

NAAC ACCREDITED



तेजस्वि नावधीतमस्तु
ISO 9001:2008 & 14001:2004

FAIRFIELD

Institute of Management & Technology

'A' Grade Institute by DHE, Govt. of NCT Delhi and Approved by the Bar Council of India and NCTE

Reference Material for Three Years

Bachelor of Economics (Hons.)

Code: 216

Semester – I

FIMT Campus, Kapashera, New Delhi-110037, Phones : 011-25063208/09/10/11, 25066256/ 57/58/59/60
Fax : 011-250 63212 Mob. : 09312352942, 09811568155 E-mail : fimtoffice@gmail.com Website : www.fimt-ggsipu.org

DISCLAIMER: FIMT, ND has exercised due care and caution in collecting the data before publishing this Reference Material. In spite of this, if any omission, inaccuracy or any other error occurs with regards to the data contained in this reference material, FIMT, ND will not be held responsible or liable. FIMT, ND will be grateful if you could point out any such error or your suggestions which will be of great help for other readers.

INDEX

Three Years Bachelor of Economics (Hons.)

Code: 216

Semester – I

S.NO.	SUBJECTS	CODE	PG.NO.
1	<i>PRIN. OF MICRO ECONOMICS</i>	101	03-39
2	<i>STATISTICAL METHOD-I</i>	103	40-102
3	<i>MATHEMATICS FOR ECONOMICS-I</i>	105	103-132
4	<i>BUSINESS ENGLISH-I</i>	107	133-175

PRINCIPLE OF MICRO ECONOMICS (101)

UNIT-1

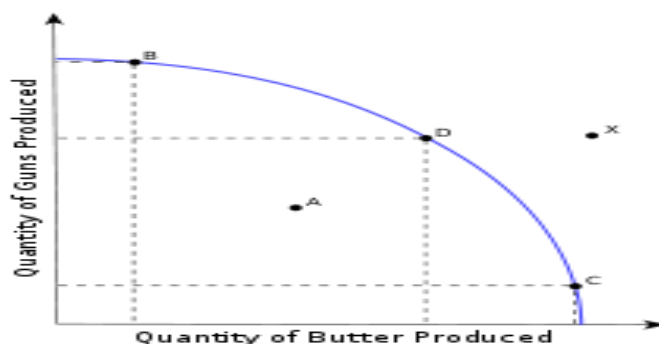
Opportunity Cost

The cost of an alternative that must be forgone in order to pursue a certain action is called opportunity cost. Put another way, the benefits you could have received by taking an alternative action. The difference in return between a chosen investment and one that is necessarily passed up, Say you invest in a stock and it returns a paltry 2% over the year. In placing your money in the stock, you gave up the opportunity of another investment - say, a risk-free government bond yielding 6%. In this situation, your opportunity costs are 4% ($6\% - 2\%$).

Marginalism

The concept of 'margin' is very popular in Economics. For example, in formal economic theory we learn that a business firm makes a decision to produce by equating marginal revenues with marginal costs. Marginal product is the addition made to total product (subtraction from the product) as a result of employing an additional (withdrawing the last factor of production. In economic theory, the concept of 'margin' is very useful; it renders the determination/derivation of an equilibrium solution quite simple and easy. However, in the real world of business management, marginalism should better be replaced by incrementalism, -In making economic decision, management is interested in knowing the impact of a chunk-change rather than a unit-change. Incremental reasoning involves a measurement of the impact of decision alternatives on economic variables like revenue and costs. Incremental revenues (or costs), for example, refer to the total magnitude of changes in total revenues (or costs) that result from a set of factors like change in prices, products, processes and patterns.

Production Possibility Frontier:-



A PPF (production possibility frontier) typically takes the form of the curve illustrated on the right. An economy that is operating on the PPF is said to be efficient, meaning that it would be impossible to produce more of one good without decreasing production of the other good. In contrast, if the economy is operating below the curve, it is said to be operating inefficiently because it could reallocate resources in order to produce more of both goods, or because some resources such as labor or capital are sitting idle and could be fully employed to produce more of both goods.

For example, assuming that the economy's available quantities of factors of production do not change over time and that technological progress does not occur, then if the economy is operating on the PPF production of guns would need to be sacrificed in order to produce more butter. If production is efficient, the economy can choose between combinations (i.e., points) on the PPF: *B* if guns are of interest, *C* if more butter is needed, *D* if an equal mix of butter and guns is required.^[1]

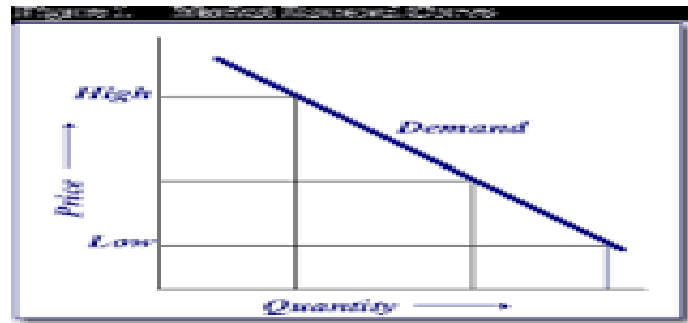
In the PPF, all points *on* the curve are points of maximum productive efficiency (i.e., no more output of any good can be achieved from the given inputs without sacrificing output of some good); all points inside the frontier (such as *A*) can be produced but are productively *inefficient*; all points outside the curve (such as *X*) cannot be produced with the given, existing resources. Not all points on the curve are Pareto efficient, however; only in the case where the marginal rate of transformation is equal to all consumers' marginal rate of substitution and hence equal to the ratio of prices will it be impossible to find any trade that will make no consumer worse off.

Any point that lies either on the production possibilities curve or to the left of it is said to be an attainable point, meaning that it can be produced with currently available resources. Points that lie to the right of the production possibilities curve are said to be unattainable because they cannot be produced using currently available resources. Points that lie strictly to the left of the curve are said to be inefficient, because existing resources would allow for production of more of at least one good without sacrificing the production of any other good. An efficient point is one that lies on the production possibilities curve. At any such point, more of one good can be produced only by producing less of the other.

Law of Demand

“Other factors remaining constant there is an inverse relationship between the price of a good and demand.”

As prices fall, we see an expansion of demand, If price rises, there will be a contraction of demand. A change in the price of a good or service causes a movement along the demand curve:
Many other factors can affect total demand - when these change, the demand curve can shift.



Consumer Equilibrium –Cardinal Utility Approach

The theory of consumer's behavior seeks to explain the determination of consumer's equilibrium. Consumer's equilibrium refers to a situation when a consumer gets maximum satisfaction out of his given resources. A consumer spends his money income on different goods and services in such a manner as to derive maximum satisfaction. Once a consumer attains equilibrium position, he would not like to deviate from it. Economic theory has approached the problem of determination of consumer's equilibrium in two different ways:

- (1) Cardinal Utility Analysis and
- (2) Ordinal Utility Analysis

Accordingly, we shall examine these two approaches to the study of consumer's equilibrium in greater detail.

Meaning of Utility:

The term utility in Economics is used to denote that quality in a good or service by virtue of which our wants are satisfied. In other words utility is defined as the want satisfying power of a commodity. According to Mrs. Robinson, "Utility is the quality in commodities that makes individuals want to buy them."

According to **Hibdon**, "Utility is the quality of a good to satisfy a want."

Concepts of Utility

There are three concepts of utility :

(1) **Initial Utility:** The utility derived from the first unit of a commodity is called initial utility. Utility derived from the first piece of bread is called initial utility. Thus, initial utility, is the utility obtained from the consumption of the first unit of a commodity. It is always positive.

(2) **Total Utility:** Total utility is the sum of utility derived from different Units of a commodity consumed by a household.

Suppose a consumer consume four units of apple. If the consumer gets 10 utils from the consumption of first apple, 8 utils from second, 6 utils from third, and 4 utils from fourth apple, then the total utility will be $10+8+6+4 = 28$

Accordingly, total utility can be calculated as :

$$TU = MU_1 + MU_2 + MU_3 + \text{_____} + MU_n$$

$$\text{Or } TU = \sum MU$$

Here TU = Total utility and $MU_1, MU_2, MU_3, + \text{_____} MU_n$ = Marginal Utility derived from first, second, third _____ and nth unit.

(3) **Marginal Utility:** Marginal Utility is the utility derived from the additional unit of a commodity consumed. The change that takes place in the total utility by the consumption of an additional unit of a commodity is called marginal utility.

According to Chapman, “Marginal utility is the addition made to total utility by consuming one more unit of commodity. Supposing a consumer gets 10 utils from the consumption of one mango and 18 utils from two mangoes, then. the marginal utility of second .mango will be $18 - 10 = 8$ utils.

Marginal utility can be measured with the help of the following formula

$$MU_{n+1} = TU_n - TU_{n-1}$$

Here MU_{n+1} = Marginal utility of nth unit,

TU_n = Total utility of ‘n’ units,

TU_{n-1} = Total utility of n-1 units,

Marginal utility can be (i) positive, (ii) zero, or (iii) negative.

(i) Positive Marginal Utility: If by consuming additional units of a commodity, total utility goes on increasing, marginal utility will be positive.

(ii) Zero Marginal Utility: If the consumption of an additional unit of a commodity causes no change in total utility, marginal utility will be zero.

(iii) Negative Marginal Utility: If the consumption of an additional unit of a commodity causes fall in total utility, the marginal utility will be negative.

Utility Analysis or Cardinal Approach

The Cardinal Approach to the theory of consumer behavior is based upon the concept of utility. It assumes that utility is capable of measurement. It can be added, subtracted, multiplied, and so on.

According to this approach, utility can be measured in cardinal numbers, like 1,2,3,4 etc. Fisher has used the term 'Util' as a measure of utility. Thus in terms of cardinal approach it can be said that one gets from a cup of tea 5 utils, from a cup of coffee 10 utils, and from a rasgulla 15 utils worth of utility.

Laws of Utility Analysis

Utility analysis consists of two important laws

1. Law of Diminishing Marginal Utility.
2. Law of Equi-Marginal Utility.

1. Law of Diminishing Marginal Utility:

Law of Diminishing Marginal Utility is an important law of utility analysis. This law is related to the satisfaction of human wants. All of us experience this law in our daily life. If you are set to buy, say, shirts at any given time, then as the number of shirts with you goes on increasing, the marginal utility from each successive shirt will go on decreasing. It is the reality of a man's life which is referred to in economics as law of Diminishing Marginal Utility. This law is also known as Gossen's First Law.

According to Chapman, “*The more we have of a thing, the less we want additional increments of it or the more we want not to have additional increments of it.*”

In short, the law of Diminishing Marginal Utility states that, other things being equal, when we go on consuming additional units of a commodity, the marginal utility from each successive unit of that commodity goes on diminishing.

Assumptions:

Every law is subject to clause “other things being equal” This refers to the assumption on which a law is based. It applies in this case as well. Main assumptions of this law are as follows:

1. Utility can be measured in cardinal number system such as 1, 2, 3_____ etc.
2. There is no change in income of the consumer.
3. Marginal utility of money remains constant.
4. Suitable quantity of the commodity is consumed.
5. There is continuous consumption of the commodity.
6. Marginal Utility of every commodity is independent.
7. Every unit of the commodity being used is of same quality and size.
8. There is no change in the tastes, character, fashion, and habits of the Consumer.
9. There is no change in the price of the commodity and its substitutes.

2. Law of Equi-Marginal Utility

This law states that the consumer maximizing his total utility will allocate his income among various commodities in such a way that his marginal utility of the last rupee spent on each commodity is equal. The consumer will spend his money income on different goods in such a way that marginal utility of each good is proportional to its price.

Limitations of Law of Equi-Marginal Utility

- It is difficult for the consumer to know the marginal utilities from different commodities because utility cannot be measured.

- Consumers are ignorant and therefore are not in a position to arrive at equilibrium.
- It does not apply to indivisible and inexpensive commodity.

Indifference Curve

An indifference curve is a geometrical presentation of a consumer's scale of preferences. It represents all those combinations of two goods which will provide equal satisfaction to a consumer. A consumer is indifferent towards the different combinations located on such a curve. Since each combination yields the same level of satisfaction, the total satisfaction derived from any of these combinations remains constant.

An indifference curve is a locus of all such points which shows different combinations of two commodities which yield equal satisfaction to the consumer. Since the combination represented by each point on the indifference curve yields equal satisfaction, a consumer becomes indifferent about their choice. In other words, he gives equal importance to all the combinations on a given indifference curve.

According to Ferguson, "An indifference curve is a combination of goods, each of which yield the same level of total utility to which the consumer is indifferent."

Indifference Schedule

An indifference schedule refers to a schedule that indicates different combinations of two commodities which yield equal satisfaction. A consumer, therefore, gives equal importance to each of the combinations.

Assumptions:

Indifference curve approach has the following main assumptions:

1. Rational Consumer: It is assumed that the consumer will behave rationally. It means the consumer would like to get maximum satisfaction out of his total income.

2. Diminishing Marginal rate of Substitution: It means as the stock of a commodity increases with the

consumer, he substitutes it for the other commodity at a diminishing rate.

3.Ordinal Utility: A consumer can determine his preferences on the basis of satisfaction derived from different goods or their combinations. Utility can be expressed in terms of ordinal numbers, i.e., first, second etc.

4. Independent Scale of Preference: It means if the income of the consumer changes or prices of goods fall or rise in the market, these changes will have no effect on the scale of preference of the consumer. It is further assumed that scale of preference of a consumer is not influenced by the scale of preference of another consumer.

5. Non-Satiety: A consumer does not possess any good in more than the required quantity. He does not reach the level of satiety. Consumer prefers more quantity of a good to less quantity.

6. Consistency in Selection: There is a consistency in consumer's behavior. It means that if at any given time a consumer prefers A combination of goods to B combination, then at another time he will not prefer B combination to A combination.

7. Transitivity: It means if a consumer prefers A combination to B combination, and B Combination to C Combination, he will definitely prefer A combination to C combination. Likewise; if a consumer is indifferent towards A and B and he is also indifferent towards Band C, then he will also he indifferent towards A and C.

Properties of Indifference Curves

1. Indifference curve slopes downward from left to right, or an indifference curve has a Negative slope
2. Indifference curve is convex to the point of origin:
3. Two Indifference Curves never cut each other:
4. Higher Indifference Curves represent more satisfaction
5. Indifference Curve touches neither x-axis nor y-axis;
6. Indifference curves need not be parallel to each other:

Indifference Curve & Price Effect

A **price effect** represents change in consumer's optimal consumption combination on account of change in the price of a good and thereby changes in its quantity purchased, price of another good and consumer's income remaining unchanged. The consumer is better-off when optimal consumption combination is located on a higher indifference curve and vice versa.

Understand that a consumer's responses to a price change differ depending upon the nature of the good, viz. a normal good, inferior good or a neutral good.

Type Of price effect;

1. Positive
2. Negative
3. Zero

These are summarized in chart.1:

Impact of fall in price of good X on its quantity demanded		
Type of Price Effect	Nature of Good X	Quantity Demanded of Good X
Positive	Normal	↑
Negative	Inferior (including Giffen Goods)	↓
Zero	Neutral	No Change in Quantity Demanded

Thus, a price effect is positive in case of normal goods. There is inverse relationship between price and quantity demanded. It is negative in case of inferior goods (including Giffen goods) where we find direct relationship between price and quantity demanded. Finally, price effect is zero in case of neutral goods where consumer's quantity demanded is fixed.

Indifference Curve & Substitution Effect

The substitution effect relates to the change in the quantity demanded resulting from a change in the price of good due to the substitution of relatively cheaper good for a dearer one, while keeping the price of the other good and real income and tastes of the consumer as constant. Prof. Hicks has explained the substitution effect independent of the income effect through compensating variation in income. “The substitution effect is the increase in the quantity bought as the price of the commodity falls, after adjusting income so as to keep the real purchasing power of the consumer the same as before. This adjustment in income is called compensating variations and is shown graphically by a parallel shift of the new budget line until it become tangent to the initial indifference curve.”

Thus on the basis of the methods of compensating variation, the substitution effect measure the effect of change in the relative price of a good with real income constant. The increase in the real income of the consumer as a result of fall in the price of, say good X, is so withdrawn that he is neither better off nor worse off than before.

The substitution effect is explained in Figure 12.17 where the original budget line is PQ with equilibrium at point R on the indifference curve I_1 . At R, the consumer is buying OB of X and BR of Y. Suppose the price of X falls so that his new budget line is PQ_1 . With the fall in the price of X, the real income of the consumer increases. To make the compensating variation in income or to keep the consumer's real income constant, take away the increase in his income equal to PM of good Y or Q_1N of good X so that his budget line PQ_1 shifts to the left as MN and is parallel to it.

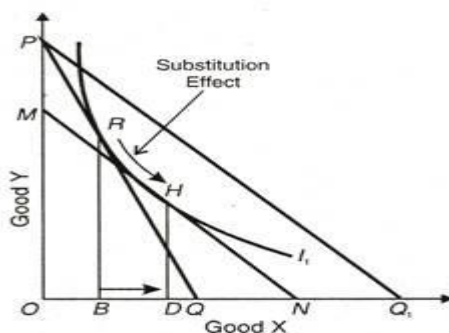


Fig. 12.17

At the same time, MN is tangent to the original indifference curve I_1 but at point H where the consumer buys OD of X and DH of Y. Thus PM of Y or Q_1N of X represents the compensating variation in income, as shown by the line MN being tangent to the curve I_1 at point H. Now the consumer substitutes X for Y and moves from point R to H or the horizontal distance from B to D. This movement is called

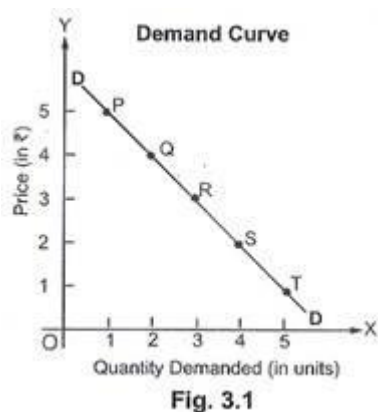
the substitution effect. The substitution effect is always negative because when the price of a good falls (or rises), more (or less) of it would be purchased, the real income of the consumer and price of the other good remaining constant. In other words, the relation between price and quantity demanded being inverse, the substitution effect is negative.

Unit-2

Individual demand curve

Individual demand curve refers to a graphical representation of individual demand schedule. With the help of Table 3.1 (Individual demand schedule), the individual demand curve can be drawn as shown in Fig. 3.1.

As seen in the diagram, price (independent variable) is taken on the vertical axis (Y-axis) and quantity demanded (dependent variable) on the horizontal axis (X-axis). At each possible price, there is a quantity, which the consumer is willing to buy. By joining all the points (P to T), we get a demand curve 'DD'.

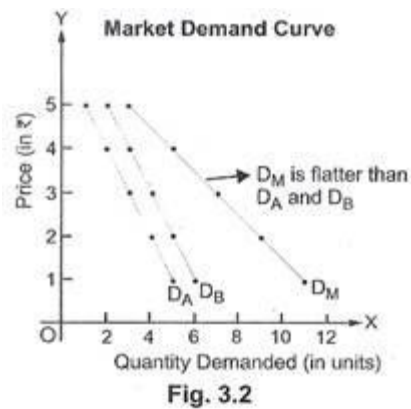


The demand curve 'DD' slopes downwards due to inverse relationship between price and quantity demanded.

Market Demand Curve:

Market demand curve refers to a graphical representation of market demand schedule. It is obtained by horizontal summation of individual demand curves.

The points shown in Table 3.2 are graphically represented in Fig. 3.2. D_A and D_B are the individual demand curves. Market demand curve (D_M) is obtained by horizontal summation of the individual demand curves (D_A and D_B).



Market demand curve ' D_M ' also slope downwards due to inverse relationship between price and quantity demanded.

Market Demand Curve is Flatter:

Market demand curve is flatter than the individual demand curves. It happens because as price changes, proportionate change in market demand is more than proportionate change in individual demand.

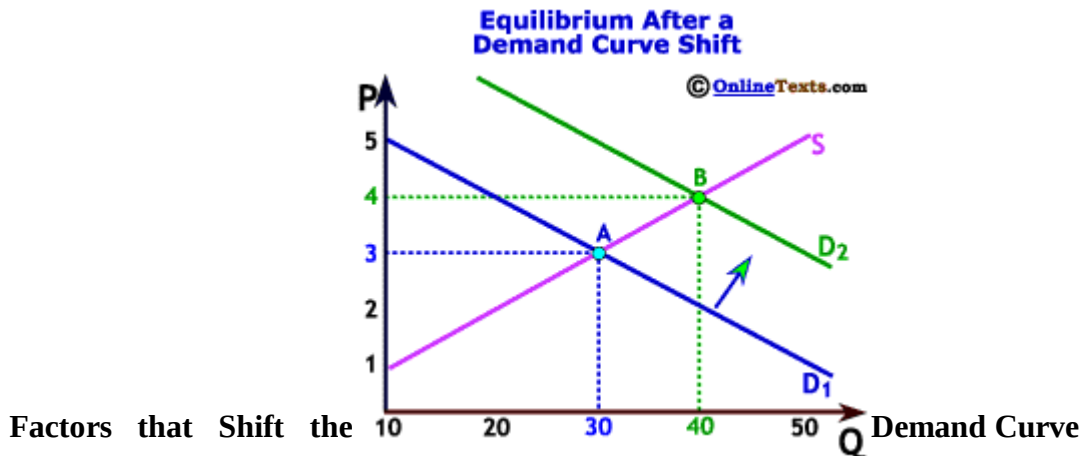
Movement along Vs shift market demand curve

It is essential to distinguish between a *movement along* a demand curve and a *shift* in the demand curve. A change in price results in a movement along a fixed demand curve. This is also referred to as a *change in quantity demanded*. For example, an increase in video rental prices from \$3 to \$4 reduces quantity demanded from 30 units to 20 units. This price change results in a movement along a given demand curve. A change in any other variable that influences quantity demanded produces a shift in the demand curve or a *change in demand*. The terminology is subtle but extremely important. The majority of the confusion that students have with supply and demand concepts involves understanding the differences between shifts and movements along curves.

TABLE 4 Change in Demand for Videos after Incomes Rise			
Price	Initial Quantity Demanded	New Quantity Demanded	Quantity Supplied
\$5	10	30	50
\$4	20	40	40
\$3	30	50	30
\$2	40	60	20
\$1	50	70	10

Suppose that incomes in a community rise because a factory is able to give employees overtime pay. The higher incomes prompt people to rent more videos. For the *same rental price*, quantity demanded is now *higher* than before. Table 4 and the figure titled "Shift in the Demand Curve" represent that scenario. As incomes rise, the quantity demanded for videos priced at \$4 goes from 20 (point A) to 40 (point A'). Similarly, the quantity demanded for videos priced at \$3 rises from 30 to 50. The entire demand curve shifts to the right.

A shift in the demand curve changes the equilibrium position. As illustrated in the figure titled "Equilibrium After a Demand Curve Shift" the shift in the demand curve moves the market equilibrium from point A to point B, resulting in a higher price (from \$3 to \$4) and higher quantity (from 30 to 40 units). Note that if the demand curve shifted to the left, both the equilibrium price and quantity would decline.



1. **Change in consumer incomes:** As the previous video rental example demonstrated, an increase in income shifts the demand curve to the right. Because a consumer's demand for goods and services is constrained by income, higher income levels relax somewhat that constraint, allowing the consumer to purchase more products. Correspondingly, a decrease in income shifts the demand curve to the left. When the economy enters a recession and more people become unemployed, the demand for many goods and services shifts to the left.
2. **Population change:** An increase in population shifts the demand curve to the right. Imagine a college town bookstore in which most students return home for the summer. Demand for books shifts to the left while the students are away. When they return, however, demand for books increases even if the prices are unchanged. As another example, many communities are experiencing "urban sprawl" where the metropolitan boundaries are pushed ever wider by new housing developments. Demand for gasoline in these new communities increases with population. Alternatively, demand for gasoline falls in areas with declining populations.
3. **Consumer preferences:** If the preference for a particular good increases, the demand curve for that good shifts to the right. Fads provide excellent examples of changing consumer preferences. Each Christmas season some new toy catches the fancy of kids, and parents scramble to purchase the product before it is sold out. A few years ago, "Tickle Me Elmo" dolls were the rage. In the year 2000 the toy of choice was a scooter. For a given price of a scooter, the demand curve shifts to the right as more consumers decide that they wish to purchase that product for their children. Of course, demand curves can shift leftward just as quickly. When fads end suppliers often find themselves with a glut of merchandise that they discount heavily to sell.
4. **Prices of related goods:** If prices of related goods change, the demand curve for the original

good can change as well. Related goods can either be substitutes or complements.

- **Substitutes** are goods that can be consumed in place of one another. If the price of a substitute increases, the demand curve for the original good shifts to the right. For example, if the price of Pepsi rises, the demand curve for Coke shifts to the right. Conversely, if the price of a substitute decreases, the demand curve for the original good shifts to the left. Given that chicken and fish are substitutes, if the price of fish falls, the demand curve for chicken shifts to the left.
- **Complements** are goods that are normally consumed together. Hamburgers and french fries are complements. If the price of a complement increases, the demand curve for the original good shifts to the left. For example, if McDonalds raises the price of its Big Mac, the demand for french fries shifts to the left because fewer people walk in the door to buy the Big Mac. In contrast, If the price of a complement decreases, the demand curve for the original good shifts to the right. If, for example, the price of computers falls, then the demand curve for computer software shifts to the right.

Elasticity of Demand

The elasticity of demand (E_d), also referred to as the price elasticity of demand, measures how responsive demand is to changes in a price of a given good. More precisely, it is the percent change in quantity demanded relative to a one percent change in price, holding all else constant. Demand of goods can be classified as either perfectly elastic, elastic, unitary elastic, inelastic, or perfectly inelastic based on the elasticity of demand. This table shows the values of elasticity of demand that correspond to the different categories.

The graph illustrates the demand curves and places along the demand curve that correspond to the table. The elasticity of demand changes as one moves along the demand curve. This is an important concept - the elasticity of demand for a good changes as you evaluate it at different price points.

1. Percentage method or Arithmetic method
2. Total Expenditure method
3. Graphic method or point method.

1. Percentage method:-

According to this method price elasticity is estimated by dividing the percentage change in amount

demand by the percentage change in price of the commodity. Thus given the percentage change of both amount demanded and price we can derive elasticity of demand. If the percentage change in amount demanded is greater than the percentage change in price, the coefficient thus derived will be greater than one.

If percentage change in amount demanded is less than percentage change in price, the elasticity is said to be less than one. But if percentage change of both amount demanded and price is same, elasticity of demand is said to be unit.

2. Total expenditure method

Total expenditure method was formulated by Alfred Marshall. The elasticity of demand can be measured on the basis of change in total expenditure in response to a change in price. It is worth noting that unlike percentage method a precise mathematical coefficient cannot be determined to know the elasticity of demand.

By the help of total expenditure method we can know whether the price elasticity is equal to one, greater than one, less than one. In such a method the initial expenditure before the change in price and the expenditure after the fall in price are compared. By such comparison, if it is found that the expenditure remains the same, elasticity of demand is One ($ed=1$).

If the total expenditure increases the elasticity of demand is greater than one ($ed>1$). If the total expenditure diminished with the change in price elasticity of demand is less than one ($ed<1$). The total expenditure method is illustrated by the following diagram.

3. Graphic method:

Graphic method is otherwise known as point method or Geometric method. This method was popularized by method. According to this method elasticity of demand is measured on different points on a straight line demand curve. The price elasticity of demand at a point on a straight line is equal to the lower segment of the demand curve divided by upper segment of the demand curve.

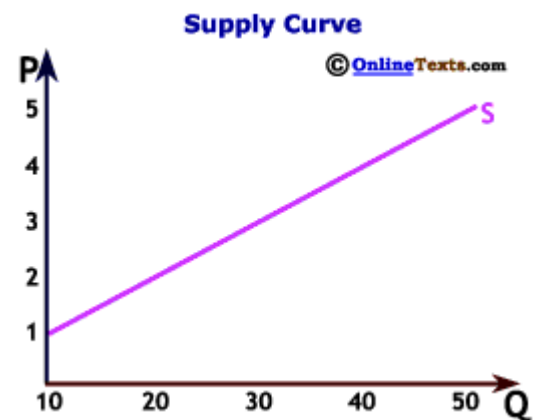
Thus at mid point on a straight-line demand curve, elasticity will be equal to unity; at higher points on the same demand curve, but to the left of the mid-point, elasticity will be greater than unity, at lower points on the demand curve, but to the right of the mid-point, elasticity will be less than unity.

The theory & Law of Supply and the Supply Curve

Supply is slightly more difficult to understand because most of us have little direct experience on the supply side of the market. Supply is derived from a producer's desire to maximize profits. When the price of a product rises, the supplier has an incentive to increase production because he can justify higher costs to produce the product, increasing the potential to earn larger profits. Profit is the difference between revenues and costs. If the producer can raise the price and sell the same number of goods while holding costs constant, then profits increase.

The *law of supply* holds that other things equal, as the price of a good rises, its quantity supplied will rise, and vice versa. Table 2 lists the quantity supplied of rental videos for various prices. At \$5, the producer has an incentive to supply 50 videos. If the price falls to \$4 quantity supplied falls to 40, and so on. The figure titled "Supply Curve" plots this positive relationship between price and quantity supplied.

TABLE 2	
Supply of Videos	
Price	Quantity Supplied
\$5	50
\$4	40
\$3	30
\$2	20
\$1	10



A *supply curve* is a graphical depiction of a supply schedule plotting price on the vertical axis and quantity supplied on the horizontal axis. The supply curve is upward-sloping, reflecting the law of supply.

Equilibrium Supply Curve & determination of Price and Quantity

What price should the seller set and how many videos will be rented per month? The seller could legally set any price she wished; however, market forces penalize her for making poor choices.

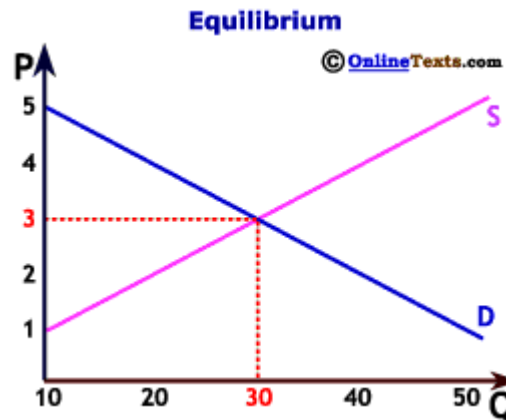
Suppose, for example, that the seller prices each video at \$20. Odds are good that few videos will be rented. On the other hand, the seller may set a price of \$1 per video. Consumers will certainly rent more videos with this low price, so much so that the store is likely to run out of videos. Through trial and error or good judgement, the store owner will eventually settle on a price that equates the forces of supply and demand.

In economics, an *equilibrium* is a situation in which:

- there is no inherent tendency to change,
- quantity demanded equals quantity supplied, and
- The market just clears.

At the market equilibrium, every consumer who wishes to purchase the product at the market price is able to do so, and the supplier is not left with any unwanted inventory. As Table 3 and the figure titled "Equilibrium" demonstrate, equilibrium in the video example occurs at a price of \$3 and a quantity of 30 videos.

TABLE 3 Video Market Equilibrium		
Price	Quantity Demanded	Quantity Supplied
\$5	10	50
\$4	20	40
\$3	30	30
\$2	40	20
\$1	50	10



Elasticity of Supply

Price elasticity of supply (PES or E_s) is a measure used in economics to show the responsiveness, or

elasticity, of the quantity supplied of a good or service to a change in its price.

The elasticity is represented in numerical form, and is defined as the percentage change in the quantity supplied divided by the percentage change in price.

When the coefficient is less than one, the supply of the good can be described as *inelastic*; when the coefficient is greater than one, the supply can be described as *elastic*. An elasticity of zero indicates that quantity supplied does not respond to a price change: it is "fixed" in supply. Such goods often have no labor component or are not produced, limiting the short run prospects of expansion. If the coefficient is exactly one, the good is said to be *unitary elastic*.

The quantity of goods supplied can, in the short term, be different from the amount produced, as manufacturers will have stocks which they can build up or run down.

Shift in the Supply Curve

1. **Change in input costs:** An increase in input costs shifts the supply curve to the left. A supplier combines raw materials, capital, and labor to produce the output. If a furniture maker has to pay more for lumber, then her profits decline, all else equal. The less attractive profit opportunities force the producer to cut output. Alternatively, car manufacturer may have to pay higher labor costs. The higher labor input costs reduces profits, all else equal. For a given price of a car, the manufacturer may trim output, shifting the supply curve to the left. Conversely, if input costs decline, firms respond by increasing output. The furniture manufacturer may increase production if lumber costs fall. Additionally, chicken farmers may boost chicken output if feed costs decline. The reduction in feed costs shifts the supply curve for chicken to the right.
2. **Increase in technology:** An increase in technology shifts the supply curve to the right. A narrow definition of technology is a cost-reducing innovation. Technological progress allows firms to produce a given item at a lower cost. Computer prices, for example, have declined radically as technology has improved, lowering their cost of production. Advances in communications technology have lowered the telecommunications costs over time. With the advancement of technology, the supply curve for goods and services shifts to the right.
3. **Change in size of the industry:** If the size of an industry grows, the supply curve shifts to the right. In short, as more firms enter a given industry, output increases even as the price remains

steady. The fast-food industry, for example, exploded in the latter half of the twentieth century as more and more fast food chains entered the market. Additionally, on-line stock trading has increased as more firms have begun delivering that service. Conversely, the supply curve shifts to the left as the size of an industry shrinks. For example, the supply of manual typewriters declined dramatically in the 1990s as the number of producers dwindled.

Unit-3

Meaning and Concept of Production

Production

Production is transformation of tangible inputs (raw materials, semi-finished goods, Sub assemblies) and intangible inputs (ideas, information, knowledge) into output(goods or services) in a specific period of time at given state of technology. Resources are used in this process to create an output that is suitable for use or has exchange value.

Factor of production

In economics, factors of production are the inputs to the production process. Finished goods are the output. Input determines the quantity of output i.e. output depends upon input. Input is the starting point and output is the end point of production process and such input-output relationship is called a production function. The product of one industry may be used in another industry.

For E.G., wheat is a output for a framer; but when it is used to produce bread it becomes a factor of production.

There are three basic factors of production:

- Land

- Labor
- Capital
- Entrepreneur

All three of these are required in combination at a time to produce a commodity.

'Factors of production' may also refer specifically to the 'primary factors', which are stocks including land, labor (the ability to work), and capital goods applied to production. Materials and energy are considered secondary factors in classical economics because they are obtained from land, labor and capital. The primary factors facilitate production but neither become part of the product (as with raw materials) nor become significantly transformed by the production process (as with fuel used to power machinery).

Four factors of production

- 1. Land**
- 2. Labor**
- 3. Capital**
- 4. Entrepreneur**

According to **Prof. Benham**, "Anything that contributes towards output is a factor of production."

Cooperation among factors is essential to produce anything because production is not a job of single factor.

Production Function

Production is transformation of tangible inputs (raw materials, semi-finished goods, Sub assemblies) and intangible inputs (ideas, information, knowledge) into output(goods or services) in a specific period of time at given state of technology. Output is thus, a function of inputs. Technical relation between inputs and outputs is depicted by production function. It denotes effective combination of inputs.

In economics, a **production function** relates physical output of a production process to physical inputs

or factors of production. In macroeconomics, aggregate production functions are estimated to create a framework in which to distinguish how much of economic growth to attribute to changes in factor allocation (e.g. the accumulation of capital) and how much to attribute to advancing technology. Some non-mainstream economists, however, reject the very concept of an aggregate production function.

Concept of production functions

In general, economic output is not a (mathematical) function of input, because any given set of inputs can be used to produce a range of outputs. To satisfy the mathematical definition of a function, a production function is customarily assumed to specify the maximum output obtainable from a given set of inputs. A production function can be defined as the specification of the minimum input requirements needed to produce designated quantities of output, given available technology. In the production function, itself, the relationship of output to inputs is non-monetary; that is, a production function relates physical inputs to physical outputs, and prices and costs are not reflected in the function. In the decision frame of a firm making economic choices regarding production—how much of each factor input to use to produce how much output—and facing market prices for output and inputs, the production function represents the possibilities afforded by an exogenous technology. Under certain assumptions, the production function can be used to derive a marginal product for each factor. The profit-maximizing firm in perfect competition (taking output and input prices as given) will choose to add input right up to the point where the marginal cost of additional input matches the marginal product in additional output. This implies an ideal division of the income generated from output into an income due to each input factor of production, equal to the marginal product of each input. The inputs to the production function are commonly termed factors of production and may represent primary factors, which are stocks. Classically, the primary factors of production were Land, Labor and Capital. Primary factors do not become part of the output product, nor are the primary factors, themselves, transformed in the production process.

Production function differs from firm to firm, industry to industry. Any change in the state of technology or managerial ability disturbs the original production function. Production function can be represented by schedules, graph, tables, mathematical equations, TP, AP & MP Curves, isoquant and so on.

Specifying the production function

A production function can be expressed in a functional form as the right side of

$$Q = f(K, L, I, O)$$

Where:

Q = quantity of output

K, L, I, O stand for quantities of factors of production (capital, labour, land or organization respectively) used in production

Stages of production

To simplify the interpretation of a production function, it is common to divide its range into 3 stages as follows:

Stage 1 (from the origin to point B): the variable input is being used with increasing output per unit, the latter reaching a maximum at point B (since the average physical product is at its maximum at that point). Because the output per unit of the variable input is improving throughout stage 1, a price-taking firm will always operate beyond this stage.

Stage 2: output increases at a decreasing rate, and the average and marginal physical product are declining. However, the average product of fixed inputs (not shown) is still rising, because output is rising while fixed input usage is constant. In this stage, the employment of additional variable inputs increases the output per unit of fixed input but decreases the output per unit of the variable input. The optimum input/output combination for the price-taking firm will be in stage 2, although a firm facing a downward-sloped demand curve might find it most profitable to operate in Stage 1.

Stage 3: too much variable input is being used relative to the available fixed inputs. Variable inputs are over-utilized in the sense that their presence on the margin obstructs the production process rather than enhancing it. The output per unit of both the fixed and the variable input declines throughout this stage. **At the boundary between stage 2 and stage 3, the highest possible output is being obtained from the fixed input.**

ISO QUANTS

An isoquant (iso product) is a curve on which the various combinations of labour and capital show the same output. According to Cohen and Cyert, “An isoproduct curve is a curve along which the maximum achievable rate of production is constant.” It is also known as a production indifference curve or a constant product curve. Just as indifference curve shows the various combinations of any two commodities that give the consumer the same amount of satisfaction (iso-utility), similarly an isoquant indicates the various combinations of two factors of production which give the producer the same level of output per unit of time. Table 24.1 shows a hypothetical isoquant schedule of a firm producing 100 units of a good.

TABLE 24.1: Isoquant Schedule:

Combination	Units of Capital	Units of Labour	Total Output (in units)
A	9	5	100
B	6	10	100
C	4	15	100
D	3	20	100

This Table 24.1 is illustrated on Figure 24.1 where labour units are measured along the X-axis and capital units on the K-axis. The first, second, third and the fourth combinations are shown as A, S, C and D respectively. Connect all these points and we have a curve IQ.

This is an isoquant. The firm can produce 100 units of output at point A on this curve by having a combination of 9 units of capital and 5 units of labour. Similarly, point B shows a combination of 6 units of capital and 10 units of labour; point C, 4 units of capital and 15 units of labour; and point D, a combination of 3 units of capital and 20 units of labour to yield the same output of 100 units.

An isoquant map shows a number of isoquants representing different amounts of output. In Figure 24.1, curves IQ, IQ1 and IQ2 show an isoquant map. Starting from the curve IQ which yields 100 units of product, the curve IQ1, shows 200 units and the IQ2 curve 300 units of the product which can be produced with altogether different combinations of the two factors.

Properties of Isoquants:

Isoquants possess certain properties which are similar to those of indifference curves.

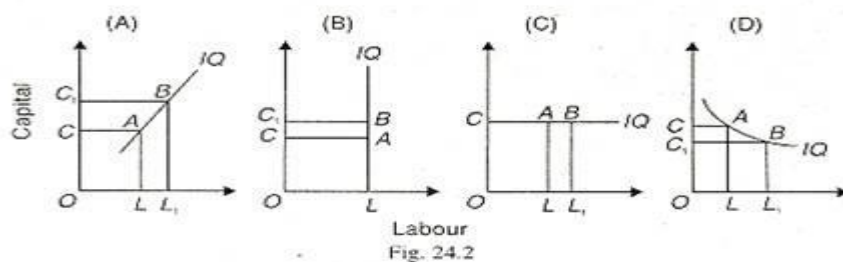
(1) Isoquants are negatively inclined:

If they do not have a negative slope, certain logical absurdities follow. If the isoquant slopes upward to the right, it implies that both capital and labour increase but they produce the same output. In Figure 24.2 (A), combination B on the IQ curve having a larger amount of both capital and labour ($OC_1 + OL_1 > OC + OL$) will yield more output than before. Therefore, point A and B on the IQ curve cannot be of equal product.

Suppose the isoquant is vertical as shown in Figure 24.2 (B), which implies a given amount of labour is combined with different units of capital. Since OL of labour and OC_1 of capital will produce a larger amount than produced by OL of labour and OC of capital, the isoquant IQ cannot be a constant product curve.

Take Figure 24.2 (C) where the isoquant is horizontal which means combining more of labour with the same quantity of capital. Here OC of capital and OL_1 of labour will produce a larger or smaller amount than produced by the combination OC of capital and OL of labour. Therefore, a horizontal isoquant cannot be an equal product curve.

Thus it is clear that an isoquant must slope downward to the right as shown in Figure 24.2 (D) where points A and B on the IQ curve are of equal quantity. As the amount of capital decreases from OC to OC_1 and that of labour increases from OL to OL_1 so that output remains constant.



(2) An Isoquant lying above and to the right of another represents a higher output level. In Figure 24.3 combination B on IQ₁ curve shows larger output than point A on the curve IQ. The combination of OC of capital and OL of labour yields 100 units of product while OC_1 of capital and OL_1 of labour produce 200 units. Therefore, the isoquant IQ₁ which lies above and to the right of the isoquant IQ, represents a larger output level.

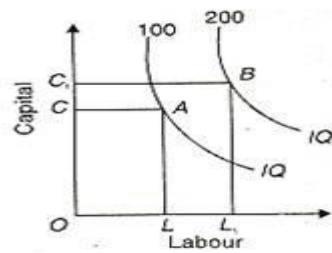


Fig. 24.3

(3) No two isoquants can intersect each other. The absurd conclusion that follows when two isoquants cut each other is explained with the aid of Figure 24.4. On the isoquant IQ , combination $A = B$. And on the isoquant IQ_1 combination $R = S$. But combination S is preferred to combination B , being on the higher portion of isoquant IQ_1 . On the other hand, combination A is preferred to R , the former being on the higher portion of the isoquant IQ . To put it algebraically, it means that $S > B$ and $R < A$. But this is logically absurd because S combination is as productive as R and A combination produces as much as B . Therefore, the same combination cannot both be less and more productive at the same time. Hence two isoquants cannot intersect each other.

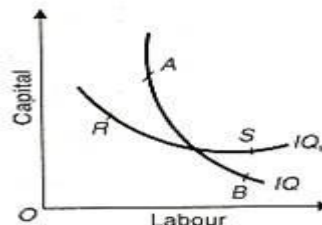


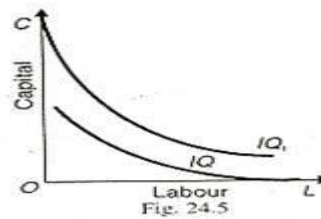
Fig. 24.4

(4) Isoquants need not be parallel because the rate of substitution between two factors is not necessarily the same in all the isoquant schedules.

(5) In between two isoquants there can be a number of isoquants showing various levels of output which the combinations of the two factors can yield. In fact, in between the units of output 100, 200, 300, etc. represented on isoquants there can be innumerable isoquants showing 120, 150, 175, 235, or any other higher or lower unit.

(6) Units of output shown on isoquants are arbitrary. The various units of output such as 100, 200, 300, etc., shown in an isoquant map are arbitrary. Any units of output such as 5, 10, 15, 20 or 1000, 2000, 3000, or any other units can be taken.

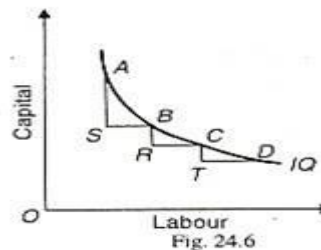
(7) No isoquant can touch either axis. If an isoquant touches X-axis, it would mean that the product is being produced with the help of labour alone without using capital at all. This is a logical absurdity for OL units of labour alone are incapable of producing anything. Similarly, OC units of capital alone cannot produce anything without the use of labour. Therefore IQ and IQ_1 cannot be isoquants, as shown in Figure 24.5.



(8) Each isoquant is convex to the origin:

As more units of labour are employed to produce 100 units of the product, lesser and lesser units of capital are used. This is because the marginal rate of substitution between two factors diminishes. In Figure 24.6, in order to produce 100 units of the product, as the producer moves along the isoquant from combination A to B and to C and D, he gives up smaller and smaller units of capital for additional units of labour. To maintain the same output of 100 units, BR less of capital and relatively RC more of labour is used.

If he were producing this output with the combination D, he would be employing CT less of capital and relatively TD more of labour. Thus the isoquants are convex to the origin due to diminishing marginal rate of substitution. This fact becomes clear from successively smaller triangles below the IQ curve $\Delta ASB > \Delta BRC > \Delta CTD$.



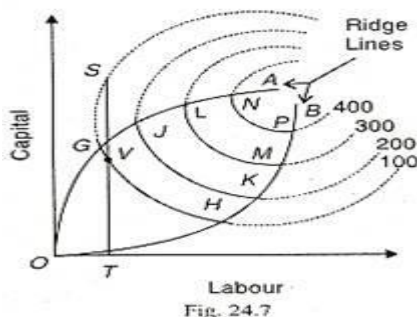
(9) Each isoquant is oval-shaped:

It is elliptical which means that at some point it begins to recede from each axis. This shape is a consequence of the fact that if a producer uses more of capital or more of labour or more of both than is necessary, the total product will eventually decline.

The firm will produce only in those segments of the isoquants which are convex to the origin and lie between the ridge lines.

This is the economic region of production. In Figure 24.7, oval-shaped isoquants are shown. Curves OA and OB are the ridge lines and in between them economically feasible units of capital and labour can be employed to produce 100, 200, 300 and 400 units of the product. For example, OT units of labour and ST units of the capital can produce 100 units of the product, but the same output can be obtained by using the same quantity of labour OT and less quantity of capital VT.

Thus only an unwise entrepreneur will produce in the dotted region of the isoquant 100. The dotted segments of an isoquant are the waste- bearing segments. They form the uneconomic regions of production. In the upper dotted portion, more capital and in the lower dotted portion more labour than necessary is employed. Hence GH, JK, LM, and NP segments of the elliptical curves are the iso- quants.



Isocost Curves:

Having studied the nature of isoquants which represent the output possibilities of a firm from a given combination of two inputs, we pass on to the prices of the inputs as represented on the isoquant map by the isocost curves. These curves are also known as outlay lines, price lines, input-price lines, factor-cost lines, constant-outlay lines, etc. Each isocost curve represents the different combinations of two inputs that a firm can buy for a given sum of money at the given price of each input.

Figure, 24.8 (A) shows three isocost curves AB, CD and EF, each represents a total outlay of 50, 75 and 100 respectively. The firm can hire OC of capital or OD of labour with Rs. 75. OC is $\frac{2}{3}$ of OD which means that the price of a unit of labour is $1\frac{1}{2}$ times less than that of a unit of capital. The line CD represents the price ratio of capital and labour. Prices of factors remaining the same, if the total outlay is raised, the isocost curve will shift upward to the right as EF parallel to CD, and if the total outlay is reduced it will shift downwards to the left as AB. The isocosts are straight lines because factor prices remain the same whatever the outlay of the firm on the two factors. The isocost curves represent the locus of all combinations of the two input factors which result in the same total cost. If the unit cost of labour (L) is w and the unit cost of capital (C) is r , then the total cost: $TC = wL + rC$. The slope of the isocost line is the ratio of prices of labour and capital i.e., w/r .

The point where the isocost line is tangent to an isoquant represents the least cost combination of the two factors for producing a given output. If all points of tangency like LMN are joined by a line, it is known as an output- factor curve or least-outlay curve or the expansion path of a firm. Salvatore defines expansion path as “the locus of points of producer’s equilibrium resulting from changes in total outlays while keeping factor prices constant.” It shows how the proportions of the two factors used might be changed as the firm expands.

For example, in Figure 24.8 (A) the proportions of capital and labour used to produce 200 (IQ₁) units of the product are different from the proportions of these factors used to produce 300 (IQ₂) units or 100 (OQ) units at the lowest cost.

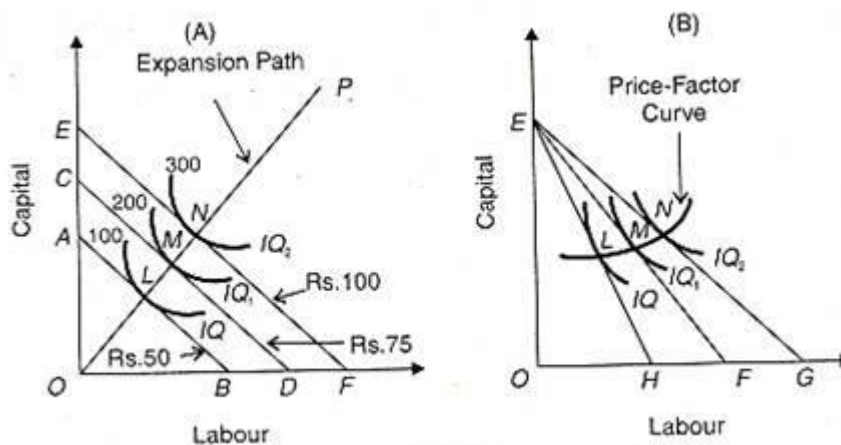


Fig. 24.8

Like the price-income line in the indifference curve analysis, a relative cheapening of one of the factors to that of another will extend the isocost line to the right. If one of the factors becomes relatively dearer, the isocost line will contract inward to the left. Given the price of capital, if the price of labour falls, the isocost line EF in Panel (B) will extend to the right as EG and if the price of labour rises, the isocost line EF will contract inward to the left as EH. If the equilibrium points L, M, and N are joined by a line, it is called the price-factor curve.

The Principle of Marginal Rate of Technical Substitution:

The principle of marginal rate of technical substitution (MRTS or MRS) is based on the production function where two factors can be substituted in variable proportions in such a way as to produce a constant level of output.

The marginal rate of technical substitution between two factors C (capital) and L (labour), MRTSLC is the rate at which L can be substituted for C in the production of good X without changing the quantity of output. As we move along an isoquant downward to the right, each point on it represents the substitution of labour for capital.

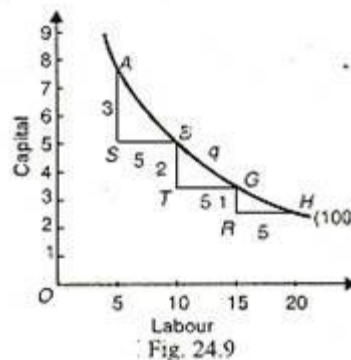
MRTS is the loss of certain units of capital which will just be compensated for by additional units of labour at that point. In other words, the marginal rate of technical substitution of labour for capital is the slope or gradient of the isoquant at a point. Accordingly, $\text{slope} = \text{MRTSLC} = - \Delta C / \Delta L$. This can be understood with the aid of the isoquant schedule, in Table 24.2.

TABLE 24.2: Isoquant Schedule:

Combination	Labour	Capital	MRTSLC	Output
1	5	9	—	100
2	10	6	3:5	100
3	15	4	2:5	100
4	20	3	1:5	100

The above table shows that in the second combination to keep output constant at 100 units, the reduction of 3 units of capital requires the addition of 5 units of labour, $MRTSLC = 3:5$. In the third combination, the loss of 2 units of capital is compensated for by 5 more units of labour, and so on.

In Figure 24.9 at point B, the marginal rate of technical substitution is AS/SB , at point G, it is BT/TG and at H, it is GR/ RH .



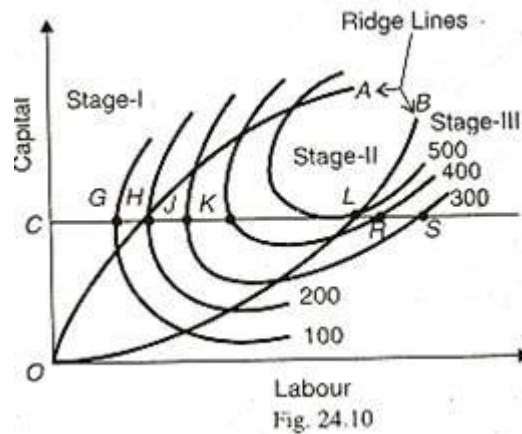
The isoquant reveals that as the units of labour are successively increased into the factor-combination to produce 100 units of good X, the reduction in the units of capital becomes smaller and smaller. It means that the marginal rate of technical substitution is diminishing. This concept of the diminishing marginal rate of technical substitution (DMRTS) is parallel to the principle of diminishing marginal rate of substitution in the indifference curve technique.

This tendency of diminishing marginal substitutability of factors is apparent from Table 24.2 and Figure 24.9. The MRTSLC continues to decline from 3:5 to 1:5 whereas in the Figure 24.9 the vertical lines below the triangles on the isoquant become smaller and smaller as we move downward so that $GR < BT < AS$. Thus, the marginal rate of technical substitution diminishes as labour is substituted for capital. It means that the isoquant must be convex to the origin at every point.

The Law of Variable Proportions:

The behaviour of the law of variable proportions or of the short-run production function when one factor is constant and the other variable can also be explained in terms of the isoquant analysis. Suppose capital is a fixed factor and labour is a variable factor. In Figure 24.10., OA and OB are the ridge lines and it is

in between them that economically feasible units of labour and capital can be employed to produce 100, 200, 300, 400 and 500 units of output.



It implies that in these portions of the isoquants, the marginal product of labour and capital is positive. On the other hand, where these ridge lines cut the isoquants, the marginal product of the inputs is zero. For instance, at point H the marginal product of capital is zero, and at point L the marginal product of labour is zero. The portion of the isoquant that lies outside the ridge lines, the marginal product of that factor is negative. For instance, the marginal product of capital is negative at G and that of labour at R. The law of variable proportions says that, given the technique of production, the application of more and more units of a variable factor, say labour, to a fixed factor, say capital, will, until a certain point is reached, yield more than proportional increases in output, and thereafter less than proportional increases in output.

Since the law refers to increases in output, it relates to the marginal product. To explain the law, capital is taken as a fixed factor and labour as a variable factor. The isoquants show different levels of output in the figure. OC is the fixed quantity of capital which therefore forms a horizontal line CD. As we move from C to D towards the right on this line, the different points show the effects of the combinations of successively increasing quantities of labour with fixed quantity of capital OC.

To begin with, as we move from C to G to H, it shows the first stage of increasing marginal returns of the law of variable proportions. When CG labour is employed with OC capital, output is 100. To produce 200 units of output, labour is increased by GH while the amount of capital is fixed at OC.

The output has doubled but the amount of labour employed has not increased proportionately. It may be observed that $GH < CG$, which means that smaller additions to the labour force have led to equal increment in output. Thus C to H is the first stage of the law of variable proportions in which the marginal product increases because output per unit of labour increases as more output is produced.

The second stage of the law of variable proportions is the portion of the isoquants which lies in between

the two ridge lines O A and OB. It is the stage of diminishing marginal returns between points H and L. As more labour is employed, output increases less than proportionately to the increase in the labour employed. To raise output to 300 units from 200 units, HJ labour is employed. Further, JK quantity of labour is required to raise output from 300 to 400 and KL of labour to raise output from 400 to 500.

So, to increase output by 100 units successively, more and more units of the variable factor (labour) are required to be applied along with the fixed factor (capital), that is $KL > JK > HJ$. It implies that the marginal product of labour continues to decline with the employment of larger quantities to it. Thus as we move from point H to K, the effect of increasing the units of labour is that output per unit of labour diminishes as more output is produced. This is known as the stage of diminishing returns.

If labour is employed further, we are outside the lower ridge line OB and enter the third stage of the law of variable proportions. In this region which lies beyond the ridge line OB there is too much of the variable factor (labour) in relation to the fixed factor (capital). Labour is thus being overworked and its marginal product is negative. In other words when the quantity of labour is increased by LR and RS, the output declines from 500 to 400 and to 300. This is the stage of negative marginal returns.

We arrive at the conclusion that a firm will find it profitable to produce only in the second stage of the law of variable proportions for it will be uneconomical to produce in the regions to the left or right of the ridge lines which form the first stage and the third stage of the law respectively.

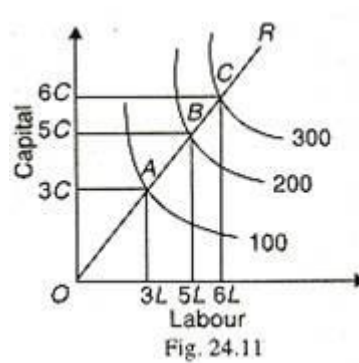
The Laws of Returns to Scale:

The laws of returns to scale can also be explained in terms of the isoquant approach. The laws of returns to scale refer to the effects of a change in the scale of factors (inputs) upon output in the long-run when the combinations of factors are changed in some proportion. If by increasing two factors, say labour and capital, in the same proportion, output increases in exactly the same proportion, there are constant returns to scale. If in order to secure equal increases in output, both factors are increased in larger proportionate units, there are decreasing returns to scale. If in order to get equal increases in output, both factors are increased in smaller proportionate units, there are increasing returns to scale.

The returns to scale can be shown diagrammatically on an expansion path “by the distance between successive ‘multiple-level-of-output’ isoquants, that is, isoquants that show levels of output which are multiples of some base level of output, e.g., 100, 200, 300, etc.”

Increasing Returns to Scale:

Figure 24.11 shows the case of increasing returns to scale where to get equal increases in output, lesser proportionate increases in both factors, labour and capital, are required.



It follows that in the figure:

100 units of output require $3C + 3L$

200 units of output require $5C + 5L$

300 units of output require $6C + 6L$

So that along the expansion path OR, $OA > AB > BC$. In this case, the production function is homogeneous of degree greater than one.

The increasing returns to scale are attributed to the existence of indivisibilities in machines, management, labour, finance, etc. Some items of equipment or some activities have a minimum size and cannot be divided into smaller units. When a business unit expands, the returns to scale increase because the indivisible factors are employed to their full capacity.

Increasing returns to scale also result from specialisation and division of labour. When the scale of the firm expands there is wide scope for specialisation and division of labour. Work can be divided into small tasks and workers can be concentrated to narrower range of processes. For this, specialized equipment can be installed. Thus with specialization, efficiency increases and increasing returns to scale follow.

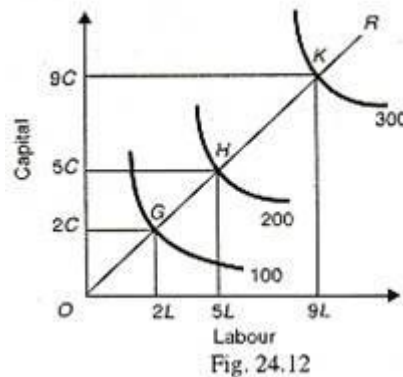
Further, as the firm expands, it enjoys internal economies of production. It may be able to install better machines, sell its products more easily, borrow money cheaply, procure the services of more efficient manager and workers, etc. All these economies help in increasing the returns to scale more than proportionately.

Not only this, a firm also enjoys increasing returns to scale due to external economies. When the industry itself expands to meet the increased 'long-run demand for its product, external economies appear which are shared by all the firms in the industry. When a large number of firms are concentrated at one place, skilled labour, credit and transport facilities are easily available. Subsidiary industries crop up to help the main industry. Trade journals, research and training centres appear which help in increasing the productive efficiency of the firms. Thus these external economies are also the cause of

increasing returns to scale.

Decreasing Returns to Scale:

Figure 24.12 shows the case of decreasing returns where to get equal increases in output, larger proportionate increases in both labour and capital are required.



It follows that:

100 units of output require $2C + 2L$

200 units of output require $5C + 5L$

300 units of output require $9C + 9L$

So that along the expansion path OR , $OG < GH < HK$.

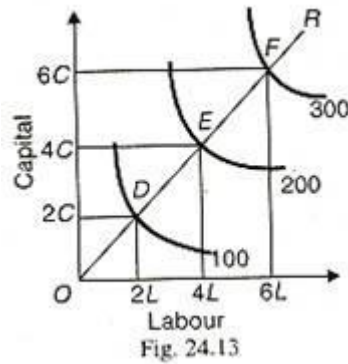
In this case, the production function is homogeneous of degree less than one.

Returns to scale may start diminishing due to the following factors. Indivisible factors may become inefficient and less productive. The firm experiences internal diseconomies. Business may become unwieldy and produce problems of supervision and coordination. Large management creates difficulties of control and rigidities. To these internal diseconomies are added external diseconomies of scale. These arise from higher factor prices or from diminishing productivities of the factors.

As the industry continues to expand the demand for skilled labour, land, capital, etc. rises. There being perfect competition, intensive bidding raises wages, rent and interest. Prices of raw materials also go up. Transport and marketing difficulties emerge. All these factors tend to raise costs and the expansion of the firms leads to diminishing returns to scale so that doubling the scale would not lead to doubling the output.

Constant Returns to Scale:

Figure 24.13 shows the case of constant returns to scale. Where the distance between the isoquants 100, 200 and 300 along the expansion path OR is the same, i.e., $OD = DE = EF$. It means that if units of both factors, labour and capital, are doubled, the output is doubled. To treble output, units of both factors are trebled.



It follows that:

100 units of output require $1(2C + 2L) = 2C + 2L$

200 units of output require $2(2C + 2L) = 4C + 4L$

300 units of output require $3(2C + 2L) = 6C + 6L$

The returns to scale are constant when internal economies enjoyed by a firm are neutralised by internal diseconomies so that output increases in the same proportion. Another reason is the balancing of external economies and external diseconomies. Constant returns to scale also result when factors of production are perfectly divisible, substitutable, homogeneous and their supplies are perfectly elastic at given prices.

That is why, in the case of constant returns to scale, the production function is homogeneous of degree one.

Relation between Returns to Scale and Returns to a Factor (Law of Returns to Scale and Law of Diminishing Returns):

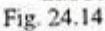
Returns to a factor and returns to scale are two important laws of production. Both laws explain the relation between inputs and output. Both laws have three stages of increasing, decreasing and constant returns. Even then, there are fundamental differences between the two laws.

Returns to a factor relate to the short period production function when one factor is varied keeping the other factor fixed in order to have more output, the marginal returns of the variable factor diminish. On the other hand, returns to scale relate to the long period production function when a firm changes its scale of production by changing one or more of its factors.

We discuss the relation between the returns to a factor (law of diminishing returns) and returns to scale (law of returns to scale) on the assumptions that:

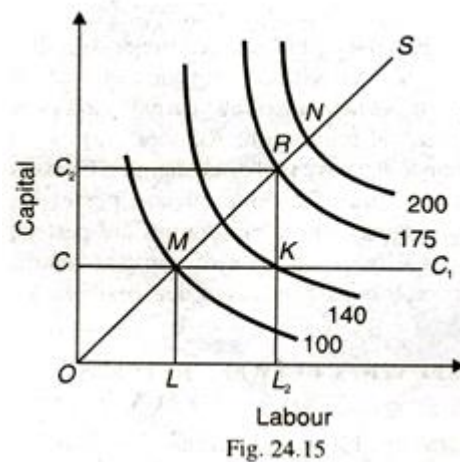
- (1) There are only two factors of production, labour and capital.
- (2) Labour is the variable factor and capital is the fixed factor.
- (3) Both factors are variable in returns to scale.

Given these assumptions, we first explain the relation between constant return to scale and returns to a variable factor in terms of Figure 24.14 where OS is the expansion path which shows constant returns to scale because the difference between the two isoquants 100 and 200 on the expansion path is equal i.e., $OM = MN$. To produce 100 units, the firm uses $OC + OL$ quantities of capital and labour and to double the output to 200 units, double the quantities of labour and capital are required so that $OC1 + OL2$ lead to this output level at point N. Thus there are constant returns to scale because $OM = MN$.



As pointed out by Stonier and Hague, “So, if production function were always homogeneous of the first degree and if returns to scale were always constant, marginal physical productivity (returns) would always fall.”

COPYRIGHT FIMT 2020



Alternatively, if both factors are doubled to $OC_2 + OL_2$ they lead to the lower output level isoquant 175 at point R than the isoquant 200 which shows diminishing returns to scale. If C is kept constant and the amount of variable factor, labour, is doubled by LL_2 we reach point K which lies on a still lower level of output represented by the isoquant 140. This proves that the marginal returns (or physical productivity) of the variable factor, labour, have diminished.

3. Now we take the relation between increasing returns to scale and returns to a variable factor. This is explained in terms of Figure 24.16 (A) and (B). In Panel (A), the expansion path OS depicts increasing returns to scale because the segment $OM > MN$. It means that in order to double the output from 100 to 200, less than double the amounts of both factors will be required. If C is kept constant and the amount of variable factor, labour, is doubled by LL_2 the level of output is reached at point K which shows diminishing marginal returns as represented by the lower isoquant 160 than the isoquant 200 when returns to scale are increasing.

In case the returns to scale are increasing strongly, that is, they are highly positive they will offset the diminishing marginal returns of the variable factor, labour. Such a situation leads to increasing marginal returns. This is explained in Panel (B) of Figure 24.16 where on the expansion path OS, the segment $OM > MN$, thereby showing increasing returns to scale. When the amount of the variable factor, labour, is doubled by LL_2 while keeping C as constant, we reach the output level K represented by the isoquant 250 which is at a higher level than the isoquant 200. This shows that the marginal returns of the variable factor, labour, have increased even when there are increasing returns to scale.

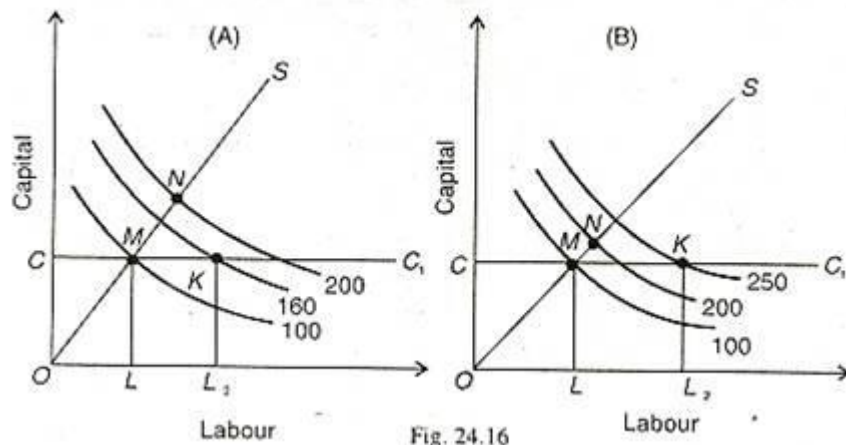


Fig. 24.16

Conclusion:

It can be concluded from the above analysis that under a homogeneous production function when a fixed factor is combined with a variable factor, the marginal returns of the variable factor diminish when there are constant, diminishing and increasing returns to scale. However, if there are strong increasing returns to scale, the marginal returns of the variable factor increase instead of diminishing.

Choice of Optimal Factor Combination or Least Cost Combination of Factors or Producer's Equilibrium:

A profit maximisation firm faces two choices of optimal combination of factors (inputs): First, to minimise its cost for a given output; and second, to maximise its output for a given cost. Thus the least cost combination of factors refers to a firm producing the largest volume of output from a given cost and producing a given level of output with the minimum cost when the factors are combined in an optimum manner. We study these cases separately.

Cost-Minimisation for a Given Output:

In the theory of production, the profit maximisation firm is in equilibrium when, given the cost-price function, it maximises its profits on the basis of the least cost combination of factors. For this, it will choose that combination which minimises its cost of production for a given output. This will be the optimal combination for it.

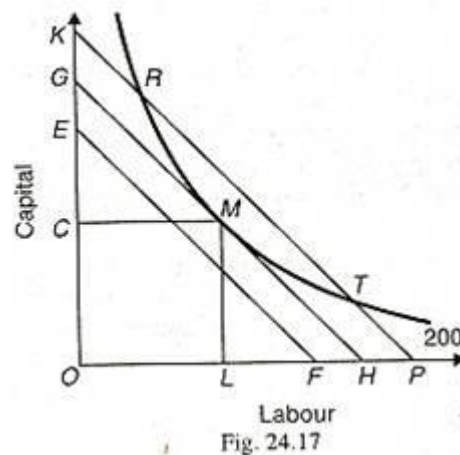
Assumptions:

This analysis is based on the following assumptions:

1. There are two factors, labour and capital.
2. All units of labour and capital are homogeneous.
3. The prices of units of labour (w) and that of capital (r) are given and constant.

4. The cost outlay is given.
5. The firm produces a single product.
6. The price of the product is given and constant.
7. The firm aims at profit maximisation.
8. There is perfect competition in the factor market.

Given these assumptions, the point of least-cost combination of factors for a given level of output is where the isoquant curve is tangent to an isocost line. In Figure 24.17, the isocost line GH is tangent to the isoquant 200 at point M. The firm employs the combination of OC of capital and OL of labour to produce 200 units of output at point M with the given cost- outlay GH. At this point, the firm is minimising its cost for producing 200 units. Any other combination on the isoquant 200, such as R or T, is on the higher isocost line KP which shows higher cost of production. The isocost line EF shows lower cost but output 200 cannot be attained with it. Therefore, the firm will choose the minimum cost point M which is the least-cost factor combination for producing 200 units of output. M is thus the optimal combination for the firm.

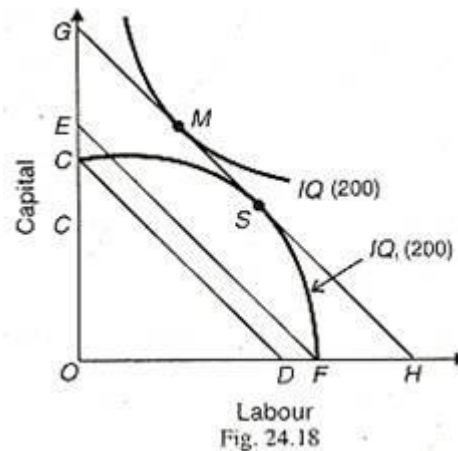


The point of tangency between the isocost line and the isoquant is an important first order condition but not a necessary condition for the producer's equilibrium. There are two essential or second order conditions for the equilibrium of the firm.

1. The first condition is that the slope of the isocost line must equal the slope of the isoquant curve. The Slope of the isocost line is equal to the ratio of the price of labour (w) to the price of capital (r) i.e., w/r . The slope of the isoquant curve is equal to the marginal rate of technical substitution of labour and capital (MRTSLC) which is, in turn, equal to the ratio of the marginal product of labour to the marginal product of capital (MPL/MPC). Thus the equilibrium condition for optimality can be written as:

The second condition is that at the point of tangency, the isoquant curve must be convex to the origin. In

other words, the marginal rate of technical substitution of labour for capital (MRTSLC) must be diminishing at the point of tangency for equilibrium to be stable. In Figure 24.18, S cannot be the point of equilibrium, for the isoquant IQ_1 , is concave where it is tangent to the isocost line GH. At point S, the marginal rate of technical substitution between the two factors increases if move to the right or left on the curve IQ_1 .



Moreover, the same output level can be produced at a lower cost CD or EF and there will be a corner solution either at C or F. If it decides to produce at EF cost, it can produce the entire output with only OF labour. If, on the other hand, it decides to produce at a still lower cost CD, the entire output can be produced with only OC capital. Both the situations are impossibilities because nothing can be produced either with only labour or only capital. Therefore, the firm can produce the same level of output at point M where the isoquant curve IQ is convex to the origin and is tangent to the isocost line GH. The analysis assumes that both the isoquants represent equal level of output, $IQ = IQ_1$.

Output-Maximisation for a Given Cost:

The firm also maximises its profits by maximising its output, given its cost outlay and the prices of the two factors. This analysis is based on the same assumptions, as given above. The conditions for the equilibrium of the firm are the same, as discussed above.

1. The firm is in equilibrium at point P where the isoquant curve 200 is tangent to the isocost line CL. At this point, the firm is maximising its output level of 200 units by employing the optimal combination of OM of capital and ON of labour, given its cost outlay CL. But it cannot be at points E or F on the isocost line CL, since both points give a smaller quantity of output, being on the isoquant 100, than on the isoquant 200. The firm can reach the optimal factor combination level of maximum output by moving along the isocost line CL from either point E or F to point P. This movement involves no extra cost

because the firm remains on the same isocost line. The firm cannot attain a higher level of output such as isoquant 300 because of the cost constraint.

Thus the equilibrium point has to be P with optimal factor combination OM + ON. At point P, the slope of the isoquant curve 200 is equal to the slope of the isocost line CL. It implies that $w/r = MPL/MPN = MRTSLC$

2. The second condition is that the isoquant curve must be convex to the origin at the point of tangency with the isocost line, as explained above in terms of Figure 24.18.

Unit -4

Concept of Cost

The term 'cost' means the amount of expenses [actual or national] incurred on or attributable to specified thing or activity. A producer requires various factors of production or inputs for producing commodity. He pays them in a form of money. Such money expenses incurred by a firm in the production of a commodity are called cost of production.

As per Institute of cost and work accounts (ICWA) India, Cost is 'measurement in monetary terms of the amount of resources used for the purpose of production of goods or rendering services. To get the results we make efforts. Efforts constitute cost of getting the results. It can be expressed in terms of money; it means the amount of expenses incurred on or attributable to some specific thing or activity.

Short run cost & long run cost

Short-run cost

Short run cost varies with output, when unlike long run cost all the factors are not variable. This cost becomes relevant, when a firm has to decide whether or not to produce more in the immediate future. This cost can be divided into two components of fixed and variable cost on the basis of variability of factors of production.

1. **Fixed cost:** In the short period the expenses incurred on fixed factors are called the fixed cost. These costs don't change with changes in level of output.

“The fixed cost is those cost that don't vary with the size of output.”

2. **Variable cost:** VC are those costs which are incurred on the use of variable factors of production. They directly change with production. The rate of increase of total variable cost is determined by the law of returns.

3. **Total cost:** TC of a firm for various levels of output are the sum of total fixed cost and total variable cost.

4. **Average cost:** per unit cost of a good is called its average cost. Average cost is total cost divided by output.

$$AC = TC/Q$$

$$AC = AFC + AVC$$

a) **Average fixed cost:** AFC is total fixed cost /total output. AFC is the per unit cost of the fixed factor of production.

b) **Average variable cost:** AVC is found by dividing the total variable cost by the total unit of output.

5. **Marginal cost:** MC is the addition made to the total cost by the production of one more unit of a commodity.

$$MC = TC_n - (TC_{n-1})$$

$$MC = \frac{\Delta TC}{\Delta Q}$$

Long-run cost

In the long run, all factors of production are variable. Hence there is no distinction between fixed and variable cost. All cost are variable cost and there is nothing like fixed cost.

a) Long run average cost

b) Long run marginal cost

a) **Long run avg. cost:** LRAC refers to minimum possible per unit cost of producing different quantities of output of a good in the long period.

b) **Long run marginal cost:** change in the total cost in the long run, due to the production of one more unit, is called LRMC.

Economies and Diseconomies of Scale

- Economies of scale are the **cost advantages** that a business can exploit by **expanding the scale of production**
- The effect is to **reduce the long run average (unit) costs of production**.
- These lower costs are an improvement in **productive efficiency** and can benefit consumers in the form of lower prices. But they give a business a competitive advantage tool.

Internal Economies of Scale.

Internal economies of scale arise from **the growth of the business itself**. Examples include:

1. Technical economies of scale:

a. Large-scale businesses can afford to invest in **expensive and specialist capital machinery**. For example, a supermarket chain such as Tesco or Sainsbury can invest in technology that improves stock control. It might not, however, be viable or cost-efficient for a small corner shop to buy this technology.

b. **Specialization of the workforce:** Larger businesses split complex production processes into separate tasks to boost productivity. The **division of labour** in mass production of motor vehicles and in manufacturing electronic products is an example.

c. **The law of increased dimensions.** This is linked to the **cubic law** where doubling the height and width of a tanker or building leads to a more than proportionate increase in the cubic capacity – this is an

important scale economy in distribution and transport industries and also in travel and leisure sectors.

2. Marketing economies of scale and monopsony power: A large firm can spread its advertising and marketing budget over a large output and it can purchase its inputs in bulk at negotiated discounted prices if it has **monopsony (buying) power** in the market.

A good example would be the ability of the electricity generators to negotiate lower prices when negotiating coal and gas supply contracts. The big food retailers have monopsony power when purchasing supplies from farmers.

3. Managerial economies of scale: This is a form of division of labour. Large-scale manufacturers employ specialists to supervise production systems and oversee human resources.

4. Financial economies of scale: Larger firms are usually rated by the financial markets to be more 'credit worthy' and have access to credit facilities, with favorable rates of borrowing. In contrast, smaller firms often face higher rates of interest on overdrafts and loans. Businesses quoted on the stock market can normally raise fresh money (i.e. extra financial capital) more cheaply through the issue of equities. They are also likely to pay a lower rate of interest on new company bonds issued through the capital markets.

5. Network economies of scale: *This is a demand-side economy of scale.* Some networks and services have huge potential for economies of scale. That is, as they are more widely used they become more valuable to the business that provides them. The classic examples are the expansion of a common language and a common currency. We can identify networks economies in areas such as **online auctions, air transport networks**. Network economies are best explained by saying that the **marginal cost of adding one more user** to the network is close to zero, but the resulting benefits may be huge because each new user to the network can then interact, trade with **all** of the existing members or parts of the network. The expansion of **e-commerce** is a great example of network economies of scale

External economies of scale

- External economies of scale occur **within an industry** and from the expansion of it

- Examples include the development of **research and development facilities in local universities** that several businesses in an area can benefit from and spending by a local authority on improving the transport network for a local town or city.
- Likewise, the **relocation of component suppliers** and other support businesses close to the main centre of manufacturing are also an external cost saving.

Diseconomies of scale

A firm may eventually experience a rise in average costs caused by diseconomies of scale.

Diseconomies of scale a firm might be caused by:

1. **Control** – monitoring the productivity and the quality of output from thousands of employees in big corporations is imperfect and costly.
2. **Co-operation** - workers in large firms may feel a sense of alienation and subsequent loss of morale. If they do not consider themselves to be an integral part of the business, their productivity may fall leading to wastage of factor inputs and higher costs. A fall in productivity means that workers may be less productively efficient in larger firms.
3. **Loss of control over costs** – big businesses may lose control over fixed costs such as expensive head offices, management expenses and marketing costs. There is also a risk that very expensive capital projects involving new technology may prove ineffective and leave the business with too much under-utilized capital.

Evaluation: Do economies of scale always improve the welfare of consumers?

- **Standardization of products:** Mass production might lead to a **standardization of products** – limiting the amount of consumer choice.
- **Lack of market demand:** Market demand may be insufficient for economies of scale to be fully exploited leaving businesses with a lot of spare capacity.

- **Developing monopoly power:** Businesses may use economies of scale to build up monopoly power and this might lead to higher prices, a reduction in consumer welfare and a loss of allocative efficiency.
- **Protecting monopoly power:** Economies of scale might be used as a **barrier to entry** –whereby existing firms can drive prices down if there is a threat of the entry of new suppliers.

Explicit Cost and Implicit Cost

Explicit cost:

All those expenses that a firm incurs to make payment to others are called explicit cost. An explicit cost is a direct payment made to others in the course of running a business, such as wage, rent and materials. Explicit costs are taken into account along with implicit ones when considering economic profit. Accounting profit only takes explicit costs into account.

Implicit cost:

Implicit cost is the cost of entrepreneur's own factors or resources. These includes the rewards for the entrepreneurs self owned land, building, labour & capital. In economics, an **implicit cost**, also called an **imputed cost**, **implied cost**.

Implicit costs also represent the divergence between economic profit (total revenues minus total costs, where total costs are the sum of implicit and explicit costs) and accounting profit (total revenues minus only explicit costs). Since economic profit includes these extra opportunity costs, it will always be less than or equal to accounting profit.

Private and Social Cost

Private cost refers to the cost of production incurred & provided for by an individual firm engaged in the production of a commodity. It is found out to get private profits. It includes both explicit as well as implicit cost. A firm is interested in minimizing private cost. Social cost refers to the cost of producing a commodity to the society as a whole. It takes into consideration of all those costs which were borne by the society directly or indirectly. It is a sum of private cost & external cost. for example, from pollution of the atmosphere.

SOCIAL COST = PRIVATE COST + EXTERNALITY

For example: - a chemical factory emits wastage as a by-product into nearby rivers and into the atmosphere. This creates negative externalities which impose higher social costs on other firms and consumers. e.g. clean up costs and health costs.

Another example of higher social costs comes from the problems caused by traffic congestion in towns, cities and on major roads and motor ways. It is important to note though that the manufacture, purchase and use of private cars can also generate external benefits to society. This why **cost-benefit analysis** can be useful in measuring and putting some monetary value on both the social costs and benefits of production.

Key Concepts:

Private Costs + External Costs = Social Costs

If external costs > 0, then private costs < social costs.

Then society tends to:

- Price the good or service too low and Produces or consumes too much of the good or service.

Different Costs Matter:

Private costs for a producer of a good, service, or activity include the costs the firm pays to purchase capital equipment, hire labor, and buy materials or other inputs. While this is straightforward from the business side, it also is important to look at this issue from the consumers' perspective. Field, in his 1997 text, *Environmental Economics* provides an example of the private costs a consumer faces when driving a car:

The private costs of this (driving a car) include the fuel and oil, maintenance, depreciation, and even the drive time experienced by the operator of the car.

Private costs are paid by the firm or consumer and must be included in production and consumption decisions. In a competitive market, considering only the private costs will lead to a socially efficient rate of output only if there are no external costs.

External costs, on the other hand, are not reflected on firms' income statements or in consumers' decisions. However, external costs remain costs to society, regardless of who pays for them. Consider a firm that attempts to save money by not installing water pollution control equipment. Because of the

firm's actions, cities located down river will have to pay to clean the water before it is fit for drinking, the public may find that recreational use of the river is restricted, and the fishing industry may be harmed. When external costs like these exist, they must be added to private costs to determine social costs and to ensure that a socially efficient rate of output is generated.

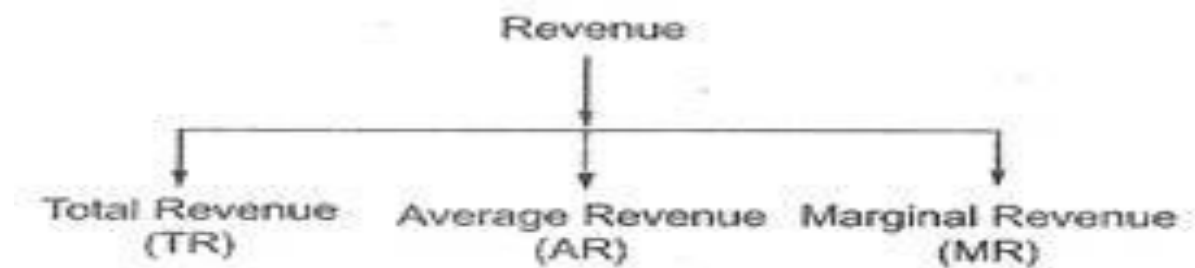
Social costs include both the private costs and any other external costs to society arising from the production or consumption of a good or service. Social costs will differ from private costs, for example, if a producer can avoid the cost of air pollution control equipment allowing the firm's production to impose costs (health or environmental degradation) on other parties that are adversely affected by the air pollution. Remember too, it is not just producers that may impose external costs on society. Let's also view how consumers' actions also may have external costs using Field's previous example on driving: The social costs include all these private costs (fuel, oil, maintenance, insurance, depreciation, and operator's driving time) and also the cost experienced by people other than the operator who are exposed to the congestion and air pollution resulting from the use of the car. The key point is that even if a firm or individual avoids paying for the external costs arising from their actions, the costs to society as a whole (congestion, pollution, environmental cleanup, visual degradation, wildlife impacts, etc.) remain. Those external costs must be included in the social costs to ensure that society operates at a socially efficient rate of output.

Revenue Concept

Revenue refers to the amount received by a firm from the sale of a given quantity of a commodity in the market.

Revenue is a very important concept in economic analysis. It is directly influenced by sales level, i.e., as sales increases, revenue also increases.

The concept of revenue consists of three important terms; Total Revenue, Average Revenue and Marginal Revenue



Total Revenue (TR):

Total Revenue refers to total receipts from the sale of a given quantity of a commodity. It is the total income of a firm. Total revenue is obtained by multiplying the quantity of the commodity sold with the price of the commodity.

Total Revenue = Quantity × Price

For example, if a firm sells 10 chairs at a price of Rs. 160 per chair, then the total revenue will be: 10 Chairs × Rs. 160 = Rs 1,600

Average Revenue (AR):

Average revenue refers to revenue per unit of output sold. It is obtained by dividing the total revenue by the number of units sold.

Average Revenue = Total Revenue/Quantity

For example, if total revenue from the sale of 10 chairs @ Rs. 160 per chair is Rs. 1,600, then:

$$\text{Average Revenue} = \text{Total Revenue/Quantity} = 1,600/10 = \text{Rs } 160$$

AR and Price are the Same:

AR is equal to per unit sale receipts and price is always per unit. Since sellers receive revenue according to price, price and AR are one and the same thing.

This can be explained as under:

$$\text{TR} = \text{Quantity} \times \text{Price} \dots (1)$$

$$\text{AR} = \text{TR/Quantity} \dots\dots (2)$$

Putting the value of TR from equation (1) in equation (2), we get

$$\text{AR} = \text{Quantity} \times \text{Price} / \text{Quantity}$$

$$\text{AR} = \text{Price}$$

AR Curve and Demand Curve are the Same:

A buyer's demand curve graphically represents the quantities demanded by a buyer at various prices. In other words, it shows the various levels of average revenue at which different quantities of the good are sold by the seller. Therefore, in economics, it is customary to refer AR curve as the Demand Curve of a firm.

Marginal Revenue (MR):

Marginal revenue is the additional revenue generated from the sale of an additional unit of output. It is the change in TR from sale of one more unit of a commodity.

$$\text{MR}_n = \text{TR}_n - \text{TR}_{n-1}$$

Where:

MR_n = Marginal revenue of nth unit;

TR_n = Total revenue from n units;

TR_{n-1} = Total revenue from (n – 1) units; n = number of units sold For example, if the total revenue realised from sale of 10 chairs is Rs. 1,600 and that from sale of 11 chairs is Rs. 1,780, then MR of the 11th chair will be:

$$MR_{11} = TR_{11} - TR_{10}$$

$$MR_{11} = \text{Rs. } 1,780 - \text{Rs. } 1,600 = \text{Rs. } 180$$

One More way to Calculate MR:

We know, MR is the change in TR when one more unit is sold. However, when change in units sold is more than one, then MR can also be calculated as:

$$MR = \text{Change in Total Revenue} / \text{Change in number of units} = \Delta TR / \Delta Q$$

Let us understand this with the help of an example: If the total revenue realised from sale of 10 chairs is Rs. 1,600 and that from sale of 14 chairs is Rs. 2,200, then the marginal revenue will be:

$$MR = TR \text{ of 14 chairs} - TR \text{ of 10 chairs} / 14 \text{ chairs} - 10 \text{ chairs} = 600/4 = \text{Rs. } 150$$

TR is summation of MR:

Total Revenue can also be calculated as the sum of marginal revenues of all the units sold.

$$\text{It means, } TR_n = MR_1 + MR_2 + MR_3 + \dots \dots \dots MR_n$$

$$\text{or, } TR = \sum MR$$

The concepts of TR, AR and MR can be better explained through Table 7.1.

Table 7.1: TR, AR and MR:

Units Sold (Q)	Price (Rs.) (P)	Total Revenue (Rs.) $TR = Q \times P$	Average Revenue (Rs.) AR = $TR/Q = P$	Marginal Revenue (Rs.) $MR_n = TR_n - TR_{n-1}$
1	10	$10 = 1 \times 10$	$10 = 10 / 1$	$10 = 10 - 0$
2	9	$18 = 2 \times 9$	$9 = 18 / 2$	$8 = 18 - 10$
3	8	$24 = 3 \times 8$	$8 = 24 / 3$	$6 = 24 - 18$
4	7	$28 = 4 \times 7$	$7 = 28 / 4$	$4 = 28 - 24$
5	6	$30 = 5 \times 6$	$6 = 30 / 5$	$2 = 30 - 28$
6	5	$30 = 6 \times 5$	$5 = 30 / 6$	$0 = 30 - 30$
7	4	$28 = 7 \times 4$	$4 = 28 / 7$	TRUE

STATISTICAL METHOD-I (103)

Define ‘Statistics’ and give characteristics of ‘Statistics’.

„Statistics“ means numerical presentation of facts. Its meaning is divided into two forms - in plural form and in singular form.

Statistics“ means a collection of numerical facts or data example price statistics, agricultural statistics, production statistics, etc. In singular form, the word means the statistical methods with the help of which collection, analysis and interpretation of data are accomplished.

Characteristics of Statistics - a) Aggregate of facts/data

- b) Numerically expressed
- c) Affected by different factors
- d) Collected or estimated
- e) Reasonable standard of accuracy
- f) Predetermined purpose
- g) Comparable
- h) Systematic collection.

Therefore, the process of collecting, classifying, presenting, analyzing and interpreting the numerical facts, comparable for some predetermined purpose are collectively known as “Statistics”.

What is meant by ‘Data’?

Data refers to any group of measurements that happen to interest us. These measurements provide information the decision maker uses. Data are the foundation of any statistical investigation and the job of collecting data is the same for a statistician as collecting stone, mortar, cement, bricks etc. is for a builder.

Discuss the Scope of Statistics.

Ans.: The scope of statistics is much extensive. It can be divided into two parts – **(i)** Statistical Methods such as Collection, Classification, Tabulation, Presentation, Analysis, Interpretation and Forecasting.

(ii) Applied Statistics – It is further divided into three parts:

- a) Descriptive Applied Statistics: Purpose of this analysis is to provide descriptive information.
- b) Scientific Applied Statistics: Data are collected with the purpose of some scientific research and with the help of these data some particular theory or principle is propounded.

c) Business Applied Statistics: Under this branch statistical methods are used for the study, analysis and solution of various problems in the field of business.

State the limitation of statistics?

Scope of statistics is very wide. In any area where problems can be expressed in qualitative form, statistical methods can be used. But statistics have some limitations

1. Statistics can study only numerical or quantitative aspects of a problem.
2. Statistics deals with aggregates not with individuals.
3. Statistical results are true only on an average.
4. Statistical laws are not exact.
5. Statistics does not reveal the entire story.
6. Statistical relations do not necessarily bring out the cause and effect relationship between phenomena.
7. Statistics is collected with a given purpose.
8. Statistics can be used only by experts.

What do you mean by Collection of Data? Differentiate between Primary and Secondary Data.

Collection of data is the basic activity of statistical science. It means collection of facts and figures relating to particular phenomenon under the study of any problem whether it is in business economics, social or natural sciences. Such material can be obtained directly from the individual units, called **primary sources** or from the material published earlier elsewhere known as the **secondary sources**.
Difference between Primary & Secondary Data. Primary Data Secondary Data **Basis nature** Primary data are original and are collected for the first time. SECONDARY Data which are collected earlier by someone else, and which are now in published or unpublished state. Collecting Agency these data are collected by the investigator himself. Secondary data were collected earlier by some other person. Post collection alterations these data do not need alteration as they are according to the requirement of the investigation. These have to be analyzed and necessary changes have to be made to make them useful as per the requirements of investing.

What is the meaning of Classification? Give objectives of Classification and essentials of an ideal classification.

Classification is the process of arranging data into various groups, classes and subclasses according to

some common characteristics of separating them into different but related parts. Main objectives of Classification: -

- (i) To make the data easy and precise
- (ii) To facilitate comparison
- (iii) Classified facts expose the cause-effect relationship.
- (iv) To arrange the data in proper and systematic way
- (v) The data can be presented in a proper tabular form only.

Essentials of an Ideal Classification :- (i) Classification should be so exhaustive and complete that every individual unit is included in one or the other class.

(ii) Classification should be suitable according to the objectives of investigation.

(iii) There should be stability in the basis of classification so that comparison can be made.

The two values which determine a class are known as class limits. First or the smaller one is known as lower limit (L1) and the greater one is known as upper limit (L2)

How many types of Series are there on the basis of Quantitative Classification? Give the difference between Exclusive and Inclusive Series.

(i) **Individual Series:** In individual series, the frequency of each item or value is only one for example; marks scored by 10 students of a class are written individually.

(ii)

Discrete Series: A discrete series is that in which the individual values are different from each other by a different amount. For example: Daily wages 5 10 15 20 No. of workers 6 9 8 5

(iii) **Continuous Series:** When the number of items is placed within the limits of the class, the series obtained by classification of such data is known as continuous series.

Exclusive Series Inclusive Series LUpper limit of one class is equal to the lower limit of next class. The two limits are not equal. Inclusion The value equal to the upper limit is included in the next class. Both upper & lower limits are included in the same class. Conversion It does not require any conversion. Inclusive series is converted into exclusive series for calculation purpose. Statistical Methods 21 Suitability It is suitable in all situations. It is suitable only when the values are in integers.

Depicting of statistical data in the form of attractive shapes such as bars, circles, and rectangles is called diagrammatic presentation. A diagram is a visual form of presentation of statistical data, highlighting their basic facts and relationship. There are geometrical figures like lines, bars, squares, rectangles, circles, curves, etc. Diagrams are used with great effectiveness in the presentation of all types of data. When properly constructed, they readily show information that might otherwise be lost amid the details of numerical tabulation.

Importance of Diagrams: A properly constructed diagram appeals to the eye as well as the mind since it is practical, clear and easily understandable even by those who are unacquainted with the methods of presentation. Utility or importance of diagrams will become clearer from the following points –

- (i) **Attractive and Effective Means of Presentation:** Beautiful lines; full of various colours and signs attract human sight, and do not strain the mind of the observer. A common man who does not wish to indulge in figures, get message from a well prepared diagram.
- (ii) **Make Data Simple and Understandable:** The mass of complex data, when prepared through diagram, can be understood easily. According to Shri Moraine, “Diagrams help us to understand the complete meaning of a complex numerical situation at one sight only”. Statistical Methods 25.
- (iii) **Facilitate Comparison:** Diagrams make comparison possible between two sets of data of different periods, regions or other facts by putting side by side through diagrammatic presentation.
- (iv) **Save Time and Energy:** The data which will take hours to understand becomes clear by just having a look at total facts represented through diagrams.

- (v) **Universal Utility:** Because of its merits, the diagrams are used for presentation of statistical data in different areas. It is widely used technique in economic, business, administration, social and other areas.
- (vi) **Helpful in Information Communication:** A diagram depicts more information than the data shown in a table. Information concerning data to general public becomes more easy through diagrams and gets into the mind of a person with ordinary knowledge.

What are the various types of graphs of frequency distribution?

Ans.: Frequency distribution can also be presented by means of graphs. Such graphs facilitate comparative study of two or more frequency distributions as regards their shape and pattern. The most commonly used graphs are as follows –

- (i) Line frequency diagram
 - (ii) Histogram
 - (iii) Frequency Polygon
 - (iv) Frequency curves
 - (v) cumulative frequency curves or Ogive curves
- Line Frequency Diagram :** This diagram is mostly used to depict discrete series on a graph. The values are shown on the X-axis and the frequencies on the Y axis. The lines are drawn vertically on X-axis against the relevant values taking the height equal to respective frequencies. Statistical Methods 31

Histogram : It is generally used for presenting continuous series. Class intervals are shown on X-axis and the frequencies on Y-axis. The data are plotted as a series of rectangles one over the other. The height of rectangle represents the frequency of that group. Each rectangle is joined with the other so as to give a continuous picture. Histogram is a graphic method of locating mode in continuous series. The rectangle of the highest frequency is treated as the rectangle in which mode lies. The top corner of this rectangle and the adjacent rectangles on both sides are joined diagonally. The point where two lines interact each other a perpendicular line is drawn on OX-axis. The point where the perpendicular line meets OX-axis is the value of mode.

Frequency Polygon : Frequency polygon is a graphical presentation of both discrete and continuous

series. For a discrete frequency distribution, frequency polygon is obtained by plotting frequencies on Y-axis against the corresponding size of the variables on X-axis and then joining all the points by a straight line. In continuous series the mid-points of the top of each rectangle of histogram is joined by a straight line. To make the area of the frequency polygon equal to histogram, the line so drawn is stretched to meet the base line (X-axis) on both sides

Frequency Curve : The curve derived by making smooth frequency polygon is called frequency curve. It is constructed by making smooth the lines of frequency polygon. This curve is drawn with a free hand so that its angularity disappears and the area of frequency curve remains equal to that of frequency polygon.

Cumulative Frequency Curve or Ogive Curve : This curve is a graphic presentation of the cumulative frequency distribution of continuous series. It can be of two types - (a) Less than Ogive and (b) More than Ogive.

Less than Ogive : This curve is obtained by plotting less than cumulative frequencies against the upper class limits of the respective classes. The points so obtained are joined by a straight line. It is an increasing curve sloping upward from left to right. **More than Ogive :** It is obtained by plotting „more than“ cumulative frequencies against the lower class limits of the respective classes. The points so obtained are joined by a straight line to give „more than ogive“. It is a decreasing curve which slopes downwards from left to right.

Measures of Central Tendency: Mean, Median, and Mode

A measure of central tendency is a summary statistic that represents the center point or typical value of a dataset. These measures indicate where most values in a distribution fall and are also referred to as the central location of a distribution. You can think of it as the tendency of data to cluster around a middle value. In statistics, the three most common measures of central tendency are the mean, median, and mode. Each of these measures calculates the location of the central point using a different method.

Mean

The mean is the arithmetic average, and it is probably the measure of central tendency that you are most familiar. Calculating the mean is very simple. You just add up all of the values and divide by the number of observations in your dataset.

$$\frac{x_1 + x_2 + \dots + x_n}{n}$$

Arithmetic mean is a mathematical average and it is the most popular measures of central tendency. It is frequently referred to as ‘mean’ it is obtained by dividing sum of the values of all observations in a series (ΣX) by the number of items (N) constituting the series.

Thus, mean of a set of numbers $X_1, X_2, X_3, \dots, X_n$ denoted by $\bar{x}$ and is defined as

$$\text{Mean} = \frac{\text{Sum of the items}}{\text{Number of the items}} = \frac{\Sigma X}{N}$$

Example : Calculated the Arithmetic Mean DIRC Monthly Users Statistics in the University Library

Month	No. of Working Days	Total Users	Average Users per month
Sep-2011	24	11618	484.08
Oct-2011	21	8857	421.76
Nov-2011	23	11459	498.22
Dec-2011	25	8841	353.64
Jan-2012	24	5478	228.25
Feb-2012	23	10811	470.04
Total	140	57064	

$$\text{Mean} = \frac{\text{Total number of users}}{\text{Total number of working days}}$$

$$= \frac{\sum X}{N} = \frac{57064}{140} = 407.6$$

Advantages of Mean

- It is easy to understand & simple calculate.
- It is based on all the values.
- It is rigidly defined.
- It is easy to understand the arithmetic average even if some of the details of the data are lacking.
- It is not based on the position in the series.

Median

The median is the middle value. It is the value that splits the dataset in half. To find the median, order your data from smallest to largest, and then find the data point that has an equal amount of values above it and below it. The method for locating the median varies slightly depending on whether your dataset has an even or odd number of values. I'll show you how to find the median for both cases. In the examples below, I use whole numbers for simplicity, but you can have decimal places.

Calculation of Median –Discrete series:

- Arrange the data in ascending or descending order.
- Calculate the cumulative frequencies.
- Apply the formula.

$$\text{Median}(M) = \text{Size of } \left(\frac{N + 1}{2} \right) \text{th item}$$

Calculation of median – Continuous series

For calculation of median in a continuous frequency distribution the following formula will be employed. Algebraically,

$$\text{Median}(M) = L1 + \frac{\frac{N}{2} - cf}{f} \times i$$

Example: Median of a set Grouped Data in a Distribution of Respondents by age

Age Group	Frequency of Median class(f)	Cumulative frequencies(cf)
0-20	15	15
20-40	32	47
40-60	54	101
60-80	30	131
80-100	19	150
Total	150	

$$\begin{aligned}\text{Median (M)} &= 40 + \frac{\frac{150}{2} - 47}{54} \times 20 \\ &= 40 + \frac{75 - 47}{54} \times 20 \\ &= 40 + \frac{28}{54} \times 20 \\ &= 40 + 0.52 \times 20 \\ &= 40 + 10.37 \\ &= \mathbf{50.37}\end{aligned}$$

Advantages of Median:

- Median can be calculated in all distributions.
- Median can be understood even by common people.
- Median can be ascertained even with the extreme items.
- It can be located graphically
- It is most useful dealing with qualitative data

Disadvantages of Median:

- It is not based on all the values.
- It is not capable of further mathematical treatment.
- It is affected fluctuation of sampling.
- In case of even no. of values it may not the value from the data.

Mode

The mode is the value that occurs the most frequently in your data set. On a bar chart, the mode is the highest bar. If the data have multiple values that are tied for occurring the most frequently, you have a multimodal distribution. If no value repeats, the data do not have a mode.

➤ Mode is the most frequent value or score

in the distribution.

➤ It is defined as that value of the item in

a series.

➤ It is denoted by the capital letter Z.

➤ highest point of the frequencies

Monthly rent (Rs)	Number of Libraries (f)
500-1000	5
1000-1500	10
1500-2000	8
2000-2500	16
2500-3000	14
3000 & Above	12
Total	65

$$Z = 2000 + \frac{16 - 8}{2(16) - 8 - 14} \times 500$$

$$Z = 2000 + \frac{8}{32 - 8 - 14} \times 500$$

$$Z = 2000 + \frac{8}{10} \times 500$$

$$Z = 2000 + 0.8 \times 500 = 2400$$

$$Z = 2400$$

Advantages of Mode:

- Mode is readily comprehensible and easily calculated
- It is the best representative of data
- It is not at all affected by extreme value.
- The value of mode can also be determined graphically.
- It is usually an actual value of an important part of the series.

Disadvantages of Mode:

- It is not based on all observations.
- It is not capable of further mathematical manipulation.
- Mode is affected to a great extent by sampling fluctuations.
- Choice of grouping has great influence on the value of mode.

Conclusion

• A measure of central tendency is a measure that tells us where the middle of a bunch of data lies. Mean is the most common measure of central tendency. It is simply the sum of the numbers divided by the number of numbers in a set of data. This is also known as average.

- Median is the number present in the middle when the numbers in a set of data are arranged in ascending or descending order. If the number of numbers in a data set is even, then the median is the mean of the two middle numbers.
- Mode is the value that occurs most frequently in a set of data.

Which is Best—the Mean, Median, or Mode?

When you have a symmetrical distribution for continuous data, the mean, median, and mode are equal. In this case, analysts tend to use the mean because it includes all of the data in the calculations. However, if you have a skewed distribution, the median is often the best measure of central tendency.

What is a Quartile?



Whenever we have an observation and we wish to divide it, there is a chance to do it in different ways. So, we use the **median** when a given observation is divided into two parts that are equal. Likewise, [quartiles](#) are values that divide a complete given set of observations into four equal parts.

Basically, there are three types of quartiles, first quartile, second quartile, and third quartile. The other name for the first quartile is lower quartile. The representation of the first quartile is ‘Q₁.’ The other name for the second quartile is median. The representation of the second quartile is by ‘Q₂.’ The other name for the third quartile is the upper quartile. The representation of the third quartile is by ‘Q₃.’

- **First Quartile** is generally the one-fourth of any sort of observation. However, the point to note here is, this one-fourth value is always less than or equal to ‘Q₁.’ Similarly, it goes for the values of ‘Q₂’ and ‘Q₃.’

What are Deciles?

Deciles are those values that divide any set of a given observation into a total of ten equal parts. Therefore, there are a total of nine deciles. These representation of these deciles are as follows – D₁, D₂, D₃, D₄, D₉.

D_1 is the typical peak value for which one-tenth (1/10) of any given observation is either less or equal to D_1 . However, the remaining nine-tenths (9/10) of the same observation is either greater than or equal to the value of D_1 .

What do you mean by Percentiles?

Last but not the least, comes the percentiles. The other name for percentiles is *centiles*. A centile or a percentile basically divide any given observation into a total of 100 equal parts. The representation of these percentiles or centiles is given as – $P_1, P_2, P_3, P_4, \dots, P_{99}$.

P_1 is the typical peak value for which one-hundredth (1/100) of any given observation is either less or equal to P_1 . However, the remaining ninety-nine-hundredth (99/100) of the same observation is either greater than or equal to the value of P_1 . This takes place once all the given observations are arranged in a specific manner i.e. ascending order.

So, in case the data we have doesn't have a proper classification, then the representation of p^{th} quartile is $(n + 1)p^{\text{th}}$

Here,

n = total number of observations.

$p = 1/4, 2/4, 3/4$ for different values of Q_1, Q_2 , and Q_3 respectively.

$p = 1/10, 2/10, \dots, 9/10$ for different values of $D_1, D_2, \dots, D_9$ respectively.

$p = 1/100, 2/100, \dots, 99/100$ for different values of $P_1, P_2, \dots, P_{99}$ respectively.

Formula

At times, the grouping of frequency distribution takes place. For which, we use the following formula during the computation:

$$Q = l_1 + [(N_p - N_i)/(N_u - N_i)] * C$$

Here,

l_1 = lower class boundary of the specific class that contains the median.

N_i = less than the cumulative frequency in correspondence to l_1 (Post Median Class)

N_u = less than the cumulative frequency in correspondence to l_2 (Pre Median Class)

C = Length of the median class ($l_2 - l_1$)

The symbol 'p' has its usual value. The value of 'p' varies completely depending on the type of quartile. There are different ways to find values or quartiles. We use this way in a grouped frequency distribution. The best way to do it is by drawing an *ogive* for the present frequency distribution.

Hence, all that we need to do to find one specific quartile is, find the point and draw a horizontal axis through the same. This horizontal line must pass through N_p . The next step is to draw a perpendicular. The perpendicular comes up from the same point of intersection of the ogive and the horizontal line. Hence, the value of the quartile comes from the value of 'x' of the given perpendicular line.

Solved Questions for You!

Question: Here are the wages of some laborers: Rs. 82, Rs. 56, Rs. 120, Rs. 75, Rs. 80, Rs. 75, Rs. 90, Rs. 50, Rs. 130, Rs. 65. Find the values of Q_1 , D_6 , and P_{82} .

Solution: The wages in ascending order – Rs. 50, Rs. 56, Rs. 65, Rs. 75, Rs. 75, Rs. 80, Rs. 90, Rs. 82, Rs. 90, Rs. 120, Rs. 130

So,

$$Q_1 = \frac{n}{4} \text{th value} = \frac{10}{4} \text{th value} = 2.5 \text{th value} = 2 \text{nd value} + 0.5 \times (\text{3rd value} - \text{2nd value}) = \text{Rs. 62.75}$$

$$D_6 = \frac{6n}{10} \text{th value} = \frac{6 \times 10}{10} \text{th value} = 6 \text{th value} + 0.6 \times (\text{7th value} - \text{6th value}) = \text{Rs. 81.20}$$

$$P_{82} = \frac{82n}{100} \text{th value} = \frac{82 \times 10}{100} \text{th value} = 8.2 \text{th value} = 8 \text{th value} + 0.2 \times (\text{9th value} - \text{8th value}) = \text{Rs. 120.20}$$

Measures of Variability/MEASURES OF DISPERSION

A measure of variability is a summary statistic that represents the amount of dispersion in a dataset. How spread out are the values? While a measure of central tendency describes the typical value,

measures of variability define how far away the data points tend to fall from the center. We talk about variability in the context of a distribution of values. A low dispersion indicates that the data points tend to be clustered tightly around the center. High dispersion signifies that they tend to fall further away. In statistics, variability, dispersion, and spread are synonyms that denote the width of the distribution. Just as there are multiple measures of central tendency, there are several measures of variability. In this blog post, you'll learn why understanding the variability of your data is critical. Then, I explore the most common measures of variability—the range, interquartile range, variance, and standard deviation. I'll help you determine which one is best for your data.

INTRODUCTION

The Measures of central tendency gives us a birds eye view of the entire data they are called averages of the first order, it serve to locate the centre of the distribution but they do not reveal how the items are spread out on either side of the central value. The measure of the scattering of items in a distribution about the average is called dispersion

- Dispersion measures the extent to which the items vary from some central value. It may be noted that the measures of dispersion or variation measure only the degree but not the direction of the variation. The measures of dispersion are also called averages of the second order because they are based on the deviations of the different values from the mean or other measures of central tendency which are called averages of the first order.

DEFINITION

- In the words of Bowley “Dispersion is the measure of the variation of the items”

According to Conar “Dispersion is a measure of the extent to which the individual items vary”

METHODS OF MEASURING DISPERSION

- Range
- Quartile Deviation
- Mean Deviation
- Standard Deviation

RANGE

Let's start with the range because it is the most straightforward measure of variability to calculate and the simplest to understand. The range of a dataset is the difference between the largest and smallest values in that dataset.

- It is defined as the difference between the smallest and the largest observations in a given set of data.
- Formula is $R = L - S$
- Ex. Find out the range of the given distribution: 1, 3, 5, 9, 11
- The range is $11 - 1 = 10$.

For example, in the two datasets below, dataset 1 has a range of $20 - 38 = 18$ while dataset 2 has a range of $11 - 52 = 41$. Dataset 2 has a broader range and, hence, more variability than dataset 1.

While the range is easy to understand, it is based on only the two most extreme values in the dataset, which makes it very susceptible to outliers. If one of those numbers is unusually high or low, it affects the entire range even if it is atypical.

Additionally, the size of the dataset affects the range. In general, you are less likely to observe extreme values. However, as you increase the sample size, you have more opportunities to obtain these extreme values. Consequently, when you draw random samples from the same population, the range tends to increase as the sample size increases. Consequently, use the range to compare variability only when the sample sizes are similar.

QUARTILE DEVIATION

- It is the second measure of dispersion, no doubt improved version over the range. It is based on the quartiles so while calculating this may require upper quartile (Q3) and lower quartile (Q1) and then is divided by 2. Hence it is half of the difference between two quartiles it is also a semi inter quartile range.

The formula of Quartile Deviation is

- $(Q D) = \underline{Q3 - Q1}$

2

MEAN DEVIATION

Mean Deviation is also known as average deviation. In this case deviation taken from any average especially Mean, Median or Mode. While taking deviation we have to ignore negative items and consider all of them as positive. The formula is given below

The formula of MD is given below

$$MD = \frac{\sum d}{N}$$

N (deviation taken from mean)

$$MD = \frac{\sum m}{N}$$

N (deviation taken from median)

$$MD = \frac{\sum z}{N}$$

N (deviation taken from mode)

Find the mean of all values ... use it to work out distances ... then find the mean of those distances!

In three steps:

- 1. Find the mean of all values
- 2. Find the **distance** of each value from that mean (subtract the mean from each value, ignore minus signs)
- 3. Then find the **mean of those distances**

Like this:

Example: the Mean Deviation of 3, 6, 6, 7, 8, 11, 15, 16

Step 1: Find the **mean**:

$$\text{Mean} = \frac{3 + 6 + 6 + 7 + 8 + 11 + 15 + 16}{8} = \frac{72}{8} = 9$$

Step 2: Find the **distance** of each value from that mean:

Value	Distance from 9
3	6
6	3
6	3
7	2
8	1
11	2
15	6
16	7

The Interquartile Range (IQR) . . . and other Percentiles

The interquartile range is the middle half of the data. To visualize it, think about the median value that splits the dataset in half. Similarly, you can divide the data into quarters. Statisticians refer to these quarters as quartiles and denote them from low to high as Q1, Q2, and Q3. The lowest quartile (Q1) contains the quarter of the dataset with the smallest values. The upper quartile (Q4) contains the quarter of the dataset with the highest values. The interquartile range is the middle half of the data that is in between the upper and lower quartiles. In other words, the interquartile range includes the 50% of data points that fall between Q1 and Q3.

STANDARD DEVIATION

- The concept of standard deviation was first introduced by Karl Pearson in 1893. The standard deviation is the most useful and the most popular measure of dispersion. Just as the arithmetic mean is the most of all the averages, the standard deviation is the best of all measures of dispersion.
- The standard deviation is represented by the Greek letter (σ). It is always calculated from the arithmetic mean, median and mode is not considered. While looking at the earlier measures of dispersion all of them suffer from one or the other demerit i.e.

- Range –it suffer from a serious drawback considers only 2 values and neglects all the other values of the series.
- Quartile deviation considers only 50% of the item and ignores the other 50% of items in the series.
- Mean deviation no doubt an improved measure but ignores negative signs without any basis.
- Karl Pearson after observing all these things has given us a more scientific formula for calculating or measuring dispersion. While calculating SD we take deviations of individual observations from their AM and then each squares. The sum of the squares is divided by the number of observations. The square root of this sum is knows as standard deviation.

MERITS OF STANDARD DEVIATION

- Very popular scientific measure of dispersion
- From SD we can calculate Skewness, Correlation etc
- It considers all the items of the series
- The squaring of deviations makes them positive and the difficulty about algebraic signs which was expressed in case of mean deviation is not found here.

DEMERITS OF STANDARD DEVIATION

- Calculation is difficult not as easier as Range and QD
- It always depends on AM
- Extreme items gain great importance

$$\text{The formula of SD is } = \sqrt{\frac{\sum d^2}{N}}$$

Problem: Calculate Standard Deviation of the following series

X – 40, 44, 54, 60, 62, 64, 70, 80, 90, 96

NO OF YOUNG ADULTS VISIT TO THE LIBRARY IN 10 DAYS (X)	d=X - A.M	d ²
40	-26	676

44	-22	484
54	-12	144
60	-6	36
62	-4	16
64	-2	4
70	4	16
80	14	196
90	24	596
96	30	900
N=10 $\Sigma X=660$		$\Sigma d^2 = 3048$

- $AM = \frac{\Sigma X}{N}$
- $= \frac{660}{10} = 66$ AM
- $SD = \sqrt{\frac{\Sigma d^2}{N}}$
- $SD = \sqrt{\frac{3048}{10}} = 17.46$

Overview of how to calculate standard deviation

The formula for standard deviation (SD) is

where $\sum$ means "sum of", x is a value in the data set, μ is the mean of the data set, and N is the number of data points in the population.

The standard deviation formula may look confusing, but it will make sense after we break it down. In the coming sections, we'll walk through a step-by-step interactive example. Here's a quick preview of the steps we're about to follow:

Step 1: Find the mean.

Step 2: For each data point, find the square of its distance to the mean.

Step 3: Sum the values from Step 2.

Step 4: Divide by the number of data points.

Step 5: Take the square root.

Correlation & Regression

Correlation

Finding the relationship between two quantitative variables without being able to infer causal relationships

Correlation is a statistical technique used to determine the degree to which two variables are related

Correlation refers to a process for establishing whether or not relationships exist between two variables. You learned that a way to get a general idea about whether or not two variables are related is to plot them on a “scatter plot”. While there are many measures of association for variables which are measured at the ordinal or higher level of measurement, correlation is the most commonly used approach. This section shows how to calculate and interpret correlation coefficients for ordinal and interval level scales. Methods of correlation summarize the relationship between two variables in a single number called the correlation coefficient. The correlation coefficient is usually given the symbol r and it ranges from -1 to $+1$.

Correlation Coefficient

The correlation coefficient, r , is a summary measure that describes the extent of the statistical relationship between two interval or ratio level variables. The correlation coefficient is scaled so that it is always between -1 and $+1$. When r is close to 0 this means that there is little relationship between the variables and the farther away from 0 r is, in either the positive or negative direction, the greater the relationship between the two variables.

The two variables are often given the symbols X and Y . In order to illustrate how the two variables are related, the values of X and Y are pictured by drawing the scatter diagram, graphing combinations of the two variables. The scatter diagram is given first, and then the method of determining Pearson's r is

presented. In presenting the following examples, relatively small sample sizes are given. Later, data from larger samples are given.

Scatter Diagram

A scatter diagram is a diagram that shows the values of two variables X and Y , along with the way in which these two variables relate to each other. The values of variable X are given along the horizontal axis, with the values of the variable Y given on the vertical axis. For purposes of drawing a scatter diagram, and determining the correlation coefficient, it does not matter which of the two variables is the X variable, and which is Y.

Later, when the regression model is used, one of the variables is defined as an independent variable, and the other is defined as a dependent variable. In regression, the independent variable X is considered to have some effect or influence on the dependent variable Y. Correlation methods are symmetric with respect to the two variables, with no indication of causation or direction of influence being part of the statistical consideration. A scatter diagram is given in the following example. The same example is later used to determine the correlation coefficient.

Types of Correlation

The scatter plot explains the correlation between the two attributes or variables. It represents how closely the two variables are connected. There can be three such situations to see the relation between the two variables –

- Positive Correlation – when the value of one variable increases with respect to another.
- Negative Correlation – when the value of one variable decreases with respect to another.

No Correlation – when there is no linear dependence or no relation between the two variables.

Types of correlation

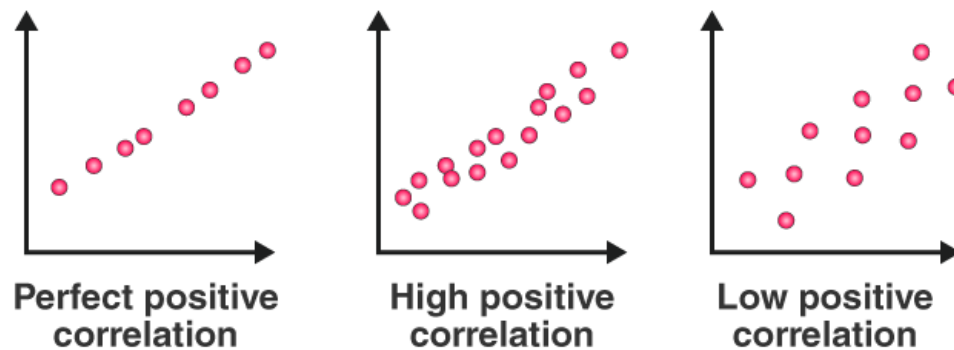
The scatter plot explains the correlation between two attributes or variables. It represents how closely the two variables are connected. There can be three such situations to see the relation between the two variables –

1. Positive Correlation
2. Negative Correlation
3. No Correlation

Positive Correlation

When the points in the graph are rising, moving from left to right, then the scatter plot shows a positive correlation. It means the values of one variable are increasing with respect to another. Now positive correlation can further be classified into three categories:

- **Perfect Positive** – Which represents a perfectly straight line
- **High Positive** – All points are nearby
- **Low Positive** – When all the points are scattered

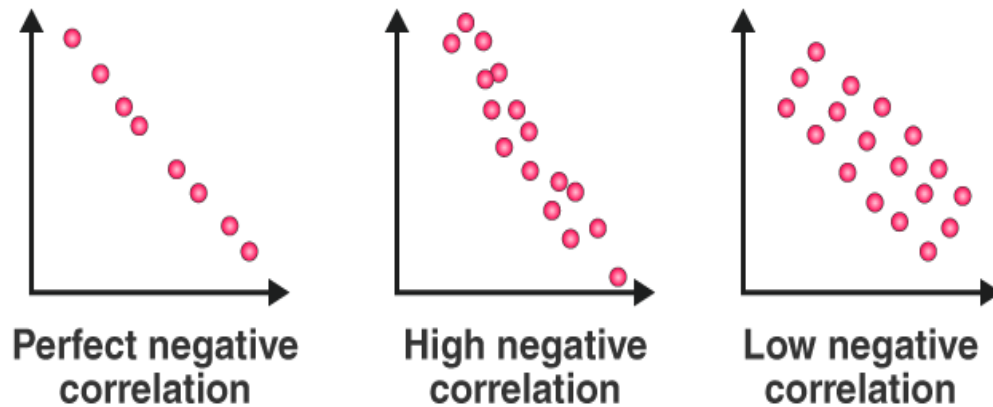


Negative Correlation

When the points in the scatter graph fall while moving left to right, then it is called a negative correlation. It means the values of one variable are decreasing with respect to another. These are also of three types:

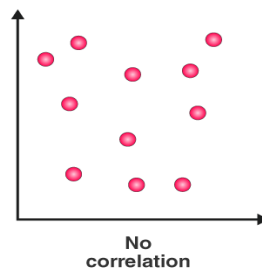
- **Perfect Negative** – Which form almost a straight line
- **High Negative** – When points are near to one another

- **Low Negative** – When points are in scattered form



No Correlation

When the points are scattered all over the graph and it is difficult to conclude whether the values are increasing or decreasing, then there is no correlation between the variables.

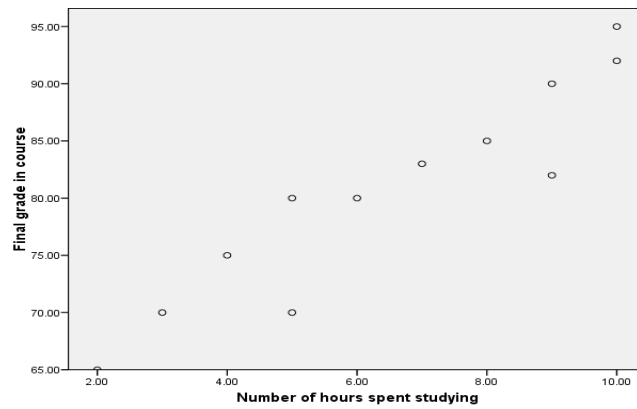


Scatter plots

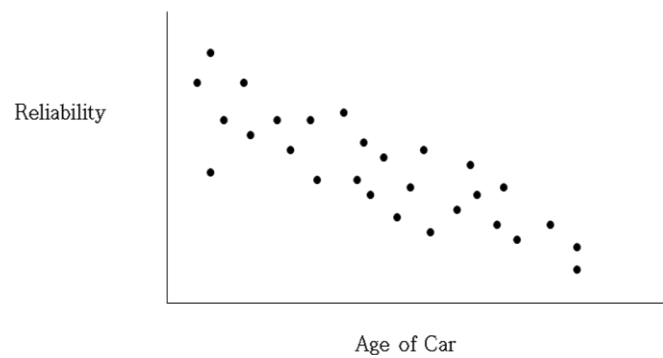
The pattern of data is indicative of the type of relationship between your two variables:

- positive relationship
- negative relationship
- no relationship

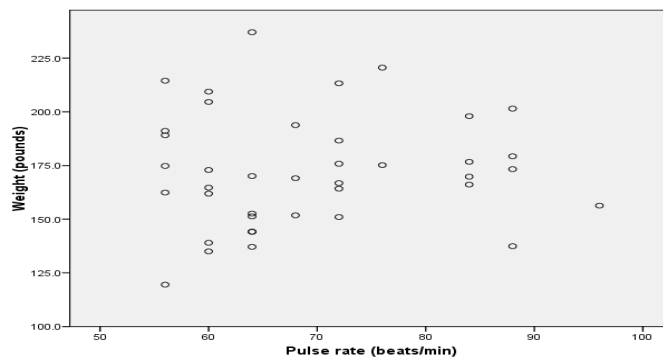
Positive relationship



Negative relationship



No relation

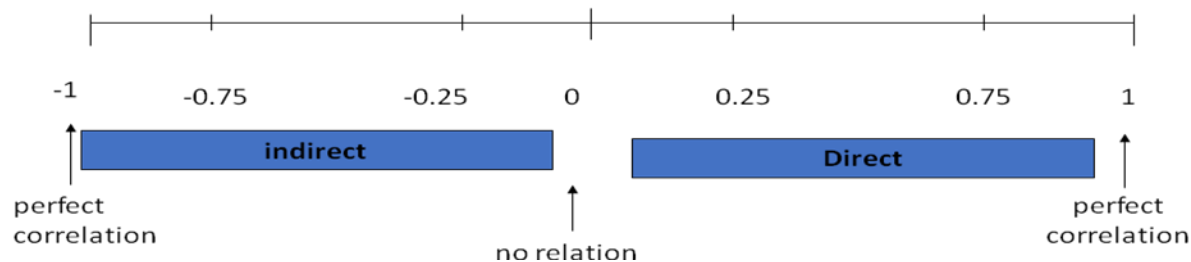


Correlation Coefficient

Statistic showing the degree of relation between two variables

Simple Correlation coefficient (r)

- It is also called Pearson's correlation or product moment correlation coefficient.
- It measures the nature and strength between two variables of the quantitative type.
- The sign of r denotes the nature of association
- while the value of r denotes the strength of association.
- If the sign is +ve this means the relation is direct (an increase in one variable is associated with an increase in the other variable and a decrease in one variable is associated with a decrease in the other variable).
- While if the sign is -ve this means an inverse or indirect relationship (which means an increase in one variable is associated with a decrease in the other).
- The value of r ranges between (-1) and (+1)
- The value of r denotes the strength of the association as illustrated by the following diagram.



- If $r = \text{Zero}$ this means no association or correlation between the two variables.
- If $0 < r < 0.25$ = weak correlation.
- If $0.25 \leq r < 0.75$ = intermediate correlation.
- If $0.75 \leq r < 1$ = strong correlation.
- If $r = 1$ = perfect correlation.

How to compute the simple correlation coefficient (r)

A sample of 6 children was selected, data about their age in years and weight in kilograms was recorded as shown in the following table. It is required to find the correlation between age and weight.

serial No	Age (years)	Weight (Kg)
1	7	12
2	6	8
3	8	12
4	5	10
5	6	11
6	9	13

These 2 variables are of the quantitative type, one variable (Age) is called the independent and denoted as (X) variable and the other (weight) is called the dependent and denoted as (Y) variables.

To find the relation between age and weight compute the simple correlation coefficient using the following formula:

Serial n.	Age (years) (x)	Weight (Kg) (y)	xy	X ²	Y ²
1	7	12	84	49	144
2	6	8	48	36	64
3	8	12	96	64	144
4	5	10	50	25	100
5	6	11	66	36	121
6	9	13	117	81	169

Total	$\Sigma x =$ 41	$\Sigma y =$ 66	$\Sigma xy = 461$	$\Sigma x^2 =$ 291	$\Sigma y^2 =$ 742
-------	--------------------	--------------------	-------------------	-----------------------	-----------------------

$$r = 0.759$$

Strong direct correlation

EXAMPLE: Relationship between Anxiety and Test Scores

Anxiety (X)	Test score (Y)	X^2	Y^2	XY
10	2	100	4	20
8	3	64	9	24
2	9	4	81	18
1	7	1	49	7
5	6	25	36	30
6	5	36	25	30
$\Sigma X = 32$	$\Sigma Y = 32$	$\Sigma X^2 = 230$	$\Sigma Y^2 = 204$	$\Sigma XY = 129$

$$r = -0.94$$

Indirect strong correlation

Spearman Rank Correlation Coefficient (r_s)

- It is a non-parametric measure of correlation.
- This procedure makes use of the two sets of ranks that may be assigned to the sample values of x and Y.

- Spearman Rank correlation coefficient could be computed in the following cases:
- Both variables are quantitative.
- Both variables are qualitative ordinal.
- One variable is quantitative and the other is qualitative ordinal.

Procedure:

1. Rank the values of X from 1 to n where n is the numbers of pairs of values of X and Y in the sample.
2. Rank the values of Y from 1 to n.
3. Compute the value of d_i for each pair of observation by subtracting the rank of Y_i from the rank of X_i
4. Square each d_i and compute $\sum d_i^2$ which is the sum of the squared values.
5. The value of r_s denotes the magnitude and nature of association giving the same interpretation as simple r.

Example - In a study of the relationship between level education and income the following data was obtained. Find the relationship between them and comment

sampl numbers	level education (X)	Income (Y)
A	Preparatory.	25
B	Primary.	10
C	University.	8
D	secondary	10
E	secondary	15
F	illiterate	50
G	University.	60

Answer:

$$\sum di^2=64$$

Comment:

There is an indirect weak correlation between level of education and income

Regression Analyses

- Regression: technique concerned with predicting some variables by knowing others
- The process of predicting variable Y using variable X

Regression

- Uses a variable (x) to predict some outcome variable (y)

Tells you how values in y change as a function of changes in values of x

Correlation and Regression

- Correlation describes the strength of a linear relationship between two variables
- Linear means “straight line”
- Regression tells us how to draw the straight line described by the correlation
- Calculates the “best-fit” line for a certain set of data

The regression line makes the sum of the squares of the residuals smaller than for any other line

Regression minimizes residuals

Exercise

A sample of 6 persons was selected the value of their age (x variable) and their weight is demonstrated in the following table. Find the regression equation and what is the predicted weight when age is 8.5 years

Serial no.	Age (x)	Weight (y)
---------------	---------	---------------

	1	7	12
2	6	8	
3	8	12	
4	5	10	
5	6	11	
6	9	13	

- Find the correlation between age and blood pressure using simple and Spearman's correlation coefficients, and comment.
- Find the regression equation?
- What is the predicted blood pressure for a man aging 25 years?

Serial	x	y	xy	x ²
1	20	120	2400	400
2	43	128	5504	1849
3	63	141	8883	3969
4	26	126	3276	676
5	53	134	7102	2809
6	31	128	3968	961
7	58	136	7888	3364
8	46	132	6072	2116
9	58	140	8120	3364
10	70	144	10080	4900

$$\frac{114486 - \frac{852 \times 2630}{20}}{41678 - \frac{852^2}{20}} = 0.4547$$

$$\hat{y} = 112.13 + 0.4547 x$$

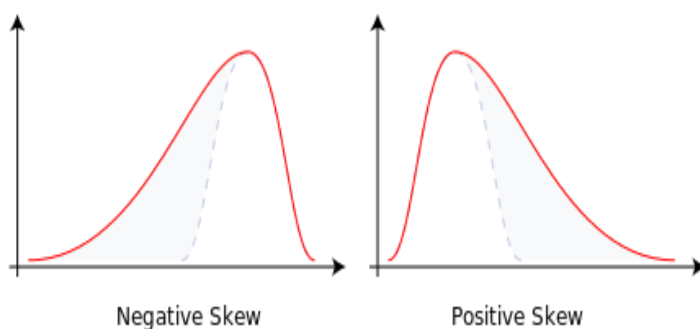
for age 25

$$\text{B.P} = 112.13 + 0.4547 * 25 = 123.49 = 123.5 \text{ mm hg}$$

Skewness

Consider the two distributions in the figure just below. Within each graph, the values on the right side of the distribution taper differently from the values on the left side. These tapering sides are called *tails*, and they provide a visual means to determine which of the two kinds of skewness a distribution has:

1. **Negative skew:** The left tail is longer; the mass of the distribution is concentrated on the right of the figure. The distribution is said to be left-skewed, left-tailed, or skewed to the left, despite the fact that the curve itself appears to be skewed or leaning to the right; left instead refers to the left tail being drawn out and, often, the mean being skewed to the left of a typical center of the data. A left-skewed distribution usually appears as a right-leaning curve



2. **Positive skew:** The right tail is longer; the mass of the distribution is concentrated on the left of the figure. The distribution is said to be right-skewed, right-tailed, or skewed to the right, despite the

fact that the curve itself appears to be skewed or leaning to the left; right instead refers to the right tail being drawn out and, often, the mean being skewed to the right of a typical center of the data. A right-skewed distribution usually appears as a left-leaning curve.

Relationship of mean and median

The skewness is not directly related to the relationship between the mean and median: a distribution with negative skew can have its mean greater than or less than the median, and likewise for positive skew.

In the older notion of nonparametric skew, defined as $\frac{\mu - \text{median}}{\sigma}$ where μ is the mean, median is the median, and σ is the standard deviation, the skewness is defined in terms of this relationship: positive/right nonparametric skew means the mean is greater than (to the right of) the median, while negative/left nonparametric skew means the mean is less than (to the left of) the median. However, the modern definition of skewness and the traditional nonparametric definition do not in general have the same sign: while they agree for some families of distributions, they differ in general, and conflating them is misleading.

If the distribution is symmetric, then the mean is equal to the median, and the distribution has zero skewness. If the distribution is both symmetric and unimodal, then the mean = median = mode. This is the case of a coin toss or the series 1,2,3,4,... Note, however, that the converse is not true in general, i.e. zero skewness does not imply that the mean is equal to the median.

A 2005 journal article points out:

Many textbooks teach a rule of thumb stating that the mean is right of the median under right skew, and left of the median under left skew. This rule fails with surprising frequency. It can fail in multimodal distributions, or in distributions where one tail is long but the other is heavy. Most commonly, though, the rule fails in discrete distributions where the areas to the left and right of the median are not equal. Such distributions not only contradict the textbook relationship between mean, median, and skew, they also contradict the textbook interpretation of the median.

Definition

Pearson's moment coefficient of skewness

The skewness of a random variable X is the third standardized moment γ_1 , defined as

Where μ is the mean, σ is the standard deviation, E is the expectation operator, μ_3 is the third central moment, and κ_t are the t -th cumulants. It is sometimes referred to as **Pearson's moment coefficient of skewness**, or simply the **moment coefficient of skewness**, but should not be confused with

Pearson's other skewness statistics (see below). The last equality expresses skewness in terms of the ratio of the third cumulant κ_3 to the 1.5th power of the second cumulant κ_2 . This is analogous to the definition of kurtosis as the fourth cumulant normalized by the square of the second cumulant. The skewness is also sometimes denoted $\text{Skew}[X]$.

If σ is finite, μ is finite too and skewness can be expressed in terms of the non-central moment $E[X^3]$ by expanding the previous formula,

Correlation

Correlation Formula

Correlation shows the relation between two variables. Correlation coefficient shows the measure of correlation. To compare two datasets we use the correlation formulas.

Pearson Correlation Coefficient Formula

The most common formula is the Pearson Correlation coefficient used for linear dependency between the data set. The value of the coefficient lies between -1 to +1. When the coefficient comes down to zero, then the data is considered as not related. While, if we get the value of +1, then the data are positively correlated and -1 has a negative correlation.

$$r = \frac{n(\sum xy) - (\sum x)(\sum y)}{\sqrt{[n \sum x^2 - (\sum x)^2][n \sum y^2 - (\sum y)^2]}}$$

Where, n = Quantity of Information

$\sum x$ = Total of the First Variable Value

$\sum y$ = Total of the Second Variable Value

$\sum xy$ = Sum of the Product of & Second Value

$\sum x^2$ = Sum of the Squares of the First Value

$\sum y^2$ = Sum of the Squares of the Second Value

Correlation coefficient formula is given and explained here for all of its types. There are various formulas to calculate the correlation coefficient and the ones covered here include Pearson's Correlation Coefficient Formula, Linear Correlation Coefficient Formula, Sample Correlation Coefficient Formula,

and Population Correlation Coefficient Formula. Before going to the formulas, it is important to understand what correlation and correlation coefficient is. A brief introduction is given below and to learn about them in detail, click the linked article.

Correlation Coefficient Formula

$$r = \frac{n(\sum xy) - (\sum x)(\sum y)}{\sqrt{[n\sum x^2 - (\sum x)^2][n\sum y^2 - (\sum y)^2]}}$$

Formula to Calculate Correlation Coefficient

About Correlation Coefficient

The correlation coefficient is a measure of the association between two variables. It is used to find the relationship is between data and a measure to check how strong it is. The formulas return a value between -1 and 1 wherein one shows -1 shows negative correlation and +1 shows a positive correlation. The correlation coefficient value is positive when it shows that there is a correlation between the two values and the negative value shows the amount of diversity among the two values.

Correlation Analysis

Correlation analysis is applied in quantifying the association between two continuous variables, for example, an dependent and independent variable or among two independent variables.

Regression Analysis

Regression analysis refers to assessing the relationship between the outcome variable and one or more variables. The outcome variable is known as the dependent or response variable and the risk elements, and cofounders are known as predictors or independent variables. The dependent variable is shown by “y” and independent variables are shown by “x” in regression analysis.

The sample of a correlation coefficient is estimated in the correlation analysis. It ranges between -1 and +1, denoted by r and quantifies the strength and direction of the linear association among two variables. The correlation among two variables can either be positive, i.e, a higher level of one variable is related to a higher level of another) or negative, i.e, a higher level of one variable is related to a lower level of the other.

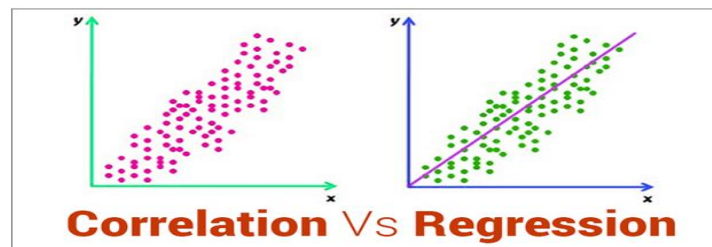
The sign of the coefficient of correlation shows the direction of the association. The magnitude of the coefficient shows the strength of the association.

For example, a correlation of $r = 0.8$ indicates a positive and strong association among two variables, while a correlation of $r = -0.3$ shows a negative and weak association. A correlation near to zero shows the non-existence of linear association among two continuous variables.

Linear Regression

Linear regression is a linear approach to modelling the relationship between the scalar components and one or more independent variables. If the regression has one independent variable, then it is known as a simple linear regression. If it has more than one independent variables, then it is known as multiple linear regression. Linear regression only focuses on the conditional probability distribution of the given values rather than the joint probability distribution. In general, all the real world regressions models involve multiple predictors. So, the term linear regression often describes multivariate linear regression.

Correlation and Regression Differences



There are some differences between Correlation and regression.

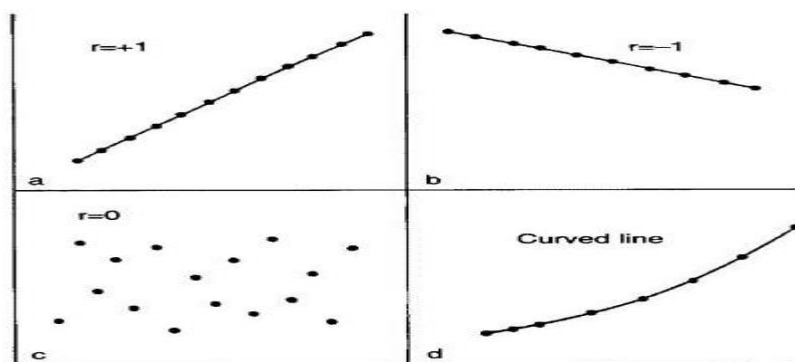
- Correlation shows the quantity of the degree to which two variables are associated. It does not fix a line through the data points. You compute a correlation that shows how much one variable changes when the other remains constant. When r is 0.0, the relationship does not exist. When r is positive, one variable goes high as the other one. When r is negative, one variable goes high as the other goes down.
- Linear regression finds the best line that predicts y from x , but Correlation does not fit a line.
- Correlation is used when you measure both variables, while linear regression is mostly applied when x is a variable that is manipulated.

Comparison Between Correlation and Regression

Basis	Correlation	Regression
Meaning	A statistical measure that defines co-relationship or association of two variables.	Describes how an independent variable is associated with the dependent variable.
Dependent and Independent variables	No difference	Both variables are different.
Usage	To describe a linear relationship between two variables.	To fit the best line and estimate one variable based on another variable.
Objective	To find a value expressing the relationship between variables.	To estimate values of a random variable based on the values of a fixed variable.

Correlation and Regression Statistics

The degree of association is measured by “r” after its originator and a measure of linear association. Other complicated measures are used if a curved line is needed to represent the relationship.



The above graph represents the correlation.

The coefficient of correlation is measured on a scale that varies from +1 to -1 through 0. The complete correlation among two variables is represented by either +1 or -1. The correlation is positive when one variable increases and so does the other; while it is negative when one decreases as the other increases. The absence of correlation is described by 0.

Index Numbers

Contents

- Introduction
- Definition
- Characteristics
- Uses
- Problems
- Classification
- Methods
- Value index numbers
- Chain index numbers.

INTRODUCTION

- An index number measures the relative change in price, quantity, value, or some other item of interest from one time period to another.
- A simple index number measures the relative change in one or more than one variable.

Characteristics, Types, Importance and Limitations

Meaning of Index Numbers:

The value of money does not remain constant over time. It rises or falls and is inversely related to the changes in the price level. A rise in the price level means a fall in the value of money and a fall in the price level means a rise in the value of money. Thus, changes in the value of money are reflected by the changes in the general level of prices over a period of time. Changes in the general level of prices can be measured by a statistical device known as 'index number.'

Index number is a technique of measuring changes in a variable or group of variables with respect to time, geographical location or other characteristics. There can be various types of index numbers, but, in the present context, we are concerned with price index numbers, which measures changes in the general price level (or in the value of money) over a period of time.

CLASSIFICATION OF INDEX NUMBERS

- *Price Index*
- *Quantity Index*
- *Value Index*
- *Composite Index*

METHODS OF CONSTRUCTING INDEX NUMBERS

- *Index Numbers*
- *Simple Aggregative*
- *Simple Average of Price Relative*
- *Unweighted*
- *Weighted*
- *Weighted Aggregated*
- *Weighted Average of Price Relatives*

SIMPLE AGGREGATIVE METHOD

It consists in expressing the aggregate price of all commodities in the current year as a percentage of the aggregate price in the base year.

P_{01} = *Index number of the current year.*

= *Total of the current year's price of all commodities.*

= *Total of the base year's price of all commodities.*

$$P_{01} = \frac{\sum P_1}{\sum P_0} \times 100$$

EXAMPLE:-

From the data given below construct the index number for the year 2007 on the base year 2008 in Rajasthan state.

Solution:-

COMMODITIES	UNITS	PRICE (Rs) 2007	PRICE (Rs) 2008
Sugar	Quintal	2200	3200
Milk	Quintal	18	20
Oil	Litre	68	71
Wheat	Quintal	900	1000
Clothing	Meter	50	60

Index Number for 2008-

$$\sum p_0 = 3236$$

$$\sum p_1 = 4351$$

$$P_{01} = \frac{\sum p_1}{\sum p_0} \times 100 = \frac{4351}{3236} \times 100 = 134.45$$

It means the prize in 2008 were 34.45% higher than the previous year

SIMPLE AVERAGE OF RELATIVES METHOD

- The current year price is expressed as a price relative of the base year price. These price relatives are then averaged to get the index number. The average used could be arithmetic mean, geometric mean or even median.

$$P_{01} = \frac{\sum \left(\frac{p_1}{p_0} \times 100 \right)}{N}$$

Where N is Numbers Of items

Example

From the data given below construct the index number for the year 2008 taking 2007 as by using arithmetic mean.

Commodities	Price (2007)	Price (2008)
P	6	10
Q	2	2
R	4	6
S	10	12
T	8	12

Solution-

Index number using arithmetic mean-

Commodities	Price (2007)	Price (2008)	Price Relative
P	6	10	166.7
Q	12	2	16.67
R	4	6	150.0
S	10	12	120.0
T	8	12	150.0

$$P_{01} = \frac{\sum \left(\frac{p_1}{p_0} \times 100 \right)}{N} = \frac{603.37}{5} = 120.63$$

$$\sum \left(\frac{p_1}{p_0} \times 100 \right)$$

Weighted index numbers

- These are those index numbers in which rational weights are assigned to various chains in an explicit fashion.

(A) Weighted aggregative index numbers-

These index numbers are the simple aggregative type with the fundamental difference that weights are assigned to the various items included in the index.

- Dorbish and bowley's method.
- Fisher's ideal method.
- Marshall-Edgeworth method.
- Laspeyres method.
- Paasche method.
- Kelly's method.

Laspeyres Method-

This method was devised by Laspeyres in 1871. In this method the weights are determined by quantities in the base.

$$p_{01} = \frac{\sum p_1 q_0}{\sum p_0 q_0} \times 100$$

Paasche's Method.

This method was devised by a German statistician Paasche in 1874. The weights of current year are used as base year in constructing the Paasche's Index number

$$p_{01} = \frac{\sum p_1 q_1}{\sum p_0 q_1} \times 100$$

Dorbish & Bowleys Method.

This method is a combination of Laspeyre's and Paasche's methods. If we find out the arithmetic average of Laspeyre's and Paasche's index we get the index suggested by Dorbish & Bowley.

$$p_{01} = \frac{\frac{\sum p_1 q_0}{\sum p_0 q_0} + \frac{\sum p_1 q_1}{\sum p_0 q_1}}{2} \times 100$$

Fisher's Ideal Index.

Fisher's ideal index number is the geometric mean of the Laspeyre's and Paasche's index numbers.

$$P_{01} = \sqrt{\frac{\sum p_1 q_0}{\sum p_0 q_0} \times \frac{\sum p_1 q_1}{\sum p_0 q_1}}$$

Marshall-Edgeworth Method.

In this index the numerator consists of an aggregate of the current years price multiplied by the weights of both the base year as well as the current year.

$$p_{01} = \frac{\sum p_1 q_0 + \sum p_1 q_1}{\sum p_0 q_0 + \sum p_0 q_1} \times 100$$

q Refers to the quantities of the year which is selected as the base.

Kelly's Method.

Kelly thinks that a ratio of aggregates with selected weights (not necessarily of base year or current year) gives the base index number

$$p_{01} = \frac{\sum p_1 q}{\sum p_0 q} \times 100$$

Example- Given below are the price quantity data, with price quoted in Rs. per kg and production in qtls.

Find- (1) Laspeyres Index (2) Paasche's Index (3) Fisher Ideal Index

.2002			2007	
ITEMS	PRICE	PRODUCTION	PRICE	PRODUCTION
BEEF	15	500	20	600
MUTTON	18	590	23	640
CHICKEN	22	450	24	500

Solution-

ITEMS	^(p₀) PRICE	^(q₀) PRODUCTION	^(p₁) PRICE	^(q₁) PRODUCTION	^(p₁q₀)	^(p₀q₀)	^(p₁q₁)	^(p₀q₁)
BEEF	15	500	20	600	10000	7500	12000	9000
MUTTON	18	590	23	640	13570	10620	14720	11520
CHICKEN	22	450	24	500	10800	9900	12000	11000
<i>TOTAL</i>					34370	28020	38720	31520

Solution-

1.Laspeyres index

$$p_{01} = \frac{\sum p_1 q_0}{\sum p_0 q_0} \times 100 = \frac{34370}{28020} \times 100 = 122.66$$

3. Paasche's Index

Fisher Ideal Index

$$p_{01} = \frac{\sum p_1 q_1}{\sum p_0 q_1} \times 100 = \frac{38720}{31520} \times 100 = 122.84$$
$$P_{01} = \sqrt{\frac{\sum p_1 q_0}{\sum p_0 q_0} \times \frac{\sum p_1 q_1}{\sum p_0 q_1}} \times 100 = \sqrt{\frac{34370}{28020} \times \frac{38720}{31520}} \times 100 = 122.69$$

Weighted average of price relative

In weighted Average of relative, the price relatives for the current year are calculated on the basis of the base year price. These price relatives are multiplied by the respective weight of items. These products are added up and divided by the sum of weights.

Weighted arithmetic mean of price relative-

$$P = \frac{P_1}{P_0} \times 100 \qquad P_{01} = \frac{\sum PV}{\sum V}$$

P=Price relative

V=Value weights= $p_0 q_0$

Value index numbers

Value is the product of price and quantity. A simple ratio is equal to the value of the current year divided by the value of base year. If the ratio is multiplied by 100 we get the value index number.

$$V = \frac{\sum p_1 q_1}{\sum p_0 q_0} \times 100$$

Chain index numbers

When this method is used the comparisons are not made with a fixed base, rather the base changes from year to year. For example, for 2007, 2006 will be the base; for 2006, 2005 will be the same and so on.

Chain index for current year-

$$= \frac{\text{Average link relative of current year} \times \text{Chain index of previous year}}{100}$$

- Example-
- From the data given below construct an index number by chain base method.

Price of a commodity from 2006 to 2008.

YEAR	PRICE
2006	50
2007	60
2008	65

solution-

YEAR	PRICE	LINK RELATIVE	CHAIN INDEX (BASE 2006)
2006	50	100	100
2007	60	$\frac{60}{50} \times 100 = 120$	$\frac{120 \times 100}{100} = 120$
2008	65	$\frac{65}{60} \times 100 = 108$	$\frac{108 \times 120}{100} = 129.60$

Important types of index numbers are discussed below:

1. Wholesale Price Index Numbers:

Wholesale price index numbers are constructed on the basis of the wholesale prices of certain important commodities. The commodities included in preparing these index numbers are mainly raw-materials and semi-finished goods. Only the most important and most price-sensitive and semi- finished goods which are bought and sold in the wholesale market are selected and weights are assigned in accordance with their relative importance.

The wholesale price index numbers are generally used to measure changes in the value of money. The main problem with these index numbers is that they include only the wholesale prices of raw materials and semi-finished goods and do not take into consideration the retail prices of goods and services generally consumed by the common man. Hence, the wholesale price index numbers do not reflect true and accurate changes in the value of money.

2. Retail Price Index Numbers:

These index numbers are prepared to measure the changes in the value of money on the basis of the retail prices of final consumption goods. The main difficulty with this index number is that the retail price for the same goods and for continuous periods is not available. The retail prices represent larger and more frequent fluctuations as compared to the wholesale prices.

3. Cost-of-Living Index Numbers:

These index numbers are constructed with reference to the important goods and services which are consumed by common people. Since the number of these goods and services is very large, only representative items which form the consumption pattern of the people are included. These index numbers are used to measure changes in the cost of living of the general public.

4. Working Class Cost-of-Living Index Numbers:

The working class cost-of-living index numbers aim at measuring changes in the cost of living of workers. These index numbers are constructed on the basis of only those goods and services which are generally consumed by the working class. The prices of these goods and index numbers are of great importance to the workers because their wages are adjusted according to these indices.

5. Wage Index Numbers:

The purpose of these index numbers is to measure time to time changes in money wages. These index numbers, when compared with the working class cost-of-living index numbers, provide information regarding the changes in the real wages of the workers.

6. Industrial Index Numbers:

Industrial index numbers are constructed with an objective of measuring changes in the industrial production. The production data of various industries are included in preparing these index numbers.

Importance of Index Numbers:

Index numbers are used to measure all types of quantitative changes in different fields.

Various advantages of index numbers are given below:

1. General Importance:

In general, index numbers are very useful in a number of ways:

- (a) They measure changes in one variable or in a group of variables.
- (b) They are useful in making comparisons with respect to different places or different periods of time,
- (c) They are helpful in simplifying the complex facts.
- (d) They are helpful in forecasting about the future,
- (e) They are very useful in academic as well as practical research.

2. Measurement of Value of Money:

Index numbers are used to measure changes in the value of money or the price level from time to time. Changes in the price level generally influence production and employment of the country as well as various sections of the society. The price index numbers also forewarn about the future inflationary tendencies and in this way, enable the government to take appropriate anti- inflationary measures.

3. Changes in Cost of Living:

Index numbers highlight changes in the cost of living in the country. They indicate whether the cost of living of the people is rising or falling. On the basis of this information, the wages of the workers can be adjusted accordingly to save the wage earners from the hardships of inflation.

4. Changes in Production:

Index numbers are also useful in providing information regarding production trends in different sectors of the economy. They help in assessing the actual condition of different industries, i.e., whether production in a particular industry is increasing or decreasing or is constant.

5. Importance in Trade:

Importance in trade with the help of index numbers, knowledge about the trade conditions and trade trends can be obtained. The import and export indices show whether foreign trade of the country is increasing or decreasing and whether the balance of trade is favourable or unfavourable.

6. Formation of Economic Policy:

Index numbers prove very useful to the government in formulating as well as evaluating economic policies. Index numbers measure changes in the economic conditions and, with this information, help the planners to formulate appropriate economic policies. Further, whether particular economic policy is good or bad is also judged by index numbers.

7. Useful in All Fields:

Index numbers are useful in almost all the fields. They are specially important in economic field.

Some of the specific uses of index numbers in the economic field are:

- (a) They are useful in analysing markets for specific commodities.
- (b) In the share market, the index numbers can provide data about the trends in the share prices,
- (c) With the help of index numbers, the Railways can get information about the changes in goods traffic.
- (d) The bankers can get information about the changes in deposits by means of index numbers.

Limitations of Index Numbers:

Index number technique itself has certain limitations which have greatly reduced its usefulness:

- (i) Because of the various practical difficulties involved in their computation, the index numbers are never cent per cent correct.

- (ii) There are no all-purpose index numbers. The index numbers prepared for one purpose cannot be used for another purpose. For example, the cost-of-living index numbers of factory workers cannot be used to measure changes in the value of money of the middle income group.
- (iii) Index numbers cannot be reliably used to make international comparisons. Different countries include different items with different qualities and use different base years in constructing index numbers.
- (iv) Index numbers measure only average change and indicate only broad trends. They do not provide accurate information.
- (v) While preparing index numbers, quality of items is not considered. It may be possible that a general rise in the index is due to an improvement in the quality of a product and not because of a rise in its.

Time series analysis

Definitions, Applications and Techniques	
Definition	Definition of Time Series: <i>An ordered sequence of values of a variable at equally spaced time intervals.</i>
Time series occur frequently when looking at industrial data	Applications: The usage of time series models is twofold: • Obtain an understanding of the underlying forces and structure that produced the observed data • Fit a model and proceed to forecasting, monitoring or even feedback and feed forward control. Time Series Analysis is used for many applications such as: • Economic Forecasting • Sales Forecasting • Budgetary Analysis • Stock Market Analysis • Yield Projections

	• Process and Quality Control • Inventory Studies • Workload Projections • Utility Studies • Census Analysis and many, many more...
<i>There are many methods used to model and forecast time series</i>	Techniques: The fitting of time series models can be an ambitious undertaking. There are many methods of model fitting including the following: • <u>Box-Jenkins ARIMA models</u> • <u>Box-Jenkins Multivariate Models</u> • <u>Holt-Winters Exponential Smoothing (single, double, triple)</u> The user's application and preference will decide the selection of the appropriate technique. It is beyond the realm and intention of the authors of this handbook to cover all these methods. The overview presented here will start by looking at some basic smoothing techniques: • Averaging Methods • Exponential Smoothing Techniques.

What Is a Time Series?

A time series is a sequence of numerical data points in successive order. In investing, a time series tracks the movement of the chosen data points, such as a security's price, over a specified period of time with data points recorded at regular intervals. There is no minimum or maximum amount of time that must be included, allowing the data to be gathered in a way that provides the information being sought by the investor or analyst examining the activity.

[Important: Time series analysis can be useful to see how a given asset, security, or economic variable changes over time.]

Understanding Time Series

A time series can be taken on any variable those changes over time. In investing, it is common to use a time series to track the price of a security over time. This can be tracked over the short term, such as the price of a security on the hour over the course of a business day, or the long term, such as the price of a security at close on the last day of every month over the course of five years.

Time Series Analysis

Time series analysis can be useful to see how a given asset, security, or economic variable changes over time. It can also be used to examine how the changes associated with the chosen data point compare to shifts in other variables over the same time period.

For example, suppose you wanted to analyze a time series of daily closing stock prices for a given stock over a period of one year. You would obtain a list of all the closing prices for the stock from each day for the past year and list them in chronological order. This would be a one-year daily closing price time series for the stock.

Delving a bit deeper, you might be interested to know whether the stock's time series shows any seasonality to determine if it goes through peaks and troughs at regular times each year. Analysis in this area would require taking the observed prices and correlating them to a chosen season. This can include traditional calendar seasons, such as summer and winter, or retail seasons, such as holiday seasons.

Alternatively, you can record a stock's share price changes as it relates to an economic variable, such as the unemployment rate. By correlating the data points with information relating to the selected economic variable, you can observe patterns in situations exhibiting dependency between the data points and the chosen variable.

Time Series Forecasting

Time series forecasting uses information regarding historical values and associated patterns to predict future activity. Most often, this relates to trend analysis, cyclical fluctuation analysis, and issues of seasonality. As with all forecasting methods, success is not guaranteed.

Unit –II PROBABILITY THEORY

Permutations and combinations

Permutations and combinations, the various ways in which objects from a [set](#) may be selected, generally without replacement, to form subsets. This selection of subsets is called a permutation when the order of selection is a factor, a combination when order is not a factor.

The concepts of and differences between permutations and combinations can be illustrated by examination of all the different ways in which a pair of objects can be selected from five distinguishable objects—such as the letters A, B, C, D, and E. If both the letters selected and the order of selection are considered, then the following 20 outcomes are possible:

AB	BA	AC	CA	AD
DA	AE	EA	BC	CB
BD	DB	BE	EB	CD
DC	CE	EC	DE	ED

Each of these 20 different possible selections is called a permutation. In particular, they are called the permutations of five objects taken two at a time, and the number of such permutations possible is denoted by the symbol ${}_5P_2$, read “5 permute 2.” In general, if there are n objects available from which to select, and permutations (P) are to be formed using k of the objects at a time, the number of different permutations possible is denoted by the symbol ${}_nP_k$. A formula for its evaluation is ${}_nP_k = n!/(n - k)!$. The expression $n!$ —read “ n [factorial](#)”—indicates that all the consecutive positive integers from 1 up to and including n are to be multiplied together, and $0!$ is defined to equal 1.

(For $k = n$, ${}_nP_k = n!$. Thus, for 5 objects there are $5! = 120$ arrangements.)

For combinations, k objects are selected from a set of n objects to produce subsets without ordering. Contrasting the previous permutation example with the corresponding combination, the AB and BA subsets are no longer distinct selections; by eliminating such cases there remain only 10 different possible subsets—AB, AC, AD, AE, BC, BD, BE, CD, CE, and DE.

The number of such subsets is denoted by ${}_nC_k$, read “ n choose k .” For combinations, since k objects have $k!$ Arrangements, there are $k!$ Indistinguishable permutations for each choice of k objects; hence dividing the permutation formula by $k!$ yields the following combination formula:

$${}_nC_k = \frac{n!}{k!(n - k)!}.$$

This is the same as the (n, k) binomial coefficient (see [binomial theorem](#)). For example, the number of combinations of five objects taken two at a time is

$$\begin{aligned} {}_5C_2 &= \frac{5!}{(2)!(5-2)!} = \frac{5!}{(2)!(3)!} = \frac{(1)(2)(3)(4)(5)}{(1)(2)(1)(2)(3)} \\ &= \frac{120}{12} = 10. \end{aligned}$$

The formulas for nPk and nCk are called counting formulas since they can be used to count the number of possible permutations or combinations in a given situation without having to list them all.

How To Tell the Difference

The difference between combinations and permutations is ordering. **With permutations we care about the order of the elements, whereas with combinations we don't.**

For example, say your locker “combo” is 5432. If you enter 4325 into your locker it won't open because it is a different ordering (aka permutation).

The **permutations** of 2, 3, 4, 5 are:

- **5432**, 5423, 5324, 5342, 5234, 5243, 4532, 4523, 4325, 4352, 4253, 4235, 3542, 3524, 3425, 3452, 3254, 3245, 2543, 2534, 2435, 2453, 2354, 2345

Your locker “combo” is a specific permutation of 2, 3, 4 and 5. If your locker worked truly by combination, you could enter any of the above permutations and it would open!

Calculating Permutations with Ease

Suppose you want to know how many permutations exist of the numbers 2, 3, 4, 5 without listing them like I did above. How would you accomplish this?

Let's use a line diagram to help us visualize the problem.

We want to find how many possible 4-digit permutations can be made from four distinct numbers. Begin

by drawing four lines to represent the 4 digits.

The first digit can be any of the 4 numbers, so place a “4” in the first blank.

Now there are 3 options left for the second blank because you’ve already used one of the numbers in the first blank. Place a “3” in the next space.

$$\underline{4} \cdot \underline{3} \quad \underline{\quad} \quad \underline{\quad}$$

For the third position, you have two numbers left.

$$\underline{4} \cdot \underline{3} \cdot \underline{2} \quad \underline{\quad}$$

And there is one number left for the last position, so place a “1” there.

$$\underline{4} \cdot \underline{3} \cdot \underline{2} \cdot \underline{1}$$

The Multiplication Principle

Using the *Multiplication Principle of combinatorics*, we know that if there are x ways of doing one thing and y ways of doing another, then the total number of ways of doing both things is $x \cdot y$. That means we need to multiply to find the total permutations.

This is a great opportunity to use shorthand **factorial notation (!)**:

$$4! = 4 \cdot 3 \cdot 2 \cdot 1 = 24$$

There are 24 permutations, which matches the listing we made at the beginning of this post.

Permutations with Repetition

What if I wanted to find the total number of permutations involving the numbers 2, 3, 4, and 5 but want to include orderings such as 5555 or 2234 where not all of the numbers are used, and some are used more than once?

How many of these permutations exist?

This turns out to be a simple calculation. Again we are composing a 4-digit number, so draw 4 lines to represent the digits.

In the first position we have 4 number options, so like before place a “4” in the first blank. Since we are allowed to reuse numbers, we now have 4 number options available for the second digit, third digit, and fourth digit as well.

That’s the same as:

$$\underline{4} \cdot \underline{4} \cdot \underline{4} \cdot \underline{4} = 4^4 = 256$$

By allowing numbers to be repeated, we end up with 256 permutations!

Choosing a Subset

Let’s up the ante with a more challenging problem:

How many different 5-card hands can be made from a standard deck of cards?

In this problem the order is irrelevant since it doesn’t matter what order we select the cards.

We’ll begin with five lines to represent our 5-card hand.

Assuming no one else is drawing cards from the deck, there are 52 cards available on the first draw, so place “52” in the first blank.

Once you choose a card, there will be one less card available on the next draw. So the second blank will have 51 options. The next draw will have two less cards in the deck, so there are now 50 options, and so on.

$$\underline{52} \cdot \underline{51} \cdot \underline{50} \cdot \underline{49} \cdot \underline{48}$$

That’s 311,875,200 *permutations*.

That’s permutations, not combinations. To fix this we need to divide by the number of hands that are different permutations but the same combination.

This is the same as saying *how many different ways can I arrange 5 cards?*

$$5! = 5 \cdot 4 \cdot 3 \cdot 2 \cdot 1$$

Note: This is mathematically similar to finding the different permutations of our locker combo

So the number of five-card hands combinations is:

$$\frac{52 \cdot 51 \cdot 50 \cdot 49 \cdot 48}{5 \cdot 4 \cdot 3 \cdot 2 \cdot 1} = 2,598,960$$

Rewriting with Factorials

With a little ingenuity we can rewrite the above calculation using factorials.

We know $52! = 52 \cdot 51 \cdot 50 \cdot \dots \cdot 3 \cdot 2 \cdot 1$, but we only need the products of the integers from 52 to 48. How can we isolate just those integers?

We'd like to divide out all the integers except those from 48 to 52. To do this divide by $47!$ since it's the product of the integers from 47 to 1.

$$\frac{52!}{47!} = \frac{52 \cdot 51 \cdot 50 \cdot 49 \cdot 48 \cdot \cancel{47} \cdot \cancel{46} \cdot \dots \cdot \cancel{2} \cdot \cancel{1}}{\cancel{47} \cdot \cancel{46} \cdot \dots \cdot \cancel{2} \cdot \cancel{1}}$$

Make sure to divide by $5!$ to get rid of the extra permutations:

$$\frac{52!}{5!47!}$$

There we go!

Now here's the cool part → we have actually derived the formula for **combinations**.

Combinations Formula

If we have n objects and we want to choose k of them, we can find the total number of combinations by using the following formula:

$$\binom{n}{k} = \frac{n!}{k!(n-k)!}$$

read: “ n choose k ”

For example, we have 52 cards ($n=52$) and want to know how many 5-card hands ($k=5$) we can make.

Plugging in the values we get:

$$\binom{52}{5} = \frac{52!}{5!(52-5)!}$$

note: $(52-5)! = 47!$

Which is exactly what we found above!

Often times you'll see this formula written in parenthesis notation, like above, but some books write it with a giant C:

Sample Space

As discussed in the previous section, the outcome of a probabilistic experiment is not known with

certainty but the possible set of outcomes is very well known. This set of all possible outcomes of an experiment is known as the sample space of the experiment and is denoted by S .

The sample space of an experiment can be thought of as a set, or collection, of all the possible outcomes of the experiment and each of the outcomes in S could be understood as an element of the set, s

Examples of sample spaces

Following are the examples of sample space and their respective elements.

1. If the outcome of an experiment is to determine the sex of a newborn child, then the sample space would consist of two elements, boy (B) and girl (G).
2. If the experiment is to roll a loaded die such that likelihood of getting an even number is twice that of an odd number, then the sample space would include all the numbers of the die, but every element of S is not equally likely.
3. An experiment consists of flipping two coins. The sample space would include following cases

$$S = \{TT, HT, TH, HH\}$$

The outcome TT refers to tails on both the coins. If the first coin is heads and second is tails, then the outcome is (HT) and if the first outcome is tails and second is heads, then it is (TH). The last possible outcome is head on both the coins (HH).

4. If two dice are rolled simultaneously and all the possible outcomes are noted, then the sample space becomes following

$$S = \{(i,j): i, j = 1, 2, 3, 4, 5, 6\}$$

In this case, 'i' refers to the outcome on the first die and 'j' is the outcome on the second die. The total number of outcomes will be 36 and they are tabled below

(1, 1) (1, 2) (1, 3) (1, 4) (1, 5) (1, 6)
(2, 1) (2, 2) (2, 3) (2, 4) (2, 5) (2, 6)
(3, 1) (3, 2) (3, 3) (3, 4) (3, 5) (3, 6)
(4, 1) (4, 2) (4, 3) (4, 4) (4, 5) (4, 6)
(5, 1) (5, 2) (5, 3) (5, 4) (5, 5) (5, 6)
(6, 1) (6, 2) (6, 3) (6, 4) (6, 5) (6, 6)

5. Consider the experiment of flipping 5 coins and noting the number of heads. The sample space is as such

$$S = \{0, 1, 2, 3, 4, 5\}$$

The sample space should consist of the exhaustive list of outcomes, such that no larger sample space can be more informative. In certain experiments, some of the outcomes may be implausible, but you have little to lose by choosing a large enough sample space.

6. Suppose the experiment consists of counting the number of telephone calls arriving in an exchange in an hour. In this the sample space will not be finite.

$$S = \{0, 1, 2, 3, \dots\}$$

Evidently, only a finite number of calls are possible in a specific time frame, but any cutoff number would be arbitrary, and might be too small.

7. Consider an experiment of testing an infinite lot of batteries to have a prescribed limit of voltage. If the battery has required units of voltage it is considered as a success (S) and the experiment stops there. However, if the battery fails (F) the voltage limit then another battery is tested out of the lot and the testing continues till a battery with correct voltage is obtained.

$$S = \{S, FS, FFS, FFFS, \dots\}$$

This sample space has infinite elements. First outcome depicts that success is obtained on the very first trial. However, third element in S implies that first two batteries failed and the third one was successful.

Concept of an Event

Consider the following experiments and the respective set of outcomes to understand the meaning of an event

In an experiment in which a coin is tossed 10 times, the experimenter wants to look at the outcomes in which at least four tails are obtained.

In an experiment of inspecting 100 projectors, the experimenter wants to look at the outcomes of finding more than 4 projectors defective.

In an experiment where the class strength is observed for the course statistics for 90 days, an experimenter wants to look at the outcomes (days) when more than 60% of students come to the class.

In the above examples, the word “outcome” should be understood as an event. It includes those outcomes which are to be studied by the experiments.

An event is any subset of a sample space. In other words, it includes well-defined set of possible outcomes of an experiment. Conventionally, events are denoted by upper case letters like A , B , C ... without any suffix, superfixes or other adornments. We would say that event A has occurred if the outcome of the experiment is one of the elements in A .

Examples of Events

The examples of sample spaces discussed in the previous section are revisited here one by one to understand the meaning of an event.

1. In the experiment of determining the gender of the newborn child, if the event A is the birth of a girl child, then

$$A = \{G\}$$

It is a subset of the sample space, $S = \{B, G\}$.

2. In the experiment of a roll of single die, if the event B is obtaining an even number, then

$$B = \{2, 4, 6\}$$

The sample space of the roll of die consists of all the numbers from 1 to 6. If an even number comes up in the roll of a die, then we say that event B has occurred.

3. If two coins are flipped together and E is the event that a tail appears on the first coin, then

$$E = \{(T, T), (T, H)\}$$

Out of the total four cases in the sample space, two are included in the set of event E.

4. Consider the experiment of simultaneously rolling two dice. This experiment has 36 possible outcomes. In other words, the sample space has 36 elements. Now, consider the following events pertaining to this experiment

- a. Let A represents an event such that sum of two numbers appearing on the dice is greater than 10

$$A = \{(5, 6), (6, 5), (6, 6)\}$$

Out of the 36 possible outcomes, only three cases are such where the sum of two numbers is greater than 10. If any of these three outcomes comes up in the roll of dice, then we can say that event A has occurred

Simple and Compound Events

An event of an experiment is simple or elementary if it cannot be decomposed any further.

For example, getting 6 on a roll of a die is denoted as event, $E = \{6\}$.

On the other hand, if an event can be decomposed in simple events or is formed by the combination of simple events then it is called a compound event.

For example, getting an even number in the roll of a die is denoted by event, $B = \{2, 4, 6\}$. Here, event B consists of three simple events.

Types of events

1. Certain Event/ Sure Event:

If the defined event is such that it contains every element of the sample space, S , then the event is called the certain event or the sure event. It is an event which will surely happen in the outcome of an experiment. Take for example, the experiment of rolling a die and the event A is defined to be obtaining a number greater than 0. Now, in this case, sample space contains numbers 1, 2, 3, 4, 5, and 6 and the event is defined as

$A = \{1, 2, 3, 4, 5, 6\}$ as all the numbers are greater than 0. It implies that in a roll of die, event A is bound to occur, therefore, it is a certain/ sure event.

2. Null event/ impossible event: If the defined event has no element in it, then it is called a null event and is denoted by $\emptyset$. It is called an impossible event as this event has no element and it can't occur. For example, consider an event of getting a number greater than 6 in a roll of a die. This is a null event as the numbers on the die do not exceed 6.

3. Mutually Exclusive Events: If occurrence of event A rules out the occurrence of another event B , then events A and B are said to be mutually exclusive events. In other words, it means only one of the two events can occur at a time.

Example 1. A card is drawn at random from a deck of 52 playing cards. Find the probability that the card drawn is either a king or a queen.

Solution. Let A be the event that the card drawn is king and B be the event that the card drawn is queen.

Therefore $P(A) = 4/52$, $P(B) = 4/52$, $P(A \cap B) = 0$,

since a card which is a queen cannot be a king.

$P(A \cup B) = P(A) + P(B) = 4/52 + 4/52 = 8/52 = 2/13$.

4. Independent Events: In the case of independent events, the outcome of an event does not affect the outcome of another event. For example, if a coin is tossed two times, then the result of the second toss is completely unaffected by the result of first toss. It means that the two tosses are independent.

5. Non-Mutually Exclusive Events- If two events A and B are not mutually exclusive then the probability that A or B will occur is equal to the sum of probabilities of each event minus the probability that both events occur together.

Events are not mutually exclusive if they both can occur at the same time, that is, A and B overlap. In , we have

$$P(A \cup B) = [P(A) - P(A \cap B)] + [P(A \cap B)] + [P(B) - P(A \cap B)].$$

Therefore, $P(A \cup B) = P(A) + P(B) - P(A \cap B)$

$$P(\overline{A \cap B}) = 1 - P(A \cap B)$$

Example. A card is drawn at random from a deck of 52 playing cards. Find the probability that the card drawn is either a king or a spade. Solution. Let A be the event that the card drawn is king and B be the event that the card drawn is a spade.

Therefore $P(A) = 4/52$, $P(B) = 13/52$, $P(A \cap B) = 1/52$,

since a card which king is a and a spade is 1.

$$P(A \cup B) = P(A) + P(B) - P(A \cap B) = 4/52 + 13/52 - 1/52 = 16/52 = 4/13$$

Approaches to Probability

Having understood the concept of experiments, sample space and events, we can define probability as a precise numerical measure of the likelihood of an event. The probability of event A would be denoted by $P(A)$ which is a number, called the probability of the event A. Take for example, an experiment of testing a battery to have voltage within prescribed limits. If the battery has required voltage, it is said that event S (success) has occurred and if the required voltage is not met that event F(failure) has occurred. Therefore, the sample space is

$$S = \{F, S\}$$

The probability assignment implies finding the numerical measure for $P(F)$ and $P(S)$. The concept of probability can be understood from three approaches: classical, relative frequency and axiomatic. Each of the approaches is discussed in detail below.

Classical Approach to Probability

This approach says that if an event A occurs in m different ways out of the total number of n possible ways of a random experiment, where every outcome has the same chance of occurrence and are mutually exclusive, then the probability of the event is m/n . In other words, the probability of event A occurring can be written as

$$P(A) = (\text{Favorable number of cases to A} / \text{Total number of cases})$$

Take for example, in a toss of a coin, the two possible outcomes are heads (H) and tails (T) and both are equally likely (given that the coin is unbiased/ fair i.e. not loaded in any way) and are mutually exclusive (as turning up of heads rules out the occurrence of tails).

$$S = \{H, T\}.$$

The total number of outcomes (as can be seen from the sample space) is two. So, in the event of heads, one outcome is favorable out of the two, thus, $P(H) = \frac{1}{2}$. Similarly, $P(T) = \frac{1}{2}$

Relative Frequency Approach to Probability

This interpretation is based on the notion of relative frequencies. If an experiment is repeated n number of times, where n is very large, and an event A is observed to occur ' m ' times out of total, then the probability of event A is defined as m/n . This is also called empirical approach to probability. The ratio m/n is the relative frequency of occurrence of A in the sequence of n replications of experiment. Let's take up an example to understand the approach.

Consider an example where event A is that a student reaches in the class on time. Suppose the experiment is done 50 times and the student arrives on time in 20 classes, then the relative frequency approach suggests that the probability of student reaching on time is $20/50$, i.e. 0.4. This relative frequency fluctuates substantially over the course of repeating the experiment and thus, the experiment is to be repeated a large number of times to get a limiting value (stabilized value).

Consider the example of tossing a coin. If a coin is tossed 100 times and the heads occurs 53 times, then $P(H) = 0.53$. Now, if the same experiment is repeated 1000 times and heads turns up on 509 times, then $P(H) = 0.509$. With increase in the number of trials, the probability converges (stabilizes) to the value of 0.5.

As the relative frequency approach of probability is based on the notion of limiting frequency, its application is limited only to the experiments which can be repeated a large number of times. Consider the case when we want to find out the probability of obtaining a contract. Now, this experiment can't be repeated again and again as the award of the contract is a one-time event and therefore, relative frequency approach fails to provide a numerical measure of getting the contract. Moreover, relative frequency approach lacks on the ground that it requires a large number of trials. This is vague, time and cost consuming.

Axiomatic Approach to Probability

Because of the difficulties attached to the classical and relative frequency approaches to probability, statisticians usually refer to the axiomatic approach to probability. In this approach, the probability of an event is required to satisfy certain axioms (basic properties) to be an appropriate measure of likelihood of the event.

An event A occurs with probability $P(A)$, if the following axioms are satisfied.

1. $P(A) \geq 0$, the probability of an event should be a non-negative number.
2. $P(S) = 1$. This axiom implies that for a certain event, the probability is 1. Here, S includes the whole sample space and the maximum possible probability to be assigned for S to occur is 1.
3. If $A_1, A_2, A_3, \dots$ is a sequence of infinite mutually exclusive events, then

$P(A_1 \cup A_2 \cup A_3 \cup \dots) = P(A_1) + P(A_2) + P(A_3) + \dots$. This axiom holds true for the finite sequence of mutually exclusive events also.

This axiom states that for any set of mutually exclusive events the probability that at least one of them will occur is equal to the sum of their respective probabilities.

Definition: Probability, on a sample space S , is a specification of numerical value $P(A)$ for all events A that satisfy the axioms 1, 2 and 3.

From the definition of axiomatic approach involving these three axioms, we can derive a bunch of properties which will help in assigning probability.

P1. If event A_1 is a subset of event A_2 , then $P(A_1) \leq P(A_2)$

P2. For every event A , $0 \leq P(A) \leq 1$. It means that the probability of any event lies between 0 and

1. If the event is a sure event, the probability of its occurrence will be 1 (100%) and if it is a null event then the probability will be 0, $P(\emptyset) = 0$.

P3. If A' is the complement of event A , then $P(A') = 1 - P(A)$. This holds true as A and A'

together constitute the entire sample space which has probability 1.

P4. If $A = A_1 \cup A_2 \cup A_3$, where A_1, A_2, A_3 are mutually exclusive events and $A = S$, then $P(A_1) + P(A_2) + P(A_3) = 1$

P5. For every two events A_1 and A_2 , $P(A \cap B') = P(A) - P(A \cap B)$

The axiomatic approach to probability is better than the classical and relative frequency approach as it subsumes their drawbacks. This approach implies to choose nonnegative numbers between 0 and 1 to assign probabilities which satisfy the axioms and properties. In particular, the classical approach to probability is used for the simple events which are equally likely.

Example 1: If a balanced die is rolled, then all the sides are equally likely to appear and in this case $P(\{1\}) = P(\{2\}) = P(\{3\}) = P(\{4\}) = P(\{5\}) = P(\{6\}) = 1/6$. The total number of cases is 6.

The probability $1/6$ satisfies axiom 1 as it is a non-negative number. The probability of all the events occurring is

$$P(S) = P(\{1\}) + P(\{2\}) + P(\{3\}) + P(\{4\}) + P(\{5\}) + P(\{6\}) = 1$$

Thus, axiom 2 is satisfied. Now consider events A and B , such that A implies getting an even number and B is getting an odd number. Both are mutually exclusive and equally likely.

$$P(A) = P(\{2, 4, 6\}) = P(\{2\}) + P(\{4\}) + P(\{6\})$$

according to Axiom 3 = $\frac{1}{2}$

Similarly, $P(B) = \frac{1}{2}$. This implies that there is 50% chance of getting an even number and 50% chance of getting an odd number.

Example 2: Events are not equally likely. Consider the experiment of choosing a student in a class of 50 students, where 20 are girls (G) and 30 are boys (B). Here the sample space of the event is {G, B}, but both the outcomes are not equally likely. A boy would be selected with the probability 0.6 (30/50) and a girl would be selected with probability 0.4 (20/50). Here, the probabilities are assigned as the ratio of favorable cases out of total cases of an experiment. The axioms of probability are satisfied.

Joint and Conditional Probability

Joint probability is a probability describing the probabilities of occurrence two or more random events simultaneously. It is the probability of the intersection of two or more events. The probability of the intersection of A and B may be written $P(A \cap B)$. In particular if we have only two random variables, then it is known as bivariate probability.

Example: Find the probability that a card is a four and black

Solution : $P(\text{four and black}) = \frac{2}{52} = \frac{1}{26}$. (Since there are 2 black fours in a deck of 52, the 4 of clubs and the 4 of spades)

Note: Joint probability cannot occur by happening of a single event, but by the occurrence of atleast two events simultaneously.

Conditional Probability

Conditional Probability of an event is its probability of occurring, given that another event has already occurred. If A and B are events in sample space S, then the conditional probability of A occurring given B has already occurred is denoted by

$$P(A/B) = \frac{P(A \cap B)}{P(B)} \text{ for } P(B) \neq 0 \quad 3.3$$

$$P(A \cap B) = P(A/B) \cdot P(B) = P(B/A) \cdot P(A)$$

states the basic rule of probability multiplication. According to this theorem, the probability of any two

events A and B occurring jointly ($P(A \cap B)$) is found by multiplying the probability of one of the events by the probability of the other, given the condition that the first event has occurred or will occur. When the occurrence of one event is conditional upon the occurrence of another event, the events are said to be dependent events.

Example. Food Plaza issued food vouchers to its regular 2000 customers. Of these customers, 1500 hold breakfast vouchers, 500 hold dinner vouchers and 40 hold both breakfast and dinner vouchers. Find the probability that a customer chosen at random holds a breakfast voucher given that the customer holds a dinner voucher.

Solution. Let B be the event of customer holding breakfast vouchers and D be the event of holding dinner vouchers. Therefore $P(B) = 1500/2000$, $P(D) = 500/2000$ and $P(B \cap D) = 40/2000$

Therefore probability of customer holding breakfast voucher given that the customer holding a dinner voucher is $P(B/D) = P(B \cap D) / P(D) = 40/2000 / 500/2000 = 2/25$

In the above example of the Conditional Probability, the sample sizes of the experiment are reduced to D, and hence the probability of B is not unconditional but conditional on given D.

Introduction- In probability a random variable can be Discrete or Continuous with their individual probability distribution which provides a complete description of a random variable or of the population over which the random variable is defined. Apart from studying these probability distributions it is often necessary to characterize the distribution by a few measures like mean and variance which summarizes the whole set of data and highlights the important characteristics of the population.

This module will focus on the probability and cumulative distributions of random variables and the different measures which helps in studying the distribution more conveniently.

.

Random Variable

Let us suppose we have a sample space S of some experiment, then a random variable (rv) is any rule that associates a number with each outcome in S. Mathematically, we can define a random variable as a function whose sample space is the domain and whose set of real numbers is the range.

Random variables are of two types

1. Discrete Random Variable
2. Continuous Random Variable

Following the usual convention in statistics, we denote random variables by upper-case letters (X) and particular values of the random variables by the corresponding lower-case letters($x_1, x_2, x_3 \dots x_n$)

Discrete Random Variable

A discrete random variable is a type random variable which can take only some specific values either finite or countable: Suppose we flip a coin and count the number of heads. The number of heads could be any integer value between zero and positive infinity. However, it could not be any number between zero and positive infinity. We could not for example get 2.5 heads. Therefore, the number of heads must be a discrete variable.

Continuous Random Variable - Continuous variable can take any numerical value in an interval or collection of intervals. A continuous random variable is a set of possible values which consists of an entire interval on a number line and is characterized by (infinitely) uncountable values within any interval.

Example: Let us suppose that the fire department makes it compulsory that all of its fire fighters must weigh between 150 to 250 pounds. Thus we can see that the weight of a fire fighter is an example of a continuous variable, since a fire fighters weight could take on any value between 150 and 250 pounds. For a random variable, a distribution can be defined which is explained below.

Probability Distribution of Random Variables

The probability distribution of a random variable is a list of all possible outcomes of some experiment and the probability associated with each outcome in a given sample space S. It shows how the total probability of 1 is distributed over the possible values of an outcome of some experiment. Mathematically, the probability function or the probability distribution of the random variable X assigns a real number to each value x_i that can be assumed by that random variable X such that

$$0 \leq P(x_i) \leq 1 \text{ for every event } x_i$$

- $P(S) = 1$
- The symbol P in $P(x_i)$ denotes a function rule according to which probabilities values are assigned to each and every event in S.

A probability distribution may be represented in tabular form, as a graph, or in the form of a

mathematical equation

Discrete Random Variable

The distribution which shows all the possible numerical values (within the range of a discrete random variable pertaining to a particular experiment) and their respective probabilities is called a Discrete Probability Distribution or Probability mass function. Discrete probability distribution assigns a probability to each of the possible outcomes X of a random experiment, survey, or procedure of statistical inference.

The Discrete Probability Distribution satisfies the following two properties

Probability of X lies between 0 and 1 inclusive., ie,

$$0 \leq P(x) \leq 1$$

- $\sum P(x) = 1$

Different types of discrete probability distributions are:

Binomial distribution

- Poisson distribution
- Geometric distribution
- Hypergeometric distribution
- Multinomial distribution
- Negative binomial distribution

The first four distributions have been discussed in detail in the sixth module.

Example: Suppose you flip a coin two times. This simple statistical experiment can have four outcomes HH, HT, TH, TT. Now the random variable X represent the number of heads that result from this experiment. The random variable X can only take the values 0,1 or 2, so it is discrete random variable.

The table depicts a discrete probability distribution because it relates each value of a discrete random variable with its probability of occurrence.

Example: Suppose we have a randomly selected family which has 3 children. What will be the probability distribution of number of boys.

Let us assume that X represent the number of boys in the family. Then $x = 0, 1, 2$, or 3 .

The column labeled "Probability" identifies the probability of the particular outcome, assuming a boy and girl are equally likely and independent i.e. $P(G)=P(B)=0.5$. Also, the possible 3-child families are mutually exclusive outcomes. Thus the probabilities of the outcomes are as follows

X= Number of boys	OUTCOME	Probability
x=0	GGG	$0.5 \times 0.5 \times 0.5 = 0.125$
x=1	GGB	$0.5 \times 0.5 \times 0.5 = 0.125$
x=1	BGG	$0.5 \times 0.5 \times 0.5 = 0.125$
x=1	GBG	$0.5 \times 0.5 \times 0.5 = 0.125$
x=2	BBG	$0.5 \times 0.5 \times 0.5 = 0.125$
x=2	BGB	$0.5 \times 0.5 \times 0.5 = 0.125$
x=2	GBB	$0.5 \times 0.5 \times 0.5 = 0.125$
x=3	BBB	$0.5 \times 0.5 \times 0.5 = 0.125$

Probability Mass Functions $P(x)$ summarizes the probability that the above discrete random variable X will take on certain values x $[0,1,2,3]$ which can be described by a formula as:

$$P(x) = P(X = x)$$

So:

$$P(X=0) = P(GGG) = 0.125$$

$$P(X=1) = P(BGG \text{ or } GBG \text{ or } GGB) = P(BGG)+P(GBG) +P(GGB) = 0.375$$

$$P(X=2) = P(BBG \text{ or } GBB \text{ or } BGB) = P(BBG)+P(GBB) +P(BGB) = 0.375$$

$$P(X=3) = P(BBB) = 0.125$$
 Probability distribution function of number of boys

X	0	1	2	3
p(x)	0.125	0.375	0.375	0.125

It can be explained by the following figure.

Continuous Probability Distribution

If X is a continuous random variable then a function $f(x)$ describing the probabilities of X is called probability distribution function of X or Probability Density Function.

A **continuous random variable** is a random variable where the data can take infinitely many values. For example, a random variable measuring the time taken for something to be done is continuous since there are an infinite number of possible times that can be taken.

For **any** continuous random variable with probability density function $f(x)$, we have that:

$$\int_{\text{all } x} f(x) dx = 1$$

This is a useful fact.

Example

X is a continuous random variable with probability density function given by $f(x) = cx$ for $0 \leq x \leq 1$, where c is a constant. Find c .

If we integrate $f(x)$ between 0 and 1 we get $c/2$. Hence $c/2 = 1$ (from the useful fact above!), giving $c = 2$.

Cumulative Distribution Function (c.d.f.)

If X is a continuous random variable with p.d.f. $f(x)$ defined on $a \leq x \leq b$, then the **cumulative distribution function** (c.d.f.), written $F(t)$ is given by:

$$F(t) = P(X \leq t) = \int_a^t f(x) dx$$

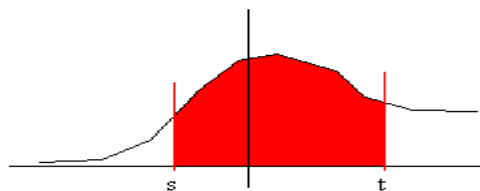
So the c.d.f. is found by integrating the p.d.f. between the minimum value of X and t .

Similarly, the probability density function of a continuous random variable can be obtained by differentiating the cumulative distribution.

The c.d.f. can be used to find out the probability of a random variable being between two values:

$P(s \leq X \leq t)$ = the probability that X is between s and t . But this is equal to the probability that $X \leq t$ minus the probability that $X \leq s$.

[We want the probability that X is in the red area:]



Hence:

- $P(s \leq X \leq t) = P(X \leq t) - P(X \leq s) = F(t) - F(s)$

Expectation and Variance

With discrete random variables, we had that the expectation was $\sum x \cdot P(X = x)$, where $P(X = x)$ was the p.d.f.. It may come as no surprise that to find the expectation of a continuous random variable, we integrate rather than sum, i.e.:

$$E(X) = \int_{\text{all } x} x f(x) dx$$

As with discrete random variables, $\text{Var}(X) = E(X^2) - [E(X)]^2$

Discrete and Continuous Random Variables:

A **variable** is a quantity whose value changes.

A **discrete variable** is a variable whose value is obtained by counting.

Examples: number of students present
 number of red marbles in a jar
 number of heads when flipping three coins
 students' grade level

A **continuous variable** is a variable whose value is obtained by measuring.

Examples: height of students in class
 weight of students in class
 time it takes to get to school
 distance traveled between classes

A **random variable** is a variable whose value is a numerical outcome of a random phenomenon.

- A random variable is denoted with a capital letter
- The probability distribution of a random variable X tells what the possible values of X are and

how probabilities are assigned to those values

- A random variable can be discrete or continuous

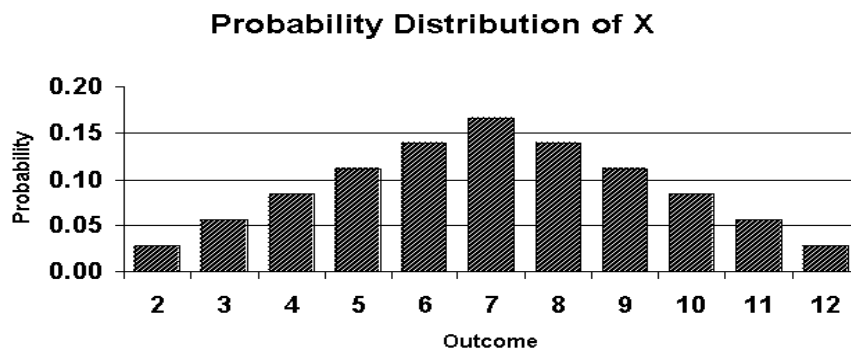
A **discrete random variable** X has a countable number of possible values.

Example: Let X represent the sum of two dice.

Then the probability distribution of X is as follows:

X	2	3	4	5	6	7	8	9	10	11	12
$P(X)$	$\frac{1}{36}$	$\frac{2}{36}$	$\frac{3}{36}$	$\frac{4}{36}$	$\frac{5}{36}$	$\frac{6}{36}$	$\frac{5}{36}$	$\frac{4}{36}$	$\frac{3}{36}$	$\frac{2}{36}$	$\frac{1}{36}$

To graph the probability distribution of a discrete random variable, construct a **probability histogram**.



A **continuous random variable** X takes all values in a given interval of numbers.

- The probability distribution of a continuous random variable is shown by a **density curve**.
- The probability that X is between an interval of numbers is the area under the density curve between the interval endpoints
- The probability that a **continuous random variable** X is exactly equal to a number is zero

Means and Variances of Random Variables:

The mean of a discrete random variable, X , is its weighted average. Each value of X is weighted by its probability.

To find the mean of X , multiply each value of X by its probability, then add all the products.

$$\begin{aligned}\mu_X &= x_1 p_1 + x_2 p_2 + \cdots + x_k p_k \\ &= \sum x_i p_i\end{aligned}$$

The mean of a random variable X is called the **expected value** of X .

Law of Large Numbers:

As the number of observations increases, the mean of the observed values, $\bar{x}$, approaches the mean of the population, μ .

The more variation in the outcomes, the more trials are needed to ensure that $\bar{x}$ is close to μ .

Rules for Means:

If X is a random variable and a and b are fixed numbers, then

$$\mu_{a+bX} = a + b\mu_X$$

If X and Y are random variables, then

$$\mu_{X+Y} = \mu_X + \mu_Y$$

Example:

Suppose the equation $Y = 20 + 100X$ converts a PSAT math score, X, into an SAT math score, Y. Suppose the average PSAT math score is 48. What is the average SAT math score?

$$\mu_X = 48$$

$$\mu_{a+bX} = a + b\mu_X$$

$$\begin{aligned}\mu_{20+100X} &= 20 + 100\mu_X \\ &= 20 + 100(48) \\ &= 500\end{aligned}$$

Example:

Let $\mu_X = 625$ represent the average SAT math score.

Let $\mu_Y = 590$ represent the average SAT verbal score.

$\mu_{X+Y} = \mu_X + \mu_Y$ represents the average combined SAT score.

Then $\mu_{X+Y} = \mu_X + \mu_Y = 625 + 590 = 1215$ is the average combined total SAT score.

The Variance of a Discrete Random Variable:

If X is a discrete random variable with mean μ , then the variance of X is

$$\begin{aligned}\sigma_X^2 &= (x_1 - \mu_X)^2 p_1 + (x_2 - \mu_X)^2 p_2 + \cdots + (x_k - \mu_X)^2 p_k \\ &= \sum (x_i - \mu_X)^2 p_i\end{aligned}$$

The standard deviation (σ_X) is the square root of the variance.

Rules for Variances:

If X is a random variable and a and b are fixed numbers, then

$$\sigma_{a+bX}^2 = b^2 \sigma_X^2$$

If X and Y are independent random variables, then

$$\sigma_{X+Y}^2 = \sigma_X^2 + \sigma_Y^2$$

$$\sigma_{X-Y}^2 = \sigma_X^2 + \sigma_Y^2$$

Example:

Suppose the equation $Y = 20 + 100X$ converts a PSAT math score, X , into an SAT math score, Y . Suppose the standard deviation for the PSAT math score is 1.5 points. What is the standard deviation for the SAT math score?

$$\begin{aligned}\sigma_X^2 &= (1.5)^2 = 2.25 \\ \sigma_{a+bX}^2 &= b^2 \sigma_X^2 \\ \sigma_{20+100X}^2 &= (100)^2 \sigma_X^2 \\ &= (100)^2 (2.25) \\ &= 22,500 \\ \sigma_X &= 150\end{aligned}$$

Suppose the standard deviation for the SAT math score is 150 points, and the standard deviation for the SAT verbal score is 165 points. What is the standard deviation for the combined SAT score?

*** Because the SAT math score and SAT verbal score are not independent, the rule for adding variances does not apply!

Continuous Random Variables - Probability Density Function (PDF)

The **probability density function** or PDF of a continuous random variable gives the relative likelihood of any outcome in a continuum occurring. Unlike the case of discrete random variables, for a continuous random variable any single outcome has probability zero of occurring. The probability density function gives the probability that any value in a *continuous set* of values might occur. Its magnitude therefore encodes the likelihood of finding a continuous random variable near a certain point.

Heuristically, the probability density function is just the distribution from which a continuous random variable is drawn, like the normal distribution, which is the PDF of a normally-distributed continuous random variable.

Definition of the Probability Density Function

The probability that a random variable X takes a value in the (open or closed) interval $[a,b]$ is given by the integral of a function called the **probability density function** $f_X(x)$:

$$P(a \leq X \leq b) = \int_a^b f_X(x) dx. P(a \leq X \leq b) = \int_a^b f_X(x) dx.$$

If the random variable can be any real number, the probability density function is normalized so that:

$$\int_{-\infty}^{\infty} f_X(x) dx = 1. \int_{-\infty}^{\infty} f_X(x) dx = 1.$$

This is because the probability that X takes some value between $-\infty$ and ∞ is one: X does take a value!

These formulas may make more sense in comparison to the discrete case, where the function giving the probabilities of events occurring is called the probability mass function $p(x)$. In the discrete case, the probability of outcome x occurring is just $p(x)$ itself. The probability $P(a \leq X \leq b)$ is given in the discrete case by:

$$P(a \leq X \leq b) = \sum_{a \leq x \leq b} p(x),$$

and the probability mass function is normalized to one so that:

$$\sum_x p(x) = 1, \sum_x p(x) = 1,$$

where the sum is taken over all possible values of x . One can see that the analogous formulas for continuous random variables are identical with the sums promoted to integrals.

The non-normalized probability density function of a certain continuous random variable X is:

$$f(x) = \frac{1}{1+x^2}. f(x) = \frac{1}{1+x^2}.$$

Find the probability that X is greater than one, $P(X > 1)$.

Solution:

First, the probability density function must be normalized. This is done by multiplying by a constant to make the total integral one. Computing the integral:

$$\int_{-\infty}^{\infty} \frac{1}{1+x^2} dx = \arctan(x) \Big|_{-\infty}^{\infty} = \pi.$$

So the normalized PDF is:

$$\tilde{f} = \frac{1}{\pi(1+x^2)}. \tilde{f} = \frac{1}{\pi(1+x^2)}.$$

Computing the probability that X is greater than one,

$$P(X > 1) = \int_1^{\infty} \frac{1}{\pi(1+x^2)} dx = \frac{1}{\pi} \left[\arctan(x) \right]_1^{\infty} = \frac{1}{\pi} \left(\frac{\pi}{2} - \arctan(1) \right) = \frac{1}{\pi} \left(\frac{\pi}{2} - \frac{\pi}{4} \right) = \frac{1}{4}. P(X > 1) = \int_1^{\infty} \frac{1}{\pi(1+x^2)} dx = \frac{1}{\pi} \left[\arctan(x) \right]_1^{\infty} = \frac{1}{\pi} \left(\frac{\pi}{2} - \arctan(1) \right) = \frac{1}{\pi} \left(\frac{\pi}{2} - \frac{\pi}{4} \right) = \frac{1}{4}.$$

Mean and Variance of Continuous Random Variables

Recall that in the discrete case the mean or **expected value** $E(X)$ of a discrete random variable was the weighted average of the possible values x of the random variable:

$$E(X) = \sum x p(x).$$

This formula makes intuitive sense. Suppose that there were n outcomes, equally likely with probability $\frac{1}{n}$ each. The expected value is just the [arithmetic mean](#), $E(X) = \frac{x_1 + x_2 + \dots + x_n}{n}$. In the cases where some outcomes are more likely than others, these outcomes should contribute more to the expected value.

In the continuous case, the generalization is again found just by replacing the sum with the integral and $p(x)$ with the PDF:

$$E(X) = \int_{-\infty}^{\infty} x f(x) dx,$$

assuming the possible values X are the entire real line. If X is constrained instead to $[0, \infty]$ or some other continuous interval, the integral limits should be changed accordingly.

The variance is defined identically to the discrete case:

$$\text{Var}(X) = E(X^2) - E(X)^2.$$

Computing $E(X^2)$ only requires inserting an x^2 instead of an x in the above formulae:

$$E(X^2) = \int_{-\infty}^{\infty} x^2 f(x) dx,$$

The mean and the variance of a continuous random variable need not necessarily be finite or exist. Cauchy distributed continuous random variable is an example of a continuous random variable having both mean and variance undefined.

MATHEMATICS FOR ECONOMICS-I (105)

Unit -1

Definition OF SETS

A set is a well defined collection of distinct objects. The objects that make up a set (also known as the elements or members of a set) can be anything: numbers, people, letters of the alphabet, other sets, and so on. Georg Cantor, the founder of set theory, gave the following definition of a set at the beginning of his *Beiträge zur Begründung der transfiniten Mengenlehre*

A set is a gathering together into a whole of definite, distinct objects of our perception or of our thought – which are called elements of the set.

Sets are conventionally denoted with capital letters. Sets A and B are equal if and only if they have precisely the same elements

As discussed below, the definition given above turned out to be inadequate for formal mathematics; instead, the notion of a "set" is taken as an undefined primitive in axiomatic set theory, and its properties are defined by the Zermelo–Fraenkel axioms. The most basic properties are that a set "has" elements, and that two sets are equal (one and the same) if and only if every element of one is an element of the other.

Describing sets

There are two ways of describing, or specifying the members of, a set. One way is by intensional definition, using a rule or semantic description:

A is the set whose members are the first four positive integers.

B is the set of colors of the French flag.

The second way is by extension – that is, listing each member of the set. An extensional definition is denoted by enclosing the list of members in curly brackets:

$C = \{4, 2, 1, 3\}$

$D = \{\text{blue, white, red}\}.$

Every element of a set must be unique; no two members may be identical. (A multi set is a generalized concept of a set that relaxes this criterion.) All set operations preserve this property. The order in which the elements of a set or multi set are listed is irrelevant (unlike for a sequence or tuple). Combining these two ideas into an example

$\{6, 11\} = \{11, 6\} = \{11, 6, 6, 11\}$

because the extensional specification means merely that each of the elements listed is a member of the

set.

For sets with many elements, the enumeration of members can be abbreviated. For instance, the set of the first thousand positive integers may be specified extensionally as:

$$\{1, 2, 3, \dots, 1000\},$$

where the ellipsis ("...") indicates that the list continues in the obvious way. Ellipses may also be used where sets have infinitely many members. Thus the set of positive even numbers can be written as $\{2, 4, 6, 8, \dots\}$.

The notation with braces may also be used in an intentional specification of a set. In this usage, the braces have the meaning "the set of all ...". So, $E = \{\text{playing card suits}\}$ is the set whose four members are ♠, ♦, ♥, and ♣. A more general form of this is set-builder notation, through which, for instance, the set F of the twenty smallest integers that are four less than perfect squares can be denoted:

$$F = \{n^2 - 4 : n \text{ is an integer; and } 0 \leq n \leq 19\}.$$

In this notation, the colon (":") means "such that", and the description can be interpreted as " F is the set of all numbers of the form $n^2 - 4$, such that n is a whole number in the range from 0 to 19 inclusive." Sometimes the vertical bar ("|") is used instead of the colon.

One often has the choice of specifying a set intensionally or extensionally. In the examples above, for instance, $A = C$ and $B = D$.

Membership

The key relation between sets is *membership* – when one set is an element of another. If a is a member of B , this is denoted $a \in B$, while if c is not a member of B then $c \notin B$. For example, with respect to the sets $A = \{1, 2, 3, 4\}$, $B = \{\text{blue, white, red}\}$, and $F = \{n^2 - 4 : n \text{ is an integer; and } 0 \leq n \leq 19\}$ defined above,

$$4 \in A \text{ and } 12 \in F; \text{ but}$$

$$9 \notin F \text{ and } \text{green} \notin B.$$

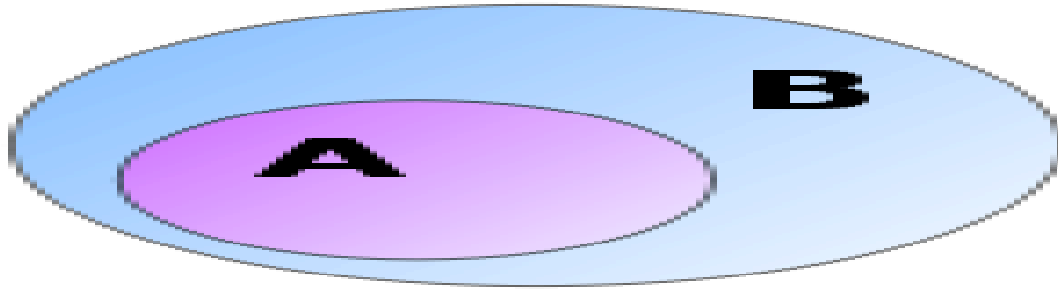
Subsets

If every member of set A is also a member of set B , then A is said to be a *subset* of B , written $A \subseteq B$ (also pronounced *A is contained in B*). Equivalently, we can write $B \supseteq A$, read as *B is a superset of A*, *B includes A*, or *B contains A*. The relationship between sets established by $\subseteq$ is called *inclusion* or *containment*.

If A is a subset of, but not equal to, B , then A is called a *proper subset* of B , written $A \subsetneq B$ (*A is a proper subset of B*) or $B \supsetneq A$ (*B is a proper superset of A*).

Note that the expressions $A \subset B$ and $B \supset A$ are used differently by different authors; some authors use

them to mean the same as $A \subseteq B$ (respectively $B \supseteq A$), whereas other use them to mean the same as $A \subsetneq B$ (respectively $B \supsetneq A$).



A is a **subset** of B

Example:

- The set of all men is a proper subset of the set of all people.
- $\{1, 3\} \subsetneq \{1, 2, 3, 4\}$.
- $\{1, 2, 3, 4\} \subseteq \{1, 2, 3, 4\}$.

The empty set is a subset of every set and every set is a subset of itself:

- $\emptyset \subseteq A$.
- $A \subseteq A$.

An obvious but useful identity, which can often be used to show that two seemingly different sets are equal:

- $A = B$ if and only if $A \subseteq B$ and $B \subseteq A$.

A partition of a set S is a set of nonempty subsets of S such that every element x in S is in exactly one of these subsets.

Power sets

The power set of a set S is the set of all subsets of S , including S itself and the empty set. For example, the power set of the set $\{1, 2, 3\}$ is $\{\{1, 2, 3\}, \{1, 2\}, \{1, 3\}, \{2, 3\}, \{1\}, \{2\}, \{3\}, \emptyset\}$. The power set of a set S is usually written as $P(S)$.

The power set of a finite set with n elements has 2^n elements. This relationship is one of the reasons for the terminology *power set*. For example, the set $\{1, 2, 3\}$ contains three elements, and the power set shown above contains $2^3 = 8$ elements.

The power set of an infinite (either countable or uncountable) set is always uncountable. Moreover, the power set of a set is always strictly "bigger" than the original set in the sense that there is no way to pair the elements of a set S with the elements of its power set $P(S)$ such that every element of S set is paired with exactly one element of $P(S)$, and every element of $P(S)$ is paired with exactly one element of S .

(There is never a bijection from S onto $P(S)$.)

Every partition of a set S is a subset of the power set of S .

Cardinality

The cardinality $|S|$ of a set S is "the number of members of S ." For example, if $B = \{\text{blue, white, red}\}$, $|B| = 3$.

There is a unique set with no members and zero cardinality, which is called the *empty set* (or the *null set*) and is denoted by the symbol $\emptyset$ (other notations are used; see empty set). For example, the set of all three-sided squares has zero members and thus is the empty set. Though it may seem trivial, the empty set, like the number zero, is important in mathematics; indeed, the existence of this set is one of the fundamental concepts of axiomatic set theory.

Some sets have infinite cardinality. The set $\mathbf{N}$ of natural numbers, for instance, is infinite. Some infinite cardinalities are greater than others. For instance, the set of real numbers has greater cardinality than the set of natural numbers. However, it can be shown that the cardinality of (which is to say, the number of points on) a straight line is the same as the cardinality of any segment of that line, of the entire plane, and indeed of any finite-dimensional Euclidean space.

Special sets

There are some sets that hold great mathematical importance and are referred to with such regularity that they have acquired special names and notational conventions to identify them. One of these is the empty set, denoted $\{\}$ or $\emptyset$. Another is the unit set $\{x\}$, which contains exactly one element, namely x . Many of these sets are represented using blackboard bold or bold typeface. Special sets of numbers include:

- $\mathbf{P}$ or $\mathbb{P}$, denoting the set of all primes: $\mathbf{P} = \{2, 3, 5, 7, 11, 13, 17, \dots\}$.
- $\mathbf{N}$ or $\mathbb{N}$, denoting the set of all natural numbers: $\mathbf{N} = \{1, 2, 3, \dots\}$ (sometimes defined containing 0).
- $\mathbf{Z}$ or $\mathbb{Z}$, denoting the set of all integers (whether positive, negative or zero): $\mathbf{Z} = \{\dots, -2, -1, 0, 1, 2, \dots\}$.
- $\mathbf{Q}$ or $\mathbb{Q}$, denoting the set of all rational numbers (that is, the set of all proper and improper fractions): $\mathbf{Q} = \{a/b : a, b \in \mathbf{Z}, b \neq 0\}$. For example, $1/4 \in \mathbf{Q}$ and $11/6 \in \mathbf{Q}$. All integers are in this set since every integer a can be expressed as the fraction $a/1$ ($\mathbf{Z} \subseteq \mathbf{Q}$).
- $\mathbf{R}$ or $\mathbb{R}$, denoting the set of all real numbers. This set includes all rational numbers, together with all irrational numbers (that is, numbers that cannot be rewritten as fractions, such as $\sqrt{2}$, as well as transcendental numbers such as π , e and numbers that cannot be defined).

- $\mathbf{C}$ or $\mathbb{C}$, denoting the set of all complex numbers: $\mathbf{C} = \{a + bi : a, b \in \mathbf{R}\}$. For example, $1 + 2i \in \mathbf{C}$.
- $\mathbf{H}$ or $\mathbb{H}$, denoting the set of all quaternions: $\mathbf{H} = \{a + bi + cj + dk : a, b, c, d \in \mathbf{R}\}$. For example, $1 + i + 2j - k \in \mathbf{H}$.

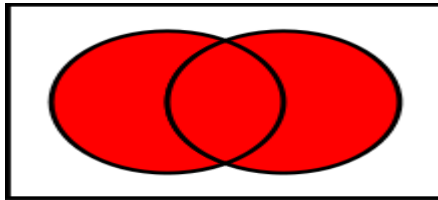
Positive and negative sets are denoted by a superscript - or +, for example: $\mathbb{Q}^+$ represents the set of positive rational numbers.

Each of the above sets of numbers has an infinite number of elements, and each can be considered to be a proper subset of the sets listed below it. The primes are used less frequently than the others outside of number theory and related fields.

Basic operations

There are several fundamental operations for constructing new sets from given sets.

Unions



The **union** of A and B , denoted $A \cup B$

Two sets can be "added" together. The *union* of A and B , denoted by $A \cup B$, is the set of all things that are members of either A or B .

Examples:

- $\{1, 2\} \cup \{1, 2\} = \{1, 2\}$.
- $\{1, 2\} \cup \{2, 3\} = \{1, 2, 3\}$.

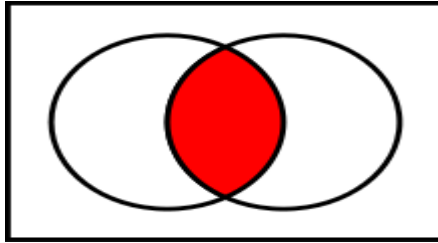
Some basic properties of unions:

- $A \cup B = B \cup A$.
- $A \cup (B \cup C) = (A \cup B) \cup C$.
- $A \subseteq (A \cup B)$.
- $A \cup A = A$.
- $A \cup \emptyset = A$.
- $A \subseteq B$ if and only if $A \cup B = B$.

Intersections

A new set can also be constructed by determining which members two sets have "in common". The

intersection of A and B , denoted by $A \cap B$, is the set of all things that are members of both A and B . If $A \cap B = \emptyset$, then A and B are said to be *disjoint*.



The **intersection** of A and B , denoted $A \cap B$.

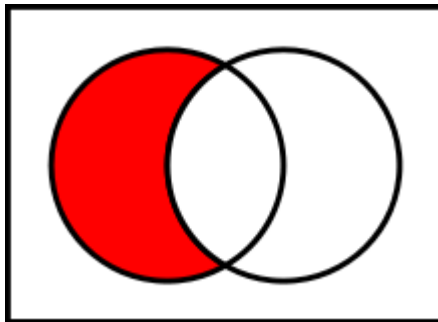
Examples:

- $\{1, 2\} \cap \{1, 2\} = \{1, 2\}$.
- $\{1, 2\} \cap \{2, 3\} = \{2\}$.

Some basic properties of intersections:

- $A \cap B = B \cap A$.
- $A \cap (B \cap C) = (A \cap B) \cap C$.
- $A \cap B \subseteq A$.
- $A \cap A = A$.
- $A \cap \emptyset = \emptyset$.
- $A \subseteq B$ if and only if $A \cap B = A$.

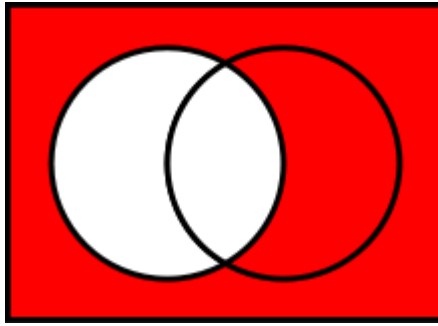
Complements



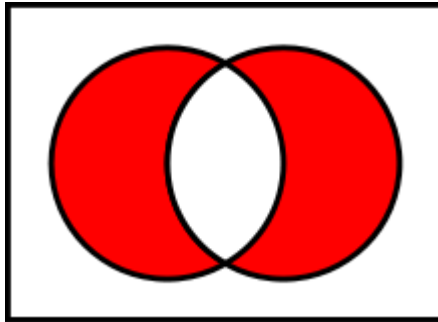
The
of B in A

relative

complement



The **complement** of A in U



The **symmetric difference** of A and B

Two sets can also be "subtracted". The *relative complement* of B in A (also called the *set-theoretic difference* of A and B), denoted by $A \setminus B$ (or $A - B$), is the set of all elements that are members of A but not members of B . Note that it is valid to "subtract" members of a set that are not in the set, such as removing the element *green* from the set $\{1, 2, 3\}$; doing so has no effect.

In certain settings all sets under discussion are considered to be subsets of a given universal set U . In such cases, $U \setminus A$ is called the *absolute complement* or simply *complements* of A , and is denoted by A' .

Examples:

- $\{1, 2\} \setminus \{1, 2\} = \emptyset$.
- $\{1, 2, 3, 4\} \setminus \{1, 3\} = \{2, 4\}$.
- If U is the set of integers, E is the set of even integers, and O is the set of odd integers, then $U \setminus E = E' = O$.

Some basic properties of complements:

- $A \setminus B \neq B \setminus A$ for $A \neq B$.
- $A \cup A' = U$.
- $A \cap A' = \emptyset$.
- $(A')' = A$.

- $A \setminus A = \emptyset$.
- $U' = \emptyset$ and $\emptyset' = U$.
- $A \setminus B = A \cap B'$.

An extension of the complement is the symmetric difference, defined for sets A, B as

$$A \Delta B = (A \setminus B) \cup (B \setminus A).$$

For example, the symmetric difference of $\{7,8,9,10\}$ and $\{9,10,11,12\}$ is the set $\{7,8,11,12\}$.

Cartesian product

A new set can be constructed by associating every element of one set with every element of another set. The *Cartesian product* of two sets A and B , denoted by $A \times B$ is the set of all ordered pairs (a, b) such that a is a member of A and b is a member of B .

Examples:

- $\{1, 2\} \times \{\text{red, white}\} = \{(1, \text{red}), (1, \text{white}), (2, \text{red}), (2, \text{white})\}$.
- $\{1, 2\} \times \{\text{red, white, green}\} = \{(1, \text{red}), (1, \text{white}), (1, \text{green}), (2, \text{red}), (2, \text{white}), (2, \text{green})\}$.
- $\{1, 2\} \times \{1, 2\} = \{(1, 1), (1, 2), (2, 1), (2, 2)\}$.

Some basic properties of cartesian products:

- $A \times \emptyset = \emptyset$.
- $A \times (B \cup C) = (A \times B) \cup (A \times C)$.
- $(A \cup B) \times C = (A \times C) \cup (B \times C)$.

Let A and B be finite sets. Then

- $|A \times B| = |B \times A| = |A| \times |B|$.

Applications

Set theory is seen as the foundation from which virtually all of mathematics can be derived. For example, structures in abstract algebra, such as groups, fields and rings, are sets closed under one or more operations.

One of the main applications of naive set theory is constructing relations. A relation from a domain A to a [codomain](#) B is a subset of the Cartesian product $A \times B$. Given this concept, we are quick to see that the set F of all ordered pairs (x, x^2) , where x is real, is quite familiar. It has a domain set $\mathbf{R}$ and a codomain set that is also $\mathbf{R}$, because the set of all squares is subset of the set of all reals. If placed in functional notation, this relation becomes $f(x) = x^2$. The reason these two are equivalent is for any given value, y that the function is defined for, its corresponding ordered pair, (y, y^2) is a member of the set F .

Definition of a relation.

We still have not given a formal definition of a relation between sets X and Y . In fact the above way of thinking about relations is easily formalized, as was suggested in class by Adam Osborne: namely, we can think of a relation R as a function from $X \times Y$ to the two-element set $\{\text{TRUE}, \text{FALSE}\}$. In other words, for $(x, y) \in X \times Y$,

we say that xRy if and only if $f((x, y)) = \text{TRUE}$.

Properties of relations.

Let X be a set. We now consider various properties that a relation R on X – i.e., $R \subseteq X \times X$ may or may not possess.

(R1) Reflexivity: for all $x \in X$, $(x, x) \in R$. In other words, each element of X bears relation R to itself. Another way to say this is that the relation R contains the equality relation. Exercise X.X: Go back and decide which of the relations in Examples X.X above are reflexive. For instance, set membership is certainly not necessarily reflexive: $1 \notin 1$ (and in more formal treatments of set theory, a set containing itself is usually explicitly prohibited), but $\subseteq$ is reflexive

(R2) Symmetry: for all $x, y \in X$, if $(x, y) \in R$, then $(y, x) \in R$. Again, this has a geometric interpretation in terms of symmetry across the diagonal $y = x$. For instance, the relation associated to the function $y = 1/x$ is symmetric since interchanging x and y changes nothing, whereas the relation associated to the function $y = x^2$ is not. (Looking ahead a bit, a function $y = f(x)$ is symmetric iff it coincides with its own inverse function.) Exercise X.X: Which of the relations in Examples X.X above are symmetric?

(R3) Anti-Symmetry: for all $x, y \in X$, if $(x, y) \in R$ and $(y, x) \in R$, then $x = y$.

For instance, $\subseteq$ satisfies anti-symmetry.

Exercise X.X: Which of the relations in Examples X.X above are anti-symmetric?

(R4) Transitivity: for all $x, y, z \in X$, if $(x, y) \in R$ and $(y, z) \in R$, then $(x, z) \in R$.

For instance, “being a parent of” is not transitive, but “being an ancestor of” is transitive.

Definition: An equivalence relation on a set X is a relation on X which is reflexive, symmetric and

transitive.

Examples of equivalence relations.

Let n be a positive integer. Then there is a natural partition of $\mathbb{Z}$ into n parts which generalizes the partition into even and odd. Namely, we put $Y_1 = \{\dots, -2n, -n, 0, n, 2n, \dots\} = \{kn \mid k \in \mathbb{Z}\}$ the set of all multiples of n , $Y_2 = \{\dots, -2n+1, -n+1, 1, n+1, 2n+1, \dots\} = \{kn+1 \mid k \in \mathbb{Z}\}$, and similarly, for any $0 \leq d \leq n-1$, we put $Y_d = \{\dots, -2n+d, -n+d, d, n+d, 2n+d, \dots\} = \{kn+d \mid k \in \mathbb{Z}\}$. That is, Y_d is the set of all integers which, upon division by n , leave a remainder of d . Earlier we showed that the remainder upon division by n is a well-defined integer in the range $0 \leq d < n$. Here by “well-defined”, I mean that for $0 \leq d_1 \neq d_2 < n$, the sets Y_{d_1} and Y_{d_2} are disjoint. Recall why this is true: if not, there exist k_1, k_2 such that $k_1n + d_1 = k_2n + d_2$, so $d_1 - d_2 = (k_2 - k_1)n$, so $d_1 - d_2$ is a multiple of n . But $-n < d_1 - d_2 < n$, so the only multiple of n it could possibly be is 0, i.e., $d_1 = d_2$. It is clear that each Y_d is nonempty and that their union is all of $\mathbb{Z}$, so $\{Y_d\}_{d=0}^{n-1}$ gives a partition of $\mathbb{Z}$. The corresponding equivalence relation is called congruence modulo n , and written as follows: $x \equiv y \pmod{n}$. What this means is that x and y leave the same remainder upon division by n .

definition

A (non-strict) **partial order** is a binary relation " $\leq$ " over a set P which is reflexive, anti symmetric, and transitive, i.e., which satisfies for all a, b , and c in P :

- $a \leq a$ (reflexivity);
- if $a \leq b$ and $b \leq a$ then $a = b$ (anti symmetry);
- if $a \leq b$ and $b \leq c$ then $a \leq c$ (transitivity).

In other words, a partial order is an anti symmetric preorder.

A set with a partial order is called a **partially ordered set** (also called a **poset**). The term *ordered set* is sometimes also used for posets, as long as it is clear from the context that no other kinds of orders are meant. In particular, totally ordered sets can also be referred to as "ordered sets", especially in areas where these structures are more common than posets.

For a, b , elements of a partially ordered set P , if $a \leq b$ or $b \leq a$, then a and b are **comparable**. Otherwise they are **incomparable**. In the figure on top-right, e.g. $\{x\}$ and $\{x,y,z\}$ are comparable, while $\{x\}$ and $\{y\}$ are not. A partial order under which every pair of elements is comparable is called a **total order** or **linear order**; a totally ordered set is also called a **chain** (e.g., the natural numbers with their standard order). A subset of a poset in which no two distinct elements are comparable is called an **antichain** (e.g.

the set of singletons $\{\{x\}, \{y\}, \{z\}\}$ in the top-right figure). An element a is said to be **covered** by another element b , written $a <: b$, if a is strictly less than b and no third element c fits between them; formally: if both $a \leq b$ and $a \neq b$ are true, and $a \leq c \leq b$ is false for each c with $a \neq c \neq b$. A more concise definition will be given below using the strict order corresponding to " $\leq$ ". For example, $\{x\}$ is covered by $\{x, z\}$ in the top-right figure, but not by $\{x, y, z\}$.

Standard examples of posets arising in mathematics include:

- The real numbers ordered by the standard *less-than-or-equal* relation $\leq$ (a totally ordered set as well).
- The set of subsets of a given set (its power set) ordered by inclusion (see the figure on top-right). Similarly, the set of sequences ordered by subsequence, and the set of strings ordered by substring.
- The set of natural numbers equipped with the relation of divisibility.
- The vertex set of a directed acyclic graph ordered by [reach ability](#).
- The set of subspaces of a vector space ordered by inclusion.
- For a partially ordered set P , the sequence space containing all sequences of elements from P , where sequence a precedes sequence b if every item in a precedes the corresponding item in b . Formally, $(a_n)_{n \in \mathbb{N}} \leq (b_n)_{n \in \mathbb{N}}$ if and only if $a_n \leq b_n$ for all n in $\mathbb{N}$.
- For a set X and a partially ordered set P , the function space containing all functions from X to P , where $f \leq g$ if and only if $f(x) \leq g(x)$ for all x in X .
- A fence, a partially ordered set defined by an alternating sequence of order relations $a < b > c < d$
- ...

Function

Consider the relation

$$f : \{(a, 1), (b, 2), (c, 3), (d, 5)\}$$

In this relation we see that each element of A has a unique image in B . This relation f from set A to B where every element of A has a unique image in B is defined as a function from A to B . So we observe that ***in a function no two ordered pairs have the same first element.***

Domain and Range:-

We also see that $\square \square$ an element $\square \square B$, i.e., 4 which does not have its preimage in A . Thus here:

(i) the set B will be termed as co-domain and

(ii) the set $\{1, 2, 3, 5\}$ is called the range.

From the above we can conclude that **range is a subset of co-domain**. Symbolically, this function can be written as

$$f : A \rightarrow B \text{ or } A \rightarrow f \rightarrow B$$

Types of functions :-

One-to-one:- Let f be a function from A to B . If every element of the set B is the image of at least one element of the set A i.e. if there is no unpaired element in the set B then we say that the function f maps the set A onto the set B . Otherwise we say that the function maps the set A into the set B . Functions for which each element of the set A is mapped to a different element of the set B are said to be **one-to-one**.

Many-to-one:- A function can map more than one element of the set A to the same element of the set B . Such a type of function is said to be **many-to-one**.

Reciprocal Function/Inverse function:-

Functions of the type $y = 1/x$, $x \neq 0$, called a reciprocal function.

Composite function:- Consider the two functions given below:

$$y = 2x + 1, x = 1, 2, 3$$

$$z = y + 1, y = 3, 5, 7$$

Then z is the composition of two functions x and y because z is defined in terms of y and y in terms of x .

Graphically one can represent this as given below :

HASHING FUNCTION

To save space and time, each record stored in a computer is assigned an address (memory location) in the computer's memory. The task of assigning the address is performed by the Hashing function (or Hash function) $H : K \rightarrow L$, which maps the set K of keys to the set L of memory addresses. Thus a Hashing function provides means of identifying records stored in a table. The function H should be one-to-one. In fact, if $k_1 \neq k_2$ implies $H(k_1) \neq H(k_2)$, then two keys will have same address and we say that

collision occurs. To resolve collisions, the following methods are used to define the hash function.

1. **Division Method.** In this method, we restrict the number of addresses to a fixed number (generally a prime) say m and define the hash function $H : K \rightarrow L$ by

$$H(k) = k \pmod{m}, k \in K,$$

where $k \pmod{m}$ denotes the remainder when k is divided by m .

2. **Midsquare Method.** As the name suggest, we square the key in this method and define hash function $H : K \rightarrow L$ by $H(k) = l$, where l is obtained by deleting digits from both ends of k^2 .
3. **Folding Method.** In this method the key k is partitioned into a number of parts, where each part, except possibly the last, has the same numbers of digits as the required memory address. Thus, if $k = k_1 + k_2 + \dots + k_n$, then the hash function $H : K \rightarrow L$ is defined by

$$H(k) = k_1 + k_2 + \dots + k_n, \text{ where the last carry has been ignored.}$$

Definition of recursive functions:- The class of **recursive functions** is defined as follows: The functions s and z are recursive, and so are all projections p^k . Functions built from recursive ones by using composition C_n or primitive recursion Pr are also recursive. Functions built from recursive ones by minimization M_n are also recursive

UNIT -2 (PARTIAL ORDER RELATIONS AND LATTICES)

Partial Order Relations on a Lattice:-

A partial order relation on a lattice (L) follows as a consequence of the axioms for the binary operations $\vee$ and $\wedge$.

PARTIALLY ORDERED SETS

A relation R on a set X is said to be anti symmetric if $a R b$ and $b R a$ imply $a = b$. Relation R on a set X is called a partial order relation if it is reflexive, anti-symmetric and transitive. A set X with the partial order R is called a partially ordered set or poset and is denoted by (X, R)

EXAMPLE

Let $\tilde{A}$ be a collection of subsets of a set S . The relation $\subseteq$ of set inclusion is a partial order relation on $\tilde{A}$.

In fact, if $A, B, C \in \tilde{A}$, then,

$A \subseteq A$, that is, A is a subset of itself which is true.

If $A \subseteq B, B \subseteq A$, then $A = B$

If $A \subseteq B, B \subseteq C$, then A is a subset of C , that is, $A \subseteq C$.

HASSE DIAGRAM

Let A be a finite set. By the theorem proved above, the digraph of a partial order on A has only cycles of length 1. In fact, since a partial order is reflexive, every vertex in the digraph of the partial order is contained in a cycle of length 1. To simplify the matter, we shall delete all such cycles from the digraph.

We also eliminate all edges that are implied by transitivity property. Thus, if $a \leq b, b \leq c$, it follows that $a \leq c$. In this case, we omit the edge from a to c . We also agree to draw the digraph of partial order with all edges pointing upward, omit the arrows and to replace then the circles by dots

“The diagram of a partial order obtained from its digraph by omitting cycles of length 1, the edges implied by transitivity and the arrows (after arranging them pointing upward) is called Hasse diagram of the partial order of the poset”.

EXAMPLE

Let $A = \{1, 2, 3, 4, 12\}$. Consider the partial order of divisibility on A . That is, if a and b are in A , $a \leq b$ if and only if $a \mid b$

Therefore, the Hasse diagram of the poset $(A, \leq)$ is as shown in Figure 1.18.

EXAMPLE

Let $S = \{a, b, c\}$ and $\tilde{A} = P(S)$ (power set of S).

Consider the partial order of set inclusion ($\subseteq$). We note that

$$\tilde{A} = P(S) = \{ \phi, \{a\}, \{b\}, \{c\}, \{a, b\}, \{a, c\}, \{b, c\}, \{a, b, c\} \}$$

Then the Hasse diagram of the poset $(\tilde{A}, \subseteq)$

Hasse diagram of a finite linearly ordered set is always of the form and thus consists of simply one path. Hence diagram of a totally order set is a chain.

Hasse Diagram of Dual Poset

If $(A, \leq)$ is a poset and $(A, \geq)$ is the dual poset, then the Hasse diagram of $(A, \geq)$ is just the Hasse diagram of $(A, \leq)$ turned upside down.

For example, let $A = \{a, b, c, d, e, f\}$ and let be the Hasse diagram of poset $(A, \leq)$. Then the Hasse diagram of dual poset $(A, \geq)$ is which can be constructed by turning the Hasse diagram of $(A, \leq)$ upside down.

EXAMPLE Let $A = \{a, b, c, d, e\}$. Then the Hasse diagram defines a partial order on B in the natural way. That is, $d \leq b$, $d \leq a$, $e \leq c$ and so on.

EXAMPLE Let n be a positive integer and D_n denote the set of all divisors of n . Considering the partial order of divisibility in D_n , draw Hasse diagram D_{24} , D_{30} and D_{36} .

Solution.

We know that

$$D_{24} = \{1, 2, 3, 4, 6, 8, 12, 24\},$$

$$D_{30} = \{1, 2, 3, 5, 6, 10, 15, 30\},$$

$$D_{36} = \{1, 2, 3, 4, 6, 9, 12, 18, 36\}.$$

Therefore, the Hasse diagram of D_{24} , D_{30} and D_{36}

$$\{5\}, \{3, 2\}, \{2, 2, 1\}, \{1, 1, 1, 1, 1\}, \{4, 1\}, \{3, 1, 1\}, \{2, 1, 1, 1\}.$$

We order the partitions of an integer m as follows:

A partition P_1 precedes a partition P_2 if the integers in P_1 can be added to obtain integers in P_2 or we can say that if the integers in P_2 can be further subdivided to obtain the integers in P_1 . For example, $\{1, 1, 1, 1\}$ precedes $\{2, 1, 1, 1\}$. On the other hand, $\{3, 1, 1\}$ and $\{2, 2, 1\}$ are non-comparable.

The Hasse diagram of the partition of $m = 5$ is

Let A be a (non-empty) linearly ordered alphabet. Then Kleene closure of A consists of all words w on A and is denoted by A^* .

Also then $|w|$ denotes the length of w .

We have following two order relations on A^* .

Alphabetical (Lexicographical) order: In this order we have

$\lambda < w$, where λ is empty word and w is any non-empty word.

Suppose $u = a u'$ and $v = b v'$ are distinct non-empty words where $a, b \in A$ and $u', v' \in A^*$. Then,

$u < v$ if $a < b$ or if $a = b$ but $u' < v'$ Short-lex order: Here A^* is ordered first by length and then

alphabetically, that is, for any distinct words u, v , in A^* , $u < v$ if $|u| < |v|$ or if $|u| = |v|$ but u precedes v alphabetically. For example, “to” proceeds “and” since $|to| = 2$ but $|and| = 3$. Similarly, “an” precedes “to” since they have the same length but “an” precedes “to” alphabetically.

This order is also called free semi-group order.

Let A be a partially ordered set with respect to a relation $\leq$. An element a in A is called a maximal element of A if and only if for all b in A , either $b \leq a$ or b and a are not comparable.

An element a in A is called greatest element of A if and only if for all b in A , $b \leq a$.

An element a in A is called minimal element of A if and only if for all b in A , either $a \leq b$ or b and a are not comparable.

An element a in A is called a least element of A if and only if for all b in A , $a \leq b$.

A greatest element is maximal but a maximal element need not be greatest element. Similarly, a least element is minimal but a minimal element need not be a least element.

The elements a_1, a_2 and a_3 are maximal elements of A , and the elements b_1, b_2 and b_3 are minimal elements. Observe that since there is no line between b_2 and b_3 we can conclude that neither $b_3 \leq b_2$ nor $b_2 \leq b_3$ showing that b_2 and b_3 are not comparable.

Let A be the poset of non-negative real numbers with usual partial order $\leq$ (read as “less than or equal to”). Then 0 is the minimal element of A . There is no maximal element of A .

Let A be a finite non-empty poset with partial order $\leq$. Then A has at least one maximal element and at least one minimal element.

Let $(A, \leq)$ be a poset and B a subset of A . An element $a \in A$ is called a least upper bound (supremum) of B if

a is an upper bound of B , that is, $b \leq a \forall b \in B$

$a \leq a'$ whenever a' is an upper bound of B .

An element $a \in A$ is called a greatest lower bound (infimum) of B if

a is a lower bound of B , that is, $a \leq b \forall b \in B$

$a' \leq a$ whenever a' is a lower bound of B .

Further, upper bounds in the poset $(A, \leq)$ correspond to lower bounds in the dual poset $(A, \geq)$ and the lower bounds in $(A, \leq)$ correspond to upper bound in $(A, \geq)$.

Similar statements also hold for greatest lower bounds and least upper bounds.

Consider Example 1.69 above:

Since $B_1 = \{a, b\}$ has no lower bound, it has no greatest lower bound. However,

$$\text{lub}(B_1) = c$$

Since the lower bounds of $B_2 = \{c, d, e\}$ are c, a and b , we have

$$\text{glb}(B_2) = c$$

The upper bounds of B_2 are f, g, h . Since f and g are not comparable, we conclude that B_2 has no least upper bound.

(Here d and e are not upper bounds of $\{c, d, e\}$ because $d \not\geq e \in B_2$ and $e \not\geq d \in B_2$)

EXAMPLE

Let $A = \{1, 2, 3, 4, 5, \dots, 11\}$ be the poset whose Hasse diagram is given (Figure 1.32).

Figure Find lub and glb of $B = \{6, 7, 10\}$, if they exist.

Solution.

Exploring all upward paths from $6, 7$ and 10 we find that $\text{lub}(B) = 10$. Similarly, by examining all downward paths from $6, 7$ and 10 , we find that $\text{glb}(B) = 4$.

EXAMPLE Let D_n denote the set of factors of a positive integer n partially ordered by divisibility. Then,

Let $\leq$ and $\leq'$ be two partial order relations on a set A . Then $\leq'$ is said to be compatible with $\leq$ if and only if

$$a \leq b \Rightarrow a \leq' b \text{ for all } a, b \in A.$$

The process of constructing a linear order (total order) which is compatible to a given partial order on a given set is called topological sorting.

The construction of a topological sorting for a general finite partially order set is based on the fact that any partially ordered set that is finite and non-empty has a minimal element.

To create a total order for a partially ordered set $(A, \leq)$, we proceed as follows:

Pick any minimal element and make it number one. Let this element be a .

Consider $A - \{a\}$. It is a subset of A and so is partially ordered. If it is empty, stop the process. If not,

pick a minimal element from it and call it element number 2. Let it be b.

Consider $A - \{\{a\}, \{b\}\}$. If this set is empty, stop the process. If not, pick a minimal element and call it number 3. Continue in this way until all the elements of the set have been used up.

We now give algorithm to construct a topological sorting for a relation $\leq$ defined on a non-empty finite set A.

Then, remove one of 4 and 18. If we remove 18, we get

Figure

total order : $3 \leq 2 \leq 6^* \leq 18$

Then,

$A = (A - \{3, 2, 6, 18\})$

Now minimal element is 4. We remove it and we get

$A = A - \{3, 2, 6, 18, 4\} = \{24\}$

total order : $3 \leq 2 \leq 6 \leq 18 \leq 4 \leq 24$

(Hasse diagram of $A - \{d\}$)

The minimal element of $A - \{d\}$ is e and we put e in sort [2]. The Hasse diagram of $A - \{d\}$, (Hasse diagram of $A - \{d, c\}$).

The minimal element of $A - \{d, e\}$ is c and we put e in sort [3]. The Hasse diagram of $A - \{d, e, c\}$ is as shown below .The minimal element of $A - \{d, e, c\}$ are a and b. We pick b and put it in sort [4]. The Hasse diagram of $A - \{d, e, c, b\}$ is shown below: (Hasse diagram of $A - \{d, e, c, b\}$).

The minimal element of $A - \{d, e, c, b\}$ is a and we put it in sort [5]. The topological sorting of $(A, \leq)$ is therefore $(A, <)$, where total order : $d < e < c < b < a$ and the Hasse diagram of $(A, <)$ is as shown in the

Let N be the set of positive integers. Then the usual relation $\leq$ (read “less than or equal to”) is a partial order on N.

Similarly, $\geq$ (read “greater than or equal to”) is a partial order on N.

But the relation $<$ (read “less than”) is not a partial order on N. In fact, this relation is not reflexive.

EXAMPLE Let N be the set of positive integers. Then the relation of divisibility is a partial order on N.

We say that “a divides b” written as $a \mid b$, if there exists an integer c such that $a \mid b$. We note that for a, b, c $\in$ N.

$$a \mid a$$

$$a \mid b, b \mid a \Rightarrow a = b$$

$$a \mid b, b \mid c \Rightarrow a \mid c.$$

Thus, relation of divisibility is a partial order on $\mathbb{N}$.

EXAMPLE

The relation of divisibility is not a partial order on the set of integers. For example, $3 \mid -3$, $-3 \mid 3$ but $3 \neq -3$, that is, the relation is not anti-symmetric and so cannot be partial order.

EXAMPLE

If R is a partial order on A , then R^{-1} (inverse relation) is also a partial order.

We know that if R is a relation on A , then

$$R^{-1} = \{(b, a) : (a, b) \in R\}, a, b \in A.$$

Since R is a partial order relation,

$$a R a \quad \forall a \in A$$

$$a R b, b R a \Rightarrow a = b$$

$$a R b, b R c \Rightarrow a R c.$$

We observe that

(i) Since R is a relation, $(a, a) \in R \quad \forall a \in A$

$$\Rightarrow (a, a) \in R^{-1}$$

$$\Rightarrow a R^{-1} a.$$

Thus the relation R^{-1} is reflexive.

(ii) If $(b, a) \in R^{-1}$ and $(a, b) \in R^{-1}$, then

$$(a, b) \in R \text{ and } (b, a) \in R,$$

$$\Rightarrow a R b \text{ and } b R a,$$

$$\Rightarrow a = b, \text{ since } R \text{ is anti-symmetric.}$$

Hence, R^{-1} is anti-symmetric.

(iii) If $(b, a) \in R^{-1}$ and $(c, b) \in R^{-1}$, then

$$(a, b) \in R \text{ and } (b, c) \in R,$$

$$\Rightarrow (a, c) \in R, \text{ since } R \text{ is transitive}$$

$$\Rightarrow (c, a) \in R^{-1}.$$

Thus $(c, b) \in R^{-1}$ and $(b, a) \in R^{-1} \Rightarrow (c, a) \in R^{-1}$ and so R^{-1} is transitive.

Hence R^{-1} is a partial order.

The poset (A, R^{-1}) is called the dual of the poset (A, R) and the partial order R^{-1} is called the dual of the partial order R .

A relation R on a set A is said to be quasi order if

R is irreflexive, that is, $(a, a) \notin R$ for any $a \in A$

R is transitive, that is, $a R b, b R c \Rightarrow a R c$ for $a, b, c \in A$.

Let (A, R) be a poset. The elements a and b of A are said to be comparable if $a R b$ or $b R a$.

We know that the relation of divisibility is a partial order on the set of natural numbers. But we see that

$3 \nmid 7$ and $7 \nmid 3$.

Thus, 3 and 7 are positive integers in $\mathbb{N}$ which are not comparable (In such a case we write $3 \parallel 7$).

If every pair of elements in a poset (A, R) is comparable, we say that A is linearly ordered (totally ordered or a chain). The partial order is then called linear order or total ordering relation. The number of elements in a chain is called the length of the chain.

Let A be a set with two or more elements and let $\subseteq$ (set inclusion) be taken as the relation on the subsets of A . If a and b are two elements of A , then $\{a\}$ and $\{b\}$ are subsets of A but they are not comparable. Hence $P(A)$ is not a chain. A subset of A is called Antichain if no two distinct elements in the subsets are related.

But, if we consider the subsets ϕ , $\{a\}$ and A of A , then this collection (subsets $\{\phi, \{a\}, A\}$ of $P(A)$) is a chain because $\phi \subseteq \{a\} \subseteq A$. Similarly, ϕ , $\{b\}$ and A form a chain.

.

Let $(a, b) R'' (a', b')$ and $(a', b') R'' (a, b)$. Then, by definition,

	$a R a'$,	$a' R a$ in A	(i)
	$b R' b'$,	$b' R' b$ in B .	(ii)

Since (A, R) and (B, R') are posets, (i) and (ii) respectively imply

$a = a'$

and

$b = b'$.

Thus, $(a, b) R'' (a', b')$ and $(a', b') R'' (a, b)$ imply

$$(a, b) = (a', b').$$

Hence R'' is anti-symmetric.

Let $(a, b) R'' (a', b')$ and $(a', b') R'' (a'', b'')$,
where $a, a', a'' \in A$ and $b, b', b'' \in B$. Then

	$a R a'$ and $a' R a''$	(iii)
	$b R' b'$ and $b' R' b''$.	(iv)

By transitivity of R , (iii) gives

$$a R a'',$$

while (iv) yields

$$b R' b''.$$

Hence,

$$(a, b) R'' (a'', b'').$$

Hence R'' is transitive and so $(A \times B, R'')$ is a poset.

The partial order R'' defined on the Cartesian product $A \times B$, as above, is called the Product Partial Order.

Definition 1.51

A relation R on a set A is called asymmetric if $a R b$ and $b R a$ do not both hold for any $a, b \in A$.

Definition 1.52

A transitive, asymmetric relation R is called a Strict Partial Ordering.

Theorem 1.27

If $\leq$ is a partial order of the set A , then a relation $<$ defined by

$$a < b \text{ if } a \leq b \text{ and } a \neq b$$

is a strict partial order of A .

Proof. We shall show that $<$ is transitive and asymmetric.

(i) Transitivity: Suppose that $a < b$ and $b < c$. Then, by definition,

$$a \leq b \text{ and } b \leq c, a \neq b, b \neq c. \quad (1)$$

Since $\leq$ is partial order, it is transitive and so $a \leq c$. It remains to show that $a \neq c$. Suppose on the contrary that $a = c$. Then,

$$c \leq b \text{ (using } a \leq b \text{ from (1)).} \quad (2)$$

From (1) and (2), we have $b \leq c$ and $c \leq b$ and so $b = c$ which is contradiction. Hence,

$$a < b \text{ and } b < c \text{ implies } a < c.$$

Proving that $<$ is transitive.

(ii) Asymmetry: Suppose that $x < y$ and $y < x$ both holds. Therefore,

$$x \leq y \text{ and } y \leq x.$$

Since $\leq$ is anti-symmetric, it follows that $x = y$, which contradicts $x < y$. Hence $x < y$ and $y < x$ cannot both hold. Thus $<$ is asymmetric.

Hence $<$ is strict partial order of A .

Remark 1.1 If $<$ is a strict partial order of A , then the relation $\leq$ defined by $x \leq y$ if $x < y$ or $x = y$ is a partial order of A (can be proved using the definitions).

Definition

A sequence of letters or other symbols, written without commas is called a string. Further,

A string of length p may be considered as an ordered p -tuple.

An infinite string such as $abababab \dots$ may be regarded as the infinite sequence $a, b, a, b, a, b, \dots$

If S is any set with a partial order relation, then the set of strings over S is denoted by S^* .

Definition Let $(A, \leq)$ and $(B, \leq)$ be chains (linearly ordered sets). Then the order relation (which is in fact totally ordered) $<$ on the Cartesian product $A \times B$ defined by

$$(a, b) < (a', b') \text{ if } a < a' \text{ or if } a = a' \text{ and } b \leq b'$$

is called Lexicographic order or Dictionary order.

EXAMPLE

Consider the plane $R^2 = R \times R$. It is linearly ordered by lexicographic order. In fact, each vertical line has usual order (less than or equal to) and points on a line (e.g., $x = a_1$ in Fig. 1.14) are less than any point on a line farther to the right (e.g. $x = a_2$ in Fig. 1.14). Thus the point $p_1(a_1, b_1) < p_2(a_2, b_2)$ because $a_1 < a_2$. Further, $p_2(a_2, b_2) < p_3(a_2, b_3)$ because in this case $a_2 = a_2$ and $b_2 \leq b_3$.

Theorem :-The digraph (directed graph) of a partial order has no cycle of length greater than 1.

Proof. Suppose on the contrary that the digraph of the partial order $\leq$ on the set A contains a cycle of length $n \geq 2$. Then there exist distinct elements $a_1, a_2, \dots, a_n$ such that

$$a_1 \leq a_2, a_2 \leq a_3, \dots, a_{n-1} \leq a_n, a_n \leq a_1.$$

By the transitivity of partial order, used $n - 1$ times, $a_1 \leq a_n$. Thus we have

$$a_1 \leq a_n \text{ and } a_n \leq a_1.$$

Since $\leq$ is partial order, anti-symmetry implies $a_1 = a_n$, which is a contradiction to the assumption that $a_1, a_2, \dots, a_n$ are distinct. Hence the result.

Definition The Transitive closure of a relation R is the smallest transitive relation containing R . It is denoted by R^∞ .

We note that from vertex 1, we have paths to the vertices 2, 3, 4 and 1. Note that path from 1 to 1 proceeds from 1 to 2 to 1. Thus we see that the ordered pairs $(1, 1)$, $(1, 2)$, $(1, 3)$ and $(1, 4)$ are in R^∞ . Starting from vertex 2, we have paths to vertices 2, 1, 3 and 4 so the ordered pairs $(2, 1)$, $(2, 2)$, $(2, 3)$ and $(2, 4)$ are in R^∞ . The only other path is from vertex 3 to 4, so we have

$$R^\infty = \{(1, 1), (1, 2), (1, 3), (1, 4), (2, 1), (2, 2), (2, 3), (2, 4), (3, 4)\}$$

A **Lattice** is an algebraic system $(L, \vee, \wedge)$ with two binary operations $\vee$ and $\wedge$, called **join** and **meet**, respectively, on a non-empty set L which satisfies the following axioms for $a, b, c \in L$:

1. Commutative Law:

$$a \vee b = b \vee a \text{ and } a \wedge b = b \wedge a.$$

2. Associative Law:

$$(a \vee b) \vee c = a \vee (b \vee c)$$

and

$$(a \wedge b) \wedge c = a \wedge (b \wedge c).$$

3. Absorption Law:

- i. $a \vee (a \wedge b) = a,$
- ii. $a \wedge (a \vee b) = a.$

We note that **Idempotent Law follows from axiom 3** above. In fact,

$$\begin{aligned} a \vee a &= a \vee [a \wedge (a \vee b)] \text{ using 3 (ii)} \\ &= a \text{ using 3 (i).} \end{aligned}$$

The proof of $a \wedge a = a$ follows by the principle of duality.

LATTICE

Definition

A **lattice** is a partially ordered set $(L, \leq)$ in which every subset $\{a, b\}$ consisting of two elements has a least upper bound and a greatest lower bound.

We denote $\text{LUB}(\{a, b\})$ by $a \vee b$ and call it **join** or **sum of a and b** . Similarly, we denote $\text{GLB}(\{a, b\})$ by $a \wedge b$ and call it **meet** or **product of a and b** .

Other symbols used are

LUB: $\oplus, +, \cup,$

GLB: $*, \cdot, \cap.$

Thus **Lattice** is a mathematical structure with **two binary operations, join and meet**.

A totally ordered set is obviously a lattice but not all partially ordered sets are lattices.

EXAMPLE Let A be any set and $P(A)$ be its power set. The partially ordered set $(P(A), \subseteq)$ is a lattice in which the meet and join are the same as the operations $\cap$ and $\cup$, respectively. If A has single element, say a , then $P(A) = \{\phi, \{a\}\}$

PROPERTIES OF LATTICES

Let $(L, \leq)$ be a lattice and let $a, b, c \in L$. Then, from the definition of $\vee$ (join) and $\wedge$ (meet) we have

- i. $a \leq a \vee b$ and $b \leq a \vee b$; $a \vee b$ is an upper bound of a and b .
- ii. If $a \leq c$ and $b \leq c$, then $a \vee b \leq c$; $a \vee b$ is the least upper bound of a and b .
- iii. $a \wedge b \leq a$ and $a \wedge b \leq b$; $a \wedge b$ is a lower bound of a and b .
- iv. If $c \leq a$ and $c \leq b$, then $c \leq a \wedge b$; $a \wedge b$ is the greatest lower bound of a and b .

Theorem

Let L be a lattice. Then for every a and b in L ,

- i. $a \vee b = b$ if and only if $a \leq b$,
 - ii. $a \vee b = a$ if and only if $a \leq b$,
 - iii. $a \wedge b = a$ if and only if $a \vee b = b$.
- I. BOUNDED, COMPLEMENTED AND DISTRIBUTIVE LATTICES
 - ii. Let $(L, \vee, \wedge)$ be a lattice and let $S = \{a_1, a_2, \dots, a_n\}$ be a finite subset of L . Then,
 - iii.
 - iv. LUB of S is represented by $a_1 \vee a_2 \vee \dots \vee a_n$,
GLB of S is represented by $a_1 \wedge a_2 \wedge \dots \wedge a_n$.
 - v. **Definition** A lattice is called **complete** if each of its non-empty subsets has a least upper bound and a greatest lower bound.
 - vi. Obviously, every finite lattice is complete.
 - vii. Also, every complete lattice must have a least element, denoted by 0 and a greatest element, denoted by I .
 - viii. The least and greatest elements if exist are called **bound (units, universal bounds)** of the lattice.
 - ix. **Definition** A lattice L is said to be **bounded** if it has a greatest element I and a least element 0 .
 - x. For the lattice $(L, \vee, \wedge)$ with $L = \{a_1, a_2, \dots, a_n\}$,

UNIT -3

DEFINITIONS AND BASIC CONCEPTS

Definition

A **graph** $G = (V, E)$ is a mathematical structure consisting of two finite sets V and E . The elements of V are called **vertices (or nodes)** and the elements of E are called **edges**. Each edge is associated with a set consisting of **either one or two vertices** called its **endpoints**.

The correspondence from edges to endpoints is called **edge-endpoint function**. This function is generally denoted by γ . Due to this function, some authors denote graph by $G = (V, E, \gamma)$.

Definition

A graph consisting of one vertex and no edges is called a **trivial graph**.

Definition

A graph whose vertex and edge sets are empty is called a **null graph**.

Definition

An edge with just one endpoint is called a **loop** or a **self-loop**.

SPECIAL GRAPHS

Definition

A graph G is said to **simple** if it has no parallel edges or loops. In a simple graph, an edge with endpoints v and w is denoted by $\{v, w\}$.

Definition

For each integer $n \geq 1$, let D_n denote the graph with n **vertices** and **no edges**. Then D_n is called the **discrete graph on n vertices**.

Definition

Let $n \geq 1$ be an integer. Then a simple graph with n vertices in which there is an edge between each pair of distinct vertices is called the **complete graph** on n vertices. It is denoted by K_n .

For example, the complete graphs K_2 , K_3 and K_4 are shown in

Definition

If each vertex of a graph G has the same degree as every other vertex, then G is called a **regular graph**.

SUBGRAPHS

Definition

A graph H is said to be a subgraph of a graph G if and only if every vertex in H is also a vertex in G , every edge in H is also an edge in G and every edge in H has the same endpoints as in G .

We may also say that G is a super graph of H

Definition

A sub graph H is said to be a **proper sub graph** of a graph G if vertex set V_H of H is a proper subset of the vertex set V_G of G or edge set E_H is a proper subset of the edge set E_G .

For example, the sub graphs in the above examples are proper sub graphs of the given graphs.

ISOMORPHISMS OF GRAPHS

We know that shape or length of an edge and its position in space are not part of specification of a graph. For example, the represent the same graph.

Definition

Let G and H be graphs with vertex sets $V(G)$ and $V(H)$ and edge sets $E(G)$ and $E(H)$, respectively. Then G is said to be **isomorphic to** H if there exist one-to-one correspondences $g: V(G) \rightarrow V(H)$ and $h: E(G) \rightarrow E(H)$ such that for all $v \in V(G)$ and $e \in E(G)$,

v is an endpoint of $e \Leftrightarrow g(v)$ is an endpoint of $h(e)$.

Definition

The property of mapping endpoints to endpoints is called **preserving incidence** or **the continuity rule** for graph mappings.

As a consequence of this property, a self-loop must map to a self-loop.

Thus, two isomorphic graphs are same except for the labelling of their vertices and edges.

WALKS, PATHS AND CIRCUITS

Definition

In a graph G , a **walk from vertex v_0 to vertex v_n** is a finite alternating sequence $\{v_0, e_1, v_1, e_2, \dots, v_{n-1}, e_n, v_n\}$ of vertices and edges such that v_{i-1} and v_i are the endpoints of e_i .

The **trivial walk** from a vertex v to v consists of the single vertex v .

Definition

In a graph G , a **path** from the vertex v_0 to the vertex v_n is a walk from v_0 to v_n that does not contain a repeated edge.

Thus a **path** from v_0 to v_n is a walk of the form

$$\{v_0, e_1, v_1, e_2, v_2, \dots, v_{n-1}, e_n, v_n\},$$

where all the edges e_i are distinct.

Definition

In a graph, a **simple path** from v_0 to v_n is a path that does not contain a repeated vertex.

Thus a simple path is a walk of the form

$$\{v_0, e_1, v_1, e_2, v_2, \dots, v_{i-1}, e_n, v_n\},$$

HAMILTONIAN CIRCUITS

Definition

A **Hamiltonian path** for a graph G is a sequence of adjacent vertices and distinct edges in which every

vertex of G appears exactly once.

Definition

A **Hamiltonian circuit** for a graph G is a sequence of adjacent vertices and distinct edges in which every vertex of G appears exactly once, except for the first and the last which are the same.

Definition

A graph is called **Hamiltonian** if it admits a Hamiltonian circuit.

MATRIX REPRESENTATION OF GRAPHS

A graph can be represented inside a computer by using the adjacency matrix or the incidence matrix of the graph.

Definition

Let G be a graph with n ordered vertices $v_1, v_2, \dots, v_n$. Then the **adjacency matrix of G** is the $n \times n$ matrix $A(G) = (a_{ij})$ over the set of non-negative integers such that

a_{ij} = the number of edges connecting v_i and v_j for all $i, j = 1, 2, \dots, n$.

We note that if G has no loop, then there is no edge joining v_i to v_i , $i = 1, 2, \dots, n$. Therefore, in this case, all the entries on the main diagonal will be 0.

Further, if G has no parallel edge, then the entries of $A(G)$ are either 0 or 1.

It may be noted that adjacent matrix of a graph is symmetric.

Conversely, given a $n \times n$ symmetric matrix $A(G) = (a_{ij})$ over the set of non-negative integers, we can associate with it a graph G , whose adjacency matrix is $A(G)$, by letting G have n vertices and joining v_i to vertex v_j by a_{ij} edges

COLOURING OF GRAPH

Definition

Let G be a graph. The assignment of colours to the vertices of G , one colour to each vertex, so that the adjacent vertices are assigned different colours is called **vertex colouring** or **colouring of the graph G** .

Definition

A graph G is **n -colourable** if there exists a colouring of G which uses n colours.

Definition

The minimum number of colours required to paint (colour) a graph G is called the **chromatic number of G** and is denoted by $\chi(G)$.

UNIT-4 Propositional Logic:

INTRODUCTION TO LOGIC

Logic is essentially the study of arguments. For example, someone may say, suppose A and B are true, can we conclude that C is true? Logic provides rules by which we can conclude that certain things are true given other things are true. Here is a simple example: A tells B that “if it rains, then the grass will get wet. It is raining”; B can then conclude that the grass is wet, if what A has told B is true. Logic provides a mechanism for showing arguments like this to be true or false.

The section starts by showing how to translate English sentences into a logical form, specifically into something called “propositions”. In fact, our study of logic starts with *propositional calculus*.

Propositional calculus is the calculus of propositions and we plan to study propositions. Most of us may associate the term calculus with integrals and derivatives, but if we check out the definition of calculus in the dictionary, we will see that calculus just means “a way of calculating”, so differential calculus, for instance, is how to calculate with derivatives and integral calculus is how to calculate with integrals.

However, before going into how to translate English sentences into propositions, we are going to introduce Boolean expressions (that is, propositions), and then discuss about translation. Hence, we will see the same ideas in two different forms.

BOOLEAN EXPRESSIONS

Definition A Boolean variable is a variable that can either be assigned true or false. We have programmed in C++ and know about types such as integers, floats, and character pointers. However, C++ also has a Boolean type, as do Java and Pascal. We can declare variables to be of Boolean type, which means that they can only take on two values: true and false. Throughout the chapter, we shall generally use the letters, p , q , and r as Boolean variables. However, in some cases we will allow these letters to have subscripts. For example, p_0 , q_{1492} and r_{1776} are all Boolean variables.

Definition A Boolean expression is either

1. a Boolean variable, or
2. it has the form $\neg\phi$, where ϕ is a Boolean expression, or
3. it has the form $(\phi * \psi)$, where ϕ and ψ are Boolean expressions and $*$ is one of the following:

$\wedge$, $\vee$, $\rightarrow$, or $\leftrightarrow$.

CONSTRUCTION OF BOOLEAN EXPRESSIONS

Suppose we are given a Boolean expression and asked to prove that it is a Boolean expression. How do we proceed? There are two different ways of doing it. The first is to build a Boolean expression from its constituent parts. Let us start off with an example. We want to show that $((p \wedge q) \vee \neg r)$ is a Boolean expression. To do so, we will take a bottom-up approach.

<i>Expression</i>	<i>Reason</i>
1. p	Boolean variable
2. q	Boolean variable
3. r	Boolean variable
4. $\neg r$	3, $\neg\phi$
5. $(p \wedge q)$	1, 2, $(\phi \wedge \psi)$
6. $((p \wedge q) \vee \neg r)$	5, 4, $(\phi \wedge \psi)$

Notice that we start with the smallest Boolean expression (namely, Boolean variables) and work our way up. Look at line number 4. The reason is “3, $\neg\phi$ ”. This means that we are using the rule $\neg\phi$ to create line 4, where ϕ is from line 3. This is just the second part of the definition being applied. And we use lines 1 and 2 and rule $(\phi \wedge \psi)$ to create the expression in line 5. Again, this rule comes from part 3 of the definition.

Definition A construction of a Boolean expression is a list of steps, where each line is either a Boolean variable or it uses a connective (e.g., $\neg$, $\wedge$, $\vee$, $\rightarrow$ or $\leftrightarrow$) to connect two other Boolean expressions (they may be the same), with line numbers that are less than itself. Each line is a valid Boolean expression.

For example, look at the construction above. If we have to add a 7th step to the construction, we would have two choices. Either we could introduce a Boolean variable (we could always do this) or we could use a connective and find a Boolean expression that is already on the list and add it. For example, we could place a $\neg$ in front of $\neg r$ (from step 4) and produce $\neg\neg r$.

The point of this exercise was to explain how to convince someone else that $((p \wedge q) \vee \neg r)$ is a Boolean expression. It is a kind of proof and uses the definition of Boolean expressions as reasons or justifications for each step.

The only difficulty with using this (and it is a small one) is that it is sometimes easier seeing an expression top-down than bottom-up. That is, it is intuitively simpler to take a complicated expression

like $((p \wedge q) \vee \neg r)$ and try to break it down to its two parts, $(p \wedge q)$ and $\neg r$.

MEANING OF BOOLEAN EXPRESSIONS

One use of logic is as a means of deciding what things are true, given that certain facts are already true. Logic provides us a framework for deducing new things that are true. However, this deduction is based on form. For example, we might say that either $x > 0$ or $x \leq 0$ and also that x is not greater than 0. Given these two facts, we should be able to conclude that $x \leq 0$.

Now both arguments actually have the same form that is, we basically said ϕ or ψ is true and then ϕ is not true, therefore we concluded ψ is true. In the first example, ϕ was $x > 0$ and ϕ was $x \leq 0$, while in the second example ϕ was “the capital of India is Mumbai” and ψ was “the capital of India is New Delhi”. In both examples, there was a similar form of the argument and the conclusion that we drew was purely based on the form.

This is actually at the heart of logic. (Some) English arguments can be translated into Boolean expressions, and then we can apply rules of logic to determine whether the arguments make sense, atleast, based on their form.

We know that Boolean variables are the building blocks of Boolean expressions. Boolean variables like p generally stand for either English or mathematical propositions.

Definition A proposition is something that is either true or false but not both.

Not all English sentences are propositions. For example, the sentence “Run away” cannot really be said to be true nor false. Not all mathematical “sentences” are propositions either. For example, $x > y$ is neither true nor false. We would need to know the values of x and y before we could draw the conclusion. Actually, we do not have to be this strict, $x > y$ is either true or false, so in some sense, we can consider it a proposition.

Once the translation has been made from English sentences or mathematical sentences into Boolean expressions, then we generally do not care what the original sentence means. We can make conclusions based on the Boolean expressions. We shall get into the details soon.

1.4.1 Conjunctions

The most basic Boolean expression is a Boolean variable, which is either true or false. Throughout this section, we shall refer to two propositions: p and q .

p I own a cat.

q I own a dog.

One way to make a more complicated sentence is to connect two sentences with “and”. For example,

“I own a cat” AND “I own a dog”. In propositional calculus (which is what we are studying now), our purpose is to determine when expressions or sentences are true. So, when is the entire expression “I own a cat AND I own a dog” true? Intuitively, we would say it is true if both parts are true.

Now let’s look at the Boolean expression equivalent of that same sentence. It happens to be $(p \wedge q)$ (again, notice the use of parentheses). The symbol for AND is $\wedge$, which we can pronounce as AND. If it helps us to remember, the $\wedge$ symbol looks sort of like an “A”, which is the first letter of AND. Sometimes, mathematicians say that $(p \wedge q)$ is a *conjunction* of p and q and that p and q are conjuncts of the conjunction. Despite the fancy name, the work “conjunction” does come up often enough and hence we ought to remember it.

So, when is $(p \wedge q)$ true? When both p and q are true. If either is false, then the whole expression is false. We can actually summarize this in a *truth table*. A truth table tells us the “truth” of a Boolean expression given that we assign either true or false to each of the Boolean variables.

Given two different Boolean variables, p and q , there are four different ways to assign truth values to them. Each of the four ways is listed below. T stands for true, while F stands for false. If we look at the column for variable q , we will see that it alternates T, then F, then T, then F, whereas the column with p to its immediate left alternates, T, T, then F, F. If we had another variable, r and placed it to the left of p , it would alternate T, T, T, T then F, F, F, F.

There is a pattern. Starting from the rightmost Boolean variable, we will alternate every turn T, F, T, F. The next column to its left will alternate T, T, F, F, T, T, F, F, etc. The next one to its left will alternate T, T, T, T, F, F, F, F. As we move to the left, we repeat the Ts twice as many times as the previous column and twice as many Fs. This pattern actually covers all possible ways of combining truth values for n Boolean variables.

Now, look at line 1 in the truth table. Look at the last column. We will notice that the entry has the value T, which means that when p is assigned T and q is assigned T, then $(p \wedge q)$ is true as well. This is just a formalization of what we said before, $(p \wedge q)$ is only true when p and q are both true (that is, both assigned to true).

The key point is to notice that we can find out the truth value of a complicated expression by knowing the truth value of the parts that make it up. This is really no different from arithmetic expressions. For example, if we had the expression $(x + y) - z$, then we could tell that the value for this expression, provided if we knew the values for each of the variables. It is the same with Boolean expressions. If we know the truth values of the Boolean variables, then using truth tables, we can determine the truth value

of a Boolean expression.

Now, we used p and q for the truth table above. However, there was nothing special about using those two Boolean variables. Any two different Boolean variables would have worked. In fact, any two Boolean *expressions* would have worked. We could have replaced p with ϕ and q with ψ and $(p \wedge q)$ with $(\phi \wedge \psi)$ and the truth table would still have been fine.

1.4.2 Disjunctions

Instead of saying “I own a cat” AND “I own a dog”, we could connect the two statements with OR, as in, “I own a cat” OR “I own a dog”. When would this statement be true? It would be true if I owned either a cat or a dog. That is, only one of the two statements has to be true. What happens if both are true? Then is the *entire* statement “I own a cat OR I own a dog” true? Going by propositional logic, we will say yes. That is, the entire OR statement is true if one or the other statement or both are true.

We use the symbol $\vee$ to represent OR. So $(p \vee q)$ (again, notice the parentheses) is the same as p OR q and the whole expression is true if p is true or q is true or both are true. It is false if both p and q are false. Sometimes the expression $(p \vee q)$ is called a *disjunction* (with a $\wedge$, it was a conjunction) and p and q are the disjuncts of the disjunction. Any Boolean expression (or subexpression) can be called a disjunction if it has the pattern $(\phi \vee \psi)$.

The use of “OR” in propositional calculus actually contrasts with the way we normally use it in English. For example, if A said, “I will go to the Cinema OR I will go to the Garden”. Usually it means, I will go to one or the other, but NOT both. This kind of “OR” is called as *exclusive* OR, while the one we use in propositional calculus is called an *inclusive* OR. We will almost always use the inclusive OR (the exclusive OR can be defined using inclusive ORs and negations, which will be introduced in the next section).

Here is the truth table for OR.

Notice the rightmost column of this truth table and compare it to the rightmost column of the truth table of $(p \wedge q)$. In the case of conjunction (i.e., AND), $(p \wedge q)$ is true in only one case, namely, when p and q are both true. It is false in all other cases. However, for $(p \vee q)$ (read p OR q), it is true in all cases except when both p and q are false. In other words, only one of either p or q has to be true for the entire expression $(p \vee q)$ to be true.

The main point covered so far:

The symbols we have seen: $\{\wedge, \vee, \rightarrow, \leftrightarrow, \neg\}$ are often called *connectives* because they connect two Boolean expressions (although in the case of $\neg$, it’s only attached to a single Boolean expression).

1.4.3 Negations

The symbol $\neg$, (pronounced “not”) is like a negative sign in arithmetic. So, if we have p (“I own a cat”), the $\neg p$ (notice there are *no* parentheses) can be read as “Not I own a cat”, or “It is not the case that I own a cat” (in which case the $\neg$ could be read as “it is not the case that”). Both of these sound awkward, but the idea is to use the original sentence and attach something before it, just like the connective. In English, it sounds more correct to say “I do not own a cat”.

Unlike $\wedge$ and $\vee$, $\neg$ only attaches to a single Boolean expression. So, the truth table is actually smaller for $\neg$ since there is only one Boolean variable to worry about.

Line	p	$\neg p$
1	T	F
2	F	T

This should be an easy truth table to understand. If p is true, then $\neg p$ is false. The reverse holds as well. If p is false, then $\neg p$ is true.

CONSTRUCTION OF TRUTH TABLES

Suppose we are given a Boolean expression, say, $((p \wedge q) \vee \neg p)$. We want to know whether this expression is true or false.

We introduce a function, v . This function will be called a *valuation*. If we were to write this function signature in pseudo-C++ code, it would look like

```
boolean v( boolExpr x );
```

In words, this function takes a Boolean expression as input (think of it as a class) and returns back a Boolean value, that is, it returns either true or false.

Let us formally define a valuation.

Definition A valuation (also called, a truth value function or a truth value assignment) is a function which assigns a truth value (that is, true or false) for a Boolean expression, under the following restrictions.

$v((\phi \wedge \psi))$	$= \min(v(\phi), v(\psi))$
$v((\phi \vee \psi))$	$= \max(v(\phi), v(\psi))$
$v(\neg \phi)$	$= 0, \text{ if } v(\phi) = 1$
	$= 1, \text{ if } v(\phi) = 0$

$v((\phi \rightarrow \psi))$	$= 1$, if $v(\phi) = F$ or $v(\psi) = T$
	$= 0$, otherwise
$v((\phi \leftarrow \psi))$	$= 1$, if $v(\phi) = v(\psi)$
	$= 0$, otherwise

The interesting thing is that because of these restrictions, once a truth value function has been defined for all the Boolean variables in a Boolean expression, the truth value for the Boolean expression (and all its subexpressions) are defined as well.

Let us take a closer look at the definition. We shall take it line by line. In the first line, we have

$$v((\phi \wedge \psi)) = \min(v(\phi), v(\psi))$$

This says that if we want to find the truth value of $(\phi \wedge \psi)$, then we have to find the truth value of ϕ (that is, $v(\phi)$) and the truth value of ψ (that is, $v(\psi)$). We take the “minimum” of $v(\phi)$ and $v(\psi)$. How does one take the minimum of the two? If we treat false as the number 0 and true as the number 1, then taking the minimum of two numbers makes sense. But is it an accurate translation of AND?

Let us think about this for a moment. When is $(\phi \wedge \psi)$ true? When ϕ is true AND when ψ is true. If we think of true as being the number 1, then we are asking what is the minimum of 1 and 1. And the minimum of those two numbers is 1. If we translate it back, we get true. That seems to work.

Now, when is $(\phi \wedge \psi)$ false? When either ϕ or ψ is false, that is, when $v(\psi) = F$ or (and this is an inclusive or) when $v(\psi) = F$. So, let’s think about this. If one of the two is false, then it has a value of 0. The minimum of 0 and anything else is 0. Why? Well, since truth values are either 0 (for false) or 1 (for true), we can only take the minimum of 0 and some other number. That number could be 0, in which case the minimum is 0, or it could be 1, in which case the minimum is still 0. So, the minimum of 0 and any other number (restricted to 0 or 1) is 0. And that makes sense too because we want $(v(\phi \wedge \psi))$ to be false (i.e., 0) when either $v(\phi)$ or $v(\psi)$ is false.

If we treat false as 0 and true as 1, then we can show

$$v((\phi \vee \psi)) = \max(v(\phi), v(\psi))$$

makes sense too. *max* is the function that takes the maximum of two numbers (in this case, we need to treat the truth values like numbers).

The real point of this is to show that to find out the truth value of a Boolean expression (i.e., to find out the value of $v(\phi)$, we need to find out the value of the smaller subexpressions.) For example, to find $v(\neg\phi)$, we need to know the value of $v(\phi)$. And to find out $v(\phi)$ we need to see what pattern ϕ follows (is it a negation, conjunction, or disjunction?) and recursively apply the definition. At each step, we break

down the equation into smaller and smaller subexpressions until we reach the smallest subexpression, which happens to be a Boolean variable.

To find the truth value of a Boolean expression, we just need to know the truth value of the Boolean variables in that expression.

Back to Derivations

Based on the insight of the previous section, we now return to our problem. We want to construct a truth table for $((p \wedge q) \vee \neg p)$. To do so, we need to find the Boolean variables in this expression. This is easy as there are only p and q .

This is how we will derive the Boolean expression. The reason will become clear, but we intend to use it to construct a truth table.

Here is the derivation.

<i>Expression</i>	<i>Reason</i>
1. p	Boolean variable
2. q	Boolean variable
3. $\neg p$	1, $\neg\phi$
4. $(p \wedge q)$	1, 2, $(\phi \wedge \psi)$
5. $((p \wedge q) \vee \neg p)$	4, 3, $(\phi \wedge \psi)$

We will create one column in the truth table for each line in the derivation. How many rows do we use? If n is the number of Boolean variables (in this case, 2), then 2^n is the number of rows (in this case, $2^2 = 4$ rows).

Now we just fill in the columns one after another. We know that column 3 is derived from column 1. So, we can just get column 3 by “negating” the values in column 1. Compare the values of column 1 and column 3. We are basically applying the definition of $\vee$ from the last section.

Column 4 is based on columns 1 and 2. We get the values by adding them.

Finally, column 5 is constructed from columns 4 and 3, respectively.

And that is how it is done!

Valuations and Truth Tables

If we know the truth values for the variables, then the value of the Boolean expression is known as well. There is a simple analog to arithmetic expressions. In an expression like $(x + y) - z$, we know the value once we know what *value* each variable has been assigned. For example, if $x = 4$, $y = 3$, and $z = 5$, the expression's value is 2. If we change the values of the variables, the result of the expression changes too. However, the result is completely defined by the values assigned to the variables.

How does v relate to truth tables? Essentially each row of a truth table corresponds to a separate v . For example, consider the last truth table in the previous section. Row 1 defines $v(p) = T$ and $v(q) = T$. Given this, $v((p \wedge q) \vee \neg p)$ must have the value *true*. Row 2 defines a different v , one where $v(p) = T$, but $v(q) = F$, and where $v((p \wedge q) \vee \neg p)$ consequently has the value *false*.

How Many Rows

If there are n Boolean variables, then there will be 2^n rows in the truth table. Why?

We start with one Boolean variable, say this is r . r can either be true or false. So, a truth table with just one variable has two lines. One when the value is true and one when it is false.

Line	r
1	T
2	F

Now we add a second variable, q . q can also be either true or false. So, suppose q is true. Then, we get a truth table that looks like:

	1	2
Line	q	r
1	T	T
2	T	F

This still gives us two rows. Notice we have done nothing to the r column. r still has the two values, T and F . Meanwhile q has been set to true. However, what about the case when q is false? We should get two more rows for when q is false because r can still be true or false in that case.

So notice that rows 3 and 4 have F in column 1, (this is when q is false) and T in rows 1 and 2.

	1	2
Line	q	r
1	T	T
2	T	F
3	F	F
4	F	F

At this point, we have shown 4 different ways in which q and r can be assigned truth values. Now, let's take it one step further. We add a third variable p . Suppose $p = T$ (or more accurately, $v(p) = T$). How many different combinations of truth values are there for q and r ? Four, right? There were four combinations of truth values for q and r when we did not have p , so why should there be any more or less when $v(p) = T$? So, there should be 4 ways when p is true.

Now how many combinations of q and r are there when p is false? It is still 4, for the same reason. Now, we have a total of 8 rows. 4 rows when p is true (since there are 4 combinations of values for q and r) and 4 more when p is false. It will give a total of 8. Does this follow our formula? There are 3 variables, so there should be 2^3 rows and 2^3 equals 8, so it matches.

But the formula for the number of rows actually makes sense. Let us pretend we have n variables. This ought to create 2^n rows or equivalently, 2^n different ways to assign truth values to n variables. Now, we want to add one more variable. We shall call this new variable, p . That will make a total of $n + 1$ variables. We expect there to be 2^{n+1} rows.

We know there are 2^n different combinations for n rows. Now, there are still 2^n combinations when p is true. After all, why should p being true (in effect, it is a constant) affect the number of combinations of the other n variables? It should not and it does not. Then, there are 2^n combinations when p is false for the same reason. So there are 2^n combinations when p is true and 2^n combinations when p is false. Add the two together and we get $2^n + 2^n = 2^{n+1}$. So, adding a new variable doubles the number of rows. We get one set of rows when the new variable is true and another set when it is false.

Making Truth Tables

Just because we know there are 2^n rows in a truth table, does not mean that it is easy to figure out a quick way to fill one up. So, here is the quick way to do it. We shall use three variables p , q , and r .

To start, fill in the right-most row with a Boolean variable. In this case, it is column 3 (the column with r). Start alternating T, F, T, F all the way down.

Go on to the column over to the left. This is column 2. Go down alternating T, T then F, F then repeat. So, 2 Ts, followed by two Fs and repeat.

Finally, move to column 1 and repeat 4 Ts, followed by 4 Fs. That is, T, T, T, T then F, F, F, F and repeat, if necessary.

In doing this we see a general pattern. Suppose we are in a particular column and are repeating the pattern of n Ts, followed by n Fs. If we move one column left, we double the number of Ts followed by double the number of Fs; that is, $2n$ Ts, followed by $2n$ Fs.

Why does this work? Well, it is related to the previous section. Think of T as being 1 and F as being 0. Look at row 1. There are three t's in this row. Convert it to a binary number. This will be 111. Now, we go to row 2. Reading across, we read T, T, F. Converting to 1s and 0s, we get 110. In binary, this is 6. Now, as we progress down from row 1 to row 8, we will be counting backwards in binary, from 7 down to 0. So, essentially a truth table for n variables has one row for each n bit binary number. We shall learn about binary numbers later on, but this is the basic idea of how filling out a truth table works.

LOGICAL EQUIVALENCE

The idea of logical equivalence is deceptively simple. We are given two Boolean expressions. How do we determine if they are the same? The first question we need to ask is, “What do we mean by ‘the same’?” The two expressions are same if they generate the same truth table. We can be a little bit more formal. As we have said earlier, once we have defined the truth values of all the Boolean variables in a Boolean expression, we can know (or calculate) the truth value of the Boolean expression.

Now, suppose we had two different Boolean expressions, a (pronounced “alpha”) and b (pronounced “beta”). Write down all the Boolean variables possible for both of them. Suppose a has m Boolean variables and b has n Boolean variables (usually, $m = n$, but this does not have to be the case). Now suppose that between the two, there are k unique variables. For example, a might have variables p , q , and r , while b has variables, p and r_2 . Between the two, there are 4 unique variables. In this example, k is 4. Thus, there are 2^k different ways of assigning truth values to these k unique variables.

For every function, v (there are 2^k of these, one for each row of the truth table), if $v(a) = v(b)$, then the a and b are logically equivalent.

Let us consider a simple example. We want to show that p is logically equivalent to $\neg\neg p$. In arithmetic, this is the equivalent of saying $x = - - x$. While this is intuitively obvious, it is better to have a method to determine when expressions are logically equivalent.

There is only one Boolean variable in both expressions, namely p . So, there are $2^1 = 2$ ways to assign truth values to p . We can either assign it true or false. Now, all we do is construct the truth table for both expressions. The result looks like:

Columns 1 and 3 are identical, which means the two are logically equivalent.

There is a stranger example. Many times, both expressions we wish to show as logically equivalent have the same Boolean variables. This makes sense, because if the variables were different, say, p as one Boolean expression and q as another, then we could find a function v , which sets $v(p) = T$ and $v(q) = F$. Hence, there would exist some function v such that $v(p) \neq v(q)$ and thus the two would not be logically equivalent.

However, sometimes, even though a Boolean expression contains a variable, it does not really depend on that variable. For example, consider $(p \wedge \neg p)$. Even if we do not know the value of p , we know $(p \vee \neg p)$ is true. Why? For $(p \vee \neg p)$ to be true, either p or $\neg p$ must be true. That is, if one is true, the other must be false. So, one of the two is always guaranteed to be true. In other words, it does not matter whether p is true or false, $(p \vee \neg p)$ is always true. This means that this expression does not really depend on p .

For a Boolean expression to depend on a variable, the truth value of the expression must change at some point, if p changes its value. That is, there must be some assignment of truth values to the other variables in an expression, a , such that $v(a) \neq v'(a)$. In this case, v might represent the valuation where p is true and v' would then represent the valuation where p is false. For all other variables, v and v' give the same truth values.

So, let us show that $(p \vee \neg p)$ and $(q \vee \neg q)$ are logically equivalent. There are two unique variables between the two of them: p and q . Now, we construct the truth table for both expressions.

Notice that columns 5 and 6 are identical. When two columns are identical in this manner, they are logically equivalent.

To show two expressions are not logically equivalent, we just need a single valuation, v , where $v(a)$ does not equal $v(b)$. Or equivalently, we need to find a single line in a truth table, where one expression's truth value is T, while the other's is false.

Now we show that $(p \wedge q)$ and $(p \vee q)$ are not logically equivalent. Again, we write up the truth tables for both.

Just by inspection, we should be able to tell that columns 3 and 4 are not the same. For example, we can look at row 2 where $v(p) = T$ and $v(q) = F$, so the result is that $v((p \wedge q)) = F$, but $v((p \vee q)) = T$. Row 3 also gives a case where $v((p \wedge q)) \neq v((p \vee q))$. The two expressions agree in rows 1 and 4. Thus, we conclude that $(p \wedge q)$ and $(p \vee q)$ are not logically equivalent.

TAUTOLOGIES AND CONTRADICTIONS

Definition A tautology is a Boolean expression which always results in a true result, regardless of what the Boolean variables in the expression are assigned to.

We saw this earlier on, with the expression $(p \vee \neg p)$. If $v(p) = T$, then the whole expression is true. If $v(p) = F$, then also the whole expression is true. Basically, a tautology means that the result of the expression is independent of whatever Boolean variables have been assigned the result is always true.

Definition A contradiction is a Boolean expression which always results in a false result, regardless of what the Boolean variables in the expression are assigned to.

A contradiction is just the opposite of a tautology, in fact, given any Boolean expression that is a tautology, we just have to negate it to get a contradiction. For example, $(p \vee \neg p)$ is a tautology. The negation of that, $\neg(p \vee \neg p)$, is a contradiction. We can apply DeMorgan's law and get $(\neg p \wedge \neg \neg p)$ and by using the simplification for double negation result in $(\neg p \wedge p)$. So, since all of these are logically equivalent, then $(\neg p \wedge p)$ is also a contradiction.

CONTRADICTION RULE

$\neg p$ is true and then deduce a contradiction, then p is true. The idea runs something like this: one generally believes math is consistent that is, we do not derive contradictions using the rules of logic (the most common contradiction is to derive $\neg q$ when we also know that q happens to be true). So, when we try to prove p , we try to assume $\neg p$ and if this leads to a contradiction, then we know that $\neg p$ cannot be true, since our system avoids contradiction and thus if $\neg p$ is not true, it is false, and if it is false, then p must be true. This is often the line of reasoning used in a proof by contradiction.

UNDERSTANDING QUANTIFIERS

What does $\forall x P(x)$ mean? We can read $\forall$ as “for all”, so the entire statement can be read as “for all x , P of x ”. If we add a domain, things can be made clearer. So, let, D , our domain, be the set $\{2, 4, 6, 8, 10\}$.

Then the statement $\forall x \in D P(x)$ makes more sense. We expect every single x picked from D to have the property P . In fact, a very convenient way to think about $\forall x \in D P(x)$ is to expand it out using ANDs.

For example, the expansion of $\forall x \in D P(x)$ would be:

$$P(2) \wedge P(4) \wedge P(6) \wedge P(8) \wedge P(10).$$

Every x in D having property P means $P(2)$ is true AND $P(4)$ is true AND $P(6)$ is true AND $P(8)$ is true AND $P(10)$ is true. In general, we will not be able to write it all out like this because D could be infinite, but it is a correct way to think about what $\forall$ really means.

What does the expansion for $\exists x \in D P(x)$ mean? $\exists$ is read as “there exists”, so the entire statement reads as “there exists an x in D , P of x (or such that $P(x)$) holds true. The phrase “There exists” should give the hint that at least one x from D has the property P . And if the previous example for $\forall$ used ANDs, what do we think $\exists$ uses? If we said ORs, that would be correct. Using the same set D as before,

$\exists x \in D P(x)$ can be expanded to:

$$P(2) \vee P(4) \vee P(6) \vee P(8) \vee P(10).$$

This disjunction is true when at least one of the P s is true and that should make sense based on our intuition of what “there exists” means (that is, it means at least one exists).

We will normally not be able to expand out universal quantifiers ($\forall$) or existential quantifiers ($\exists$) in this manner because D , the domain, may be infinite. Nevertheless, it is useful to think about an expansion when trying to get a feel of what quantifiers mean.

Unit I

PROJECT AND REPORT WRITING

MEANING OF PROJECT

Project means a “Proposed Plan of Action”. It is a planned set of interrelated tasks to be executed over a fixed period and within certain cost and other limitations. It is a sequence of tasks which is planned from the beginning to the end and also it is restricted by time, required results & resources. There has to be a definite deadline, budget & outcomes.

BASICS OF PROJECT WRITING

HOW TO WRITE WELL

1. Precision

There should be exact meaning to the word you intend to say.

2. Vigour

In addition to be precise, you also need to assign some weight & meaning to the words.

3. Spelling and grammar

The spelling & grammar should be such used, which does not embarrass you if done incorrectly.

STRUCTURE OF A REPORT

Format

Use 1.5 line spacing, Font-Times New Roman and Font-Size-12.

Title

Give your document a brief descriptive title.

Introduction

Introduce your topic over here.

Body

The body is the content of your document where you present your data and make your points.

Carefully organize the body of your report so that similar topics are included together, and the logic of your report flows smoothly.

Break up your writing with heading, subheadings, tables, figures, and lists.

Conclusions or Summary

Use a summary or conclusions paragraph to wrap-up your document and provide closure.

It should give a final idea of the whole report.

References

The reference list is a list of the sources of information used and cited in the document. Cite references to show the source(s) of information and data included in your document. An easy citation format is simply to include the author's last name (or publication name) and year of publication in parentheses.

BASICS OF WRITING A REPORT

1. Identify the audience.
2. Decide on the length of the report in advance.
3. Divide the contents of the report into clearly labeled categories.
 - ☐ Provide an executive summary.
 - ☐ Write an introduction.
 - ☐ Include a methodology section and describe it.
 - ☐ Elaborate on the findings of the project.
 - ☐ Explain preliminary project successes.
 - ☐ Describe project challenges and obstacles.
 - ☐ Suggest recommendations and solutions.
 - ☐ Write the report conclusion.
4. Use formatting techniques & sub-headings to guide the attention of readers.
5. Review the report for errors.

Report Writing:-

A report can be defined as a testimonial or account of some happening. It is purely based on observation and analysis. A report gives an explanation of any circumstance. In today's corporate world, reports play a crucial role. They are a strong base for planning and control in an organization, i.e., reports give information which can be utilized by the management team in an organization for making plans and for solving complex issues in the organization.

A report discusses a particular problem in detail. It brings significant and reliable information to the limelight of top management in an organization. Hence, on the basis of such information, the management can make strong decisions. Reports are required for judging the performances of various departments in an organization.

How to write an effective report

1. Determine the objective of the report, i.e., identify the problem.
2. Collect the required material (facts) for the report.
3. Study and examine the facts gathered.
4. Plan the facts for the report.
5. Prepare an outline for the report, i.e., draft the report.
6. Edit the drafted report.
7. Distribute the draft report to the advisory team and ask for feedback and recommendations.

The essentials of good/effective report writing are as follows-

1. Know your objective, i.e., be focused.
2. Analyze the niche audience, i.e., make an analysis of the target audience, the purpose for which audience requires the report, kind of data audience is looking for in the report, the implications of report reading, etc.
3. Decide the length of report.
4. Disclose correct and true information in a report.
5. Discuss all sides of the problem reasonably and impartially. Include all relevant facts in a report.
6. Concentrate on the report structure and matter. Pre-decide the report writing style. Use vivid structure of sentences.
7. The report should be neatly presented and should be carefully documented.
8. Highlight and recap the main message in a report.
9. Encourage feedback on the report from the critics. The feedback, if negative, might be useful if properly supported with reasons by the critics. The report can be modified based on such feedback.
10. Use graphs, pie-charts, etc to show the numerical data records over years.
11. Decide on the margins on a report. Ideally, the top and the side margins should be the same (minimum 1 inch broad), but the lower/bottom margins can be one and a half times as broad as others.
12. Attempt to generate reader's interest by making appropriate paragraphs, giving bold headings for each paragraph, using bullets wherever required, etc.

PARAGRAPH WRITING

First of all, WHAT IS A PARAGRAPH? Paragraphs comprises of related sentences. The focus is on only “One-Idea”.

HOW TO WRITE A PARAGRAPH

The paragraph should be made while keeping the following four-points in view:-

1. Unity
2. Order
3. Coherence
4. Completeness
5. Well developed

PROPOSAL WRITING

What is a Proposal?

It is basically a persuasive offer to complete some work, to sell a product, to provide some service or to provide a solution to a problem. A **business proposal** is a written offer from a seller to a prospective buyer. Business proposals are often a key step in the complex sales process—i.e., whenever a buyer considers more than price in a purchase. A proposal puts the buyer's requirements in a context that favors the sellers products and services, and educates the buyer about the capabilities of the seller in satisfying their needs. A successful proposal results in a sale, where both parties get what they want, a win-win situation.

WRITING A PROPOSAL

Proposal will either be accepted or rejected. Obviously, you want your proposal to be accepted. To help make this possible, follow the six steps listed below.

1. *Your proposal should define the problem and state how you plan to solve the problem.*
2. *Do not assume that your readers will believe your solution is the best.*
3. *Your proposal should be researched thoroughly.*
4. *Your proposal should prove that your solution works.*
5. *Your proposal should be financially feasible.*
6. *Your finished proposal should look attractive.*

PAPER READING (HOW TO READ A RESEARCH PAPER)

1. Read critically
2. Read creatively
3. Make notes as you read the paper
4. After the first read-through, try to summarize the paper in one or two sentences
5. If possible, compare the paper to other works
6. Summarizing the paper

VOICE MODULATION

Voice Modulation is the adjustment of the pitch or tone of voice to become enough to be clearly heard and understood by the audience.

COMPONENTS OF VOICE MODULATION:-

Pace or Speech speed: speed should be such that you are able to understand what is said.

Pitch or Depth of voice: Keep it at a level that is comfortable for you.

Pause: Pauses should be given at required intervals.

Power: One should speak from inside the abdomen to make it commanding by generating intensity in your voice.

Volume: Try and match your listener's speech volume.

Emphasis: Put emphasis by putting some pressure or focus on the key words.

Inflection and pausing effectively: Inflection means ups and downs of words. Inflection links meaning and feeling with your words.

Tone

Rhythm and Melody

Identifying your optimal pitch

BASICS OF PROJECT PRESENTATION

WHAT IS PRESENTATION?

When we talk about presentation, we mean that there is a person or group of persons who will showcase his/her work before an audience. Now, when there is Project presentation, then the person who has made the project will give a brief summary including the following things:-

1. Introduction about the Topic.
2. Objective of the Study done by the person.
3. Research methodology used in the Project.
4. Conclusion
5. Limitations
6. Suggestions
7. References of the sources used for making the project.

Unit II

How to make Presentations

Be neat and avoid trying to cram too much into one slide.

Be brief use keywords rather than long sentences

Have a very clear introduction, to motivate what you do and to present the problem you want to solve. The introduction is **not technical** in nature, but strategic (i.e. why this problem, big idea).

Use only **one idea per slide**. Have a good conclusions slide: put there the main ideas, the ones you really want people to remember. Use **only one "conclusions" slide**.

The conclusion slide should be the last one. Do not put other slides after conclusions, as this will weaken their impact.

Don't count on the audience to remember any detail from one slide to another (like color-coding, applications you measure, etc.). If you need it remembered, re-state the information a second time.

Try to cut out as much as possible; less is better.

Use a good presentation-building tool, like MS PowerPoint.

Humor is very useful; prepare a couple of puns and jokes beforehand (but not epic jokes, which require complicated setup). However, if you're not good with jokes, better avoid them altogether. Improvising humor is very dangerous.

The more you rehearse the talk, the better it will be. A rehearsal is most useful when carried out loud. 5 rehearsals is a minimum for an important presentation.

☐ ☐ Not everything has to be written down; speech can and should complement the information on the slides.

☐ Be enthusiastic.

☐ Give people time to think about the important facts by slowing down, or even stopping for a moment.

Text

- ☐ Slides should have short titles. A long title shows something is wrong.
- ☐ Use uniform capitalization rules.
- ☐ All the text on one slide should have the same structure (e.g. complete phrases, idea only, etc.).
- ☐ Put very little text on a slide; avoid text completely if you can. Put no more than one idea per slide (i.e. all bullets should refer to the same thing). If you have lots of text, people will read it faster than you talk, and will not pay attention to what you say.
- ☐ Use very few formulas (one per presentation). The same goes for program code (**at most one code fragment** per presentation).
- ☐ Do not put useless graphics on each slide: logos, grids, affiliations, etc.
- ☐ Spell-check. A spelling mistake is an attention magnet.

Illustrations

- ☐ Use suggestive graphical illustrations as much as possible. **Prefer an image to text.**

Do not "waste" information by using unnecessary colors. Each different color should signify something different, and something important. Color-code your information if you can, but don't use too many different colors. Have high-contrast colors.

A **few** real photos related to your subject look very cool (e.g. real system, hardware, screen-shots, automatically generated figures, etc.). Real photos are much more effective during the core of the talk than during the intro.

- ☐ Sometimes a matte pastel background looks much better than a white one.
- ☐ Use strong colors for important stuff, pastel colors for the unimportant.
- ☐ Encode information cleverly: e.g. make arrow widths showing flows proportional to the flow capacity.

VARIOUS PRESENTATION TOOLS

To make your presentations entertaining and at the same time educational you should consider the best presentation tools below.

1. Slide Rocket

It is an online presentation which allows you to engage with people and be able to deliver results. It helps you to come up with a presentation that will wow your audience. It gives you the power to the internet and you can also integrate with the public from free web resources like the You Tube or flicker. This online software allows you to create and share your work publically or privately.

2. Presentation

It is like a community where you can create and show your presentation. You can work alone or with others while it allows you to make it private or public. With Prezentit your works are always available, anytime and anywhere for viewing or edit. It is an online software and it is based on the JavaScript allowing real time collaboration.

3. Author Stream

It enables presenters to share PowerPoint either publically or privately. It is the best tool for marketing purposes and it is also an easy process to make videos from online PowerPoint presentations. This app is all about viewing and sharing PowerPoint presentations either with friends, family, or the public.

4. Empressr

It is the best original app with based rich media. And now you can share information the way you prefer and it allows you to make it public or private. It is built with tools which record either videos or audio which is added to the slides.

5. Google Docs

It allows the users to access your documents anytime and anywhere easily. It is a web-based office suite which is offered by Google within its Google drive services. The user is allowed to collaborate with coworkers in real time on different aspects including spreadsheets, documents, presentation and many others.

6. Casmio

This is the best solution for the multimedia presentations for videos and photos. It is the easy and yet very powerful platform for your work. It works and supports playback for the computer browser, mobile phone, and iPad. It allows you to upload PDF, MP3, images and many more. You can download presentation as videos or even auto-upload to YouTube.

7. Zoho

It is online software which manage your sales, markets your products and offer customer support. Therefore it allows the user to focus on the business and leave the rest to the app. It has a wide function for storing and sharing of information which also allows you to collaborate with others.

Guidelines for Effective Presentations

We need to make presentations to a wide variety of audiences, for example, Board members, employees, community leaders and groups of customers. Usually there is a lot that can be quickly gained or quickly lost from a presentation.

1. **Background:** The purpose of a presentation is communication. Poorly prepared displays (slides or overhead transparencies) and poor delivery plague many technical sessions at statistical meetings. The speaker often speaks too quickly or too quietly, or uses displays that cannot be read clearly.

2. **Content organization:** Your presentation will be most effective when the audience walks away understanding the five things any listener to a presentation really cares about:

- What is the problem and why is it a problem?
- What has been done about it before?
- What is the presenter doing (or has done) about it?
- What additional value does the presenter's approach provide?
- Where do we go from here?

3. **Planning:** Who is your audience? Be really clear about who your audience is and about why is it important for them to be in the meeting. Members of your audience will want to know right away why they were the ones chosen to be in your presentation. Be sure that your presentation makes this clear to them right away. What do you want to accomplish? List and prioritize the top three goals that you want to accomplish with your audience. Inform? Persuade? Be clear about the tone that you want to set for your presentation, for example, hopefulness, celebration, warning, teamwork, etc. Consciously identifying the tone to yourself can help you cultivate that mood to your audience.

4. **Level of audience and knowledge** - List the major points of information that you want to convey to your audience. When you are done making that list, then ask yourself, if everyone in the audience understands all of those points, then will I have achieved the goal that I set for this meeting? What medium will you use for your presentation?

5. **Design a brief opening** (about 5-10% of your total time presentation time) that:

- a. Presents your goals for the presentation.
- b. Clarifies the benefits of the presentation to the audience.

- c.Explains the overall layout of your presentation.
- d.Prepare the body of your presentation (about 70-80% of your presentation time).
- e.Design a brief closing (about 5-10% of your presentation time) that summarizes the key points from your presentation.
- f.Design time for questions and answers (about 10% of the time of your presentation).

6.Organizing your presentation: Beginning - middle - end . Time allocation Begin with an overview Build up your middle and emphasize your points End quickly with a summary

7.Supporting your presentation Hard copy of your presentation - slides/overheads Notes to support your presentation if you provide the supplemental information during your presentation, then your audience will very likely read that information during your presentation, rather than listening to you. Therefore, hand out this

information after you have completed your presentation. Or, hand it out at the beginning of your presentation and ask them not to read it until you have completed your presentation.

8.Don't crowd presentation with detail Bullets - short, five to six words per line Five or six lines per overhead/slide Colors Brightest colors jump forward - text Dark colors for background Consistent color scheme throughout to link ideas together

9.Typefaces (fonts) Simple styles one or two styles per slide Retain styles throughout Style variations - italics, bold, shadow, color Layout Consistent Balance

10.Emphasis Reveal information one idea at a time - buildups Use transitions appropriately DON'T read your information-- use the words in your own presentation Charts Accurate Easy to interpret Have impact

11.Basic Guidelines about Your Delivery

- ☐ If you are speaking to a small group (for example, 2-15 people), then try to accomplish eye contact with each person for a few seconds throughout your delivery.
- ☐ Look up from your materials, or notes, every 5-10 seconds, to look into the audience.

Speak a little bit louder and a little bit slower than you normally would do with a friend. A good way to practice these guidelines is to speak along with a news anchor when you are watching television.

Vary the volume and rate of your speech. A monotone voice is absolutely toxic to keeping the attention of an audience.

- ☐ Stand with your feet at shoulder-length apart.
- ☐ Keep your hands relatively still.

Boredom factors during Presentation

- ☐ A flat or monotone voice.
- ☐ A static presentation of facts and figures without real-life examples.
- ☐ Over-use of words such as: like, um, uh, ok, you know.
- ☐ Not giving the audience a chance to warm up to the speaker. A speaker can avoid this by making a joke or telling a
- ☐ Funny story to break the ice.
- ☐ No enthusiasm for the subject matter.
- ☐ Lack of facial expression.
- ☐ Being unprepared or unorganized.
- ☐ Too many people in a small room which will make the room too warm and uncomfortable.
- ☐ Room temperature that is too warm or too cold.
- ☐ Lack of handouts or informational literature.

How to overcome boredom factors during Presentation

☐ ☐ *Do your homework.* Remember that failing to plan is planning to fail. So, commit the time and effort to properly prepare your audience-centric presentation.

☐ Part of your audience analysis process is *anticipating their needs, reactions and potential questions*. Be prepared to deal with their reactions and respond to their questions with succinct and focused answers ... just in case.

Also accept the fundamental difference between *being an expert and having expertise*. You don't need to be an expert, just have more expertise on the topic than the audience does so you can accomplish your outcomes. An old Sicilian proverb comes to mind here – 'In the land of the blind, the one-eyed man is king.'

□ Rehearse your presentation several times, out loud, standing up, working with your slides. Audio or video tape it for self-critique. For really important presentations, rehearse with a small audience of colleagues who can relate to the topic and provide focused, objective constructive feedback.

INTERACTIVE PRESENTATION

Interactive presentation is a dynamically different way of using PowerPoint that gives presenters fast access to individual ideas while speaking. Rather than merely scrolling through slides, you'll make choices along the way. It provides a **jumping-off point** to more than 200 hours worth of visual explanations and examples allowing a presenter efficient, easy access to his or her **entire collection of materials at any point in his presentation**. Whether the presenter thinks of a relevant example on the spot or decides to reference material in response to a question, he or she can do it seamlessly in a way that standard linear, bullet point presentations simply do not allow.

In order to give an effective presentation, there are a few key steps to take. This goes for a student giving a class presentation or a professional giving a business presentation. Keeping these elements in mind, a successful presentation will be given.

1. Know your topic inside and out. By the time you get up in front of your audience, you must be comfortable talking about the topic and be able to answer any questions that may come your way. It is not a bad thing to have note cards or an outline in front of you while presenting.

2. Know to whom you are presenting. Learn what intrigues the audience and use that to your advantage.

3. Know your limits. Just because you think you are humorous doesn't mean the audience thinks you're humorous. Present in a way that you are most comfortable with, while still keeping it professional.

4. Know the purpose of the presentation. What exactly is it that you are trying to do? Sell? Educate? Entertain? Keep this in mind while preparing.

5. Use proper visuals. Handouts are an effective visual. Audience members can take notes and follow along. PowerPoints and videos also work effectively. People tend to remember visuals better than text. When using PowerPoint, Prezi, or other slideshow programs, keep it professional looking yet intriguing.

Use your own design elements and themes. Be creative.

6.PRACTICE! You may feel comfortable after just reading through your presentation, but that does not necessarily mean you are comfortable speaking it to an audience. So, practice the presentation out loud in front of mirror multiple times.

7.Allow time for Q&A at the end of your presentation. This is the time when your audience can ask for clarification or provide their input on certain elements within the topic.

Presentation as part of a job interview

Delivered well, with the desired impact, a presentation can certainly enhance your chances of success. However, for many it is an obstacle that can have the opposite effect – due to nerves, lack of preparation and focus. So it's very important to do the right sort of preparation to ensure that you get it right on the day. There are different types of presentations that you can be asked to deliver at an interview. You may be asked to prepare a presentation on a certain topic or on a topic of your choice.

Planning the Presentation

If asked to prepare a presentation, careful planning beforehand will help you to deliver it with greater confidence and success as does preparing for an interview in all other situations too. Do your research on the company to get a good understanding for the corporate style and culture of the company. This will help you to tailor your presentation to the needs of your interviewer(s). Check out the company website for information. You may also be able to use the site's search facility to discover more about the person or people who will be interviewing you.

Structuring your message

Decide on one main key message that you want to get across. This is like the spinal cord of the presentation – it helps hold the presentation together. It also provides a strong motivation for your audience to listen to you. Your presentation should follow a clear structure – with a strong opening, main body and ending - as this will help you stay focused and avoid losing track of your thoughts if you are nervous.

Your opening should capture attention at the start. It should clearly communicate your key message to your audience. Keep it succinct and punchy, using short sentences. A long rambling opening gives the

impression that it is going to be a long rambling presentation. Structure the main body of your presentation to three main sections. Three is a powerful number that people tend to remember things in. By restricting your presentation to three main sections it will keep a strong focus. You can then have sub-sections within each of the three main sections if you need to expand on your points. Your ending should also be memorable. Use the opportunity to re-emphasise your key message and leave a lasting impression.

The impromptu Presentation

You may be put on the spot and asked to give a presentation without prior warning. For these situations you need some form of standard structures in your head that you can call upon at very short notice.

Power of THREE

Using the Power of Three is a helpful tool as well here. Use three key words/phrases to help you create a quick structure in your head. One example structure with three areas that you can use quickly if asked to present on a problem solving or strategic issue is:

1. What was the issue?
2. What did you do to resolve it?
3. What was the outcome?

Chronological structures

Another structure you can use for impromptu presentations is:-

1. Past
2. Present
3. Future

This is a useful structure for the “Tell me about yourself” presentation (or question) that is commonly asked at interview by describing your personal history under these three titles.

The STAR structure

Another similar structure to consider is S.T.A.R. This has four steps to it.

Situation - describe the context / background

Task - what was your responsibility?

Action - what did you do?

Result - what were the results of your actions?

The STAR structure is often used to help formulate responses to competency based interview questions - but can also be used to help provide you with a higher level structure for a slightly longer impromptu presentation.

Delivering with impact

Nerves can take over and hinder your performance. Also – when you are nervous you are more likely to rush and this will make you feel even more nervous. To help control your nerves, take two deep breaths before you start, breathing out for as long as possible to help release tension and encourage you to slow down when you start to speak.

Art of Listening

Meaning of Listening

Listening is considered to be the one of the most important part of the oral communication. The term is used in order to make oral communication effective. Poor listening skills of an individual may affect the individual very badly specially in an organization where the maximum number of time a person spent in communication therefore it is very much important if will talk from organizational prospective because a effective and active listening by an individual plays a very important role in contributing towards the success of the business.

Difference between listening and hearing

Hearing is one of the five senses of a person and it is the ability to perceive sound by detecting vibrations through an organ such as the ear. According to Merriam-Webster, hearing is “the process, function, or power of perceiving sound; specifically: the special sense by which noises and tones are received as stimuli.” In hearing, vibrations are detected by the ear and then converted into nerve impulses and sent to the brain. A person who is unable to hear has a condition known as deafness. Hearing occurs even in sleep, where the ear processes the sounds and passes them on to the brain, but the brain does not always react to the sound. Listening also known as ‘active listening’ is a technique used in communication which requires a person to pay attention to the speaker and provide feedback. Listening is a step further than hearing, where after the brain receives the nerve impulses and deciphers it, it then sends feedback. Listening requires concentration, deriving meaning from the sound that is heard and reacting to it.

The listening process

- 1.Sensing or Selection
- 2.Interpretation
- 3.Evaluation
- 4.Response
- 5.Memory

Types of Listening

1.**Passive Listening:** It involves Physical presence but mental absence of the listener. The listener is merely hearing out not absorbing the message.

.**Marginal Listening:** In marginal listening small pieces of the message are listened to and assimilated. He allows information to sink only in bits and pieces.

3.**Sensitive Listening:** Listener tries to understand the viewpoint of the speaker. If taken in isolation, sensitive listening results in onesided sympathetic stand.

4.**Active Listening:** The listener absorbs all that is being said and moves in accordance with the intent of the speaker. The listener asks questions to understand the viewpoint of the speaker.

UNIT – III

The deal is **nobody gets a job unless they first have a job interview**. That's pretty obvious, right? So how do you get a job interview? There are a few ways, but the focus for today will be getting out the old resume and preparing to be interview bait. Some sticking points to remember are that everybody else applying for the job has a resume also.

Keep in mind that any resume should be:

- 1.**Visual**
- 2.**Relevant**
- 3.**Concise**
- 4.**Exciting**
- 5.**Truthful**

Visual Impact

Take any standard resume that you could produce using the Resume Wizard in Microsoft Word. My resume is shown here (with some information changed for privacy concerns). Comparing these two side-

by-side, I think most people would agree that mine looks most visually distinct. Later, we'll look at some of the really simple tools available in Microsoft Word for making print media look really cool.

Relevant

Nothing says low-class as much as an impersonal email or resume. I take that back, a misaddressed or irrelevant message is lower than that. So put forth a little bit of effort. Personalize the information. If you're applying for a highschool social studies position, then your mad finger-painting skills may not be the most impressive thing to the principal. By the same token, your experience working at Burger King ten years ago is not as impressive as your babysitting service twenty years ago. This is the reason that I wrote about some things you can do while still in school that will make getting a job that much easier.

Concise

Deleting extraneous information makes the powerful stuff that much more powerful. There is nothing that says resumes have to be limited to one page in today's world. However, I have consciously trimmed mine to one page simply to make it more powerful. We'll go into some of the tricks I use tomorrow, but *Selected Work Experience* seems more impressive to me than *Teaching Experience*. Be careful that you don't leave off important information. Do not be so zealous to trim to one page that you create a gap of 3 years here and 2 years there. That looks like you are inconsistent and, while it could be answered by a simple phone call from the interviewer, it could also be answered by a second page. Save trees, but also save time for those who are weeding through applicants.

Exciting

Market yourself. **Tell them a story** that they want to hear. Show them the benefits of hiring you. If you can convey who you are and what it is that you are looking for in your life, then they are more inclined to consider you and give you the interview. Remember, everybody else applying for the job has a resume. But most of them won't tell the interviewer a story about why they should meet. You will, right?

Truthful

In telling your story, you have to be truthful. Don't make stuff up just because you think it's what they want to hear. Start with where you are. Find a way to make that relevant and helpful to the betterment of the school environment. Exaggeration will merely bite you in the end. Thursday, we'll look at how to

make a visually effective resume that jumps off the page and into the good graces of the interviewer. I know you won't want to miss that. Selling yourself to an employer is your first challenge, and your resume will be your sales pitch. Sales resumes need to be results-oriented, emphasizing how you will contribute to your employer's bottom line. Start by creating a profile or career summary that highlights your relevant skills and value to potential employers. Include the main reasons an employer should call you for an interview, and clearly show your areas of expertise and industry knowledge. For example, if you are pursuing a pharmaceutical sales representative position, those keywords and your supporting knowledge should be in the profile. This section is perfect for exhibiting the drive, energy and enthusiasm that is so important in the sales profession.

Document Your Achievements

A need to continually achieve is key to sales success. Prove you are an achiever. Document your three biggest victories, and be prepared to reel off a list of at least seven other significant wins in your life from school, sports, music, class politics, etc. You will achieve again for the employer, because past behavior is the best predictor of future behavior. You may not have sales success, but you have had success in other areas. Success leaves clues.

It's deceptively easy to make mistakes on your resume and exceptionally difficult to repair the damage once an employer gets it. So prevention is critical, especially if you've never written one before. Here are the most common pitfalls and how you can avoid them.

Avoid Common Resume Mistakes

1. Typos and Grammatical Errors

Your resume needs to be grammatically perfect. If it isn't, employers will read between the lines and draw not-so-flattering conclusions about you, like: "This person can't write," or "This person obviously doesn't care."

2. Lack of Specifics

Employers need to understand what you've done and accomplished. For example:

A. Worked with employees in a restaurant setting.

B. Recruited, hired, trained and supervised more than 20 employees in a restaurant with \$2 million in annual sales.

Both of these phrases could describe the same person, but details and specifics in example B will more likely grab an employer's attention.

3. Attempting One Size Fits All

Whenever you try to develop a one-size-fits-all resume to send to all employers, you almost always end up with something employers will toss in the recycle bin. Employers want you to write a resume specifically for them. They expect you to clearly show how and why you fit the position in a specific organization.

4. Highlighting Duties Instead of Accomplishments

It's easy to slip into a mode where you simply start listing job duties on your resume. For example:

- Attended group meetings and recorded minutes.
- Worked with children in a day-care setting.
- Updated departmental files.

Employers, however, don't care so much about what you've done as what you've accomplished in your various activities. They're looking for statements more like these:

- Used laptop computer to record weekly meeting minutes and compiled them in a Microsoft Word-based file for future organizational reference.
- Developed three daily activities for preschool-age children and prepared them for a 10-minute holiday program performance.

5. Going on Too Long or Cutting Things Too Short

Despite what you may read or hear, there are no real rules governing the length of your resume. Why? Because human beings, who have different preferences and expectations where resumes are concerned, will be reading it. That doesn't mean you should start sending out five-page resumes, of course. Generally speaking, you usually need to limit yourself to a maximum of two pages. But don't feel you have to use two pages if one will do. Conversely, don't cut the meat out of your resume simply to make it conform to an arbitrary one-page standard.

6. A Bad Objective

Employers do read your resume's objective statement, but too often they plow through vague pufferies

like, “Seeking a challenging position that offers professional growth.” Give employers something specific and, more importantly, something that focuses on their needs as well as your own. Example: “A challenging entry-level marketing position that allows me to contribute my skills and experience in fund-raising for nonprofits.”

7. No Action Verbs

Avoid using phrases like “responsible for.” Instead, use action verbs: “Resolved user questions as part of an IT help desk serving 4,000 students and staff.”

8. Leaving Off Important Information

You may be tempted, for example, to eliminate mention of the jobs you’ve taken to earn extra money for school. Typically, however, the soft skills you’ve gained from these experiences (e.g., work ethic, time management) are more important to employers than you might think.

9. Visually Too Busy

If your resume is wall-to-wall text featuring five different fonts, it will most likely give the employer a headache. So show your resume to several other people before sending it out. Do they find it visually attractive? If what you have is hard on the eyes, revise.

10. Incorrect Contact Information

I once worked with a student whose resume seemed incredibly strong, but he wasn’t getting any bites from employers. So one day, I jokingly asked him if the phone number he’d listed on his resume was correct. It wasn’t.

Guidelines for a good resume

Here are five rules to help you write a resume that does its job:

- 1. Summarize Your Unique Value**
- 2. Communicate with Confidence**
- 3. Watch Your Language**
- 4. Key in on Keywords**
- 5. Keep it Concise**

What do these really mean?

1. Summarize Your Unique Value

A resume should begin with a **Summary** (or if you're a student, new grad, or career changer, an **Objective**). Use this space to tell employers who you are and how your skills and qualifications meet their needs. Although your real objective may be to get away from your micro-managing boss or shorten your commute, *don't say that on your resume!*

Your Summary or Objective is where you explain how and why you are uniquely qualified to contribute to the company.

Bonus: *Once you've crafted a solid message that summarizes your value, you can use it as the basis for your response to every hiring manager's favorite line: "Tell me about yourself."*

2. Communicate with Confidence

Tell the potential employer what you've accomplished in your current and

previous roles to show how you made a difference. This is not the time to be humble or modest, or to assume the employer will read between the lines. For instance, if your resume just states the facts, without context (e.g., "Sold 50,000 widgets between January and June"), the reader won't know if that's better, worse, or the same as what the company had achieved in the past. But a confident statement like "Boosted widget sales 35% in the first six months" or "Increased widget sales from 40K to 50K within six months" is bound to jump off the page.

3. Watch Your Language

Don't start your sentences with **I** or **We** or **Our**. In fact, **don't even use full sentences.**

4. Key in on Keywords Here's an awful truth: Resumes, in many cases, are not even read. Rather, they're scanned (either by a machine or by someone who is not the hiring manager). What they're scanning for is keywords or phrases that match their hiring criteria.

Not sure what keywords to put in your resume? Read the job description for a position that interests you, as well as descriptions for similar jobs. Then read your target companies' web sites.

HOW TO FACE AN INTERVIEW BOARD

New job opportunities are arising fast in every part of the world. Anyone can see mushrooming call-centers and branch offices of multinationals. Everyday new institutes are opening to prepare aspiring young men and women for lucrative jobs in public and private sector. Interview for an anticipated post

has become a very important step in the professional life of a young person. Despite trying their best for this moment things go wrong for most of the candidates.

BASIC PREPARATION

Don't bang your head with an interview board like an enthusiastic teenager (they have a passion for driving fast bikes mindlessly and get bruises). Keep your mind cool and prepare well. Refresh your general knowledge, rehearse answering the expected questions (in front of a mirror now and then). Arrange your certificates properly in an attractive folder. Avoid taking irrelevant documents. You generally need: copies of your educational and experience certificates, bio-data, and application (sent to the company). Be sure you fulfill the requirements mentioned in the advertisements of the company otherwise you may become a laughing stock as soon as your folder is examined. Never push your folder for the perusal of the board until it is demanded. Don't forget to pray to God for success before leaving for an interview. It will definitely boost your confidence. Take light breakfast/meals as it will keep your mind light and you will feel less nervous. A heavy stomach can cloud your mind, leading to nervousness.

BE WELL DRESSED

Wear a jacket, ride your bike, and appear before the board, like your favourite hero but you end up becoming a villain. Although this may be the right step for many to attract the attention of some people but you need a formal dress to woo the board. You must feel comfortable and confident in your dress. It should be clean and well ironed. Avoid gaudy or bright-coloured suits. You have to wear a moderate dress according to the occasion. Polish your shoes, have a hair cut, if needed.

CONTROLLING NERVOUSNESS

Nervousness is natural. Even the best of the orators, businessmen or politicians used to get nervous. So there is nothing extraordinary about your nervousness. Moreover, you can control it easily. In fact, nervousness is a form of stored mental energy. The best solution is to use it positively:

1. Concentrate in preparing,
2. Shine your personality (good dress, shoes, haircut etc.) ,
3. Eat light food,
4. Pray to god,

SPEAK WELL

A candidate must have command over language, pronounce the words clearly and have a good store of vocabulary. Hissentence-making should be grammatically correct. Answers should be in brief and to the point. Lack of fluency, bad grammar, incorrect pronunciation or answering in a hurry without listening to the board properly can surely land you in trouble. Such candidates never get good posts in big companies. Join a good institute to improve your sentence-making,pronunciation and fluency.

EXPECTED QUESTIONS

Interview is basically a series of questions asked from the interviewee to test his ability, wisdom and personality.{Interview-boards of many big companies also have an expert who can understand human psychology, and who is capable enough to read the mind of a candidate by studying his body language). We can divide the expected question in three categories: 1. Questions relating to personal information of a person (family background, interests, education, experience etc.); 2. Questions relating to his knowledge about the work he will be responsible for in the company; 3. Questions to check the personality of a person – his nature, ideology, decision-making & problem-solving ability etc. Companies may have different set of questions according to their work culture. However most of the questions are related to the categories of the questions given above. Some irrational questions are also asked by some interview boards. Don't panic in such a situation. Maintain your self-confidence and answer in a simple and straightforward way. Interview board may be checking your psychological structure. If you get irritated or try to be over smart you will definitely be discarded from the list of expected winners.

DON'T BE CLEVER

Many candidates try to act cleverly to show their intelligence but it always misfires. Remember that every company needs sincere and hard-working employees. Clever cats are kept at a distance. Listen to the board members carefully. Whenever a question is asked from you answer in brief. If you don't understand a question request the board politely to repeat it. Talk like you are obeying your seniors, giving them all the respect. Speak confidently and clearly. Your voice must be loud enough so that the board may hear and understand it easily. Don't give any unnecessary information or you may become a big bore for the board. In case, you are unable to think answer of a question, just say sorry. Avoid repeating the same sentence or phrase. Never contradict the board even if you have a vast knowledge of the subject. Don't criticize any other community, company or person while giving your answers. Control

your emotions. Imposing or an egoist personality is never liked by a board. Humility and politeness are your winning edges. At the end of the interview don't forget to say the board 'thank you very much'. Before leaving move a few steps backward respectfully and come out.

IMPROVE BODY-LANGUAGE

The importance of body language is yet to be recognized fully. Sometimes body language of a candidate even proves weightier than his/her ability. It will solve this puzzle of your mind: why do some candidates are selected although they are less capable than others in every way? If your body language is negative it will create an artificial distance between you and the interview board. The board will be confused and unable to judge your abilities. A good body language carries your confidence to the board, assuring and convincing them about your ability. They think you will learn fast even if you are less capable at the time. Your good dress, smile, confidence and good manners convey a language, which affect the interview greatly. Don't sit on the chair like a hard rock. Sit easily, ready to face the bombardment of questions. If you panic at the first salvo you will lose the game. Instead, feel relaxed and think that you are here to gain experience and learn.

PROPER BODY POSTURE

Interview skills and communication skills are not just about **speech techniques and structures**. You may have come across studies or statistics which state that up to 60% of the impression that you make is through your body language. Whatever the reality behind this statement, it is undoubtable that the way you dress and behave at an interview will strongly influence the person who is looking at you, even if it is subconscious.

To make a strong impression, there are a number of rules regarding correct body language that you need to reflect upon and adopt:

Choose a good position within the room

At an interview, you will normally be directed to a specific seat (i.e. you will have no choice). However, interviews can often be conducted in oversized environments (e.g. a meeting room with a table for 8 when there are only 3 of you). Make sure you choose a seat which enables you to see everyone involved without having to rotate your head exaggatedly. In most cases, it may be best to hover around to see which chairs the interviewers are aiming for before making your selection. If there is a window, choose a chair that faces it so that your face is lit from the front, unless there is good lighting all round. If you turn your back to the window, the interviewers may see you in sepia!

Maintain a good posture

If you are being interviewed at a table, make sure that you are not too close to the table. As a rule of

thumb, your body language should be such that if you let your arms fall loosely on the table in front of you, they should fall with your elbows slightly outside of the table. If your elbows are actually on the table then you are too close. If your elbows are more than a few inches away (or you have to lean forward a lot to put your hands on the table) then you are too far. For most people, the ideal distance between chest and table is about 4 inches. Plant both feet onto the ground so that you remain stable; and put your hands on the table (people who place their hands below the table come across as having something to hide). Keep yourself upright, with a slight slant forward and relax our shoulders. Slouching is bad body language! If there is no table (or only a low table) then simply rest your hands on your lap.

Don't be afraid to "own the space"

Just because you are under observation, it does not mean that you should recoil in a corner. It is okay to stand or sit with your legs slightly apart, and in fact, it is a sign of confident body language (don't overdo it though, it would become indecent!)

Limit your hand and arm movement

It is perfectly okay in your body language to move your arms and hands around, and if that is the way that you normally behave then don't try to become someone else. Your personality and enthusiasm are as important as everything else. However make sure that such movements do not become distracting and do not take the focus away from your face. To achieve this, make sure that your movements are limited to the corridor in front of you, never higher than your chest, and never under the table. If there is no table, you can let your hands go as far down as your lap.

If your hands go outside towards the left or right, your interviewers will follow them and may stop concentrating on you. If your hands go over chest level, you will most likely obscure your lips or eyes. If you have a tendency to fidget in a very distracting manner, intertwine your fingers and rest your hands on the table. Whatever you do, never cross your arms. It will make you look unreceptive, guarded and lacking in confidence.

Smile

A nervous smile is better than no smile at all. No one wants to recruit a grumpy person or someone who looks like they are not enjoying themselves. Good interviewers will understand that you may be nervous and will make attempts to put you at your ease. Make sure you reward their efforts with an easy smile. No need to overdo it. It is not a contest for straight teeth, but simply a reasonable attempt to engage with them. Smile lightly also when you are being introduced to each member of your panel. With this body language you can build a good rapport. It is also perfectly acceptable to laugh if the situation warrants it (but avoid making jokes just for the sake of introducing a laugh into the conversation. You'll probably

end up being the only one laughing, and you'll soon be crying.)

Maintain eye contact

If you do not make eye contact, you will come across as evasive and insecure which is poor body language. If you stare at people too much, you will make *them* insecure. There are two situations here: either you are being interviewed by just one person, in which case you will have no choice but to look at them all the time; or you are being interviewed by more than one person. If this case, then look mostly at the person who is asking you the question, and occasionally glance aside to involve the others (they will be grateful that you are trying to involve them into the conversation even if they have not asked that particular question).

Beware of the props

If you have a pen with you, avoid fiddling with it. It will only end up flying in the wrong direction. Similarly, if they offer you a drink (tea, coffee, water, etc), make sure that you can cope with it and that won't need to go to the loo or start crossing your legs half-way through the interview. Generally you should avoid picking up any drink if you can. Other than the fact that it may end up down your shirt or on your lap, the movement of the water in a glass that you have just picked up will reveal just how nervous you are.

Mirror the interviewer's behaviour

Mirroring (i.e. acting similarly) to someone is an indication that there is a connection through body language. It should happen normally but you may be able to influence it too, if only to give the interviewer the feeling that you are getting on. For example, if the interviewer is sitting back then you may want to sit back a little too; if he leans forward, you may lean forward to. Be careful not to overdo it though and do not mirror instantly, otherwise it will look like some kind of Laurel and Hardy sketch.

And relax ...

At the end of the day, you can't spend all your energy focussing on body language. There is no point having a **brilliant body language** if you are talking rubbish. Bearing in mind that **body language is a reflection of your level of confidence**, it is important that you build your confidence up first through good preparation and then go to the interview relaxed. You will be surprised of how much of the above you can do naturally.

IMPORTANCE OF GESTURE

In a job interview, it's likely that your body language will have more of a positive impact on your success than anything you say. Consider the following scenarios: As you're waiting to be called in for a job interview, do you patiently check emails on your phone, or do you nervously practice answers to

tough questions? When introduced to your interviewer, do you make strong eye contact and offer a firm handshake? And as the meeting begins, do you speak passionately and expressively, or are your responses rehearsed and carefully controlled?

In each of these examples, your body language is giving off important signals about what kind of employee you would be. In fact, studies indicate that body language accounts for a full 55% of any response, while what you actually say accounts for just 7%. The remaining 38% is taken up by "paralanguage," or the intonation, pauses and sighs you give off when answering a question. In other words, even if your spoken answers convey intelligence and confidence, your body language during job interviews may be saying exactly the opposite.

"Our nonverbal messages often contradict what we say in words," says Arlene Hirsch, a Chicago career consultant. "When we send mixed messages, or our verbal messages don't agree with our body language, our credibility can crumble because most smart interviewers will believe the nonverbal over the verbal."

Unemployed job seekers, for example, are often so traumatized by their long and difficult job hunts that they appear downcast, even when discussing their strengths. Tough questions can throw them off balance, and their anxiety may cause them to fidget or become overly rigid. Since nonverbal communication is considered more accurate than verbal communication, this kind of behavior reveals your *inner* confidence, say career counselors. The words that you say during an interview can be deceiving – sometimes people don't mean what they say or say what they mean – but your job interview body language is subconscious, and thus more spontaneous and less controlled.

Still, many people discount the importance of job interview body language because they've been trained to place more emphasis on spoken words instead. To become more adept at interpreting and using body language, career advisers suggest that you heighten your awareness of nonverbal signals and learn to trust your "gut" instinct.

Once you've learned to harness your body's nonverbal forms of communication, use the following tips to accentuate your job interview body language so that you appear more professional and self-assured:

STEPS TO SUCCEED IN INTERVIEWS

There are a multitude of opinions on how to succeed in a job interview. In my short tenure as a recruiter, I've spoken to hundreds of successful candidates about their interviewing experiences. The following is a short list full of helpful information that will give you the best possible chance for that same success!

Study information about the company through its website and other forms of information that you can find out about them, such as professional networking website profiles. Learn everything that you can about the company and don't forget to take notes to study in advance of the interview (if you have

reasonable time before the interview). The more you learn the better chance you'll have to advance through to the next step in the hiring process.

Always be on time for your interview. Arrive 10-15 minutes early. You may need to fill out forms related to the job, and it will help that you are finished before the actual interview is scheduled. Make sure you have a contact phone number available in case of an emergency during your travel time to the interview. When you arrive, ask the receptionist for the name of the interviewer, and write it down! The receptionist is your first point of contact. Be genuine as you communicate, as they will have an impact with any comments they may make to the interviewer.

Always dress for professionally. If you look professional, you'll likely be treated with professional respect. Wear clean and comfortable attire. Avoid casual clothing such as blue jeans and T-Shirts (even if you would wear this type of clothing during actual work hours).

There is a difference between being cocky and confident. Confidence comes from being prepared. Cockiness comes from living the life of an idiot. Don't be so foolish to think that you are better than anyone else interviewing for the position. Just be as prepared as you can possibly be. If you're prepared, you'll come across as confident. If not...well!

Your resume is an extremely important tool. Have extra copies available to the interviewer. Make sure (in advance) that your resume is very well prepared. It should "fit" the job description and meet the minimum requirements in the description. Make sure you have indicated on your resume what each company you have worked for did, what did they produce or manufacture. This will give additional insight to the interviewer about your work history.

Be polite and considerate of the person interviewing you. Their time is valuable and you must respect this. Be prepared to ask short questions that will get their attention and more than a yes or no answer. Show them that you have come prepared! Show them that you are genuinely interested in getting an offer.

Do you ask for the position? YES! Just remember to use caution in how you ask. I can't begin to tell you how to ask. Only you will know based on how the interview has gone. If the interview went well, just be polite, avoid over confidence and ask this way...I believe that I've shown you I'm qualified for the position, I'd like to ask for the job!

You've always heard that it is a great idea to send a thank you note or email to the interviewer. This is still great advice, and one point that should never be missed! How long you should wait is a matter of how your interview went. Be realistic about the success of the interview. If it did not go well, send one

anyway, indicating your understanding of the end result and the fact that you value their time spent (this is a great method of bridge building for future interviews). If it went well, send one indicating the same things and include a reminder of something you both had in common related to the company

Practice mock interview in classrooms with presentations on self

Self-presentation is expressive. We construct an image of ourselves to claim personal identity, and present ourselves in a manner that is consistent with that image. If we feel like this is restricted, we exhibit reactance/be defiant. We try to assert our freedom against those who would seek to curtail our self-presentation expressiveness. A classic example is the idea of the "preacher's daughter", whose suppressed personal identity and emotions cause an eventual backlash at her family and community. People adopt many different impression management strategies. One of them is ingratiation, where we use flattery or praise to increase our social attractiveness by highlighting our better characteristics so that others will like us

HIGHLIGHT YOUR POSITIVE AND NEGATIVE TRAITS OF YOUR PERSONALITY

Positive Trait

Dedicated

If you are someone who takes your work seriously and consistently performs at peak levels, this is a personality trait you should highlight in your job interviews. Employers are looking for employees who have a passion for their jobs and aren't simply punching a time clock to earn a paycheck. If you take pride in your work and strive to continually produce the best product or service possible, make this personality trait a key aspect of your interview highlights.

Timely

Being on time for work and on deadline for projects is a key "must" for employers. If you are punctual, deadline- focused and prepared for every meeting an hour ahead of schedule, highlight this information for potential employers. This personality trait demonstrates that you are unlikely to make clients wait, leave customers unattended or fail to prepare an important presentation on schedule.

Personable

Kindness, graciousness and professionalism are all significant personality traits employers look for. Having a sense of humor, a calm attitude and an overall outgoing, personable nature marks you as someone who is easy to get along with and work with. Even if you are interviewing for a serious type of position, or don't want to be viewed as a pushover, a pleasant personality and genial nature can earn extra points with potential employers.

Creative

If you can brainstorm and think outside the box, emphasize the creative aspect of your personality during job interviews. Regardless of the industry, creative thinkers are valued corporate assets. Outline ways in which your creativity has been of value to employers in the past.

Loyal

Perhaps one of the greatest personality traits employees value in staffers is loyalty. If you stick with a project through thick and thin, share credit with others and never take advantage of professional relationships, this is something potential employers should know about you. Emphasize your loyal personality by outlining your longevity in past positions.

Negative Trait

Flip or Over-Confident Attitude

Presenting yourself in a confident manner is good. Coming across as a know-it-all who considers himself better than everyone else is bad. Don't exaggerate your skills or performance abilities and don't outline achievements in a boastful manner. You can set yourself apart as a highly qualified and talented candidate without resorting to bragging or coming across as an overly-aggressive person who will be hard to manage or work with.

Hostility

Even if you previously worked for the most abusive boss in the industry, avoid talking poorly about him or other colleagues during your interview. Don't complain about past working conditions or low salaries, and never malign a company or its products and services. If you had a bad work experience and you're asked why you left the job, simply explain that you decided to explore new opportunities. Chances are, if you're coming from a position in a similar industry and you were employed with a problem company, your interviewer knows about the poor working conditions. You'll be respected for holding your tongue and not putting down your old boss.

Electronics Use

Turn off your cell phone before you even walk into the interview setting. Never check your e-mail or text messages during an interview. This rule applies to all electronic devices, including laptops and tablets. If you forget to turn your phone off and it rings during the interview, apologize and silence it or let it go to voicemail. Never pick it up and begin a conversation.

Vulgar Language

Don't use slang or poor grammar during an interview and never use foul language. Try to avoid words and expressions such as, "yeah," "ya know" and too many "um's." the way you present yourself verbally says a lot about how you will interact with clients and customers, so speak clearly and authoritatively

with professionalism and respect.

DEALING WITH PEOPLE WITH FACE TO FACE

Communicating meaningfully is becoming more difficult than ever before. While technology has created an ever-increasing number of ways to communicate rapidly over great distances, many people are now so well insulated and protected by these devices we use that we are losing the skills and abilities to communicating in the most influential way—face to face. There's a real danger to the maintenance and perpetuation of meaningful communications and personal and professional relationships. If you become overly dependent on e-mail or text messages, you focus on the object, but not the person. As a result, you become uncomfortable communicating face to face.

Tweets, text messages, e-mail, and Facebook posts all transmit words over distances so they can be received without the sender's presence. The human element and context are absent.

These messages are typically short, sequential, and directed. There's no instantaneous interaction or connection that allows the other person to understand the tone, inflection, or emotion that is carried with the words. The sender cannot effectively project the elements of trust, confidence, credibility, and concern that are crucial to developing and building a relationship. That failure to convey the feelings that accompany the words so people build trust, credibility, and understanding can have a phenomenal impact on business and success, including:

- ☐ Miscommunication and understanding
- ☐ Wasted time
- ☐ Lost profits
- ☐ A minimized ability to effectively project trust, confidence, and credibility to build relationships

Meaningful communications that carry these powerful and important characteristics can only be achieved in face-to-face interactions. Communicating with impact and achieving influence is not only about what you say—it's also how you say it. You have influence on others because you see their face, observe and experience their message, and actively listen and engage their interest to build relationships. There are also certain topics of conversation where face-to-face communication will be the best way to achieve clarity and understanding needed for mutual success and beneficial action. These include:

- ☐ Negotiating salaries, vacations, and termination
- ☐ Resolving a dispute, challenge, or conflict between two or more people or organizations
- ☐ Seeking clarification after written communications have failed

Face-to-face communication is a crucial skill. It requires you to focus. You must be comfortable in the presence of other people for more than a few minutes. Communicating with impact and influence face to face also requires discipline, determination, and self-control.

To increase your impact and influence, begin applying these eight must-haves:

1. Make your moments together count. Everyone has the right to speak. Listen before you speak. Earn the right to be heard. Think about what you want to say before you say it. Make every communication moment worth your and your listener's time. Every word counts. Think before you speak. Tailor what you say to meet your listener's needs.

2. Pay attention by listening for the unspoken emotions. Concentrate on the speaker closely. Focus intently on their face. Do not let your eyes dart away and drift off, since that signals you are no longer paying attention. Do not interrupt. Wait to speak only when the person has finished what they want to say. Hear their words and read their face so you gain maximum understanding of the why behind their words.

3. Honor the other person's time. Prepare and get to the point quickly by speaking in short and concise sentences. Replace your non-words ("uh," "um," "so," "you know...") with a pause to find your thought. Avoid rambling and cluttering your message with unnecessary points. Ask for a clear and specific action. Don't take 20 minutes when you only asked for 10.

4. Prepare for your face-to-face meeting ahead of time. K.N.O.W. your listener.

☐ **K:** What does your listener know about your topic?

☐ **N:** What does your listener need to know to take the action you want them to take in the time frame you have for this conversation?

☐ **O:** What is your listener's opinion about your topic?

☐ **W:** Who is your listener? What additional information do you know about your listener to help you customize your message for them?

Tailor your agenda and message to achieve the understanding you need and to influence your listener to act on what you have to say.

5. Watch your body language. Avoid non-verbal abuse. Every movement you make counts. Control your facial expressions. Don't smile, snicker, whistle, roll your eyes, grimace, look sideways, wink, or send the evil eye. Your behavior and non-verbal cues are as important as the words you say. Don't

fidget, act nervous, express fear, or allow your posture to convey uncertainty, insincerity, lack of caring, arrogance, overconfidence, dismay, or criticism.

6.Be sincere and authentic. Speak in your authentic voice. Be sincere, be genuine, and allow others to see the real you.

7.Maintain the power of the floor. Be interesting. Watch for the signs that you are no longer the center of attention:

- ☐ Your listener begins working on their Blackberry, iPad, iPhone, etc.
- ☐ Your listener starts nodding off.
- ☐ Your listener begins to have side conversations.
- ☐ Your listener interrupts you.

Stop. Earn their attention. Get back on track.

8. Ask for feedback. Face-to-face communications is a two-way street. Balanced feedback allows people to be relaxed and comfortable. However, when people start feeling comfortable, they also may become lazy and lose their professionalism. Don't forget who you are and what you are doing. Ask for specific feedback on things such as the points you raised, the manner in which you presented, the way you responded. Give yourself feedback by asking, "What worked and what didn't work?"

UNIT –IV

QUALITIES OF A LEADER

Having a great idea and assembling a team to bring that concept to life is the first step in creating a successful business venture. While finding a new and unique idea is rare enough; the ability to successfully execute this idea is what separates the dreamers from the entrepreneurs. However you see yourself, whatever your age may be, as soon as you make that exciting first hire, you have taken the first steps in becoming a powerful leader. When money is tight, stress levels are high, and the visions of instant success don't happen like you thought, it's easy to let those emotions get to you, and thereby your team. Take a breath, calm yourself down, and remind yourself of the leader you are and would like to become. Here are some key qualities that every good leader should possess, and learn to emphasize.

Honesty

Whatever ethical plane you hold yourself to, when you are responsible for a team of people, it's important to raise the bar even higher. Your business and its employees are a reflection of yourself, and if you make honest and ethical behavior a key value, your team will follow suit.

Ability to Delegate

Finessing your brand vision is essential to creating an organized and efficient business, but if you don't learn to trust your team with that vision, you might never progress to the next stage. Its important to remember that trusting your team with your idea is a sign of strength, not weakness. Delegating tasks to the appropriate departments is one of the most important skills you can develop as your business grows. The emails and tasks will begin to pile up, and the moreyou stretch yourself thin, the lower the quality of your work will become, and the less you will produce. The key to delegation is identifying the strengths of your team, and capitalizing on them.

Communication

Knowing what you want accomplished may seem clear in your head, but if you try to explain it to someone else and are met with a blank expression, you know there is a problem. If this has been your experience, then you may want to focus on honing your communication skills. Being able to clearly and succinctly describe what you want done is extremely important. If you can't relate your vision to your team, you won't all be working towards the same goal.

Sense of Humor

If your website crashes, you lose that major client, or your funding dries up, guiding your team through the process without panicking is as challenging as it is important. Morale is linked to productivity, and it's your job as the team leader to instill a positive energy. That's where your sense of humor will finally pay off. Encourage your team to laugh at the mistakes instead of crying. If you are constantly learning to find the humor in the struggles, your work environment will become a happy and healthy space, where your employees look forward to working in, rather than dreading it.

Confidence

There may be days where the future of your brand is worrisome and things aren't going according to plan. This is true with any business, large or small, and the most important thing is not to panic. Part of your job as a leader is to put out fires and maintain the team morale. Keep up your confidence level, and assure everyone that setbacks are natural andthe important thing is to focus on the larger goal. As the leader, by staying calm and confident, you will help keep the team feeling the same. Remember, your team will take cues from you, so if you exude a level of calm damage control, your team will pick up on that feeling. The key objective is to keep everyone working and moving ahead.

Commitment

If you expect your team to work hard and produce quality content, you're going to need to lead by example. There is no greater motivation than seeing the boss down in the trenches working alongside

everyone else, showing that hard work is being done on every level. By proving your commitment to the brand and your role, you will not only earn the respect of your team, but will also instill that same hardworking energy among your staff. It's important to show your commitment not only to the work at hand, but also to your promises.

Positive Attitude

You want to keep your team motivated towards the continued success of the company, and keep the energy levels up. Whether that means providing snacks, coffee, relationship advice, or even just an occasional beer in the office, remember that everyone on your team is a person. Keep the office mood a fine balance between productivity and playfulness.

Creativity

Some decisions will not always be so clear-cut. You may be forced at times to deviate from your set course and make an on the fly decision. This is where your creativity will prove to be vital. It is during these critical situations that your team will look to you for guidance and you may be forced to make a quick decision. As a leader, it's important to learn to think outside the box and to choose which of two bad choices the best option is. Don't immediately choose the first or easiest possibility; sometimes it's best to give these issues some thought, and even turn to your team for guidance. By utilizing all possible options before making a rash decision, you can typically reach the end conclusion you were aiming for.

Ability to Inspire

Creating a business often involves a bit of forecasting. Especially in the beginning stages of a startup, inspiring your team to see the vision of the successes to come is vital. Make your team feel invested in the accomplishments of the company. Whether everyone owns a piece of equity, or you operate on a bonus system, generating enthusiasm for the hard work you are all putting in is so important. Being able to inspire your team is great for focusing on the future goals, but it is also important for the current issues. When you are all mired deep in work, morale is low, and energy levels are fading, recognize that everyone needs a break now and then.

KNOWING YOUR SKILLS AND ABILITIES

One of the most important things you can do before looking for work or an alternative career is to consider what skills and abilities you already have. These are your most valuable assets and are very important.

Three kinds of skills you need in the world of work are:

- ☐ technical;
- ☐ transferable; and

☐ Personal.

Technical skills are the specialized skills and knowledge required to perform specific duties, sometimes referred to as ‘work skills’. For example:

Each one of these skills is made up of specific skills a person must be able to do in order to complete technical tasks.

Transferable skills are the skills required to perform a variety of tasks. They are your greatest asset as they can be ‘transferred’ from one area of work to another.

These skills can be useful when you are trying to make a career change.

Personal skills are the individual attributes you have such as personality and work habits. These often describe what you are like and how you would naturally go about doing things.

☐ Working under pressure	☐ Honest and reliable	☐ Has initiative
☐ Trustworthy	☐ Fast learner	☐ Planning/organisational
☐ Self-motivated	☐ Professional	☐ Loyal

GROUP DISCUSSION TECHNIQUES

Discussions of any sort are supposed to help us develop a better perspective on issues by bringing out diverse view points. Whenever we exchange differing views on an issue, we get a clearer picture of the problem and are able to understand it. The understanding makes us better equipped to deal with the problem. This is precisely the main purpose of a discussion. The dictionary meaning of the word Group Discussion is to talk about a subject in detail. So, group discussion may refer to a communicative situation that allows its participants to express views and opinions and share with other participants. It is a systematic oral exchange of information, views and opinions about a topic, issue, problem or situation among members of a group who share certain common objectives.

GD is essentially an interactive oral process. The group members need to listen to each other and use

voice and gesture effectively, use clear language and persuasive style.

GD is structured: the exchange of ideas in a GD takes place in a systematic and structured way. Each of the participants gets an opportunity to express his/her views and comments on the views expressed by other members of the group. GD involves a lot of group dynamics, that is, it involves both -person to person as well as group to group interactions. Every group member has to develop a goal oriented or group oriented interaction. A participant needs to be aware of needs of other group members and overall objectives of the discussion.

Definition: Group discussion may be defined as – a form of systematic and purposeful oral process characterized by the formal and structured exchange of views on a particular topic, issue, problem or situation for developing information and understanding essential for decision making or problem solving. There are several types of oral group communication. In Public Speaking, the speaker is evaluated by the audience; however there is not much interaction between audience and speaker. A chairperson conducts the meeting and controls and concludes the deliberations.

Group Discussion differs from debate in nature, approach and procedure. Debates include representation of two contrasting viewpoints while GD can include multiple views. A GD may help achieve group goals as well as individual needs. The examiner observes the personality traits of several candidates who participate in the G.D.

Importance of Group Discussion skills

A Group Discussion helps problem solving, decision making and personality assessment. Whether one is a student, a job seeker, a professional engineer or a company executive one needs effective GD skills. Students need to participate in academic discussions, meetings, classroom sessions or selection GDs for admission to professional courses. A job- seeker may be required to face selection GDs as part of the selection process. Professionals have to participate in different meetings at the workplace .In all these situations, an ability to make a significant contribution to group deliberation and helping the group in the process of decision making is required. The importance of GD has increased in recent times due to its increasing role as an effective tool in a) problem solving b) decision making c) personality assessment. In any situation of problem, the perceptions of different people are discussed, possible solutions are suggested. The best option is chosen by the group. While taking a decision, the matter is discussed, analyzed, interpreted and evaluated.

Characteristics of Successful Group Discussion

For any group discussion to be successful, achieving group goal is essential. Following characteristics are necessary:

Having a clear objective : The participants need to know the purpose of group discussion so that they can concentrate during the discussion and contribute to achieving the group goal. An effective GD typically begins with a purpose stated by the initiator.

Motivated Interaction: When there is a good level of motivation among the members, they learn to subordinate the personal interests to the group interest and the discussions are more fruitful.

Logical Presentation: Participants decide how they will organize the presentation of individual views, how an exchange of the views will take place, and how they will reach a group consensus. If the mode of interaction is not decided, few of the members in the group may dominate the discussion and thus will make the entire process meaningless.

Cordial Atmosphere: Development of a cooperative, friendly, and cordial atmosphere avoids the confrontation between the group members.

Evaluation in a GD

In any kind of GD, the aim is to judge the participants based on personality, knowledge, communicative ability to present the knowledge and leadership skills. Today team players are considered more important than individual contributors. Hence the potential to be a leader is evaluated and also ability to work in a team is tested. The evaluators generally assess the oral competence of a candidate in terms of team listening, appropriate language, clarity of expression, positive speech attitudes and adjustments, clear articulation, and effective non-verbal communication.

Personality: Even before one starts communicating, impression is created by the appearance, the body language, eye- contact, mannerisms used etc. The attire of a participant creates an impression, hence it is essential to be dressed appropriately. The hairstyle also needs to suit the occasion. Other accessories also have to be suitable for the occasion. The facial expression helps to convey attitudes like optimism, self-confidence and friendliness. The body language, anon-verbal communication skill gives important clues to personality assessment. It includes the posture of a person, the eye-contact and overall manner in which one moves and acts. In the entire participation in the GD, the body language has an important role in the impact created.

Content: Content is a combination of knowledge and ability to create coherent, logical arguments on the basis of that knowledge. Also a balanced response is what is expected and not an emotional response. In a group discussion, greater the knowledge of the subject more confident and enthusiastic would be the participation. Participants need to have a fair amount of knowledge on a wide range of subjects. The discussion of the subject must be relevant, rational, convincing and appealing to the listeners. One needs

to keep abreast with national and international news, political, scientific, economic, cultural events, key newsmakers etc. This has to be supplemented by one's own personal reasoning and analysis. People with depth and range of knowledge are always preferred by dynamic companies and organizations.

Communication Skills:

First and foremost feature of communication skills is that it is a two way process. Hence the communicator has to keep in mind the listeners and their expectations. The participants need to observe the group dynamics. Since GD tests one's behavior as well as one's influence on the group, formal language and mutual respect are obvious requirements. One may not take strong views in the beginning itself but wait and analyse the pros and cons of any situation. If one needs to disagree, learn to do so politely. One can directly put forward the personal viewpoint also. One may appreciate the good points made by others can make a positive contribution by agreeing to and expanding an argument made by another participant. An idea can be appreciated only when expressed effectively. A leader or an administrator has the ability to put across the idea in an influential manner. Hence the participants in a group discussion must possess not only subject knowledge but also the ability to present that knowledge in an effective way. Since oral skills are used to put across the ideas, the ability to speak confidently and convincingly makes a participant an impressive speaker. The members of the selection committee closely evaluate the oral communication skills of the candidates. The effective communication would imply use of correct grammar and vocabulary, using the right pitch, good voice quality, clear articulation, logical presentation of the ideas and above all, a positive attitude.

Listening Skills:

Lack of active listening is often a reason for failure of communication. In the GD, participants often forget that it is a group activity and not a solo performance as in elocution. By participating as an active listener, he/she may be able to contribute significantly to the group deliberations. The listening skills are closely linked to the leadership skills as well.

Leadership Skills:

The success of any group depends to a large extent upon the leader. One of the common misconceptions about leadership is that the leader is the one who controls the group. There are different approaches to the concept of leadership. By studying the personality traits of great leaders or actual dimensions of behavior to identify leadership one can learn to cultivate essential traits of leaders. In a GD, a participant with more knowledge, one who is confident, one who can find some solution to the problem and display initiative and responsibility will be identified as the leader. A candidate's success in a GD test will depend not only on his/her subject knowledge and oral skills but also on his/her ability to provide

leadership to the group. Adaptability, analysis, assertiveness, composure, self-confidence, decision making, discretion, initiative, objectivity, patience, and persuasiveness are some of the leadership skills that are useful in proving oneself as a natural leader in a GD. The leader in a group discussion should be able to manage the group despite differences of opinion and steer the discussion to a logical conclusion within the fixed time limit. In a selection GD, the group, which may consist of six to ten persons, is given a topic to discuss within 30 to 45 minutes. After announcing the topic, the total GD time, and explaining the general guidelines and procedures governing the GD, the examiner withdraws to the background leaving the group completely free to carry on with the discussion on its own without any outside interference. In the absence of a designated leader to initiate the proceedings of the discussion, the group is likely to waste time in cross talks, low-key conversations, cross-consultations, asides, and so on. The confusion may last until someone in the group takes an assertive position and restores the chaos into order. It could be any candidate. In order to get the GD started, the assertive, natural leader will have to remind the group of its goal and request them to start the discussion without wasting time. Leadership functions during a GD include initiative, analysis, and assertiveness and so on. GD does not have a formal leader, hence one of the participants is expected to take the initiative. The leader will promote positive group interactions; point out areas of agreement and disagreement; Help keep the discussion on the right track and lead the discussion to a positive and successful conclusion within the stipulated time. The ability to analyse a situation is a quality of leadership. Analytical skills and objectivity in expressing opinions are absolute requirements for leadership. With patience and composure one can develop the analytical skills. Reaching consensus by considering the group opinion will make the GD successful. Assertiveness that is an ability to bring order to the group by handling the conflict is another desirable quality of leadership. Self confidence is a quality which helps win the agreement from other participants.

DEBATE

Debate is a formal contest of argumentation between two teams or individuals. More broadly, and more importantly, debate is an essential tool for developing and maintaining democracy and open societies. More than a mere verbal or performance skill, debate embodies the ideals of reasoned argument, tolerance for divergent points of view and rigorous self-examination. Debate is, above all, a way for those who hold opposing views to discuss controversial issues without descending to insult, emotional appeals or personal bias. A key trademark of debate is that it rarely ends in agreement, but rather allows for a robust analysis of the question at hand. Debate is not a forum for asserting absolute truths, but

rather a means of making and evaluating arguments that allows debaters to better understand their own and others' positions. This sense of a shared journey toward the truth brings debaters closer together, even when they represent opposing sides of an issue or come from vastly different cultures or social classes. In so doing, debate fosters the essential democratic values of free and open discussion.

EXTEMPORE

"Extempore" or "ex tempore" refers to a stage or theater performance that is carried out without preparation or forethought. Most often the term is used in the context of speech, singing and stage acting.

Well basically it's about projecting confidence and telling what you know in a short span of time. Firstly your body language should convey the fact that you are not shaky about coming to stage One should feel the confidence within oneself

Secondly you should know a few facts about the topic you are going to talk about. For this its a good idea if you go through the daily newspapers and have a general understanding about things. Then its all about talking effectively without stuttering and good posture. Your body language should project good confidence. At the same time you also shouldn't appear smug. Well, when it comes to content, its better if you organize your points and tell them in a systematic manner. Its good if you mention most of the points without going much deep in to any of them. Its always better to limit your speech to the time allotted for one speech, especially if it's a competition.

Difference between Debate and Extempore

Debate is two or more people speaking to each other-two way traffic, Extempore means without preparation as opposed to prepared speech. Debate is discussing about a particular issue, exploring the pros and cons as well as one's thoughts and convictions to put forth his views and ideas. Extempore, means delivering a talk or address without much back home preparation, just like that address instantaneously to the audience.

INCREASE YOUR PROFESSIONALISM

Proficient at job

First and foremost, being good at your job is what defines you as someone with a high degree of professionalism in the workplace. A fancy wardrobe or a crisp hand shake is no match for the skills that make you effective at your job. Before you consider anything else, assess all the skills you need to

perform all aspects of your job. Identify strengths and weaknesses and develop a plan to improve any areas that will help you be more effective.

Clear on the expectations of the job

When a professional is clear on what is expected from their role and act accordingly they convey competence and put others at ease. They focus on carrying out the actions that help them meet those expectations. This conveys confidence as a professional. When others learn they can count on you to be focused and competent, they can then focus on their job. When they don't have to manage you or keep you from getting distracted you become an entity they can count on.

Able to separate personal from professional

Have you ever had a bad day and behaved in a way you were not proud of? If so, you've learned what it is like to experience professionalism from someone else. If the person you were dealing with focuses on serving you, rather than getting offended by you as a person, you're likely to notice. You may even make a point to change your demeanor in acknowledgement of their professional attitude. Depending on how important the transaction is you may leave with a feeling of gratitude for the maturity that you experienced. When it comes to being in the other person's shoes, your maturity and ability to focus on doing your job will separate you from your competition.

Practice "Right Speech" Right speech is a term that comes from Buddhist teachings. It refers to how we talk about others. Do you engage in water cooler gossip or do you politely remove yourself from negative banter? If you must discuss the actions of a coworker or client, consider the reason for the discussion. Is it about problem solving? Is it designed to take a corrective action or is it simply an expression of bad feelings. Also consider who you are sharing the information with. Are you charitable and non-judgemental or are you angry and vengeful. When it comes to karma and business transactions the energy you put out will come back to you.

Demonstrate Appropriate Enthusiasm

When you enjoy your life and your job, it's obvious. Others see it in your behaviors. When you smile, when you offer to assist them and with every project you tackle, you build self-confidence and radiate enthusiasm. While it's possible to overdo enthusiasm, the problem always seems to come from a lack of it. A colleague explains, "insurance policies aren't all that interesting in and of themselves but my customers are interesting and I enjoy helping them." If you don't feel like this at your job, chances are you're not demonstrating appropriate enthusiasm because it is not something that you can fake. If professionalism is important to you and your business, consider these four ways to improve on it.

Practice these 4 key concepts. With practice you will improve your self-confidence and professionalism. Over time and the efforts you put forth will be noticed and met with appreciation. Colleagues and clients alike will respond to your increasing professionalism with positive acknowledgement. Want help implementing them? Consulting with a business coach might be the first step. Holding a staff workshop might be another solution.

Audio Video Recording

An audio video recording is a recording that contains both audio and video information, usually gained by utilizing a system that contains both a microphone and camera. There are a number of different ways in which this type of recording can be achieved, though film and video cameras that include a built in microphone are quite common. With the rise in digital technology for capturing video and audio, this type of recording has become easier to achieve and store. An audio video recording can also be created without the use of a camera or microphone, through software that can record audio and video being viewed and heard on a computer.

While many people may think of both audio and video in reference to a single recording, the term “audio video recording” specifically indicates such a recording. This distinction largely stems from cameras that were only capable of capturing video images and not audio. For this type of recording, the audio would have to be recorded using a separate microphone and recorder, and the audio and video would then need to be synchronized during playback. An audio video recording, however, includes both audio and video signals recorded together, ensuring synchronicity when the recording is played.

LEADERSHIP QUIZ WITH CASE STUDY

Aditya Birla Group's Growth Strategy

In 2008, the Aditya Birla Group (ABG) was a US\$ 28 billion corporation. It employed 100,000 people belonging to 25 nationalities and over 50% of its revenues were attributed to its overseas operations in countries like the US, the UK, China, Germany, Hungary, and Brazil, among others. The group's product portfolio comprised Aluminum (Hindalco-Indal), Copper (Birla Copper), Fertilizers (Indo Gulf Fertilizers Ltd.), Textiles and Cement (Grasim Industries Ltd.), Insulators (Birla NGK Insulators Pvt. Ltd.),

Viscose Filament Yarn (Indian Rayon and Industries Ltd.), Carbon black (Birla Carbon), Insurance (Birla Sun Life Insurance Company Ltd.), Telecommunications (Idea Cellular Ltd.)

and BPO (Minacs Worldwide Ltd.).

In 2007, the group acquired Novelis Inc., the Atlanta (US)-based aluminum producer to become

one of the largest rolled-aluminum products manufacturers in the world⁶. The group had also acquired a majority stake in Indal from Alcan of Canada in the year 2000, and this had positioned it in the value-addition chain of the business, from metal to downstream products. Birla Copper enjoyed a good market share in the country and the acquisition of mines in Australia in the year 2003⁷ elevated it to an integrated copper producer. Indo-Gulf Fertilizers possessed a brand that commanded strong cash flows and a leadership position in the fertilizer industry. The group had entered into a 50 : 50 joint venture with NGK Corporation of Japan for its insulators division in 2002. This was expected to provide ABG access to the latest in product and manufacturing technology for the insulators division and also to open up the path to global markets. In 2006, the group purchased the equity holding of NGK and made the venture its subsidiary. Group company, Birla Sun Life, offered insurance and mutual fund products in the Indian market. In 2006, the group acquired Minacs Worldwide, a BPO company, and acquired Tata's stake in Idea Cellular⁸. In 2007, the group acquired Trinethra, a chain of retail stores.

The group's strategy toward the business portfolio was to exit from those areas of business where they had a minor presence or where losses were being incurred and to consolidate and build upon operations where competencies and business strengths existed. For instance, the group's textiles division Grasim had consolidated its operations by closing down operations at its pulp and fiber plants located at Mavoor and had sold the loss making fabric operations at Gwalior in 2002.

The group also divested itself of its stake in Mangalore Refinery and Petrochemicals Ltd. to the leading Indian oil company ONGC in 2002. Analysts felt that the group's ability to grow had stemmed largely from the emphasis placed on building meritocracy in the group. Under the leadership of Kumar Mangalam Birla (Birla), several initiatives were taken with the focus on learning and relearning, performance management, and organizational renewal.

Birla also instituted steps to retire aged managers and replaced them with young managers who came in with fresh and 'out of the box' ideas. The group instituted Gyanodaya, the group's learning center, to facilitate transfer of best practices across the group companies. The training methodology comprised classroom teaching and e-learning initiatives and the training calendar was accessible to the group

employees through the group-wide intranet.

The group also put in place 'The Organizational Health Survey' aimed at tracking the satisfaction levels of the group's managers. The survey was seen as a gauge of the happiness at work index in the group. The implementation of these initiatives resulted in the group becoming one of the preferred employers in Asia

.Toward performance management, the group had instituted the Aditya Birla Sun awards to recognize the successes of the group companies. This resulted in information sharing and encouraged healthy competition among these companies.

Q1) Critically analyze the growth strategy adopted by the Aditya Birla Group. What are your views on the business portfolio adopted by the group?

Q2) Analyze the initiatives taken by the group on the personnel and culture front under the leadership of Kumar Mangalam Birla

Dialogue session on Current topics, Economy, Education System, Environment, Politics

- **Economy:-** An economy (Greek οίκος-household and νέμωμαι - manage) or economic system consists of the production, distribution or trade, and consumption of limited goods and services by different agents in a given geographical location. The economic agents can be individuals, businesses, organizations, or governments. Macroeconomics is focused on the movement and trends in the economy as a whole, while in microeconomics the focus is placed on factors that affect the decisions made by firms and individuals. The Economy of India is the seventh-largest in the world by nominal GDP and the third-largest by purchasing power parity (PPP).
- **Education System-** Education is the process of facilitating learning. Knowledge, skills, values, beliefs, and habits of a group of people are transferred to other people, through storytelling, discussion, teaching, training, or research. Education frequently takes place under the guidance of educators, but learners may also educate themselves in a process called autodidactic learning. Any experience that has a formative effect on the way one thinks, feels, or acts may be considered educational. Education is commonly and formally divided into stages such as

preschool, primary school, secondary school and then college, university or apprenticeship. The methodology of teaching is called pedagogy.

- **Environment-** Environment (biophysical), the physical and biological factors along with their chemical interactions that affect an organism or a group of organisms. Environment (systems), the surroundings of a physical system that may interact with the system by exchanging mass, energy, or other properties. The natural environment encompasses all living and non-living things occurring naturally on Earth or some region thereof. It is an environment that encompasses the interaction of all living species. Climate, weather, and natural resources that affect human survival and economic activity.
- **Politics-** Politics is the practice and theory of influencing other people. Politics involves the making of a common decision for a group of people, that is, a uniform decision applying in the same way to all members of the group. Politics in India take place within the framework of its constitution, as India is a federal parliamentary democratic republic in which the President of India is the head of state and the Prime Minister of India is the head of government. India follows the dual polity system, i.e. a double government which consists of the central authority at the centre and states at the periphery. The constitution defines the organization, powers and limitations of both central and state governments, and it is well-recognized, rigid and considered supreme; i.e. laws of the nation must conform to it. There is a provision for a bicameral legislature consisting of an Upper House, i.e. Rajya Sabha, which represents the states of the Indian federation and a lower house i.e. Lok Sabha, which represents the people of India as a whole. The Indian constitution provides for an independent Judiciary which is headed by the Supreme Court. The court's mandate is to protect the constitution, to settle disputes between the central government and the states, inter-state disputes, and nullify any central or state laws that go against the constitution.