside underinvestment concerns in others: Aghion and Howitt themselves highlighted an ambiguous welfare relationship between market competition and the incentive to innovate — while competition can spur innovation (the "escape competition" effect, where firms innovate specifically to escape neck-and-neck rivalry with close competitors), it can also reduce post-innovation monopoly rents and thus dampen the incentive to invest in R&D in the first place (the "Schumpeterian effect," where monopoly rents are needed to fund/incentivise R&D) — the net relationship between competition and growth is

theoretically non-monotonic (often inverted-U shaped), a nuance the simple statement in the question does not capture.

2. **Not the sole source of growth:** Even within endogenous growth theory, creative destruction/quality-ladder innovation is one of several mechanisms proposed (alongside, e.g., Romer's variety-expansion model, and human-capital-based models like Lucas's) — treating it as the fundamental source of growth may understate the role of capital accumulation, human capital, institutions, and diffusion of existing technology (particularly relevant for developing economies like India, where catching up to the technological frontier through diffusion/adoption may matter more than frontier-pushing innovation itself).
3. **Distributional and transitional costs are understated:** "Destruction" of incumbent firms has real transitional costs — displaced workers, stranded capital, regional economic disruption — that the aggregate growth-rate focus of the model does not fully address, an important consideration for real-world innovation and labour-market policy (e.g., debates around automation and job displacement).
4. **Measurement challenges:** Empirically isolating "creative destruction"-driven growth from other sources of productivity growth remains difficult, and the model's strong predictions (e.g., about the competition-innovation relationship) have received mixed empirical support across different country/industry contexts.

Conclusion

Aghion's statement captures a genuinely important and empirically resonant insight — that innovation-driven obsolescence of incumbent technologies and firms is a central, observable feature of long-run growth in modern economies, and the Aghion–Howitt framework provides a rigorous microfounded model of this process, with valuable implications for competition and innovation policy. However, treating creative destruction as the single fundamental source of growth somewhat overstates the case; a more balanced view recognises it as one of several complementary engines of growth — alongside capital accumulation, human capital formation, and technology diffusion — whose relative importance varies across countries and stages of development, particularly for emerging economies still catching up to the global technology frontier.

(c) The Mundell–Fleming model extends the closed-economy IS-LM framework to an open economy with international capital flows, and one of its most influential results concerns the "impossible trinity" (or trilemma): under conditions of perfect capital mobility, a country cannot simultaneously maintain a fixed exchange rate, independent monetary policy, and free capital movement — it must give up (at most) one of the three. This directly determines whether monetary policy is effective under fixed versus flexible exchange rate regimes.

Setup: Perfect Capital Mobility

Under perfect capital mobility, capital flows respond instantaneously and completely to any interest rate differential relative to the world interest rate — domestic interest rates are effectively "pinned" to the world rate, since any deviation would trigger unlimited capital inflows or outflows until the differential is eliminated.

Monetary Policy Under a Fixed Exchange Rate Regime — Ineffective

1. Suppose the central bank attempts an expansionary monetary policy (increases money supply), which would normally lower the domestic interest rate below the world rate.
2. With perfect capital mobility, this interest rate gap triggers massive capital outflows, as investors move funds abroad seeking the higher world return.
3. This capital outflow creates strong downward (depreciation) pressure on the domestic currency.
4. However, under a fixed exchange rate commitment, the central bank is obligated to defend the peg — it must sell foreign exchange reserves and buy back domestic currency, which directly contracts the money supply — exactly reversing/offsetting the initial monetary expansion.
5. Result: The money supply is forced back to its original level to maintain the peg, the domestic interest rate returns to the world rate, and output/income remain unchanged. Monetary policy is thus completely ineffective under fixed exchange rates with perfect capital mobility — the central bank sacrifices monetary policy autonomy entirely to defend the fixed rate.

Monetary Policy Under a Flexible (Floating) Exchange Rate Regime — Highly Effective

1. Under a flexible exchange rate, the central bank has no obligation to intervene in the foreign exchange market — the exchange rate is allowed to adjust freely to clear the market.
2. An expansionary monetary policy again lowers the domestic interest rate below the world rate, triggering capital outflows and downward pressure on the currency — but this time, the currency is allowed to depreciate.
3. This depreciation makes domestic goods cheaper relative to foreign goods, boosting net exports (exports rise, imports fall) — this shifts the IS curve rightward, reinforcing and amplifying the expansionary effect of the initial LM shift.
4. Result: Under flexible exchange rates, monetary policy is highly effective — indeed, more effective than in a closed economy, since the exchange-rate depreciation channel adds an additional expansionary boost (via net exports) on top of the standard domestic interest-rate/investment channel.

Illustrative Contrast

If the RBI wished to stimulate a slowing economy through monetary easing, the Mundell-Fleming framework predicts that under a fixed exchange rate regime (such as India's system before 1991, when the rupee was managed/pegged), such easing would largely leak out through reserve losses and be self-defeating, forcing the central bank to reverse the policy to defend the currency. Under India's current managed float regime (closer to flexible, though not perfectly so, given the RBI's periodic FX intervention), monetary easing has more room to work through both the domestic interest-rate channel and the exchange-rate/net-export channel, though the RBI's own FX interventions to smooth volatility mean India does not experience the textbook full-flexibility outcome exactly.

Conclusion

The Mundell-Fleming framework shows that the effectiveness of monetary policy under perfect capital mobility is fundamentally shaped by the exchange rate regime: fixed exchange rates render monetary policy entirely ineffective (since defending the peg requires exactly offsetting any attempted change in money supply), while flexible exchange rates render monetary policy highly effective (since the exchange-rate channel reinforces the interest-rate channel) — a result that underpins the well-known "impossible trinity," and which has directly informed the choice of many countries, including India post-1991, to move toward greater exchange rate flexibility partly in order to retain a meaningful degree of independent monetary policy in an increasingly capital-mobile world.