

UNIT - I

Lesson 1 - Set theory and Set Operations

Contents:

- 1.1 Aims and Objectives
- 1.2 Sets and elements
- 1.3 Further set concepts
- 1.4 Venn Diagrams
- 1.5 Operations on Sets
- 1.6 Set Intersection
- 1.7 Let – us Sum Up
- 1.8 Lesson – End Activities
- 1.9 References

1.1 Aims and Objectives

This Lesson introduces some basic concepts in Set Theory, describing sets, elements, Venn diagrams and the union and intersection of sets.

1.2 Sets and elements

Sets of objects, numbers, departments, job descriptions, etc. are things that we all deal with every day of our lives. Mathematical Set Theory just puts a structure around this concept so that sets can be used or manipulated in a logical way. The type of notation used is a reasonable and simple one.

For example, suppose a company manufactured 5 different products a, b, c, d, and e. Mathematically, we might identify the whole set of products as P , say, and write:

$$P = (a, b, c, d, e)$$

which is translated as 'the set of company products, P , consists of the *members (or elements)* a, b, c, d and e .

The elements of a set are usually put within braces (curly brackets) and the elements separated by commas, as shown for set P above.

A mathematical set is a collection of distinct objects, normally referred to as *elements or members*.

Sets are usually denoted by a capital letter and the elements by small letters.

Example 1 (Illustrations of sets)

- a) The employees of a company working in the purchase department could be written as:

$$P = (\text{Jones, Wilson, Gopan, Smith, Hari})$$

- b) The warehouse locations of a large supermarket chain could be written as:

$$W = (\text{Mumbai, Delhi, Bangalore, Chennai, Triandrum, Kochi})$$

1.3 Further set concepts

- a) *Subsets*. A subset of some set A , say, is a set which contains some of the elements of A . For example,

if $A = (h, i, j, k, l)$, then:

$X = (i, j, l)$ is a subset of A

$Y = (h, l)$ is a subset of A

$Z = (i, j)$ is a subset of A and also a subset of X .

- b) *The number of a set*. The number of a set A , written as $n[A]$, is defined as the number of elements that A contains.

For example,

if $A = (a, b, c, d, e)$, then $n[A] = 5$ (since there are 5 elements in A);

if $D = (\text{Sales, Purchasing, Inventory, Payroll})$, then $n[D] = 4$.

- c) *Set equality*. Two sets are equal only if they have identical elements. Thus, if

$A = (x, y, z)$ and $B = (x, y, z)$, then $A = B$.

- d) *The Universal Set*. In some problems in involving sets, it is necessary to consider one or more sets under consideration as belonging to some larger set that contains them. For example, if we were considering the set of skilled workers (S , say) on a production line, it might be convenient to consider the universal set (U , say) as all of the workers on the line. In other words, where a universal set has been defined, all the sets under consideration must necessarily be subsets of it.

- e) *The complement of a set*. If A is any set, with some universal set U defined, the complement of A , normally written as A' , is defined as 'all those elements that are not contained in A but are contained in U '. For the example of the workers on the production line (given in d above), S was specified as the set of skilled workers within the universal set of all workers on the line. Therefore, S' would be all the workers that were not skilled. i.e. the set of unskilled workers.

1.4 Venn Diagrams

A *Venn diagram* is a simple pictorial representation of a set. For example, if $M = (a,b,c,d,e,f,g)$ then we could represent this information in the form of a Venn diagram as in Figure 1.1

Figure 1.1

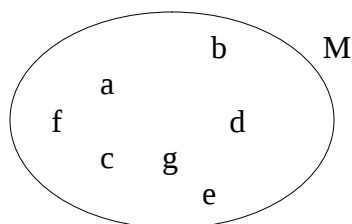
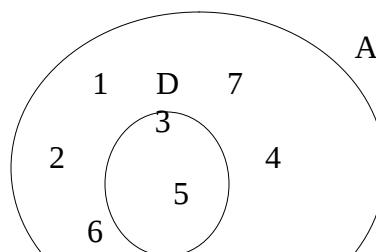


Figure 1.2



Venn diagrams are useful for demonstrating general relationships between sets.

For example, if a firm maintains a fleet of 7 cars, we might write $A = (1,2,3,4,5,6,7)$ (each car being numbered for convenience). If also it was important to identify those cars of the fleet that were being used by the directors, we might have $D = \{3,5\}$. i.e. Cars 3 and 5 are director's cars. This situation could be represented in Venn diagram form as in Figure 1.2. This diagram nicely demonstrates the fact that D is a subset of A , which normally means that $n[D] < n[A]$. In this case $n[D] = 2$ and $n[A] = 7$.

1.5 Operations on Sets

In ordinary arithmetic and algebra there are four common operations that can be performed; namely, addition, subtraction, multiplication and division. With sets, however, just two operations are defined. These are set *union* and set *intersection*. Both of these operations are described, with examples, in the following sections.

The *Union* of two sets A and B is written as $A \cup B$ and defined as that set which contains all the elements lying within *either* A or B or *both*.

For example, if $A = (c,d,f,h,j)$ and $B = (d,m,c,f,n,p)$, then the union of A and B is $A \cup B = (c,d,f,h,j,m,n,p)$, these being the elements that lie in either A or B . So that any element of A *must* be an element of $A \cup B$; similarly any element of B *must also* be an element of $A \cup B$.

Set union for three or more sets is defined in an obvious way. That is, if A , B and C are any three sets, $A \cup B \cup C$ is the set containing all the elements lying within (i) anyone of A , B or C , (ii) any two of them or (iii) all three.

Example 2 (To demonstrate set union)

If $A = (m,n,o,p)$; $B = (m,o,p,q)$; $C = (m,p,r)$; and the universal set is defined as $U = (k,l,m,n,o,p,q,r,s)$, then:

- a) $A \cup B = (m,n,o,p,q)$
- b) $A \cup C = (m,n,o,p,r)$
- c) $B \cup C = (m,o,p,q,r)$
- d) $A \cup B \cup C = (m,n,o,p,q,r)$
- e) $(A \cup B)' = (k,l,r,s)$, which is describing all the elements that are not in $A \cup B$ but are in the universal set U .

1.6 Set Intersection

The *intersection* of two sets A and B is written as $A \cap B$ and defined as that set which contains all the elements lying within *both* A or B .

For example, if $A = (a,b,c,d,f,g)$ and $B = (c,f,g,h,j)$, then the intersection of A and B is $A \cap B = (c,f,g)$, since these are the elements that lie in both sets.

The intersection of three or more sets is a natural extension of the above. If P , Q and R are any three sets then $P \cap Q \cap R$ is the set containing all the elements that lie *in all three sets*.

Any combinations of union and intersection can be used with sets. For, example, if X and Y are the sets specified above and $Z = (d,f,g,j)$. then: $(X \cap Y) \cup Z = (c,f,g) \cup (d,f,g,j) = (c,d,f,g,j)$ which can be described in words as 'the set of elements that are in either both of X and Y or in Z '.

Example 3 (To demonstrate set intersection)

If $A = (m,n,o,p)$; $B = (m,o,p,q)$; $C = (n,q,r)$; with a universal set defined as (k,l,m,n,o,p,q,r,s) . Then:

- a) $A \cap B = (m,o,p)$, since all these elements are in *both* sets.
Similarly,
- b) $A \cap C = (n)$
- c) $B \cap C = (q)$.
- d) $A \cap B \cap C$ has no elements, is sometimes called the *empty set* and can be written $A \cap B \cap C = \{\}$. Note $n[\{\}] = 0$.
- e) $(A \cap B)' = (k,l,n,q,r,s)$ is the complement of $A \cap B$ and is the set of all elements that are NOT in both A and B .
- f) $(A \cup B) \cap C = (m,n,o,p,q) \cap (n,q,r) = (n,q)$ is the set of all elements that are in A or B AND ALSO in C .

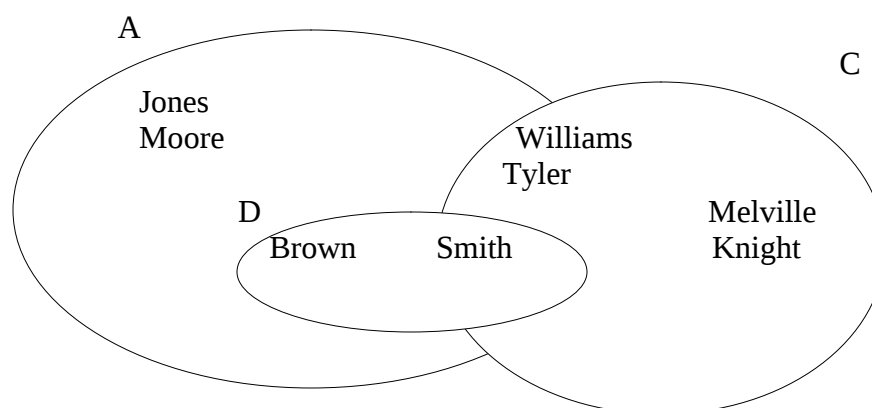
Example 4 (The union and intersection of given sets)**Question**

In a particular insurance life office, employees Smith, Jones, Williams and Brown have 'A' levels, with Smith and Brown also having a degree. Smith, Melville, Williams, Tyler, Moore and Knight are associate members of the Chartered Insurance Institute (ACII) with Tyler, and Moore having 'A' levels. Identifying set A as those employees with 'A' levels, set C as those employees who are ACII and set D as graduates:

- Specify the elements of sets A , C and D .
- Draw a Venn diagram representing sets A , C and D , together with their known elements.
- What special relationship exists between sets A and D ?
- Specify the elements of the following sets and for each set, state in words what information is being conveyed.
 - $A \cap C$
 - $D \cup C$
 - $D \cap C$
- What would be a suitable universal set for this situation?

Answer

- $A = \{\text{Smith, Jones, Williams, Brown, Tyler, Moore}\};$
 $C = \{\text{Smith, Melville, Williams, Tyler, Moore, Knight}\}; D = \{\text{Smith, Brown}\}$
- The Venn diagram is shown in Figure 1.3.

**Figure 1.3**

- From the diagram, it can be seen that D is a subset of A .
- This information can be obtained either from the Venn diagram or from the sets listed in, a) above.
 - $A \cap C = \{\text{Williams, Tyler, Smith}\}$. This set gives the employees who have both 'A' levels and are ACII.
 - $D \cup C = \{\text{Brown, Smith, Williams, Tyler, Melville, Knight}\}$. This set gives the employees who are either graduates or ACII.

- iii. $D \cap C = \{\text{Smith}\}$. This set gives the single employee who is both a graduate and ACII qualified.
- e) A suitable universal set for this situation would be the set of all the employees working in the Life office.

1.7 Let us Sum Up

This Lesson presented described about the set, set theory, Venn Diagrams, and its applications. A set is a collection of distinct objects, called elements, which are normally enclosed within brackets and separated by commas. Venn diagram is a pictorial representation of one or more sets. The Union and Intersection of sets were also discussed in detail. Some examples to understand the concept is also given in the Lesson.

1.8 Lesson – End Activities

1. Define set, subset.
2. Give the purpose of drawing Venn diagrams.

1.9 References

Navaneethan. P. – Business Mathematics.

Lesson 2 - Functions and Co-ordinate Systems

Contents:

- 2.1 Aims and Objectives**
- 2.2 Definitions**
- 2.3 Types of Functions**
- 2.4 Solution of Functions**
- 2.5 Business Applications**
- 2.6 Let us Sum Up**
- 2.7 Lesson – End Activities**
- 2.8 Reference**

2.1 Aims and Objectives

For decision problems which use mathematical tools, the first requirement is to identify or formally define all significant interactions or relationships among primary factors (also called variables) relevant to the problem. These relationships usually are stated in the form of an equation (or set of equations) or inequations. Such type of simplified mathematical relations help the decision-maker in understanding (any) complex management problems. For example, the decision-maker knows that demand of an item is not only related to price of that item but also to the price of the substitutes. Thus if he can define specific mathematical relationship (also called model) that exists, then the demand of the item in the near future can be forecasted. The main objective of this unit is to study mathematical relationships (or functions) in the context of managerial problems.

2.2 Definitions

Variable

A variable is something whose magnitude can vary or which can assume various values. The variables used in applied mathematics include: sale, price, profit, cost, etc. since magnitude of variables can vary, therefore these are represented by symbols (such as x, y, z etc) instead of a specific number. In applied mathematics a variable is represented by the first letter of its name, for example p for price or profit; q for quantity, c for cost; s for saving or sales; d for demand and so forth. When we write $X = 5$, the variable takes specific value.

Variables can be classified in a number of ways. For example, a variable can be discrete (suspect to counting, e.g. 2 houses, 3 machines etc.), or continuous (suspect to measurement, e.g. temperature, height etc.).

Constant and Parameter

A quantity that remains fixed in the context of a given problem or situation is called a constant. An absolute (or numerical) constant such as 2, π , e, etc. retains the same value in all problems whereas an arbitrary (or parametric) constant or parameter retains the same value throughout any particular problem but may assume different values in different problems, such as wage rates of different category of labourers in an industrial unit.

The Absolute or numerical value of a constant 'b' is denoted by $|b|$ and means the magnitude of 'b' regardless of its algebraic sign. Thus $|b| = |-b|$ or $|+b|$.

Functions

We come across situations in which two or more variables are related to each other. For example, demand (D) of a commodity is related to its price (p). It can be mathematically expressed as

$$D = f(p) \quad (2-1)$$

This relationship is read as "demand is function of price" or simply "f of p". it does not mean D equals f times p. This mathematical relationship has two variables, D and p. these are called variables because they can take on different numerical values.

Let us now consider a mathematical relationship that contains three variables. Assume that the demand (D) of a commodity is related to the price (p) per unit of the commodity, and the level of advertising expenditure (A). then the general relationship among these variables can be expressed as

$$D = f(p, A) \quad (2-2)$$

The functional notations of the type (2-1) and (2-2) are meant to give a general idea that certain variables are, somehow, related. However for making managerial decisions, we need a specific and explicit, not a general and implicit relationship among selected variables. For example, for the purpose of finding the value of demand (D), we make the general relationship (2-2) more specific as shown in (2-3).

$$D = 4 + 3p - 2pA + 2A^2 \quad (2-3)$$

Now for any given values of p and A, the value of D can be calculated using the relationship (2-3). This means that the value of D depends on the values of p and A. Hence D is called the dependent variable and p and A are called independent variables. In this case, it may be noted that we have established a rule of correspondence between the dependent variable and independent variable (s). That is as soon as values are assigned to the independent variables (s), the corresponding unique value for the dependent variable is determined by the given specific relationship. That is why a function is sometimes defined as a rule of correspondence between variables. The set of values given to independent variable is called

the domain of the function while the corresponding set of values of the dependent variable is called the range of the function. Other examples of functional relationships are as follows:

- i) the distance (d) covered is a function of time (T) and speed (s), i.e. $d = f(T, s)$.
- ii) Sales volume (v) of the commodity is a function of price (p), i.e. $V = f(p)$.
- iii) Total inventory cost (T) is a function of order quantity (Q), i.e. $T = f(Q)$.
- iv) The volume of the sphere (v) is a function of its radius R , i.e. $V = f(R)$ or $V = \frac{4}{3} \pi r^3$
- v) The extension (y) of a spring is proportional to the weight (m) (Hooke's law), i.e. $Y = km$ or $Y = km$.
- vi) The net present value (y) of an investment is a function of net cash flows (C_t) in different time periods, project's initial cash outlay (B), firm's cost of capital (P) and the life of the project (N), i.e. $y = f(C_t, B, P, N)$.

The following example will illustrate the meaning of these terms.

Example 1

Suppose an industrial worker gets Rs. 50 per day. If he works for 25 days in a particular month, then his total wage for this month is $50 \times 25 = \text{Rs. } 1250$. During some other month he may have worked a total of only 24 days, then he would have earned Rs. 1200. Thus, the total wages of the worker, assuming no overtime, can always be calculated as follows:

Total wages = 25 X number of days worked

Let,

T = total wages

D = number of days worked

Then,

$T = 50 D$.

This represents the relationship between total wages and number of days worked. In general, the above relationship can also be written as:

$$T = KD$$

Where K is a constant for particular class of worker (s), to be assigned or determined in a specific situation. Since the value of K can vary for a specific situation, problem or context therefore it is called a parameter, whereas constants such as π (denoted by π) which has approximate value of 3.1416 remains same from one problem context to another are called absolute constants. Quantities such as T and D which can assume various values in a given problem are called variables.

Exercise

1. Find the domain and range of each of the following functions
 - a) $Y = 1/x - 1$
 - b) $Y = x; y \geq 0$
 - c) $Y = 4 - x; y \geq 0$
2. Let $4p + 6q = 60$ be an equation containing variables p (price) and q (quantity). Identify the meaningful domain and range for the given function when price is considered as independent variable.

2.3 Types of Functions

In this Section some different types of functions are introduced.

Linear Functions:

A linear function is one in which the power of independent variable is 1, the general expression of linear function having only one independent variable is:

$$Y = f(x) = a + bx$$

Where a and b are given real numbers and x is an independent variable taking all numerical values in an interval.

A function with only one independent variable is also called single variable function.

Further, a single-variable function can be linear and non-linear. For example,

$$Y = 3 + 2x, \text{ (linear single-variable function)}$$

And

$$Y = 2 + 3x - 5x^2 + x^2, \text{ (non-linear single-variable function)}$$

A linear function with one variable can always be graphed in a two dimensional plane (or space). This graph can always be plotted by giving different values to x and calculating corresponding values of y . the graph of such functions is always a straight line.

Example 2

Plot the graph of the function, $y = 3 + 2x$

For plotting the graph of the given function, assigning various values to x and then calculating the corresponding values of y as shown in the table below:

X	0	1	2	3	4	5	...
Y	3	5	7	9	11	13	...

The graph of the given function is shown in Figure 1.4

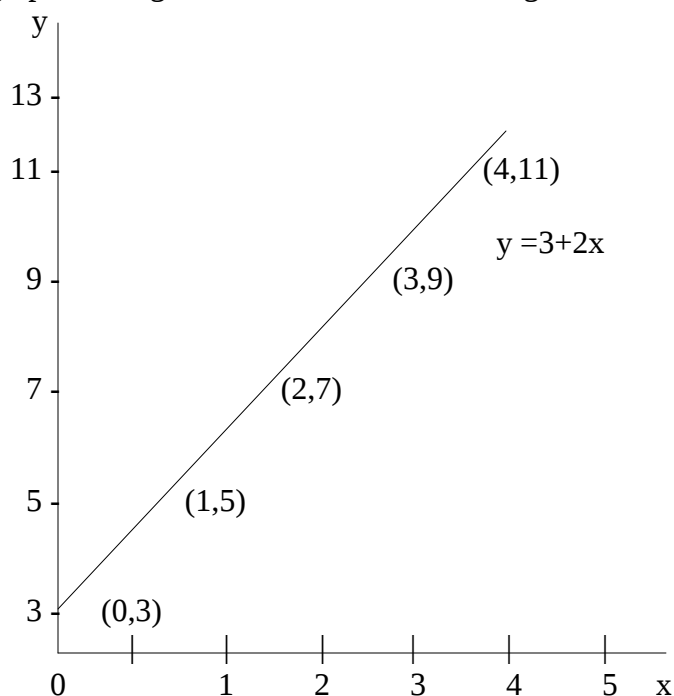


Figure 1.4

A function with more than one independent variable is defined, in general, form, as:

$$Y = f(x_1, x_2, \dots, x_n) = a_0 + a_1x + a_2x_2 + \dots + a_nx_n$$

Where $a_0, a_1, a_2, \dots, a_n$ are given real numbers and $x_1, x_2, \dots, x_n$ are independent variables taking all numerical value in the given intervals. Such functions are also called multivariable functions. A multivariable function can be linear and non-linear, for example,

$$Y = 2 + 3x_1 + 5x_2 \text{ (linear multi-variable function)}$$

and

$$Y = 3 + 4x_1 + 15x_1x_2 + 10x_2^2 \text{ (non-linear multivariable function)}$$

Multivariable functions may not be graphed easily because these require three-dimensional plane or more dimensional plane for plotting the graph. In general, a function with n variables will require $(n+1)$ dimensional plane for plotting its graph.

Polynomial Functions:

A function of the form

$$Y = f(x) = a_1x^{n-1} + \dots + a_nx^0 \quad (1-4)$$

Where $a_1, \dots, a_n$ are real numbers, $a_1 \neq 0$ and n is a positive integer is called a polynomial of degree n .

a) if $n = 1$, then the polynomial function is of degree 1 and is called a linear function.

That is, for $n = 1$, function (1-4) can be written as:

$$y = a_1x^1 + a_nx^0 \quad (a_1 \neq 0)$$

This is usually written as

$$Y = a + bx \text{ (since } x^0 = 1)$$

Where 'a' and 'b' symbolise a_n and a_1 respectively.

b) if $n = 2$, then the polynomial function is of degree 2 and is called a quadratic function.

That is, for $n = 2$, the function (1-4) can be written as:

c) $y = a_1x^2 + a_2x^1 + a_nx^0 (a_1 \neq 0)$

This is usually written as

$$Y = ax^2 + bx + c$$

where

$$a_1 = a, a_2 = b \text{ and } a_n = c$$

Absolute Value Functions

The functional relationship expressed by

$$Y = |x|$$

Is known as an absolute value function, where $|x|$ is known as magnitude (or absolute value) of x . By absolute value we mean that whether x is positive or negative, its absolute value remains positive. For example $|7|=7$ and $|-6|=6$.

Plotting of the graph of the function $y=|x|$, assigning various values to x and then calculating the corresponding values of y , is shown in the table below:

X	...	-3	-2	-1	0	1	2	3	...
Y	...	3	2	1	0	1	2	3	...

The graph of the given function is shown in Figure 1.5

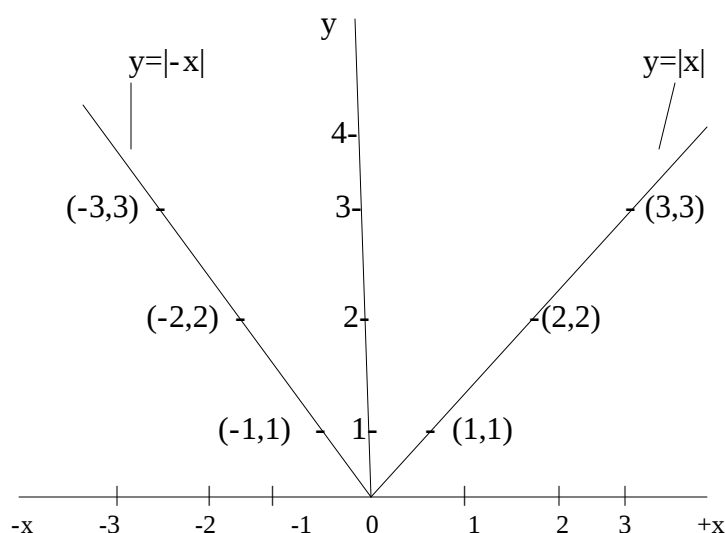


Figure 1.5

Inverse Function

Take the function $y = f(x)$. Then the value of y , can be uniquely determined for given values of x as per the functional relationship. Sometimes, it is required to consider x as a function of y , so that for given values of y , the value of x can be uniquely determined as per the functional relationship. This is called the inverse function and is also denoted by $x = f^{-1}(y)$. For example consider the linear function:

$$Y = ax + b$$

Expressing this in terms of x , we get

$$\begin{aligned} X &= y - b/a \\ &= y/a - b/a = cy + d \end{aligned}$$

where $c = 1/a$, and $d = -b/a$

This is also a linear function and is denoted by $x = f^{-1}(y)$

Step Function

For different values of an independent variable x in an interval, the dependent variable $y = f(x)$ takes a constant value, but takes different values in different intervals. In such cases the given function $y = f(x)$ is called a step function. For example

$$y = f(x) \quad \begin{aligned} &y_1, \text{ if } 0 < x < 50 \\ &y_2, \text{ if } 51 < x < 100 \\ &y_3, \text{ if } 101 \leq x \leq 150 \end{aligned}$$

The shape of the graph of this function looks as shown in Figure 1.6, for $y_3 < y_2 < y_1$

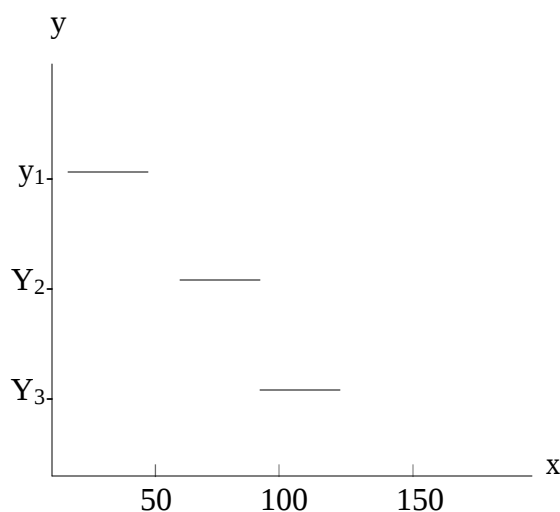


Figure 1.6

Algebraic and Transcendental Functions

Functions can also be classified with respect to the mathematical operations (addition, subtraction, multiplication, division, powers and roots) involved in the functional relationship between dependent variable and independent variable (s). When only finite number of terms are involved in a functional relationship and variables are affected only by the mathematical operations, then the function is called an algebraic function, otherwise transcendental function. The following functions are algebraic functions of x .

i) $y = 2x^3 + 5x^2 - 3x + 9$

ii) $y = x + 1/x^2$

iii) $y = x^3 - 1/x + 2$

The sub-classes of transcendental functions are follows:

a) Exponential function

If the independent variable in any functional relationship appears as an exponent (or power), then that functional relationship is called exponential function, such as

i) $y = ax, a > 1$

ii) $y = ka^x, a > 1$

iii) $y = ka^{bx}, a > 1$

iv) $y = ke^x$

where a, b, e and k are constants with 'a' taking only a positive value.

Such functions are useful for describing sharp increase or decrease in the value of dependent variable. For example, the exponential function $y = ka^x$ curve rises to the right for $a > 1, k > 0$ and falls to the left $a < 1, k > 0$ as shown in the Figure 1.7 (a) and (b).

Figure 1.7(a)

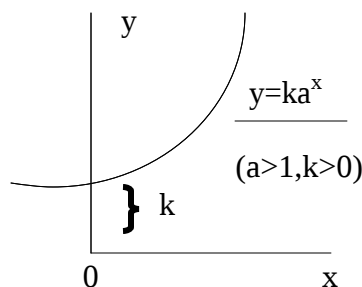
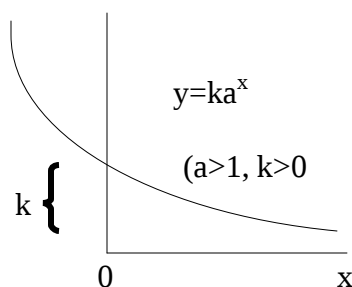


Figure 1.7(b)



b) Logarithmic functions

A logarithmic function is expressed as

$$Y = \log_a x$$

Where $a > 1$ and > 0 is the base. It is read as 'y' is the log to the base a of x'. This can also be written as

$$X = a^y$$

Thus from an exponential function $y=a^x$, we may construct the logarithmic function $x=a^y$ by interchanging the variables. This shows that the inverse of an exponential function is a logarithmic function.

The two most widely used bases for logarithms are '10' and 'e' (=2.7182).

- i) Common logarithm: It is the logarithm to the base 10 of a number x. it is written as $\log_{10}x$. if $y=\log_{10}x$, then $x=10^y$
- ii) Natural logarithm: It is the logarithm to the base 'e' of a number x. it is written as $\log_e x$ or $\ln x$. when no base is mentioned, it will be understood that the base is e.

Some important properties of the logarithmic function $y=\log_e x$ are as follows:

- i) $\log 1=0$
- ii) $\log_e e=1$
- iii) $\log (xy)=\log x+\log y$
- iv) $\log (x/y)=\log x - \log y$
- v) $\log (x^n) = n \log x$
- vi) $\log_e 10 = 1/\log_{10}e$
- vii) $\log_e a = (\log_e 10) (\log_{10} a) = \log_{10} a / \log_{10} e$
- viii) logarithm of zero and negative number is not defined.

Exercise

1. Draw the graph of the following functions

- a) $y=3x-5$
- b) $y=x^2$
- c) $y=\log_2 x$

2. The data of machine operating cost © and the age (t) of the machine are shown in the following table:

t (years)	:	1	2	3	4	5
c (in '000's)	:	5	8	13	20	29

- i) Express operating cost as a function of the machine age
- ii) Sketch the graph of the function derived in (i).

2.4 Solution of Functions

The value (s) of x at which the given function f(x) becomes equal to zero are called the roots (or zeros) of the function f(x). For the linear function

$$Y = ax + b$$

The roots are given by

$$ax + b = 0$$

$$\text{Or } x = -b/a$$

Thus if $x = -b/a$ is substituted in the given linear function $y = ax + b$ then it becomes equal to zero.

In the case of quadratic function, $Y = ax^2+bx+c$,

We have to solve the equation $ax^2+bx+c=0$; $a \neq 0$ to find the roots of y . The general value of x for which the given quadratic function will become zero is given by

$$X = \frac{-b \pm \sqrt{b^2-4ac}}{2a}$$

Thus, in general, there are two values of x for which y becomes zero. One value is

$$X = \frac{-b + \sqrt{b^2-4ac}}{2a}$$

and other value is

$$X = \frac{-b - \sqrt{b^2-4ac}}{2a}$$

It is very important to note that the number of roots of the given function are always equal to the highest power of the independent variable.

Particular Cases:

The expression b^2-4ac in the above formula is known as discriminant which determines the nature of the roots as discussed below:

- i) If $b^2-4ac > 0$, then the two roots are real and unequal
- ii) If $b^2-4ac = 0$ or $b^2 = 4ac$, then the two roots are equal and are equal to $-b/2a$
- iii) If $b^2-4ac < 0$, then the two roots are imaginary (not real) because of the square root of a negative number.
- iv) The roots of a polynomial of the form:
 $Y = (x-a)(x-b)(x-c)(x-d) \dots$ are $a, b, c, d, \dots$

2.5 Business Applications

In business applications, there are lot of situations to deal with supply and demand functions; cost functions; profit functions; revenue functions; production functions; utility functions; etc. in applied mathematics. In this section, a few examples are given by constructing such functions and obtaining their solutions:

Example 3 (Linear Functions)

A company sells x units of an item each day at the rate of Rs. 50 per unit. The cost of manufacturing and selling these units is Rs. 35 per unit plus a fixed daily overhead cost of Rs. 1000. Determine the profit function. How would you interpret the situation if the company manufactures and sells 400 units of the item a day.

Solution:

The total revenue received by the company per day is given by:

$$\begin{aligned} \text{Total revenue (R)} &= (\text{price per unit}) \times (\text{number of items sold}) \\ &= 50x \end{aligned}$$

The total cost of manufactured items per day is given by:

$$\begin{aligned} \text{Total cost (C)} &= (\text{Variable cost per unit}) \times (\text{no. of items manufactured}) + (\text{fixed daily overhead cost}) \\ &= 35x + 1000 \end{aligned}$$

$$\begin{aligned}\text{Thus, Total profit (p)} &= (\text{Total revenue}) - (\text{Total cost}) \\ &= 50.x - (35.x + 1000) = 15.x - 1000\end{aligned}$$

If 400 units of the item are manufactured and sold, then the profit is given by:

$$\begin{aligned}P &= 15 \times 400 - 1000 \\ &= -400\end{aligned}$$

The negative profit indicates loss. Thus if the company manufactures and sells 400 units of the item, it would incur a loss of Rs. 400 per day.

Example 4 (Quadratic Functions)

Let the market supply function of an item be $q = 160 + 8p$, where q denotes the quantity supplied and p denotes the market price. The unit cost of production is Rs. 4. It is felt that the total profit should be Rs. 500. What market price has to be fixed for the item so as to achieve this profit?

Solution:

Total profit function can be constructed as follows:

$$\begin{aligned}\text{Total profit (p)} &= \text{Total revenue} - \text{Total cost} \\ &= (\text{price per unit} \times \text{Quantity supplied}) - (\text{unit cost} \times \text{Quantity supplied}) \\ &= p.q - c.q \\ &= (p - c).q\end{aligned}$$

Given that $c = \text{Rs. } 4$ and $q = 160 + 8p$. Then total profit function becomes

$$\begin{aligned}P &= (p - 4)(160 + 8p) \\ &= 8p^2 + 128p - 640\end{aligned}$$

If $P = 500$, then we have

$$\begin{aligned}500 &= 8p^2 + 128p - 640 \\ 8p^2 + 128p - 1140 &= 0 \\ p &= 6.36 \text{ or } -22.37\end{aligned}$$

since negative price has no economic meaning, therefore the required price per unit should be Rs. 6.36.

Exercise

- Consider the quadratic equation $x^2 - 8x + c = 0$. For what value of c , the equation has
 - real roots,
 - equal roots, and
 - imaginary roots?
- A newsboy buys papers for p_1 paise per paper and sells them at a price of p_2 paise per paper ($p_1 > p_2$). The unsold papers at the end of the day are bought by a wastepaper dealer for p_3 paise per paper ($p_3 < p_1$).
 - Construct the profit function of the newsboy.
 - construct the opportunity loss function of the newsboy.

2.6 Let us Sum Up

The objective of this unit is to provide you exposure to functional relationship among decision variables. We started with the mathematical concept of function and defined terms such as constant, parameter, independent and dependent variable. Various examples of functional relationships are mentioned to see the concept in broad

perspective. Various types of functions which are normally used in managerial decision-making are enumerated along with suitable examples, their graphs and solution procedure. Finally, the applications of functional relationships are demonstrated through several examples.

2.7 Lesson – End Activities

1. Define functions, variable.
2. With an example, discuss the inverse function.

2.8 Reference

Navaneethan. P – Business Mathematics.

Lesson 3 - Matrices and Matrices Operations

Contents

- 3.1 Aims and Objectives**
- 3.2 Matrices : Definition and Notations**
- 3.3 Some special Matrices**
- 3.4 Matrix Representation of Data**
- 3.5 Operations on Matrices**
- 3.6 Determinant of a Square Matrix**
- 3.7 Let us Sum Up**
- 3.8 References**

3.1 Aims and Objectives

Matrices have applications in management disciplines like finance, production, marketing etc. Also in quantitative methods like linear programming, game theory, input-output models and in many statistical applications matrix algebra is used as the theoretical base. Matrix algebra can be used to solve simultaneous linear equations.

3.2 Matrices : Definition and Notations

A matrix is a rectangular array or ordered numbers. The term ordered implies that the position of each number is significant and must be determined carefully to represent the information contained in the problem. These numbers (also called elements of the matrix) are arranged in rows and columns of the rectangular array and enclosed by either square brackets, $[\]$; or parantheses $(\)$, or by pair of double vertical line $\| \ \|$.

A matrix consisting of m rows and n columns is written in the following form

A column

$$\begin{bmatrix} A_{11} & a_{12} & \dots & a_{1n} \\ A_{21} & a_{22} & \dots & a_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ A_{m1} & a_{m2} & \dots & a_{mn} \end{bmatrix}$$

Where $a_{11}, a_{12}, \dots$ denote the numbers (or elements) of the matrix. The dimension (or order) of the matrix is determined by the number of rows and columns. Here, in the given matrix, there are m rows and n columns. Therefore, it is of the dimension $m \times n$ (read as m by n). In the dimension of the given matrix the number of rows is always specified first and then the number of columns.

Boldface capital letters such as $A, B, C, \dots$ are used to denote entire matrix. The matrix is also sometimes represented as $A = [a_{ij}]_{m \times n}$ where a_{ij} denotes the i th row and the j th element of a . Some examples of the matrices are

$$A = \begin{bmatrix} -1 & 1 \\ 2 & 3 \end{bmatrix}_{2 \times 2}; B = \begin{bmatrix} 1 & 1 & 2 \\ 2 & 4 & -3 \end{bmatrix}_{2 \times 3}; C = \begin{bmatrix} 5 & 5 & 10 \\ 6 & 2 & 10 \\ 2 & 1 & 2 \end{bmatrix}_{3 \times 3}$$

The matrix A is a 2×2 matrix because it has 2 rows and 2 columns. Similarly the matrix B is a 2×3 matrix while matrix C is a 3×3 matrix.

Exercise

Tick mark the correct alternative indicating the dimension of the matrix

$$\begin{bmatrix} 2 & 3 & 4 \\ 6 & 8 & 9 \\ 3 & 5 & 7 \end{bmatrix}$$

- i) 3×4 ii) 4×3 iii) None of these

3.3 Some special Matrices

a) Square matrix

A matrix in which the number of rows equals the number of columns is called a square matrix. For example

$$\begin{bmatrix} 2 & 3 & 7 \\ 3 & 5 & 2 \\ 4 & 3 & 1 \end{bmatrix}_{3 \times 3}$$

is a square matrix of dimension 3. The elements 2, 5 and 1 in this matrix are called the diagonal elements and the diagonal is called the principal diagonal.

b) Diagonal matrix

A square matrix, in which all non-diagonal elements are zero whereas diagonal elements are non-zero, is called a diagonal matrix. For example

$$\begin{bmatrix} 2 & 0 & 0 \\ 0 & 5 & 0 \\ 0 & 0 & 1 \end{bmatrix}_{3 \times 3}$$

is a diagonal matrix of dimension 3.

c) Scalar matrix

A diagonal matrix in which all diagonal elements are equal is called a scalar matrix. For example

$$\begin{bmatrix} k & 0 & 0 \\ 0 & k & 0 \\ 0 & 0 & k \end{bmatrix}_{3 \times 3}$$

is a scalar matrix, where k is a real (or complex) number.

d) Identity (or unit) matrix

A scalar matrix in which all diagonal elements are equal to one, is called an identity (or unit) matrix and is denoted by I . Following are two different identity matrices

$$I_2 = \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix}_{2 \times 2} ; I_3 = \begin{bmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{bmatrix}_{3 \times 3}$$

An identity matrix of dimension n is denoted by I_n . It has n elements in its diagonal each equal to 1 and other elements are zero.

d) The zero (or null) matrix

A matrix is said to be the zero matrix if every element of it is zero. It is denoted as O . Following are three different zero matrices

3.4 Matrix Representation of Data

Before discussing the operations on matrices, it is necessary for you to know a few situations in which data can be represented in matrix form.

1. Transportation Problem

The unit cost of transportation of an item from each of the two factories to each of the three warehouses can be represented in a matrix as shown below:

Warehouses

		W_1	W_2	W_3
Factory	F_1	20	15	30
	F_2	25	20	15

Similarly, we can also construct a time matrix $[t_{ij}]$, where t_{ij} =time of transportation of an item from factory I to warehouse j. Note that the time of transportation is independent of the amount shipped.

2. Distance Matrix

The distance (in kms.) between given number of cities can be represented as matrix as shown below:

		City			
		A	B	C	D
City	A	-	1,470	2,158	1,732
	B	1,470	---	1,853	2,385
	C	2,158	1,853	---	1,635
	D	1,732	2,365	1,635	----

3. Diet matrix

The vitamin content of two types of foods and two types of vitamins can be represented in a matrix as shown below:

		Vitamins	
		A	B
Food	F_1	150	120
	F_2	170	100

4. Assignment Matrix

The time required to perform three jobs by three workers can be represented matrix as shown below:

		Job		
		J ₁	J ₂	J ₃
Worker	W ₁	5	3	2
	W ₂	4	5	3
	W ₃	2	4	6

5. Pay – off Matrix

Suppose two players A and B play a coin tossing game. If outcome (H,H) or (T,T) occurs, then player B loses Rs. 20 to player A, otherwise gains as shown in the matrix:

		Player B	
		H	T
Player A	H	20	-20
	T	-20	20

The minus sign with the pay off means that player A pays to B.

6. Brand Switching matrix

The proportion of users in the population surveyed switching to brand j of an item in a period, given that they were using brand I can be represented as a matrix.

To		Brand 1	Brand2	Brand 3
From	Brand 1	0.3	0.6	0.1
	Brand 2	0.6	0.3	0.1
	Brand 3	0.2	0.5	0.3

Here the sum of the elements of each row is 1 because these are proportions.

3.5 Operations on Matrices

1. Addition and Subtraction of Matrices

The sum of two matrices of same order is obtained by adding the corresponding elements of the given matrices. The difference of two matrices of same order is obtained by subtracting the corresponding elements of the given matrices.

$$\text{For example, if } A = \begin{bmatrix} 2 & -3 \\ 1 & 0 \\ -2 & -1 \end{bmatrix} \text{ and } B = \begin{bmatrix} 1 & 3 \\ 1 & 2 \\ -1 & 0 \end{bmatrix}, \text{ then,}$$

$$A+B = \begin{bmatrix} 2+1 & -3+3 \\ 1+1 & 0+2 \\ -2+(-1) & -1+0 \end{bmatrix}$$

$$= \begin{bmatrix} 3 & 0 \\ 2 & 2 \\ -3 & -1 \end{bmatrix}$$

$$\begin{aligned}\text{Also, } A-B &= \begin{bmatrix} 2-1 & -3-3 \\ 1-1 & 0-2 \\ -2-(-1) & -1-0 \end{bmatrix} \\ &= \begin{bmatrix} 1 & -6 \\ 0 & -2 \\ -1 & -1 \end{bmatrix}\end{aligned}$$

Properties of Matrix Addition

If A, B and C are three matrices of same dimension, then,

- a) matrix addition is commutative, i.e. $A + B = B + A$
- b) matrix addition is associative, i.e. $(A+B)+C = A+(B+C)$
- c) zero matrix is the additive identity, i.e. $A+\mathbf{0} = A$
- d) B is an additive inverse if $A+B = \mathbf{0}$.

2. Scalar Multiplication

If A is any matrix of dimension $m \times n$ and k is any scalar (real number), then kA is obtained by multiplying each element of A by the scalar k.

For example, if $A = \begin{bmatrix} 2 & 1 & -1 \\ 3 & 1 & 5 \\ 2 & 0 & 1 \end{bmatrix}$, then

$$\begin{aligned}3A &= \begin{bmatrix} 3 \times 2 & 3 \times 1 & 3 \times -1 \\ 3 \times 3 & 3 \times 1 & 3 \times 5 \\ 3 \times 2 & 3 \times 0 & 3 \times 1 \end{bmatrix} \\ &= \begin{bmatrix} 6 & 3 & -3 \\ 9 & 3 & 15 \\ 6 & 0 & 3 \end{bmatrix}\end{aligned}$$

3. Multiplication of Matrices

If the number of columns in the first matrix is equal to the number of rows in the second matrix, then the matrices are compatible for multiplication. That is, if there are n columns in the first matrix then the number of rows in the second matrix must be n. Otherwise the matrices are said to be incompatible and their multiplication is not defined.

The operation of multiplication

- a) The element of a row of the first matrix should be multiplied by the corresponding elements of a column of the second matrix.

- b) The products are then summed and the location of the resulting element in the new matrix determines the row from first matrix has to be multiplied with which column from second.

Example 1. Let $A = \begin{bmatrix} 1 & 0 & 3 \\ 2 & 1 & 5 \end{bmatrix}$ and $B = \begin{bmatrix} 2 & 1 \\ 1 & 0 \\ 3 & 2 \end{bmatrix}$

Since A is of order 2×3 and B is of order 3×2 , the matrices are compatible for multiplication and the resultant matrix should 2×2 .

In the first matrix R_1 is $[1 \ 0 \ 3]$ and R_2 is $[2 \ 1 \ 5]$ and

Columns of the second matrix are C_1 is $\begin{bmatrix} 2 \\ 1 \\ 3 \end{bmatrix}$ and C_2 is $\begin{bmatrix} 1 \\ 0 \\ 2 \end{bmatrix}$

$$\text{Then, } A \times B = \begin{bmatrix} R_1 \times C_1 & R_1 \times C_2 \\ R_2 \times C_1 & R_2 \times C_2 \end{bmatrix}$$

$$\begin{aligned} R_1 \times C_1 &= 1 \times 2 + 0 \times 1 + 3 \times 3 = 11 & R_1 \times C_2 &= 1 \times 1 + 0 \times 0 + 3 \times 2 = 7, \\ R_2 \times C_1 &= 2 \times 2 + 1 \times 1 + 5 \times 3 = 20 & \text{and } R_2 \times C_2 &= 2 \times 1 + 1 \times 0 + 5 \times 2 = 12 \end{aligned}$$

$$\text{Therefore } AB = A \times B = \begin{bmatrix} 11 & 7 \\ 20 & 12 \end{bmatrix}$$

Properties of multiplication

1. Matrix multiplication, in general, is not commutative. i.e, $AB \neq BA$.
2. Matrix multiplication is associative. i.e., $A(BC) = (AB)C$
3. Matrix multiplication is distributive, i.e, $A(B+C) = AB + AC$

4. Transpose of a Matrix

The matrix obtained by interchanging the rows and columns of a matrix A is called the transpose of A and is denoted by A' or A^T . Thus if A is an $m \times n$ matrix, then, A^T will be an $n \times m$ matrix.

$$\text{For example, if } A = \begin{bmatrix} 2 & -3 \\ 1 & 0 \\ -2 & -1 \end{bmatrix}, \text{ then } A^T = \begin{bmatrix} 2 & 1 & -2 \\ -3 & 0 & -1 \end{bmatrix}$$

Properties of Transpose

1. Transpose of a sum (or difference) of two matrices is the sum (or difference) of the transposes, i.e. $(A \pm B)^T = A^T \pm B^T$
2. Transpose of transpose is the original matrix. i.e. $(A^T)^T = A$
3. The transpose of a product of two matrices is the product of their transposes taken in reverse order. i.e., $(AB)^T = B^T A^T$

Exercise

1. If matrices A and B are defined as

$$A = \begin{bmatrix} 0 & 2 & 3 \\ 2 & 1 & 4 \end{bmatrix}; B = \begin{bmatrix} 7 & 6 & 3 \\ 1 & 4 & 5 \end{bmatrix}$$

then compute

- a) $A+B$
- b) $A-B$
- c) $B-A$

2. If two matrices A and B are defined as

$$A = \begin{bmatrix} 0 & 2 & 3 \\ 2 & 1 & 4 \end{bmatrix}; B = \begin{bmatrix} 7 & 6 & 3 \\ 1 & 4 & 5 \end{bmatrix}$$

then compute $2A+3B$.

3. If two matrices A and B are defined as

$$A = \begin{bmatrix} 2 & 1 & 2 \\ 2 & 4 & 0 \end{bmatrix}, B = \begin{bmatrix} 2 & 2 \\ 1 & 4 \\ 2 & 0 \end{bmatrix}$$

then verify that $(AB)^t = B^t A^t$

3.6 Determinant of a Square Matrix

The determinant of a square matrix is a scalar (i.e. a number). Determinants are possible only for square matrices. For more clarity, we shall be defining it in stages, starting with square matrix of order 1, then for matrix of order 2, etc. The determinant of a square matrix A is denoted either by $|A|$ or $\det. A$.

- i) Determinant of order 1. Let $A = [a_{11}]$ be a matrix of order 1. Then $\det A = a_{11}$
- ii) Determinant of order 2. Let

$$A = \begin{bmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{bmatrix}$$

be a square matrix of order 2, then $\det. A$ is defined as

$$\det. A = \begin{vmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{vmatrix} = a_{11}a_{22} - a_{21}a_{12}$$

For example

$$\det. A = \begin{vmatrix} 3 & 4 \\ 1 & 2 \end{vmatrix} = 3 \times 2 - 1 \times 4 = 2$$

to write the expansion of a determinant to matrices of order 3, 4, ..., let us first define two important terms:

a) Minor: Let A be a square matrix of order m . Then minor of an element a_{ij} is the determinant of the residual matrix (or submatrix) obtained from A by deleting row i and column j containing the element a_{ij} .

In the $|A|$, the minor of the element a_{ij} is denoted by M_{ij} . Thus, in the determinant of order 3

$$\begin{vmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \\ a_{31} & a_{32} & a_{33} \end{vmatrix}$$

the minor of the element a_{11} is obtained by deleting first row and first column containing element a_{11} and is written as

$$M_{11} = \begin{vmatrix} a_{22} & a_{23} \\ a_{32} & a_{33} \end{vmatrix}$$

Similarly, minor of a_{12} is

$$M_{12} = \begin{vmatrix} a_{21} & a_{23} \\ a_{31} & a_{33} \end{vmatrix}$$

b) Cofactor. The cofactor c_{ij} of an element a_{ij} is defined as

$$C_{ij} = (-1)^{i+j} M_{ij}$$

Where M_{ij} is the minor of an element a_{ij} .

Now using the concept of minor and cofactor, you can write the expansion of a determinant of order 3 as shown below:

$$\begin{vmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \\ a_{31} & a_{32} & a_{33} \end{vmatrix} = a_{11} C_{11} + a_{12} C_{12} + a_{13} C_{13}$$

$$= a_{11}(-1)^{1+1} M_{11} + a_{12}(-1)^{1+2} M_{12} + a_{13}(-1)^{1+3} M_{13}$$

$$=a_{11} \begin{vmatrix} a_{11} & a_{22} \\ a_{32} & a_{33} \end{vmatrix} - a_{12} \begin{vmatrix} a_{21} & a_{23} \\ a_{31} & a_{33} \end{vmatrix} + a_{13} \begin{vmatrix} a_{21} & a_{22} \\ a_{31} & a_{32} \end{vmatrix}$$

$$=a_{11}(a_{22} a_{33}-a_{32} a_{23}) - a_{12}(a_{21} a_{33}-a_{31} a_{23}) + a_{13}(a_{21} a_{32}-a_{31} a_{22})$$

The expansion of the given determinant can also be done by choosing elements in any row and column. In the above example expansion was done by using the elements of the first row.

Example 2

Find the value of the determinant

$$\det.A = \begin{vmatrix} 1 & 18 & 72 \\ 2 & 40 & 96 \\ 2 & 45 & 75 \end{vmatrix}$$

Solution:

If you expand the determinant by using the elements of the first column, then you will get

$$\begin{vmatrix} 1 & 18 & 72 \\ 2 & 40 & 96 \\ 2 & 45 & 75 \end{vmatrix} = 1 \begin{vmatrix} 40 & 96 \\ 45 & 75 \end{vmatrix} - 2 \begin{vmatrix} 18 & 72 \\ 45 & 75 \end{vmatrix} + 2 \begin{vmatrix} 18 & 72 \\ 40 & 96 \end{vmatrix}$$

$$= 1(3000-4300) - 2(1350-3240) + 2(1728-2880)$$

$$= 1 \times (-1320) - 2 \times (-1890) + 2 \times (-1152)$$

$$= -1320 + 3780 - 2304$$

$$= -3624 + 3780 = 156$$

Properties of determinants

Following are the useful properties of determinants of any order. These properties are very useful in expanding the determinants.

1. The value of a determinant remains unchanged. If rows are changed into column and columns into rows, i.e.
 $|A| = |A^t|$
2. If two rows (or columns) of a determinant are interchanged, then the value of the determinant so obtained is the negative of the original determinant.
3. If each element in any row or column of a determinant is multiplied by a constant number say K, then the determinant so obtained is K times the original determinant.
4. The value of a determinant in which two rows (or columns) are equal is zero.
5. If any row (or column) of a determinant is replaced by the sum of the row and a linear combination of other rows (or columns), then the value of the determinant so obtained is equal to the value of the original determinant.
6. The rows (or columns) of a determinant are said to be linearly dependent if $|A|=0$, otherwise independent.

Example 3

Verify the following result

$$\begin{vmatrix} 1 & a & a^2 \\ 1 & b & b^2 \\ 1 & c & c^2 \end{vmatrix} = (a-b)(b-c)(c-a)$$

Applying row operation (property 5)

$$\begin{array}{l} R_2 \rightarrow R_2 + (-1)R_1 \\ R_3 \rightarrow R_3 + (-1)R_1 \end{array}$$

On the given determinant, the new determinant so obtained

$$\begin{vmatrix} 1 & a & a^2 \\ 0 & b-a & b^2-a^2 \\ 0 & c-a & c^2-a^2 \end{vmatrix}$$

Expanding the new determinant by the elements of first column, you will get

$$\begin{vmatrix} b-a & b^2-a^2 \\ c-a & c^2-a^2 \end{vmatrix} = \begin{vmatrix} b-a & (b-a)(b+a) \\ c-a & (c-a)(c+a) \end{vmatrix}$$

Again performing row operation,

$$R_2 \rightarrow 1/(b-a) R_2$$

$$R_3 \rightarrow 1/(c-a) R_3$$

You will have

$$(b-a)(c-a) \begin{vmatrix} 1 & b+a \\ 1 & c+a \end{vmatrix}$$

$$= (b-a)(c-a)(c+a) - (b+a)\}$$

$$= (b-a)(c-a)(c-b)$$

$$= (a-b)(b-c)(c-a)$$

Example4 : $\begin{vmatrix} 2 & 1 & 2 \\ 0 & 5 & 4 \\ 5 & 2 & 1 \end{vmatrix} = 2 \begin{vmatrix} 5 & 4 \\ 2 & 1 \end{vmatrix} - 1 \begin{vmatrix} 0 & 4 \\ 5 & 1 \end{vmatrix} + 2 \begin{vmatrix} 0 & 5 \\ 5 & 2 \end{vmatrix}$

$$= 2(5-8) - 1(0-20) + 2(0-25)$$

$$= -36$$

Singular and Non-singular Matrices: A matrix A is said to be singular if $|A| = 0$; otherwise it is called non-singular.

Exercise

If $a+b+c = 0$, then verify the following result.

$$\begin{vmatrix} a & b & c \\ 0 & a & b \\ b & 0 & a \end{vmatrix} = c(2ab - c^2)$$

3.7 Let us Sum Up

Matrices play an important role in quantitative analysis of managerial decision. They also provide very convenient and compact methods of writing a system of linear simultaneous equations and methods of solving them. These tools have also become very useful in all functional areas of management. Another distinct advantage of matrices is that once the system of equations can be set up in matrix form, they can be solved quickly using a computer. A number of basic matrix operations (such as matrix addition, subtraction, multiplication) were discussed in this Lesson.

3.8 Lesson – End Activities

1. Define matrix, square matrix, diagonal matrix, scalar matrix.
2. Mention the properties of transpose of a matrix.
3. List the properties of determinants.

3.9 References

P.R. Vittal – Business Mathematics and Statistics.

Lesson 4 - Inverse of Matrix

Contents

- 4.1 Aims and Objectives
- 4.2 Inverse of a Matrix
- 4.3 Let us Sum up
- 4.4 Lesson – End Activities
- 4.5 References

4.1 Aims and Objectives

In the last Lesson, Matrix algebra, matrix operations and applications of matrix theory, etc., were discussed in details. This Lesson exclusively describes inverse of matrix, which is another important operation of matrix algebra.

4.2 Inverse of a Matrix

If for a given square matrix A, another square matrix B of the same order is obtained such that

$$AB = BA = I$$

Then matrix B is called the inverse of A and is denoted by $B=A^{-1}$

Before start discussing the procedure of finding the inverse of a matrix, it is important to know the following results:

1. The matrix $B=A^{-1}$ is said to be the inverse of matrix A if and only if

$$AA^{-1}=A^{-1}A=I.$$

2. That is, if the inverse of a square matrix multiplied by the original matrix, then result is an identity matrix. The inverse A^{-1} does not mean I/A or $I \cdot A$. This is simply a notation to denote the inverse of A

3. Every square matrix may not have an inverse. For example, zero matrix has no inverse. Because, inverse of square matrix exists only if the value of its determinant is non-zero, i.e. A^{-1} exists if and only if $|A| \neq 0$.

For example, let B be the inverse of the matrix A, then

$$AB=BA=I$$

Or $|AB|=|I|$

Or $|A||B|=1(|I|=1)$

Hence $|A| \neq 0$.

4. If a square matrix A has an inverse, then it is unique. It can also be proved by letting two inverse B and C of A.

We then have

$$AB = BA = I \dots(i)$$

And

$$AC = CA = I \dots(ii)$$

Pre-multiplying (i) by C, we get

$$CAB = CI$$

$$IB = CI$$

or

$$B = C (CA = I)$$

This implies that the inverse of a square matrix is unique.

Singular Matrix

A matrix is said to be singular if its determinant is equal to zero; Otherwise non-singular.

Properties of the inverse

- i) The inverse of the inverse is the original matrix, i.e. $(A^{-1})^{-1} = A$.
- ii) The inverse of the transpose of a matrix is the transpose of its inverse, i.e. $(A^t)^{-1} = (A^{-1})^t$
- iii) The identity matrix is its own inverse, i.e. $I^{-1} = I$
- iv) The inverse of the product of two non-singular matrices is equal to the products of two inverse in the reverse order, i.e. $(AB)^{-1} = B^{-1} A^{-1}$

Methods of finding inverse of a matrix

The procedure of finding inverse of a square matrix $A = [a_{ij}]$ of order n can be summarized in the following steps:

3. Construct the matrix of co-factors of each element a_{ij} in $|A|$ as follows:

$$\begin{bmatrix} C_{11} & C_{12} & \dots & C_{1n} \\ C_{21} & C_{22} & \dots & C_{2n} \\ \vdots & \vdots & & \vdots \\ \vdots & \vdots & & \vdots \\ C_{m1} & C_{m2} & \dots & \end{bmatrix} C_{mn}$$

In this case cofactors are the elements of the matrix

2. Take the transpose of the matrix of cofactors constructed in step 1. It is called adjoint of A and is denoted by $\text{Adj. } A$.

3. Find the value of $|A|$

4. Apply the following formula to calculate the inverse of A

$$A^{-1} = \frac{\text{Adj. } A}{|A|}, |A| \neq 0$$

Example 1

Find the inverse of the matrix

$$A = \begin{bmatrix} 1 & 3 & 0 \\ -2 & 3 & 3 \\ 1 & 1 & 4 \end{bmatrix}$$

Solution

The determinant of matrix A is expanded with respect to the elements of first row:

$$|A| = \begin{vmatrix} 1 & 3 & 0 \\ -2 & 3 & 3 \\ 1 & 1 & 4 \end{vmatrix} = 1 \begin{vmatrix} 3 & 3 \\ 1 & 4 \end{vmatrix} - 3 \begin{vmatrix} -2 & 3 \\ 1 & 4 \end{vmatrix} + 0 \begin{vmatrix} -2 & 3 \\ 1 & 1 \end{vmatrix}$$

$$= 9 - 3(-11) = 42$$

Since $|A| \neq 0$, therefore the inverse of A exists. The matrix of cofactor of elements A is:

$$C_{11} = (-1)^{1+1} M_{11} = \begin{vmatrix} 3 & 3 \\ 1 & 4 \end{vmatrix} = 9$$

$$C_{12} = (-1)^{1+2} M_{12} = \begin{vmatrix} -2 & 3 \\ 1 & 4 \end{vmatrix} = 11$$

$$C_{13} = (-1)^{1+3} M_{13} = \begin{vmatrix} -2 & 3 \\ 1 & 1 \end{vmatrix} = -5$$

$$C_{21} = (-1)^{2+1} M_{21} = \begin{vmatrix} -3 & 0 \\ 1 & 4 \end{vmatrix} = -12$$

$$C_{22} = (-1)^{2+2} M_{22} = \begin{vmatrix} 1 & 0 \\ 1 & 4 \end{vmatrix} = 4$$

$$C_{23} = (-1)^{2+3} M_{23} = \begin{vmatrix} -1 & 3 \\ 1 & 1 \end{vmatrix} = 2$$

$$C_{31} = (-1)^{3+1} M_{31} = \begin{vmatrix} 3 & 0 \\ 3 & 3 \end{vmatrix} = 9$$

$$C_{32} = (-1)^{3+2} M_{32} = \begin{vmatrix} -1 & 0 \\ -2 & 3 \end{vmatrix} = -3$$

$$C_{33} = (-1)^{3+3} M_{33} = \begin{vmatrix} 1 & 3 \\ -2 & 3 \end{vmatrix} = 9$$

The matrix of cofactors of elements of matrix A is

$$\begin{bmatrix} C_{11} & C_{12} & C_{13} \\ C_{21} & C_{22} & C_{23} \\ C_{31} & C_{32} & C_{33} \end{bmatrix} = \begin{bmatrix} 9 & 11 & -5 \\ -12 & 4 & 2 \\ 9 & -3 & 9 \end{bmatrix}$$

$$\begin{array}{ccc} C_{31} & C_{32} & C_{33} \\ \hline & 9 & -3 & 9 \end{array}$$

The adj. A is now constructed by taking transpose of the cofactor matrix:

$$\text{Adj. A} = (\text{Co-factor A})^t \quad \begin{bmatrix} 9 & -12 & 9 \\ 11 & 4 & -3 \\ -5 & 2 & 9 \end{bmatrix}$$

Hence

$$A^{-1} = \frac{\text{Adj A}}{|A|}$$

$$= \frac{1}{42} \begin{bmatrix} 9 & -12 & 9 \\ 11 & 4 & -3 \\ -5 & 2 & 9 \end{bmatrix}$$

Exercise

For the matrix

$$A = \begin{bmatrix} 1 & 4 & 0 \\ -1 & 2 & 0 \\ 0 & 0 & 2 \end{bmatrix}$$

- i) Calculate A^{-1}
- ii) Verify $(A^t)^{-1} = (A^{-1})^t$
- iii) Verify $(\text{adj A})^{-1} = \text{adj}(A^{-1})$

4.3 Let us Sum Up

Subsequent to the last Lesson, a discussion on matrix inversion and procedure for finding matrix inverse was discussed in this Lesson. Examples were also given in support of the inverse of a matrix. The inverse of matrix finds applications in most of the problems in matrix algebra like in business applications while solving linear equations.

4.4 Lesson – End Activities

1. How to find the Inverse of a Matrix?

4.5 Reference

Navaneethan, P. – Business Mathematics.

Lesson 5 - Matrix Methods to Solve Simultaneous Equations

Contents

- 5.1 Aims and Objectives
- 5.2 Solution of Linear Simultaneous Equations
- 5.3 Let us Sum Up
- 5.4 Lesson – End Activities
- 5.5 Reference

5.1 Aims and Objectives

Matrix theory was discussed in detail in the previous Lessons. In business applications there are several occasions in which mathematical solution are to be made using simultaneous equations. Matrix algebra is useful in solving a set of linear simultaneous equations involving more than two variables. Now the procedure for getting the solution will be demonstrated in this Lesson.

5.2 Solution of Linear Simultaneous Equations

Consider the set of linear simultaneous equations

$$\begin{aligned}x - y + z &= 4 \\ 2x + 5y - 2z &= 3\end{aligned}$$

These equations can also be solved by using ordinary algebra. However, to demonstrate the use of matrix algebra, the first step is to write the given system of equations to matrix form as follows:

$$\begin{bmatrix} 1 & 1 & 1 \\ 2 & 5 & -2 \end{bmatrix} \begin{bmatrix} X \\ Y \\ Z \end{bmatrix} = \begin{bmatrix} 4 \\ 3 \end{bmatrix}$$

or $AX=B$
 where $A = \begin{bmatrix} 1 & 1 & 1 \\ 2 & 5 & -2 \end{bmatrix}$

Is known as the coefficient matrix in which coefficients of x are written in first column, coefficients of y in second column and the coefficients of z in the third column.

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

Is the matrix of unknown variables x,y and z, and

$$B = \begin{bmatrix} 4 \\ 3 \end{bmatrix}$$

is the matrix formed with the right hand terms in equations which do not involve unknowns x, y and z .

Generalizing the situation, let us consider m linear equations in n -unknowns $x_1, x_2, \dots, x_n$;

$$A_{11} X_1 + a_{12} X_2 + \dots + a_{1n} X_n = b_1$$

$$A_{21} X_1 + a_{22} X_2 + \dots + a_{2n} X_n = b_2$$

.....

$$A_{m1} X_1 + a_{m2} X_2 + \dots + a_{mn} X_n = b_m$$

Writing this system of equations in matrix form,

$$AX = B$$

Where

$$A = \begin{bmatrix} A_{11} & a_{12} \dots a_{1n} \\ a_{21} & a_{22} \dots A_{2n} \\ \dots & \dots \\ a_{m1} & a_{m2} \dots a_{mn} \end{bmatrix} m \times n$$

$$X = \begin{bmatrix} X_1 \\ X_2 \\ \cdot \\ \cdot \\ \cdot \\ X_n \end{bmatrix} n \times 1$$

$$B = \begin{bmatrix} b_1 \\ b_2 \\ \cdot \\ \cdot \\ \cdot \\ b_m \end{bmatrix} m \times 1$$

Classification of linear Equations

If matrix B is zero matrix, i.e. $B=0$, then the system $AX=0$ is said to be homogeneous system. Otherwise, the system is said to be non-homogeneous.

Homogeneous Linear Equations

When the system is homogenous, i.e. $b_1=b_2= \dots =b_m=0$, the only possible solution is $X=0$ or $X_1=X_2=\dots X_n=0$. it is called a trivial solution. Any other solution if it exists is called non-trivial solution of the homogenous linear equations.

In order to solve the equation $Ax=0$, we perform such an elementary operations or transformations on the given coefficient matrix A which does not change the order of the matrix. An elementary operation is of any one of the following three types:

- i) The interchange of any two rows (or columns)

$$[A:B] = \begin{bmatrix} A_{11} & a_{12} \dots a_{1n} & . & b_1 \\ a_{21} & a_{22} \dots & A_{2n} & . & b_2 \\ \dots & \dots & \dots & \dots & \dots \\ a_{m1} & a_{m2} \dots a_{mn} & . & b_m \end{bmatrix}$$

$$\begin{bmatrix} 1 & 3 & -2 \\ 2 & -1 & 4 \\ 1 & -11 & 14 \end{bmatrix} \begin{bmatrix} X1 \\ X2 \\ X3 \end{bmatrix} = \begin{bmatrix} 0 \\ 0 \\ 0 \end{bmatrix} \text{ or } AX=0$$

$$[A:0]^+ \begin{bmatrix} 1 & 3 & -2 & : 0 \\ 2 & -1 & 4 & : 0 \\ 1 & -11 & 14 & : 0 \end{bmatrix}$$

$$R_3 = R_3 - R_1$$

The new equivalent matrix is:

$$\begin{bmatrix} 1 & 3 & -2 & . & 0 \\ 0 & 7 & 8 & . & 0 \\ 0 & -14 & 16 & . & 0 \end{bmatrix}$$

Again applying $R_3 \leftarrow R_3 - 2R_1$. The new equivalent matrix is:

$$\begin{bmatrix} 1 & 3 & -2 & . & 0 \\ 0 & -7 & 8 & . & 0 \\ 0 & 0 & 0 & . & 0 \end{bmatrix}$$

The equations equivalent to the given system of equations obtained by elementary row operations are:

$$X_1 + 3X_2 - 2X_3 = 0$$

$$-7X_2 + 8X_3 = 0 \text{ or } X_2 - (8/7)X_3 = 0$$

$$0 = 0$$

The last equation, though true, is redundant and the system is equivalent to

$$X_1 + 3X_2 - 2X_3 = 0$$

$$X_2 - (8/7)X_3 = 0$$

This is not in triangular form because the number of equations being less than the number of unknowns.

This system can be solved in terms of X_3 by assigning an arbitrary constant value, k to it.

The general solution to the given system is given by

$$X_3 = k$$

$$X_2 = (8/7)k$$

$$X_1 + 3X_2 = 2k_3 \text{ or } X_1 = -3(8/7)k + 2k = (-10/7)k$$

Exercise

Solve the following system of equations using Gauss-Jordan Method

i) $4X_1 + X_2 = 0$

$$-8X_1 + 2X_2 = 0$$

ii) $X_1 - 2X_2 + 3X_3 = 0$

$$2X_1 + 5X_2 + 6X_3 = 0$$

Non-homogeneous Linear Equations:

The non-homogeneous linear equations can be solved by any of the following methods

1 Matrix Inverse Method

2 Cramer's Method

3 Gauss-Jordan Method

Again, for the purpose of demonstrating above solution methods, we shall consider three equations with three unknowns.

1. Matrix Inverse Method

Let $AX = B$

Be the given system of linear equations, and also A^{-1} be the inverse of A .

Pre-multiplying both sides of the equation by A^{-1} ,

$$A^{-1}(AX) = A^{-1}B$$

$$(A^{-1}A)X = A^{-1}B$$

$$IX = A^{-1}B$$

$$X = A^{-1}B$$

Where I is the identity matrix.

The value of X gives the general solution to the given set of simultaneous equations.

This solution is thus obtained by (i) first finding A^{-1} , and (ii) post multiplying A^{-1} by B.

When the system has a solution, it is said to be consistent, otherwise inconsistent. A consistent system has either just one solution or infinitely many solutions.

Example 1

The daily cost, C of operating a hospital, is a linear function of the number of in patients I, and out-patients, P, plus a fixed cost a, i.e.,

$$C = a + bP + dI.$$

Given the following data for three days, find the values of a, b, and d by setting up a linear system of equations and using the matrix inverse.

Day	Cost (in Rs.)	No.of in-patients, I	No.of out-patients, P
1	6,950	40	10
2	6,725	35	9
3	7,100	40	12

Solution:

Based on the given daily cost equation, the system of equations for three days cost can be written as:

$$a + 10b + 40d = 6,950$$

$$a + 9b + 35d = 6,725$$

$$a + 12b + 40d = 7,100$$

This system can be written in the matrix form as follows:

$$\begin{bmatrix} 1 & 10 & 40 \\ 1 & 9 & 35 \\ 1 & 12 & 40 \end{bmatrix} \begin{bmatrix} a \\ b \\ d \end{bmatrix} = \begin{bmatrix} 6,950 \\ 6,725 \\ 7,100 \end{bmatrix}$$

Which is of the form $AX=B$, where

$$\begin{bmatrix} 1 & 10 & 40 \\ 1 & 9 & 35 \\ 1 & 12 & 40 \end{bmatrix}; \quad X = \begin{bmatrix} a \\ b \\ d \end{bmatrix} = \quad , \text{ and } B = \begin{bmatrix} 6,950 \\ 6,725 \\ 7,100 \end{bmatrix}$$

The inverse of a matrix A is obtained as follows:

$$\text{Adj.}A = \begin{bmatrix} C_{11} & C_{12} & C_{13} \\ C_{21} & C_{22} & C_{23} \\ C_{31} & C_{32} & C_{33} \end{bmatrix} = \begin{bmatrix} 60 & 5 & -3 \\ -80 & 0 & 2 \\ 10 & -5 & 1 \end{bmatrix}^t = \begin{bmatrix} 60 & -80 & 10 \\ 5 & 0 & -5 \\ -3 & 2 & 1 \end{bmatrix}$$

$$\begin{bmatrix} 1 & 10 & 40 \\ 9 & 35 \\ 1 & 35 \\ 1 & 9 \end{bmatrix}$$

$$|A| = \begin{vmatrix} 1 & 9 & 35 \\ 1 & 12 & 40 \end{vmatrix} = 1 \begin{vmatrix} 12 & 40 \\ -10 & 1 \end{vmatrix} - 40 \begin{vmatrix} 1 & 1 \\ 40 & 1 \end{vmatrix} + 40 \begin{vmatrix} 1 & 12 \\ 1 & 1 \end{vmatrix}$$

$$= (360 - 420) - 10(40 - 35) + 40(12 - 9) \\ = -10 \quad 0$$

Since $|A| \neq 0$, therefore inverse of matrix A exists and is computed as

$$A^{-1} = \frac{\text{Adj. } A}{|A|}$$

$$= -1/10 \begin{bmatrix} 60 & -80 & 10 \\ 5 & 0 & -5 \\ -3 & 2 & 1 \end{bmatrix}$$

$$\therefore X = A^{-1}B$$

$$\text{or } \begin{bmatrix} a \\ b \\ d \end{bmatrix} = -1/10 \begin{bmatrix} 60 & -80 & 10 \\ 5 & 0 & -5 \\ -3 & 2 & 1 \end{bmatrix} \begin{bmatrix} 6,950 \\ 6,725 \\ 7,100 \end{bmatrix}$$

$$= -1/10 \begin{bmatrix} 60 \times 6,950 - 80 \times 6,725 + 10 \times 7,100 \\ 5 \times 6,950 + 0 \times 6,725 - 5 \times 7,100 \\ -3 \times 6,950 + 2 \times 6,725 + 1 \times 7,100 \end{bmatrix}$$

$$= -1/10 \begin{bmatrix} -50,000 \\ -750 \\ -300 \end{bmatrix} = \begin{bmatrix} 5,000 \\ 75 \\ 30 \end{bmatrix}$$

$$\text{or } a = 5000, b = 75 \text{ and } d = 30$$

Exercise

A salesman has the following record of sales during three months for three items A, B and C, which have different rates of commission.

Months	Sales of Units			Total Commission drawn (in Rs.)
	A	B	C	
January	90	100	20	800
February	150	50	40	900
March	60	100	30	850

Find out the rates of commission on items A, B and C.

2. Cramer's Method

When the number of equations is equal to the number of unknowns and the determinant of the coefficients has non-zero value, then the system has a unique solution which can be found by using Cramer's formula.

$$X_j = \frac{D_j}{D}, j = 1, 2, \dots, n$$

Where $D = |a_{ij}|$ and determinant D_j is obtained from D by replacing column j by the column of constant terms (i.e. matrix B).

Example 2

An automobile company uses three types of steel, S_1, S_2 and S_3 for producing three different types of cars C_1, C_2 and C_3 . Steel requirements (intones) for each type of car and total available steel of all the three types is summarized in the following table.

Types of steel	Type of car			Total steel available
	C_1	C_2	C_3	
S_1	2	3	4	29
S_2	1	1	2	13
S_3	3	2	1	16

Determine the number of cars of each type which can be produced.

Solution:

Let X_1, X_2 and X_3 be the number of cars of the type C_1, C_2 and C_3 respectively which can be produced. Then system of three linear equations is:

$$2x_1 + 3x_2 + 4x_3 = 29$$

$$x_1 + x_2 + 2x_3 = 13$$

$$3x_1 + 2x_2 + x_3 = 16$$

These equations can also be represented in matrix form as shown below:

$$\begin{bmatrix} 2 & 3 & 4 \\ 1 & 1 & 2 \\ 3 & 2 & 1 \end{bmatrix} \begin{bmatrix} x_1 \\ x_2 \\ x_3 \end{bmatrix} = \begin{bmatrix} 29 \\ 13 \\ 16 \end{bmatrix}$$

The determinant of the coefficients matrix is

$$\begin{bmatrix} 2 & 3 & 4 \\ 1 & 1 & 2 \\ 3 & 2 & 1 \end{bmatrix} = 2 \begin{bmatrix} 1 & 2 \\ 2 & 1 \end{bmatrix} - 3 \begin{bmatrix} 1 & 2 \\ 3 & 1 \end{bmatrix} + 4 \begin{bmatrix} 1 & 1 \\ 3 & 2 \end{bmatrix}$$

$$= 2(1-4) - 3(1-6) + 4(2-3)$$

$$= 5(0)$$

Applying Cramer's Method

$$x_1 = \frac{D_1}{D} = \frac{1}{5} \begin{bmatrix} 29 & 3 & 4 \\ 13 & 1 & 2 \\ 16 & 2 & 1 \end{bmatrix} = 2$$

$$x_2 = \frac{D_2}{D} = \frac{1}{5} \begin{bmatrix} 2 & 29 & 4 \\ 1 & 13 & 2 \\ 3 & 16 & 1 \end{bmatrix} = 3$$

$$x_3 = \frac{D_3}{D} = \frac{1}{5} \begin{bmatrix} 2 & 3 & 29 \\ 1 & 1 & 13 \\ 3 & 2 & 16 \end{bmatrix} = 4$$

Hence, the number of cars of type C1, C2 and C3 which can be produced are 2, 3 and 4 respectively.

Total time available is 80 hours and 60 hours in department I and II respectively. Determine the number of units of product A and B which should be produced.

3. Gauss – Jordan Method

We can solve a system of linear equations by transform the augmented matrix $[A:B]$ into a triangular form.

Example 3 : Solve $x + 2y = 3$
 $2x + 5y = 2$

Solution: The system of equations can be written as $\begin{bmatrix} 1 & 2 \\ 2 & 5 \end{bmatrix} \begin{bmatrix} x \\ y \end{bmatrix} = \begin{bmatrix} 3 \\ 2 \end{bmatrix}$

That is, $A = \begin{bmatrix} 1 & 2 \\ 2 & 5 \end{bmatrix}$ and $B = \begin{bmatrix} 3 \\ 2 \end{bmatrix}$

The augmented matrix is $[A:B] = \begin{bmatrix} 1 & 2 & 3 \\ 2 & 5 & 2 \end{bmatrix}$

Apply transformation, $R_2 = R_2 - 2R_1$

$$= \begin{bmatrix} 1 & 2 & 3 \\ 0 & 1 & -4 \end{bmatrix}$$

Then the equations are $x + 2y = 3$ and $y = -4$. Substitute the value of y in the first equation we get the value of x . That is, $x = 3 - (-8) = 11$.

So the solution is $x = 11$ and $y = -4$.

Exercises

1. If $A = \begin{bmatrix} 2 & 3 & 4 \\ 1 & 1 & 2 \\ 3 & 2 & 1 \end{bmatrix}$ and $B = \begin{bmatrix} 9 & -12 & 9 \\ 11 & 4 & -3 \\ -5 & 2 & 9 \end{bmatrix}$, then find i) $A+B$ ii) $A-B$ iii) AB

and iv) BA . Also show that $AB \neq BA$.

2. Find the inverse of the matrix $A = \begin{bmatrix} 1 & -1 & 4 \\ -1 & 1 & 8 \\ 3 & -4 & 1 \end{bmatrix}$

3. Solve the equations by matrix inverse method, Cramer's method and Gauss-Jordan Method

$$x + 3y + z = 11$$

$$2x + y + 4z = 7$$

$$-x + 2y + 2z = 5$$

4. A firm makes two products A and B. Each product requires production time in each of two departments I and II as shown below:

Product	Time taken (in hrs/week)	
	Deptt.I	Deptt.II
A	5	4
B	6	2

5.3 Let us Sum Up

The methods for solving linear equations using matrix theory is described in this Lesson. The three important methods of solving the equations using Cramer's rule, matrix inverse method and Gauss – Jordan method, are described in detail in this Lesson. Number of examples were given in support of the above said operations.

5.4 Lesson – End Activities

1. How to solve a system of linear equation by Gauss – Jordan Method?

5.5 References

1. Navaneethan. P – Business Mathematics.
2. P.R. Vital – Business Mathematics and Statistics.

UNIT - II

Lesson 6 - Sequences and Series

Contents

- 6.1 Aims and Objectives
- 6.2 Sequence
- 6.3 Series
- 6.3 Arithmetic Progression (AP)
- 6.5 Geometric Progression (GP)
- 6.6 Let us Sum Up
- 6.7 Lesson – End Activities
- 6.8 References

6.1 Aims and Objectives

This Lesson deals with the concepts and applications of sequence and series. A clear understanding about sequence and series is provided. Applications of series like Arithmetic Progression and Geometric Progression and practical applications in business are also dealt with in this Lesson.

6.2 Sequence

If for every positive integer n , there corresponds a number a_n such that a_n is related to n by some rule, then the terms $a_1, a_2, \dots, a_n, \dots$ are said to form a sequence. A sequence is denoted by bracketing its n^{th} term, i.e. (a_n) or $\{a_n\}$.

Example of a few sequences are:

- i) If $a_n = n^2$, then sequence $\{a_n\}$ is $1, 4, 9, 16, \dots, a_n, \dots$
- ii) If $a_n = 1/n$, then sequence $\{a_n\}$ is $1, 1/2, 1/3, 1/4, \dots, 1/n, \dots$
- iii) If $a_n = n^2/n+1$, then sequence $\{a_n\}$ is $1/2, 4/3, 9/4, \dots, n^2/n+1, \dots$

The concept of sequence is very useful in finance. Some of the major areas where it plays a vital role are: “instalment buying”; simple and compound interest problems; ‘annuities and their present values’, mortgage payments and so on

6.3 Series

A series is formed by connecting the terms of a sequences with plus or minus sign. Thus if a_n is the n th term of a sequence, then $a_1 + a_2 + \dots + a_n$ is the given series of n terms.

6.4 Arithmetic Progression (AP)

A progression is a sequence whose successive terms indicate the growth or progress of some characteristics. An arithmetic progression is a sequence whose term increases or decreases by a constant number called common difference of an A.P. and is denoted by d . In other words, each term of the arithmetic progression after the first is obtained by adding a constant d to the preceding term. The standard form of an A.P. is written as

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

where 'a' is called the first term. Thus the corresponding standard form of an arithmetic series becomes

$$a+(a+d)+(a+2d)+(a+3d)+\dots$$

Example 1

Suppose we invest Rs. 100 at a simple interest of 15% per annum for 5 years. The amount at the end of each year is given by

$$115, 130, 145, 160, 175$$

This forms an arithmetic progression

The n^{th} Term of an A.P.

The n^{th} term of an A.P. is also called the general term of the standard A.P. it is given by.

$$T_n = a+(n-1)d; \quad n=1, 2, 3, \dots$$

Sum of the First n terms of an A.P

Consider the first n terms of an A.P.

$$a, a+d, a+2d, a+3d, \dots, a+(n-1)d$$

The sum, S_n of these terms is given by

$$\begin{aligned} S_n &= a+(a+d) + (a+2d) + (a+3d) + \dots + a+(n-1)d \\ &= (a+a+\dots+a) + d(1+2+(n-1)3+\dots+) \end{aligned}$$

$$= n.a + d \left\{ \frac{n(n-1)}{2} \right\} \quad \text{(using formula for the sum of first (n-1) natural numbers)}$$

$$= n/2 \{2a+(n-1)d\}$$

Example 2

Suppose Mr. Anil repays a loan of Rs. 3250 by paying Rs. 20 in the first month and then increases the payment by Rs. 15 every month. How long will he take to clear his loan?

Solution

Since Mr. Anil increases the monthly payment by a constant amount, Rs.15 every month, therefore $d = 15$ and first month instalment is, $a = \text{Rs. } 20$. This forms an A.P. Now if the entire amount be paid in n monthly instalments, then we have

$$S_n = n/2 \{2a+(n-1)d\}$$

$$\text{Or } 3250 = n/2 \{2 \times 20 + (n-1)15\}$$

$$6500 = n\{25+15n\}$$

$$15n^2 + 25n - 6500 = 0$$

This is a quadratic equation in n . Thus to find the values of n which satisfy this equation, we shall apply the following formula as discussed before.

$$\begin{aligned} N &= \frac{-b \pm \sqrt{b^2 - 4ac}}{2a} = \frac{-25 \pm \sqrt{(25)^2 - 4 \times 15 \times (-6500)}}{2 \times 15} \\ &= \frac{-25 \pm 625}{30} = 20 \text{ or } -21.66 \end{aligned}$$

The value, $n = -21.66$ is meaningless as n is positive integer. Hence Mr. Anil will pay the entire amount in 20 months.

Exercise

- Find the 15th term of an A.P. whose first term is 12 and common difference is 2.
- A firm produces 1500 RV sets during its first year. The total production of the firm at the end of the 15th year is 8300 TV sets, then
 - estimate by how many units, production has increased each year.
 - based on estimate of the annual increment in production, forecast the amount of production for the 10th year.

6.5 Geometric Progression (GP)

A geometric progression (GP) is a sequence whose each term increases or decreases by a constant ratio called common ratio of G.P. and is denoted by r . In other words, each term of G.P. is obtained after the first by multiplying the preceding term by a constant r . The standard form of a G.P. is written as :

$$a, ar, ar^2, \dots$$

Where 'a' is called the first term. Thus the corresponding geometric series in standard form becomes

$$a + ar + ar^2 + \dots$$

Example 3

Suppose we invest Rs. 100 at a compound interest of 12% per annum for three years. The amount at the end of each year is calculated as follows:

$$\begin{aligned} \text{i) Interest at the end of first year} &= 100 \times 12/100 = \text{Rs. } 12 \\ \text{Amount at the end of first year} &= \text{Principal} + \text{Interest} \\ &= 100 + 100(12/100) \\ &= 100(1 + 12/100) \end{aligned}$$

This shows that the principal of Rs. 100 becomes Rs. $100(1 + 12/100)$ at the end of first year.

$$\begin{aligned} \text{ii) Amount at the end of second year} &= \\ &= (\text{Principal at the beginning of second year}) \{1 + 12/100\} \\ &= 100\{1 + 12/100\} \{1 + 12/100\} \\ &= 100\{1 + 12/100\}^2 \end{aligned}$$

$$\text{iii) Amount at the end of third year} = 100\{1 + 12/100\}^2 \{1 + 12/100\}$$

$$=100\{1+12/100\}^3$$

Thus, the progression giving the amount at the end of each year is

$$100\{1+12/100\}^2 ; 100\{1+12/100\}^2 ; 100\{1+12/100\}^3 ; \dots$$

This is a G.P with common ratio $r = (1+12/100)$

In general, if P is the principal and I is the compound interest rate per annum, then the amount at the end of first year becomes $P(1+i)$. Also the amount at the end of successive years forms a G.P. is

$$P(1+i/100); P(1+i/100)^2 : \dots$$

with $r = (1+i/100)$

The n^{th} Term of G.P.

The n^{th} term of G.P. is also called the general term of the standard G.P. It is given by

$$T_n = ar^{n-1}, n=1,2,3,\dots$$

It may be noted here that the power of r is one less than the index of T_n , which denotes the rank of this term in the progression.

Sum of the First n Terms in G.P.

Consider the first n terms of the standard form of G.P.

$$a, ar, ar^2, \dots, ar^{n-1}$$

The sum, S_n of these terms is given by

$$S_n = a + ar + ar^2 + \dots + ar^{n-2} + ar^{n-1}$$

Multiplying both sides by r , we get

$$rS_n = ar + ar^2 + ar^3 + \dots + ar^{n-1} + ar^n$$

Subtracting the above equations, we have

$$S_n - rS_n = a - ar^n$$

$$S_n(1-r) = a(1-r^n)$$

$$\text{or } S_n = \frac{a(1-r^n)}{(1-r)} ; r \neq 1 \text{ and } < 1$$

Changing the signs of the numerator and denominator, we have

$$S_n = \frac{a(r^n-1)}{(r-1)} ; r \neq 1 \text{ and } > 1$$

a) If $r = 1$, G.P. becomes $a, a, a, \dots$ so that S_n in this case is $S_n = n.a.$

b) If number of terms in a G.P. are infinite, the

$$S_n = \frac{a}{1-r}, r < 1$$

$$= \frac{a}{r-1}, r > 1$$

Example 4

A car is purchased for Rs. 80,000. Depreciation at 5% per annum for the first 3 years and 10% per annum for the next 3 years. Find the money value of the car after a period of 6 years.

Solution:

i) Depreciation for the first year = $80,000 \times 5/100$. Thus the depreciated value of the car at the end of first year is:

$$= (80,000 - 80,000 \times 5/100)$$

$$= 80,000(1-5/100)$$

ii) Depreciation for the second year

$$= (\text{Depreciated value at the end of first year}) \times (\text{Rate of depreciation for the second year})$$

$$= 80,000(1-5/100) (5/100)$$

Thus the depreciated value at the end of the second year is

$$= (\text{Depreciated value after first year}) - (\text{Depreciation for the second year})$$

$$= 80,000(1-5/100) - 80,000(1-5/100) (5/100)$$

$$= 80,000(1-5/100) (1-5/100)$$

$$= 80,000(1-5/100)^2$$

Calculating in the same way, the depreciated value at the end of three years is

iii) Depreciation for fourth year

$$= 80,000(1-5/100)^3 (10/100)$$

Thus the depreciated value at the end of the fourth year is

$$= (\text{Depreciated value after three year}) \times (\text{Depreciation for fourth year})$$

$$= 80,000(1-5/100)^3 - 80,000(1-5/100)^3 (10/100)$$

$$= 80,000(1-5/100)^3 (1-10/100)$$

Calculating in the same way, the depreciated value at the end of six years becomes

$$= 80,000(1-5/100)^3 (1-10/100)^3$$

$$= \text{Rs. } 49,989.24$$

Exercise

1. Determine the common ratio of the G.P.

49, 7, 1/7, 1/49,

a) Find the sum to first 20 terms of G.P.

b) Find the sum to infinity of the terms of G.P.

2. The population of a country in 1985 was 50 crore.

Calculate the population in the year 2000 if the compounded annual rate of increase is

(a) 1% (b) 2%

6.6 Let us Sum Up

Attention is then directed to defining the Arithmetic and Geometric Progressions and subsequently to their applications. An Arithmetic Progression is a sequence whose terms increases or decreases by a constant number. Geometric Progression is a sequence whose terms increases or decreases by a constant ration. Applications of these series find in sales, calculating interest in business, population studies, other socio-economic applications. Some of the examples shown in this Lesson brings out the importance of the study.

6.7 Lesson – End Activities

1. Define Arithmetic progression

2. Define Geometric progression

6.8 Reference

1. Navaneethan. P – Business Mathematics

2. P.R. Vital – Business Mathematics and Statistics.

Lesson 7 - Mathematics for Finance

Contents

- 7.1 Aims and Objectives
- 7.2 Some terms used in business calculations
- 7.3 Difference between Simple and Compound Interest
- 7.4 Formula for Amount Accrued (simple interest)
- 7.5 Formula for Amount Accrued (compound interest)
- 7.6 Notes on previous formula
- 7.7 Formula for calculating APR
- 7.8 Depreciation
- 7.9 Formula for Reducing Balance Depreciation
- 7.10 Let us Sum Up
- 7.11 Lesson – End Activities
- 7.12 References

7.1 Aims and Objectives

So far we have discussed about various mathematical functions and theories. This Lesson deals with the applications of such theories in Finance. In financial management, lot of calculations are involved in the case of interest, depreciation values, and so on.

7.2 Some terms used in business calculations

- a) **Principal amount (P).** This is the amount of money that is initially being considered. It might be an amount about to be invested or loaned or it may refer to the initial value or cost of plant or machinery. Thus if a company was considering a bank loan value or cost of plant or machinery. Thus if a company was considering a bank loan of say Rs.20000, this would be referred to as the principal amount to be borrowed.
- b) **Accrued amount (A).** This term is applied generally to a principal amount after some time has elapsed for which interest has been calculated and added. It is quite common to qualify a precisely according to time elapsed. Thus A^1 , A^2 , etc would mean the amount accrued at the end of the first and second years and so on. The company referred to in (a) above might owe, say, an accrued amount of Rs.22000 at the end of the first year and Rs. 24200 at the end of the second year (if no repayments had been made prior to this time).
- c) **Rate of interest (i).** Interest is the name given to a proportionate amount of money which is added to some principal amount (invested or borrowed). It is normally denoted by symbol i and expressed as a percentage rate per annum. For example if Rs. 100 is invested at interest rate 5% per annum (pa), it will accrue to Rs. $100 + (5\% \text{ of Rs. } 100) = \text{Rs } 100 + \text{Rs.}5 = \text{Rs.}105$ at the end of one year. Note however, that for calculation purposes, a percentage rate is best written as a proportion. Thus, an interest rate of 10% would be written as $i = 0.1$ and 12.5% as $i = 0.125$ and so on.

- d) **Number of time periods (n).** The number of time periods over which amounts of money are being invested or borrowed is normally denoted by the symbol n . although n is usually a number of years, it could represent other time periods, such as a number of quarters or months.

7.3 Difference between Simple and Compound Interest

When an amount of money is invested over a number of years, the interest earned can be dealt with in two ways.

- a) **Simple interest.** This is where any interest earned is NOT added back to the principal amount invested.

For example, suppose that Rs. 200 is invested at 10% simple interest per annum. The following table shows the state of the investment, year by year.

Year	Amount on which interest is calculated	Interest earned	Cumulative amount accrued
1	Rs. 200	10% of Rs. 200 = Rs. 20	Rs. 220
2	Rs. 200	10% of Rs. 200 = Rs. 20	Rs. 240
3	Rs. 200	10% of Rs. 200 = Rs. 20	Rs. 260 ... etc.

- b) **Compound interest.** This is where interest earned is added back to the previous amount accrued.

For example, suppose that Rs. 200 is invested at 10% compound interest. The following table shows the state of the investment, year by year:

Year	Amount on which interest is calculated	Interest earned	Cumulative amount accrued
1	Rs. 200	10% of Rs. 200 = Rs.20	Rs. 220
2	Rs. 220	10% of Rs. 220 = Rs. 22	Rs. 242
3	Rs. 242	10% of Rs. 242 = Rs. 24. 20	Rs. 266.20

...
...etc

The difference between the two methods can easily be seen by comparing the above two tables. Notice that the amount on which simple interest is calculated is always the same, namely, the original principal.

Note that whether a principal amount is being invested [(as in a) and b)] or borrowed makes no difference to the considerations for interest.

7.4 Formula for Amount Accrued (simple interest)

If Rs. 1000 is invested at 5% simple interest (pa) then 5% of Rs. 1000 = Rs.50 will be earned every year.

In general terms, if amount P is invested at $100i\%$ simple interest (per time period), then the amount of interest earned per time period is given by $P.i$.

For the data given, $P = 1000$ (Rs.) and $100i\% = 5\%$ (i.e. $i = 0.05$).

Therefore, the interest earned per year = $P.i = 1000(0.05) = \text{Rs. } 50$.

Also, if Rs. 3000 ($P=3000$) is invested at a (simple) interest rate of 9.5% ($i=0.095$), the interest earned in any year is $P.i = 3000(0.095) = \text{Rs. } 285$. Over a period of, say, 5 years, the accrued amount of the investment would be given by: $\text{Rs.}[2000+5(285)] = \text{Rs. } 4425$. that is, the original principal plus 5 year's-worth of interest.

In general terms, if amount P is invested at $100i\%$ simple interest, then the amount accrued over n years is given by the following formula.

Accrued amount formula (simple interest) :

$$A_n = P(1+i.n)$$

Where: A_n = accrued amount at end of n-th year

P = Principal amount

i = (proportional) interest rate per year

n = number of years.

Generally, simple interest is of no great practical value in modern business and commercial situations, since in practice interest is always compounded.

7.5 Formula for Amount Accrued (compound interest)

Section 2.2.3 gave details of the year-by-year state of an investment of Rs. 200 at 10% compounded per annum (sometimes called an investment schedule). The following derives, in general terms, a formula to give the accrued amount of a principal at the end of any time period.

If principal P is invested at a (compound) interest rate of $100i\%$ over n time periods, then:

$A_1 = P + P.i$ i.e. the total amount accrued at the end of the first time period is the amount of the original principal plus the interest earned (i) for this period.

So, $A_1 = P + P.i$ the interest earned in one time period is $P.i$

i.e. $A_1 = P(1+i)$ factorizing

$A_2 = P(1+i) + P(1+i).i$ amount accrued at end of second time period is the amount at the beginning of the period plus the interest earned (I)

i.e. $A_1 = P(1+i) + P(1+i).i$ interest earned is accrued amount (at beginning) times i

$= P(1+i)(1+i)$ factorizing

i.e. $A_2 = P(1+i)^2$

Similarly, $A_3 = P(1+i)^3$, $A_4 = P(1+i)^4$ and so on.

For example, if Rs. 5000 ($P=5000$) is invested at 9% p.a ($i=0.09$), then:

Amount accrued after 1 year: $A_1 = P(1+i) = 5000(1.09) = \text{Rs. } 5450.$

Similarly: $A_2 = P(1+i)^2 = 5000(1.09)^2 = \text{Rs. } 5940.50.$

and $A_3 = P(1+i)^3 = 5000(1.09)^3 = \text{Rs. } 6475.15.$

... and so on.

In general terms, if amount P is invested at $100i\%$ compound interest, then the amount accrued over n years is given by the following formula.

Accrued amount formula (simple interest) :

$$A_n = P(1+i)^n$$

Where: A_n =accrued amount at end of n -th year

P = Principal amount

i =(proportional) interest rate per year

n =number of years.

7.6 Notes on previous formula

a) The above formula can be transposed to make P the subject as follows:

$$P = \frac{A}{(1+i)^n}$$

So that, given an interest rate and a time period, a principal can be found if accrued amount is known.

For example, if some principal amount is invested at 12% ($i=0.12$) and amounts to Rs. 4917.25 ($=A$) after 3 years ($n=3$), then P (the principal amount) can be found using the above formula as:

$$P = \frac{4917.25}{1.12^3} = \frac{4917.25}{1.4049} = \text{Rs. } 3,500$$

b) The standard time period for the calculation of interest is usually a year. Hence, from this point on, the value of n (the number of time periods) will be assumed to be a number of years unless stated otherwise.

Example 1 (using accrued amount formula to solve simple problems)

a) What will be the value of Rs. 450 compounded at 12% for 3 years?

Here, $P = 450$, $i=12/100 = 0.12$ and $n=3$

Therefore $A = P(1+i)^n = 450(1+0.12)^3 = 450 (1.12)^3 = \text{Rs. } 632.22$

c) A principal amount accrues to Rs. 8500 if it is compounded at 14.5% over 6 years.

Find the value of this original amount.

Here it is necessary to use the inverted formula, since P needs to be found.

With $A = 8500$, $i=0.145$ and $n=6$ (and using the formula in section 11(a)):

$$\text{We have: } P = \frac{A}{(1+i)^n} = \frac{8500}{(1+0.145)^6} = \frac{8500}{(1.145)^6} = 3772.12$$

That is, Rs. 37772.12 needs to be invested (at 14.5% over 6 years) in order to accrue to Rs. 8500.

7.7 Formula for calculating APR

Formula to calculate APR :

Given a nominal annual rate of interest, the effective rate or **actual percentage rate** (APR) can be calculated as:

$$APR = (1+i/n)^n - 1$$

Where: i= given nominal rate (as a proportion)

n= number of equal compounding periods in one year.

Thus, for example

- a) 10% nominal, compounded quarterly, has $APR = (1.025)^4 - 1 = 0.1038 = 10.38\%$
- b) 24% nominal, compounded monthly, has $APR = (1.02)^{12} - 1 = 0.2682 = 26.82\%$

Example 7 (To calculate principal amount and APR)

A company will have to spend Rs. 300,000 on new plant in two years from now. Currently investment rates are at a nominal 10%.

- a) What single sum should now be invested, if compounding is six-monthly?
- b) What is the APR?

Answer

- a) Since compounding is six-monthly, the investment (P, say) must accrue to a value of Rs. 300,000 after four six-monthly periods. Note also that the interest rate for each six-month period is $(10/2)\% = 5\%$.

Using, the compounding (accrued amount) formula, $300000 = P(1+0.05)^4$

And re-arranging gives:
$$P = \frac{300000}{1.05^4} = 246810.75$$

That is, the amount to be invested is Rs. 246810.75

- b) Using the previous APR formula:

$$APR = (1+0.05)^2 - 1 = (1.05)^2 - 1 = 0.1025 = 10.25\%$$

7.8 Depreciation

Depreciation is an allowance made in estimates, valuations or balance sheets, normally for 'wear and tear'.

It is normal accounting practice to depreciate the values of certain assets. There are several different techniques available for calculating depreciation, two of which are:

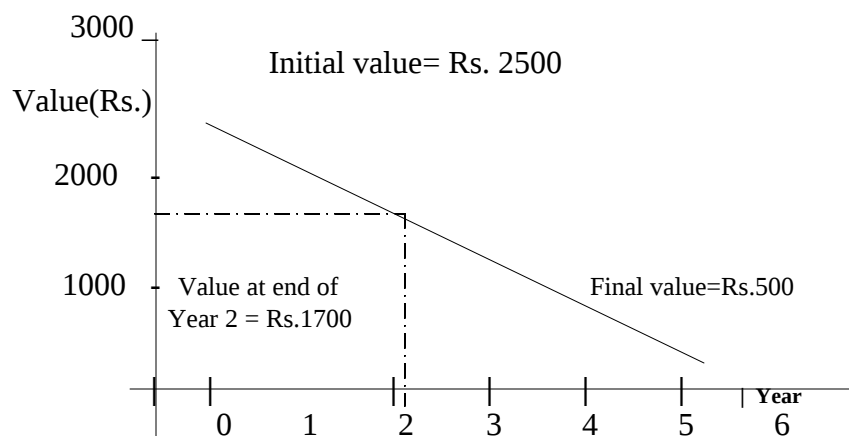
- a) Straight line (or Equal Instalment) depreciation and
- b) Reducing balance depreciation.

These two methods can be thought of as the converse of the interest techniques dealt with so far in the Lesson. That is, instead of adding value to some original principal amount (as with interest), value is taken away in order to reduce the original amount. Straight

line depreciation is the converse of simple interest with amounts being subtracted (rather than added), while the reducing balance method is the converse of compound interest.

Straight line depreciation of the value of a machine

Figure 2.1



7.8.1 Straight line depreciation

Given that the value of an asset must be depreciated, this technique involves subtracting the same amount from the original book value each year. Thus, for example, if the value of a machine is to depreciate from Rs. 2500 to Rs. 500 over a period of five years, then the annual depreciation would be Rs $(2500-500)/5 = \text{Rs. } 400$. The term 'straight line' comes from the fact that 2500 can be plotted against year 0 and 500 against year 5 on a graph, the two points joined by a straight line and then values for intermediate years can be read from the line. See Figure 2.1

7.8.2 Reducing balance depreciation

We already know that to increase some value P by 100%, we need to calculate the quantity $P(1+i)$, and if this is done n times, successively, the accrued value is given by $P(1+i)^n$ (from section 10).

A similar argument is applied if we need to decrease (or equivalently depreciate) some value B say by 100%. Here, we need to calculate $B(1-i)$. Notice that the multiplier is now $1-i$, rather than $1+i$. If the depreciation is carried out successively, n times, the accrued value will be given as $B(1-i)^n$.

For example, Rs. 2550 depreciated by 15% is $\text{Rs. } 2550(1-0.15) = \text{Rs. } 2550(0.85) = \text{Rs. } 2167.50$. Also, if Rs 2550 was successively depreciated over four time periods by 15%, the final depreciated value $= \text{Rs. } 2550(1-0.15)^4 = \text{Rs. } 2550(0.85)^4 = \text{Rs } 1331.12$.

Reducing balance depreciation is the name given to the technique of depreciating the book value of an asset by a constant percentage.

7.9 Formula for Reducing Balance Depreciation

Reducing balance depreciated value formula :

If book value B is subject to reducing balance depreciation at rate $100i\%$ over n time periods, the depreciated value at the end of the n -th time period is given by:

$$D = B(1-i)^n$$

Where: D = depreciated value at the end of the n -th time period

B = original book value

i = depreciation rate (as a proportion)

n = number of time periods (normally years).

Note that by re-arranging the above formula, any one of the variables B , i and n could be found, given the other three.

For example: $B = \frac{D}{(1-i)^n}$ (giving B in terms of D , i and n)

Thus if the depreciated value (D) of an asset was Rs. 5378.91 after three years depreciation at 25%, the original book value can be calculated as:

$$\begin{aligned} B &= \text{Rs. } \frac{5378.91}{0.75^3} \\ &= \text{Rs. } 12750. \end{aligned}$$

Also, since: $(1-i)^n = D/B$

Then: $1-i = \sqrt[n]{D/B}$ (giving $1-i$ in terms of D , V and n)

Example 8 (Reducing balance and straight line depreciation)

A mainframe computer whose cost is Rs. 220,000 will depreciate to a scrap value of Rs. 12000 in 5 years.

- If the reducing balance method of depreciation is used, find the depreciation rate.
- What is the book value of the computer at the end of the third year?
- How much more would the book value be at the end of the third year if the straight line method of depreciation had been used ?

Answer

Variables given are: $D=12000$, $B=220000$, $n=5$ and $D/B = \frac{12,000}{220,000} = 0.0545$

a) Now $1-i = \sqrt[5]{0.0545}$

Therefore, $1-i = 0.5589$
and so $i=0.4411$.

Thus the depreciation rate is 44.11%

b) The book value at the end of year 3 is given by: $B(1-i)^3 = 220000(0.5589)^3$
 $= \text{Rs. } 38408.$

c) Using the straight line method, the annual depreciation is: $\frac{\text{Rs. } (220000-12000)}{5}$

$$= \text{Rs. } 41600$$

Thus, after three years, the book value would be : $\text{Rs}[220000 - 3(41600)] = \text{Rs } 95200$.
So that the book value, using this method, would be $\text{Rs}[95200 - 38408] = \text{Rs. } 56792$ more than using the reducing balance method (at the end of the third year).

7.10 Let us Sum Up

Let us sum up the Lesson with the understanding of certain mathematical applications in Financial management. Usual applications involves calculations of different types of interest on capital, like simple interest, compound interest, accrued interest, calculation of depreciation, etc. With suitable examples and with simple illustration, the concept has been well presented in this Lesson for better understanding.

7.11 Lesson – End Activities

1. Mention the differences between simple interest and compound interest.
2. What is the need for giving depreciation?

7.12 References

1. Navaneetha. P – Business Mathematics
2. P.R. Vital – Business Mathematics and Statistics.

UNIT – III

Lesson 8 - Quantitative Techniques

Contents

- 8.1 Aims and Objectives**
- 8.2 Meaning of Quantitative Techniques**
- 8.3 Statistics**
- 8.4 Types of Statistical Data**
- 8.5 Classification of Statistical Methods**
- 8.6 Various Statistical Techniques**
- 8.7 Advantages of Quantitative Approach to Management**
- 8.8 Applications of Quantitative Techniques in Business and Management**
- 8.9 Let us Sum Up**
- 8.10 References**

8.1 Aims and Objectives

You may be aware of the fact that prior to the industrial revolution individual business was small and production was carried out on a very small scale mainly to cater to the local needs. The management of such business enterprises was very different from the present management of large scale business. The decisions were much less extensive than at present. Thus they used to make decisions based upon his past experience and intuition only. Some of the reasons for this were:

1. The marketing of the product was not a problem because customers were, for the large part, personally known to the owner of the business. There was hardly any competition in the business.
2. Test marketing of the product was not needed because the owner used to know the choice and requirement of the customers just by personal interaction.
3. The manager (also the owner) also used to work with his workers at the shopfloor. He knew all of them personally as the number was small. This reduced the need for keeping personal data.
4. The progress of the work was being made daily at the work centre itself. Thus production records were not needed.
5. Any facts the owner needed could be learnt direct from observation and most of what he required was known to him.

Now, in the face of increasing complexity in business and industry, intuition alone has no place in decision-making because basing a decision on intuition becomes highly questionable when the decision involves the choice among several courses of action each of which can achieve several management objectives simultaneously. Hence there is a need for training people who can manage a system both efficiently and creatively.

Quantitative techniques have made valuable contribution towards arriving at an effective decision in various functional areas of management-marketing, finance, production and

personnel. Today, these techniques are also widely used in regional planning, transportation, public health, communication, military, agriculture, etc.

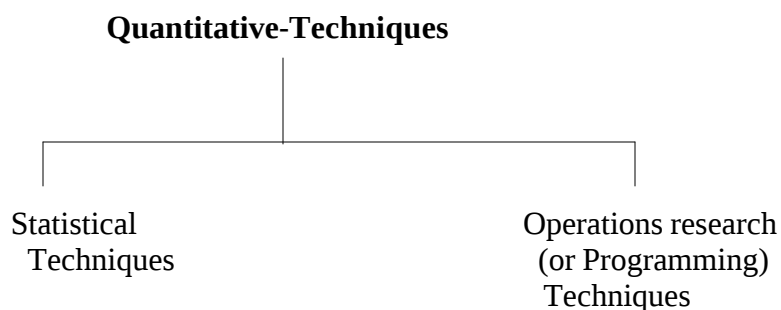
Quantitative techniques are being used extensively as an aid in business decision-making due to following reasons:

1. Complexity of today's managerial activities which involve constant analysis of existing situation, setting objectives, seeking alternatives, implementing, co-ordinating, controlling and evaluating the decision made.
2. Availability of different types of tools for quantitative analysis of complex managerial problems.
3. Availability of high speed computers to apply quantitative techniques (or models) to real life problems in all types of organisations such as business, industry, military, health, and so on. Computers have played an important role in arriving at the optimal solution of complex managerial problems.

In spite of these reasons, the quantitative approach, however, does not totally eliminate the scope of qualitative or judgment ability of the decision-maker. Of course these techniques complement the experience and knowledge of decision-maker in decision-making.

8.2 Meaning of Quantitative Techniques

Quantitative techniques refer to the group of statistical, and operations research (or programming) techniques as shown in the following chart.



The quantitative approach in decision-making requires that, problems be defined, analysed and solved in a conscious, rational, systematic and scientific manner based on data, facts, information, and logic and not on mere whims and guesses.

In other words, quantitative techniques (tools or methods) provide the decision – maker a scientific method based on quantitative data in identifying a course of action among the given list of courses of action to achieve the optimal value of the predetermined objective or goal. One common characteristic of all types of quantitative techniques is that numbers, symbols or mathematical formulae (or expressions) are used to represent the models of reality.

8.3 Statistics

Statistics

The word statistics can be used in a number of ways. Commonly it is described in two senses namely:

1. Plural Sense (Statistical Data)

The plural sense of statistics means some sort of statistical data. When it means statistical data, it refers to numerical description of quantitative aspects of things. These descriptions may take the form of counts or measurements. For example, statistics of students of a college include count of the number of students, and separate counts of number of various kinds as such, male and females, married and unmarried, or undergraduates and post-graduates. They may also include such measurements as their heights and weights.

2. Singular Sense (Statistical Methods)

The large volume of numerical information (or data) gives rise to the need for systematic methods which can be used to collect, organise or classify, present, analyse and interpret the information effectively for the purpose of making wise decisions. Statistical methods include all those devices of analysis and synthesis by means of which statistical data are systematically collected and used to explain or describe a given phenomena.

The above mentioned five functions of statistical methods are also called phases of a statistical investigation. Methods used in analysing the presented data are numerous and contain simple to sophisticated mathematical techniques. As an illustration, let us suppose that we are interested in knowing the income level of the people living in a certain city. For this we may adopt the following procedures:

- a) Data Collection: The following data is required for the given purpose:
 - Population of the city
 - Number of individuals who are getting income
 - Daily income of each earning individual
- b) Organise (or Condense) the data: the data so obtained should now be organised in different income groups. This will reduce the bulk of the data.
- c) Presentation: the organised data may now be presented by means of various types of graphs or other visual aids. Data presented in an orderly manner facilitates statistical analysis.
- d) Analysis: on the basis of systematic presentation (tabular form or graphical form) determine the average income of an individual and extent of disparities that exist. This information will help to get an understanding of the phenomenon (i.e. income of individuals.)

- e) Interpretation: All the above steps may now lead to drawing conclusions which will aid in decision-making-a policy decision for improvement of the existing situation.

Characteristics of data

It is probably more common to refer to data in quantitative form as statistical data.

It is probably more common to refer to data in quantitative form as statistical data. But not all numerical data is statistical. In order that numerical description may be called statistics they must possess the following characteristics:

- i) They must be aggregate of facts, for example, single unconnected figures cannot be used to study the characteristics of the phenomenon.
- ii) They should be affected to a marked extent by multiplicity of causes, for example, in social services the observations recorded are affected by a number of factors (controllable and uncontrollable)
- iii) They must be enumerated or estimated according to reasonable standard of accuracy, for example, in the measurement of height one may measure correct upto 0.01 of a cm; the quality of the product is estimated by certain tests on small samples drawn from a big lot of products.
- iv) They must have been collected in a systematic manner for a pre-determined purpose. Facts collected in a haphazard manner, and without a complete awareness of the object, will be confusing and cannot be made the basis of valid conclusions. For example collected data on price serve no purpose unless one knows whether he wants to collect data on wholesale or retail prices and what are the relevant commodities in view.
- v) They must be placed in relation to each other. That is, data collected should be comparable; otherwise these cannot be placed in relation to each other, e.g. statistics on the yield of crop and quality of soil are related but these yields cannot have any relation with the statistics on the health of the people.
- vi) They must be numerically expressed. That is, any facts to be called statistics must be numerically or quantitatively expressed. Qualitative characteristics such as beauty, intelligence, etc. cannot be included in statistics unless they are quantified.

8.4 Types of Statistical Data

An effective managerial decision concerning a problem on hand depends on the availability and reliability of statistical data. Statistical data can be broadly grouped into two categories:

- 1) Secondary (or published) data
- 2) Primary (or unpublished) data

The Secondary data are those which have already been collected by another organisation and are available in the published form. You must first check whether any such data is available on the subject matter of interest and make use of it, since it will save considerable time and money. But the data must be scrutinised properly since it was

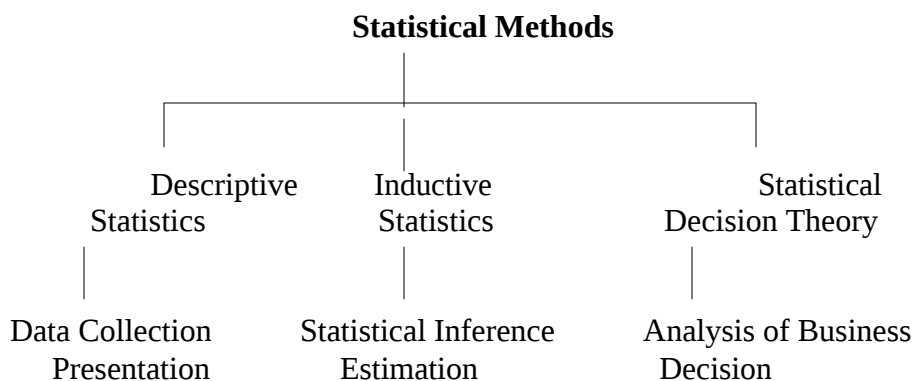
originally collected perhaps for another purpose. The data must also be checked for reliability, relevance and accuracy.

A great deal of data is regularly collected and disseminated by international bodies such as: World Bank, Asian Development Bank, International Labour Organisation, Secretariat of United Nations, etc., Government and its many agencies: Reserve Bank of India, Census Commission, Ministries-Ministry of Economics Affairs, Commerce Ministry; Private Research Organisations, Trade Associations etc.

When secondary data is not available or it is not reliable, you would need to collect original data to suit your objectives. Original data collected specifically for a current research are known as **primary data**. Primary data can be collected from customers, retailers, distributors, manufacturers or other information sources, primary data may be collected through any of the three methods: observation, survey, and experimentation. Data are also classified as micro and macro. Micro data relate to a particular unit region whereas macro data relate to the entire industry, region or economy.

8.5 Classification of Statistical Methods

The field of statistics provides the methods for collecting, presenting and meaningfully interpreting the given data. Statistical Methods broadly fall into three categories as shown in the following chart.



Descriptive Statistics

There are statistical methods which are used for re-arranging, grouping and summarising sets of data to obtain better information of facts and thereby better description of the situation that can be made. For example, changes in the price-index. Yield by wheat etc. are frequently illustrated using the different types of charts and graphs. These devices summarise large quantities of numerical data for easy understanding. Various types of averages, can also reduce a large mass of data to a single descriptive number. The

descriptive statistics include the methods of collection and presentation of data, measure of Central tendency and dispersion, trends, index numbers, etc.

Inductive Statistics

It is concerned with the development of some criteria which can be used to derive information about the nature of the members of entire groups (also called population or universe) from the nature of the small portion (also called sample) of the given group. The specific values of the population members are called 'parameters' and that of sample are called 'Statistics'. Thus, inductive statistics is concerned with estimating population parameters from the sample statistics and deriving a statistical inference.

Samples are drawn instead of a complete enumeration for the following reasons:

- i) the number of units in the population may not be known
- ii) the population units may be too many in number and/or widely dispersed. Thus complete enumeration is extremely time consuming and at the end of a full enumeration so much time is lost that the data becomes obsolete by that time.
- iii) It may be too expensive to include each population item.

Inductive statistics, includes the methods like: probability and probability distributions; sampling and sampling distribution; various methods of testing hypothesis; correlation, regression, factor analysis; time series analysis.

Statistical Decision Theory

Statistical decision theory deals with analysing complex business problems with alternative course of action (or strategies) and possible consequences. Basically, It is to provide more concrete information concerning these consequences, so that best course of action can be identified from alternative courses of action.

Statistical decision theory relies heavily not only upon the nature of the problem on hand, but also upon the decision environment. Basically there are four different states of decision environment as given below:

State of decision	Consequences
Certainty	Deterministic
Risk	Probabilistic
Uncertainty	Unknown
Conflict	Influenced by an opponent

Since statistical decision theory also uses probabilities (subjective or prior) in analysis, therefore it is also called a subjectivist approach. It is also known as Bayesian approach because Baye's theorem, is used to revise prior probabilities in the light of additional information.

8.6 Various Statistical Techniques

A brief comment on certain standard techniques of statistics which can be helpful to a decision-maker in solving problems is given below.

i) **Measures of Central Tendency:** Obviously for proper understanding of quantitative data, they should be classified and converted into a frequency distribution (number of times or frequency with which a particular data occurs in the given mass of data.). This type of condensation of data reduces their bulk and gives a clear picture of their structure. If you want to know any specific characteristics of the given data or if frequency distribution of one set of data is to be compared with another, then it is necessary that the frequency distribution help us to make useful inferences about the data and also provide yardstick for comparing different sets of data. Measures of average or central tendency provide one such yardstick. Different methods of measuring central tendency, provide us with different kinds of averages. The main three types of averages commonly used are:

a) **Mean:** the mean is the common arithmetic average. It is computed by dividing the sum of the values of the observations by the number of items observed.

b) **Median:** the median is that item which lies exactly half-way between the lowest and highest value when the data is arranged in an ascending or descending order. It is not affected by the value of the observation but by the number of observations. Suppose you have the data on monthly income of households in a particular area. The median value would give you that monthly income which divides the number of households into two equal parts. Fifty per cent of all the households have a monthly income above the median value and fifty per cent of households have a monthly income below the median income.

c) **Mode:** the mode is the central value (or item) that occurs most frequently. When the data organised as a frequency distribution the mode is that category which has the maximum number of observations. For example, a shopkeeper ordering fresh stock of shoes for the season would make use of the mode to determine the size which is most frequently sold. The advantages of mode are that (a) it is easy to compute, (b) is not affected by extreme values in the frequency distribution, and (c) is representative if the observations are clustered at one particular value or class.

ii) **Measures of Dispersion:** the measures of central tendency measure the most typical value around which most values in the distribution tend to coverage. However, there are always extreme values in each distribution. These extreme values indicate the spread or the dispersion of the distribution. The measures of this spread are called 'measures of dispersion' or 'variation' or 'spread'. Measures of dispersion would tell you the number of values which are substantially different from the mean, median or mode. The commonly used measures of dispersion are range, mean deviation and standard deviation. The data may spread around the central tendency in a symmetrical or an asymmetrical pattern. The measures of the direction and degree of symmetry are called measures of the skewness. Another characteristic of the frequency distribution is the shape of the peak, when it is plotted on a graph paper. The measures of the peakedness are called measures of Kurtosis.

iii) **Correlation:** Correlation coefficient measures the degree to which the change in one variable (the dependent variable) is associated with change in the other variable (independent one). For example, as a marketing manager, you would like to know if there is any relation between the amount of money you spend on advertising and the sales you achieve. Here, sales is the dependent variable and advertising budget is the independent variable. Correlation coefficient, in this case, would tell you the extent or relationship between these two variables, whether the relationship is directly proportional (i.e. increase or decrease in advertising is associated with decrease in sales) or it is an inverse relationship (i.e. increasing advertising is associated with decrease in sales and vice-versa) or there is no relationship between the two variables. However, it is important to note that correlation coefficient does not indicate a casual relationship, Sales is not a direct result of advertising alone, there are many other factors which affect sales. Correlation only indicates that there is some kind of association-whether it is casual or causal can be determined only after further investigation. You may find a correlation between the height of your salesmen and the sales, but obviously it is of no significance.

iv) **Regression Analysis:** For determining causal relationship between two variables you may use regression analysis. Using this technique you can predict the dependent variables on the basis of the independent variables. In 1970, NCAER (National Council of Applied and Economic Research) predicted the annual stock of scooters using a regression model in which real personal disposable income and relative weighted price index of scooters were used as independent variable.

The correlation and regression analysis are suitable techniques to find relationship between two variables only. But in reality you would rarely find a one-to-one causal relationship, rather you would find that the dependent variables are affected by a number of independent variables. For example, sales affected by the advertising budget, the media plan, the content of the advertisements, number of salesmen, price of the product, efficiency of the distribution network and a host of other variables. For determining causal relationship involving two or more variables, multi-variable statistical techniques are applicable. The most important of these are the multiple regression analysis, discriminant analysis and factor analysis.

v) **Time Series Analysis :** A time series consists of a set of data (arranged in some desired manner) recorded either at successive points in time or over successive periods of time. The changes in such type of data from time to time are considered as the resultant of the combined impact of a force that is constantly at work. This force has four components: (i) Editing time series data, (ii) secular trend, (iii) periodic changes, cyclical changes and seasonal variations, and (iv) irregular or random variations. With time series analysis, you can isolate and measure the separate effects of these forces on the variables. Examples of these changes can be seen, if you start measuring increase in cost of living, increase of population over a period of time, growth of agricultural food production in India over the last fifteen years, seasonal requirement of items, impact of floods, strikes, wars and so on.

vii) **Index Numbers:** Index number is a relative number that is used to represent the net result of change in a group of related variables that has some over a period of time. Index numbers are stated in the form of percentages. For example, if we say that the index of prices is 105, it means that prices have gone up by 5% as compared to a point of reference, called the base year. If the prices of the year 1985 are compared with those of 1975, the year 1985 would be called “given or current year” and the year 1975 would be termed as the “base year”. Index numbers are also used in comparing production, sales price, volume employment, etc. changes over period of time, relative to a base.

viii) **Sampling and Statistical Inference:** In many cases due to shortage of time, cost or non-availability of data, only limited part or section of the universe (or population) is examined to (i) get information about the universe as clearly and precisely as possible, and (ii) determine the reliability of the estimates. This small part or section selected from the universe is called the sample, and the process of selection such a section (or part) is called sampling.

Schemes of drawing samples from the population can be classified into two broad categories:

a) **Random sampling schemes:** In these schemes drawing of elements from the population is random and selection of an element is made in such a way that every element has equal chance (probability) of being selected.

b) **Non-random sampling schemes:** in these schemes, drawing of elements for the population is based on the choice or purpose of selector.

The sampling analysis through the use of various ‘tests’ namely Z-normal distribution, student’s ‘t’ distribution; F-distribution and χ^2 –distribution make possible to derive inferences about population parameters with specified level of significance and given degree of freedom.

8.7 Advantages of Quantitative Approach to Management

Executives at all levels in business and industry come across the problem of making decision at every stage in their day-to-day activities. Quantitative techniques provide the executive with scientific basis for decision-making and enhance his ability to make long-range plans and to solve every day problems of running a business and industry with greater efficiency and confidence.

Some of the advantages of the study of statistics are:

1. **Definiteness:** the study of statistics helps us in presenting general statements in a precise and a definite form. Statements of facts conveyed numerically are more precise and convincing than those stated qualitatively. For example, the statement that “literacy rate as per 1981 census was 36% compared to 29% for 1971 census” is more convincing than stating simply that “literacy in our country has increased”.

2. **Condensation:** The new data is often unwieldy and complex. The purpose of statistical methods is to simplify large mass of data and to present a meaningful information from them. For example, it is difficult to form a precise idea about the income position of the people of India from the data of individual income in the country. The data will be easy to understand and more precisely if it can be expressed in the form of per capita income.
3. **Comparison:** According to Bodding, the object of statistics is to enable comparisons between past and present results with a view to ascertaining the reasons for change which have taken place and the effect of such changes in the future. Thus, if one wants to appreciate the significance of figures, then he must compare them with other of the same kind. For example, the statement “per capita income has increased considerably” shall not be meaningful unless some comparison of figures of past is made. This will help in drawing conclusions as to whether the standard of living of people of India is improving.
4. **Formulation of policies:** Statistics provides that basic material for framing policies not only in business but in other fields also. For example, data on birth and mortality rate not only help in assessing future growth in population but also provide necessary data for framing a scheme of family planning.
5. **Formulating and testing hypothesis:** statistical methods are useful in formulating and testing hypothesis or assumption or statement and to develop new theories. For example, the hypothesis: “whether a student has benefited from a particular media of instruction”, can be tested by using appropriate statistical method.
6. **Prediction:** For framing suitable policies or plans, and then for implementation it is necessary to have the knowledge of future trends. Statistical methods are highly useful for forecasting future events. For example, for a businessman to decide how many units of an item should be produced in the current year, it is necessary for him to analyse the sales data of the past years.

8.8 Applications of Quantitative Techniques in Business and Management

Some of the areas where statistics can be used are as follows:

Management

i) Marketing:

- Analysis of marketing research information
- Statistical records for building and maintaining an extensive market
- Sales forecasting

ii) Production

- Production planning, control and analysis
- Evaluation of machine performance
- Quality control requirements
- Inventory control measures
-

iii) Finance, Accounting and Investment:

- Financial forecast, budget preparation
- Financial investment decision
- Selection of securities
- Auditing function
- Credit, policies, credit risk and delinquent accounts

iv) Personnel:

- Labour turn over rate
- Employment trends
- Performance appraisal
- Wage rates and incentive plans

Economics

- Measurement of gross national product and input-output analysis
- Determination of business cycle, long-term growth and seasonal fluctuations
- Comparison of market prices, cost and profits of individual firms
- Analysis of population, land economics and economic geography
- Operational studies of public utilities
- Formulation of appropriate economic policies and evaluation of their effect

Research and Development

- Development of new product lines
- Optimal use of resources
- Evaluation of existing products

Natural Science

- Diagnosing the disease based on data like temperature, pulse rate, blood pressure etc.
- Judging the efficacy of particular drug for curing a certain disease
- Study of plant life

Exercises

1. Comment on the following statements:
 - a) “Statistics are numerical statement of facts but all facts numerically stated are not statistics”
 - b) “Statistics is the science of averages”.
2. What is the type of the following models?
 - a) Frequency curves in statistics.
 - b) Motion films.
 - c) Flow chart in production control, and
 - c) Family of equations describing the structure of an atom.

3. List at least two applications of statistics in each, functional area of management.
4. What factors in modern society contribute to the increasing importance of quantitative approach to management?
5. Describe the major phases of statistics. Formulate a business problem and analyse it by applying these phases.
6. Explain the distinction between:
 - a) Static and dynamic models
 - b) Analytical and simulation models
 - c) Descriptive and prescriptive models.
7. Describe the main features of the quantitative approach to management.

8.9 Let us Sum Up

We have so far learned the quantitative techniques and quantitative approach to management with its characteristics.

8.10 Lesson – End Activities

1. What are the different types of statistical data available.
2. Mention the advantages of quantitative approach to management.

8.11 References

1. Gupta. S.P. – Statistical Methods.

Lesson 9 - Presentation of Data

Contents

- 9.1 Aims and Objectives**
- 9.2 Classification of Data**
- 9.3 Objectives of Classification**
- 9.4 Types of Classification**
- 9.5 Construction of a Discrete Frequency Distribution**
- 9.6 Construction of a Continuous Frequency Distribution**
- 9.7 Guidelines for Choosing the Classes**
- 9.8 Cumulative and Relative Frequencies**
- 9.9 Charting of Data**
- 9.10 Let us Sum Up**
- 9.11 Lesson – End Activities**
- 9.12 References**

9.1 Aims and Objectives

The successful use of the data collected depends to a great extent upon the manner in which it is arranged, displayed and summarized. This Lesson mainly deals with the presentation of data. Presentation of data can be displayed either in tabular form or through charts. In the tabular form, it is necessary to classify the data before the data tabulated. Therefore, this unit is divided into two section, viz., (a) classification of data and (b) charting of data.

9.2 Classification of Data

After the data has been systematically collected and edited, the first step in presentation of data is classification. Classification is the process of arranging the data according to the points of similarities and dissimilarities. It is like the process of sorting the mail in a post office where the mail for different destinations is placed in different compartments after it has been carefully sorted out from the huge heap.

9.3 Objectives of Classification

The principal objectives of classifying data are:

- i) to condense the mass of data in such a way that salient features can be readily noticed
- ii) to facilitate comparisons between attributes of variables
- iii) to prepare data which can be presented in tabular form
- iv) to highlight the significant features of the data at a glance

9.4 Types of Classification

Some common types of classification are:

- Geographical i.e., according to area or region
- Chronological, i.e., according to occurrence of an event in time.
- Qualitative, i.e., according to attributes.
- Quantitative, i.e., according to magnitudes.

Geographical Classification: In this type of classification, data is classified according to area or region. For example, when we consider production of wheat State wise, this would be called geographical classification. The listing of individual entries are generally done in an alphabetical order or according to size to emphasise the importance of a particular area or region.

Chronological Classification: when the data is classified according to the time of the occurrence, it is known as chronological classification. For example, sales figure of a company for last six years are given below:

Year	Sales (Rs. Lakhs)	Year	Sales (Rs. Lakhs)
1982-83	175	1985-86	485
1983-84	220	1986-87	565
1984-85	350	1987-88	620

Qualitative Classification: When the data is classified according to some attributes (distinct categories) which are not capable of measurement is known as qualitative classification. In a simple (or dichotomous) classification, an attribute is divided into two classes, one possessing the attribute and the other not possessing it. For example, we may classify population on the basis of employment, i.e., the employed and the unemployed. Similarly we can have manifold classification when an attribute is divided so as to form several classes. For example, the attribute education can have different classes such as primary, middle, higher secondary, university, etc.

Quantitative Classification: when the data is classified according to some characteristics that can be measured, it is called quantitative classification. For example, the employees of a company may be classified according to their monthly salaries. Since quantitative data is characterized by different numerical values, the data represents the values of a variable. Quantitative data may be further classified into one or two types: discrete or continuous. The term discrete data refers to quantitative data that is limited to certain numerical values of a variable. For example, the number of employees in an organisation or the number of machines in a factory are examples of discrete data.

Continuous data can take all values of the variable. For example, the data relating to weight, distance, and volume are examples of continuous data. The quantitative classification becomes the basis for frequency distribution.

When the data is arranged into groups or categories according to conveniently established divisions of the range of the observations, such an arrangement in tabular form is called a frequency distribution. In a frequency distribution, raw data is represented by distinct groups which are known as classes. The number of observations that fall into each of the classes is known as frequency. Thus, a frequency distribution has two parts, on its left there are classes and on its right are frequencies.

When data is described by a continuous variable it is called continuous data and when it is described by a discrete variables, it is called discrete data. The following are the two examples of discrete and continuous frequency distributions.

No.of Employees	No.of companies	Age (years)	No.of workers
110	25	20-25	15
120	35	25-30	22
130	70	30-35	38
140	100	35-40	47
150	18	40-45	18
160	12	45-50	10
Discrete frequency distribution		Continuous frequency distribution	

9.5 Construction of a Discrete Frequency Distribution

The process of preparing a frequency distribution is very simple. In the case of discrete data, place all possible values of the variable in ascending order in one column, and then prepare another column of 'Tally' mark to count the number of times a particular value of the variable is repeated. To facilitate counting, block of five 'Tally' marks are prepared and some space is left in between the blocks. The frequency column refers to the number of 'Tally' marks, a particular class will contain. To illustrate the construction of a discrete frequency distribution, consider a sample study in which 50 families were surveyed to find the number of children per family. The data obtained are:

3 2 2 1 3 4 2 1 3 4 5 0 2
 1 2 3 3 2 1 1 2 3 0 3 2 1
 4 3 5 5 4 3 6 5 4 3 1 0 6
 4 3 1 2 0 1 2 3 4 5

To condense this data into a discrete frequency distribution, we shall take the help of 'Tally' marks as shown below:

No. of Children	No. of families	Frequency
0	III	4
1	IIII III	9
2	IIII IIII	10
3	IIII IIII II	12
4	IIII II	7
5	IIII I	6
6	II	2
		Total 50

9.6 Construction of a Continuous Frequency Distribution

In constructing the frequency distribution for continuous data, it is necessary to clarify some of the important terms that are frequently used.

Class Limits: Class limits denote the lowest and highest value that can be included in the class. The two boundaries (i.e., lowest and highest) of a class are known as the lower limit and the upper limit of the class. For example, in the class 60-69, 60 is the lower limit and 69 is the upper limit or we can say that there can be no value in that class which is less than 60 and more than 69.

Class Intervals: The class interval represents the width (span or size) of a class. The width may be determined by subtracting the lower limit of one class from the lower limit of the following class (alternatively successive upper limits may be used). For example, if the two classes are 10-20 and 20-30, the width of the class interval would be the difference between the two successive lower limits of the same class, i.e., $20-10=10$.

Class Frequency: The number of observations falling within a particular class is called its class frequency or simply frequency. Total frequency (sum of all the frequencies) indicates the total number of observations considered in a given frequency distribution.

Class Mid-point: Mid-point of a class is defined as the sum of two successive lower limits divided by two. Therefore, it is the value lying halfway between the lower and upper class limits. In the example taken above the mid-point would be $(10+20)/2=15$ corresponding to the class 10-20 and 25 corresponding to the class 20-30.

Types of Class Interval: There are different ways in which limits of class intervals can be shown such as:

- i) Exclusive and Inclusive method, and
- ii) Open-end

Exclusive Method: The class intervals are so arranged that the upper limit of one class is the lower limit of the next class. The following example illustrates this point.

Sales (Rs. Thousands)	No. of firms	Sales (Rs. Thousands)	No. of firms
20-25	20	35-40	27
25-30	28	40-45	12
30-35	35	45-50	8

In the above example there are 20 firms whose sales are between Rs. 20,000 and Rs. 24,999. A firm with sales of exactly Rs. 25 thousand would be included in the next class viz. 25-30. Therefore in the exclusive method, it is always presumed that upper limit is excluded.

Inclusive Method: In this method, the upper limit of one class is included in that class itself. The following example illustrate this point.

Sales (Rs. Thousands)	No.of firms	Sales (Rs. Thousands)	No.of firms
20-24.999	20	35-39.999	27
25-29.999	28	40-44.999	12
30-34.999	35	45-49.999	8

In this example, there are 20 firms whose sales are between Rs. 20,000 and Rs. 24,999. A firm whose sales are exactly Rs. 25,000 would be included in the next class. Therefore in the inclusive method, it is presumed that upper limit is included.

It may be observed that both the methods give the same class frequencies, although the class intervals look different. Whenever inclusive method is used for equal class intervals, the width of class intervals can be obtained by taking the difference between the two lower limits (or upper limits).

Open-End: In an open-end distribution, the lower limit of the very first class and upper limit of the last class is not given. In distribution where there is a big gap between minimum and maximum values, the open-end distribution can be used such as in income distributions. The income disparities, of residents of a region may vary between Rs. 800 to Rs. 50,000 per month. In such a case, we can form classes like: Less than Rs. 1,000
1,000 - 2,000
2,000 - 5,000
5,000 - 10,000
10,000 - 25,000
25,000 and above

Remark: To ensure continuity and to get correct class intervals, we shall adopt exclusive method. However, if inclusive method is suggested then it is necessary to make an

adjustment to determine the class interval. This can be done by taking the average value of the difference between the lower limit of the succeeding class and the upper limit of the class. In terms of formula:

$$\text{Correction factor} = \frac{\text{Lower Limit of second class} - \text{Upper Limit of the first class}}{2}$$

This value so obtained is deducted from all lower limits and added to all upper limits. For instance, the example discussed for inclusive method can easily be converted into exclusive case. Take the difference between 25 and 24,999 and divide it by 2. Thus correction factor becomes $(25 - 24,999)/2 = 0.0005$. Deduct this value from lower limits and add it to upper limits.

The new frequency distribution will take the following.

Sales (Rs. Thousands)	No.of firms	Sales (Rs. Thousands)	No.of firms
19.9995-24.9995	20	34.9995-39.9995	27
24.9995-29.9995	28	39.9995-44.9995	12
29.9995-34.9995	35	44.9995-49.9995	8

9.7 Guidelines for Choosing the Classes

The following guidelines are useful in choosing the class intervals.

1. The number of classes should not be too small or too large. Preferably, the number of classes should be between 5 and 15. However, there is no hard and fast rule about it. If the number of observations is smaller, the number of classes formed should be towards the lower side of this towards the upper side of the limit.
2. If possible, the widths of the intervals should be numerically simple like 5,10,25 etc. Values like 3,7,19 etc. should be avoided.
3. It is desirable to have classes of equal width. However, in case of distributions having wide gap between the minimum and maximum values, classes with unequal class interval can be formed like income distribution.
4. The starting point of a class should begin with 0,5,10 or multiplies thereof. For example, if the minimum value is 3 and we are taking a class interval of 10, the first class should be 0-10 and not 3-13.
5. The class interval should be determined after taking into consideration the minimum and maximum values and the number of classes to be formed. For example, if the income of 20 employees in a company varies between Rs. 1100 and Rs.5900 and we want to form 5 classes, the class interval should be 1000

$$\frac{(5900-1100)}{1000} = 4.8 \text{ or } 5.$$

All the above points can be explained with the help of the following example wherein the ages of 50 employees are given:

22 21 37 33 28 42 56 33 32 59
 40 47 29 65 45 48 55 43 42 40
 37 39 56 54 38 49 60 37 28 27
 32 33 47 36 35 42 43 55 53 48
 29 30 32 37 43 54 55 47 38 62

In order to form the frequency distribution of this data, we take the difference between 60 and 21 and divide it by 10 to form 5 classes as follows:

Age(Years)	Tally Marks	Frequency
20-30	IIII II	7
30-40	IIII IIII IIII I	16
40-50	IIII IIII IIII	15
50-60	IIII IIII	9
60-70	III	3
		Total 50

9.8 Cumulative and Relative Frequencies

It is often useful to express class frequencies in different ways. Rather than listing the actual frequency opposite each class, it may be appropriate to list either cumulative frequencies or relative frequencies or both.

Cumulative Frequencies: As its name indicates, it cumulates the frequencies, starting at either the lower or highest value. The cumulative frequency of a given class interval thus represents the total of all the previous class frequencies including the class against which it is written. To illustrate the concept of cumulative frequencies consider the following example

Monthly salary (Rs.)	No.of employees	Monthly Salary (Rs.)	No.of employees
1000-1200	5	2000-2200	25
1200-1400	14	2200-2400	22
1400-1600	23	2400-2600	7
1600-1800	50	2600-2800	2
1800-2000	52		

If we keep on adding the successive frequency of each class starting from the frequency of the very first class, we shall get cumulative frequencies as shown below:

Monthly Salary(Rs.)	No. of employees	Cumulative frequency
1000-1200	5	5
1200-1400	14	19
1400-1600	23	42
1600-1800	50	92
1800-2000	52	144
2000-2200	25	169
2200-2400	22	191
2400-2600	7	198
2600-2800	2	200
Total 200		

Relative Frequencies: Very often, the frequencies in a frequency distribution are converted to relative frequencies to show the percentage for each class. If the frequency of each class is divided by the total number of observations (total frequency), then this proportion is referred to as relative frequency. To get the percentage of each class, multiply the relative frequency by 100. For the above example, the values computed for relative for relative frequency and percentage are shown below:

Monthly Salary (Rs.)	No. of employees	Relative frequency	percentage
1000-1200	5	0.025	2.5
1200-1400	14	0.070	7.0
1400-1600	23	0.115	11.5
1600-1800	50	0.250	25.0
1800-2000	52	0.260	26.0
2000-2200	25	0.125	12.5
2200 -2400	22	0.110	11.0
2400-2600	7	0.035	3.5
2600-2800	2	0.010	1.0
	200	1.000	100%

There are two important advantages in looking at relative frequencies (percentages) instead of absolute frequencies in a frequency distribution.

1. Relative frequencies facilitate the comparisons of two or more than two sets of data.
2. Relative frequencies constitute the basis of understanding the concept of probability.

9.9 Charting of Data

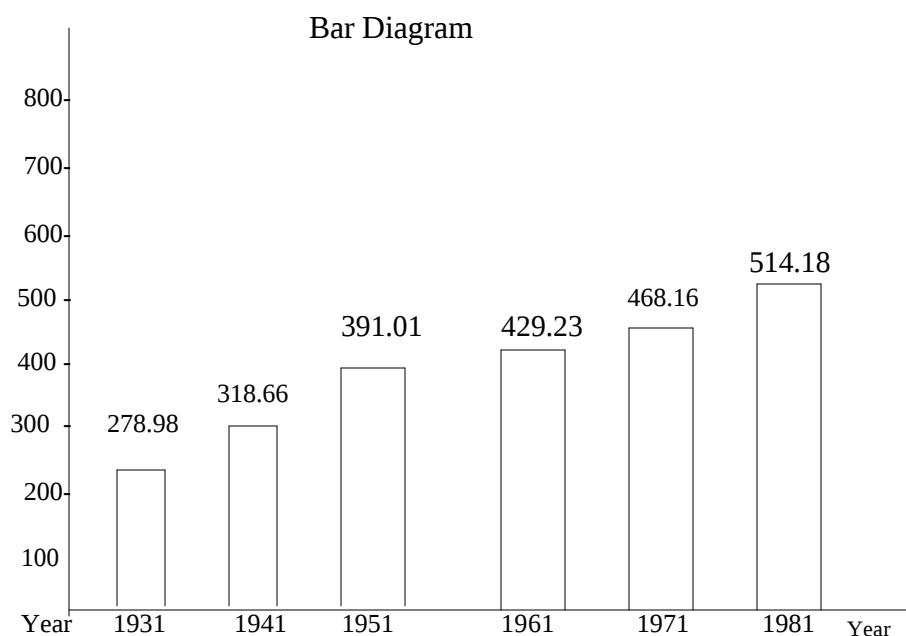
Charts of frequency distributions which cover both diagrams and graphs are useful because they enable a quick interpretation of the data. A frequency distribution can be presented by a variety of methods. In this section, the following four popular methods of charting frequency distribution are discussed in detail.

- i) Bar Diagram
- ii) Histogram
- iii) Frequency Polygon
- iv) Ogive or Cumulative Frequency Curve

Bar Diagram: Bar diagrams are most popular. One can see numerous such diagrams in newspapers, journals, exhibitions, and even on television to depict different characteristics of data. For example, population, per capita income, sales and profits of a company can be shown easily through bar diagrams. It may be noted that a bar is thick line whose width is shown to attract the viewer. A bar diagram may be either vertical or horizontal.

In order to draw a bar diagram, we take the characteristic (or attribute) under consideration on the X-axis and the corresponding value on the Y-axis. It is desirable to mention the value depicted by the bar on the top of the bar.

To explain the procedure of drawing a bar diagram, we have taken the population figures (in millions) of India which are given below:



Take the years on the X-axis and the population figure on the Y-axis and draw a bar to show the population figure for the particular year. This is shown above:

As can be seen from the diagram, the gap between one bar and the other bar is kept equal. Also the width of different bars is same. The only difference is in the length of the bars and that is why this type of diagram is also known as one dimensional.

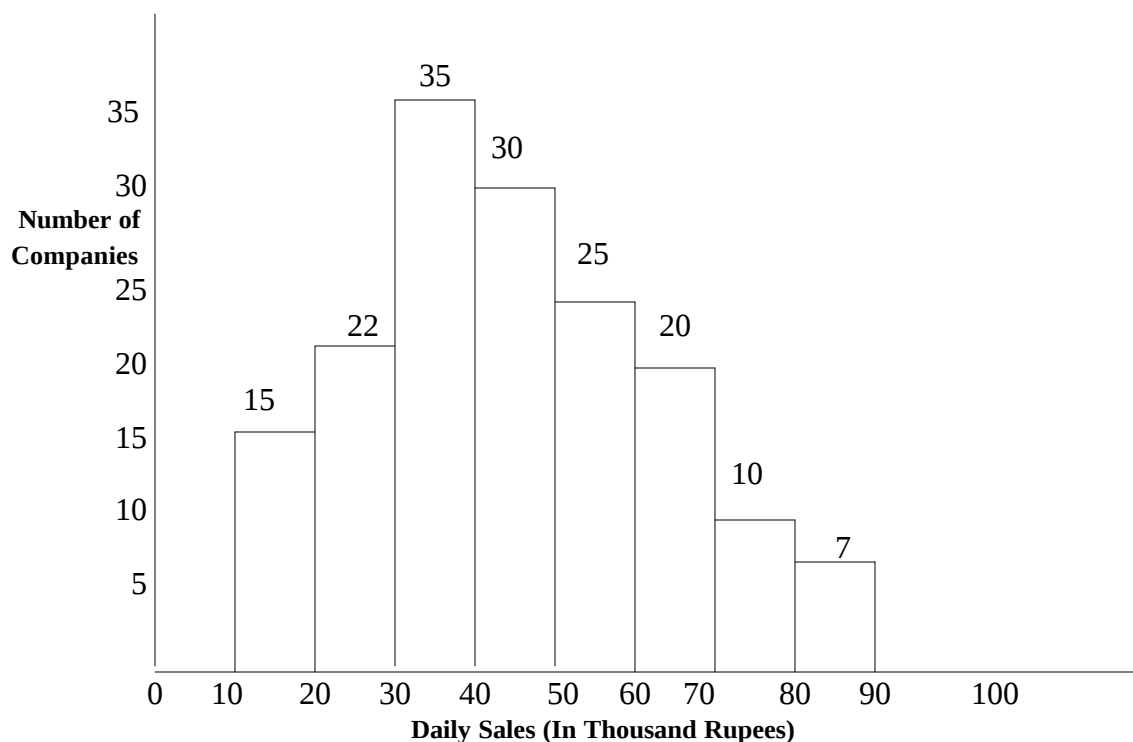
Histogram: One of the most commonly used and easily understood methods for graphic presentation of frequency distribution is histogram. A histogram is a series of rectangles having areas that are in the same proportion as the frequencies of a frequency distribution.

To construct a histogram, on the horizontal axis or X-axis, we take the class limits of the variable and on the vertical axis or Y-axis, we take the frequencies of the class intervals shown on the horizontal axis. If the class intervals are of equal width, then the vertical bars in the histogram are also of equal width. On the other hand, if the class intervals are unequal, then the frequencies have to be adjusted according to the width of the class interval. To illustrate a histogram when class intervals are equal, let us consider the following example.

Daily Sales (Rs. Thousand)	No. of companies	Daily Sales (Rs. Thousand)	No. of companies
10-20	15	50-60	25
20-30	22	60-70	20
30-40	35	70-80	16
40-50	30	80-90	7

In this example, we may observe that class intervals are of equal width. Let us take class intervals on the X-axis and their corresponding frequencies on the Y-axis. On each class interval (as base), erect a rectangle with height equal to the frequency of that class. In this manner we get a series of rectangles each having a class interval as its width and the frequency as its height as shown below :

Histogram with Equal Class Intervals



It should be noted that the area of the histogram represents the total frequency as distributed throughout the different classes.

When the width of the class intervals are not equal, then the frequencies must be adjusted before constructing the histogram.

The following example will illustrate the procedure

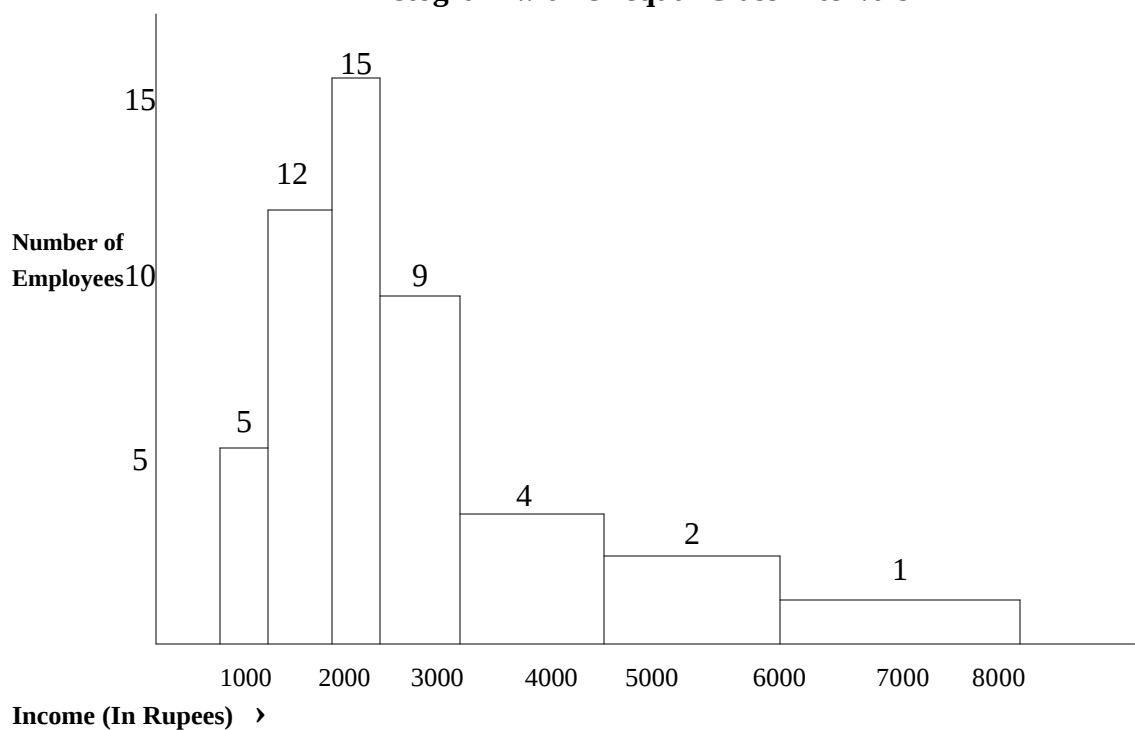
Income (Rs.)	No. of employees	Income(Rs.)	No. of employees
1000-1500	5	3500-5000	12
1500-2000	12	5000-7000	8
2000-2500	15	7000-8000	2
2500-3500	18		

As can be seen, in the above example, the class intervals are of unequal width and hence we have to find out the adjusted frequency of each class by taking the class with the lowest class interval as the basis of adjustment. For example, in the class 2500-3500, the class interval is 1000 which is twice the size of the lowest class interval, i.e., 500 and therefore the frequency of this class would be divided by two, i.e., it would be $18/2=9$. In a similar manner, the other frequencies would be obtained. The adjusted frequencies for various classes are given below:

Income (Rs.)	No. of employees	Income(Rs.)	No. of employees
1000-1500	5	3500-5000	4
1500-2000	12	5000-7000	2
2000-2500	15	7000-8000	1
2500-3500	18		

The histogram of the above distribution is shown below:

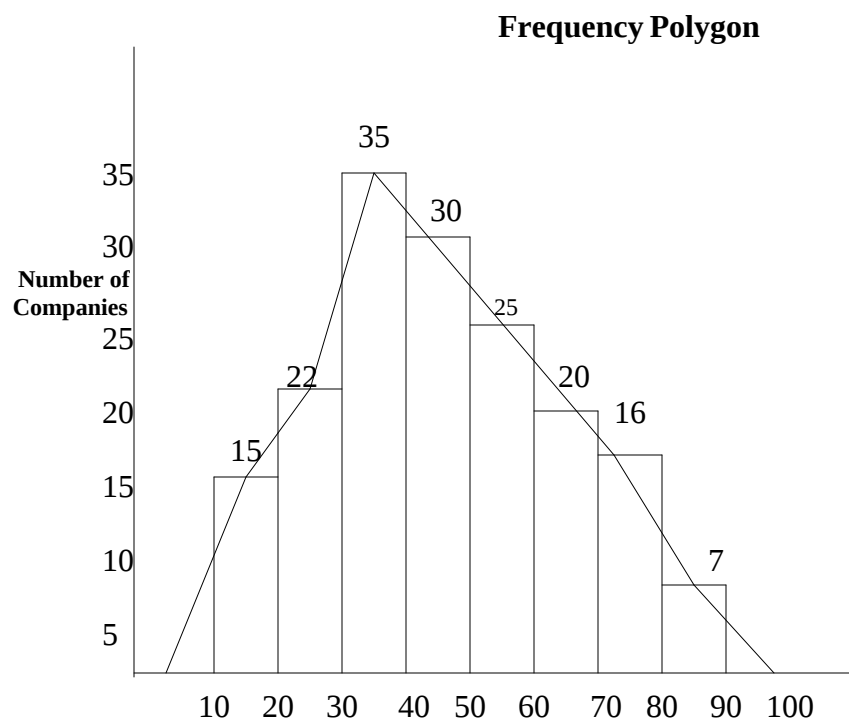
Histogram with Unequal Class Intervals



Income (In Rupees) >

It may be noted that a histogram and a bar diagram look very much alike but have distinct features. For example, in a histogram, the rectangles are adjoining and can be of different width whereas in bar diagram it is not possible.

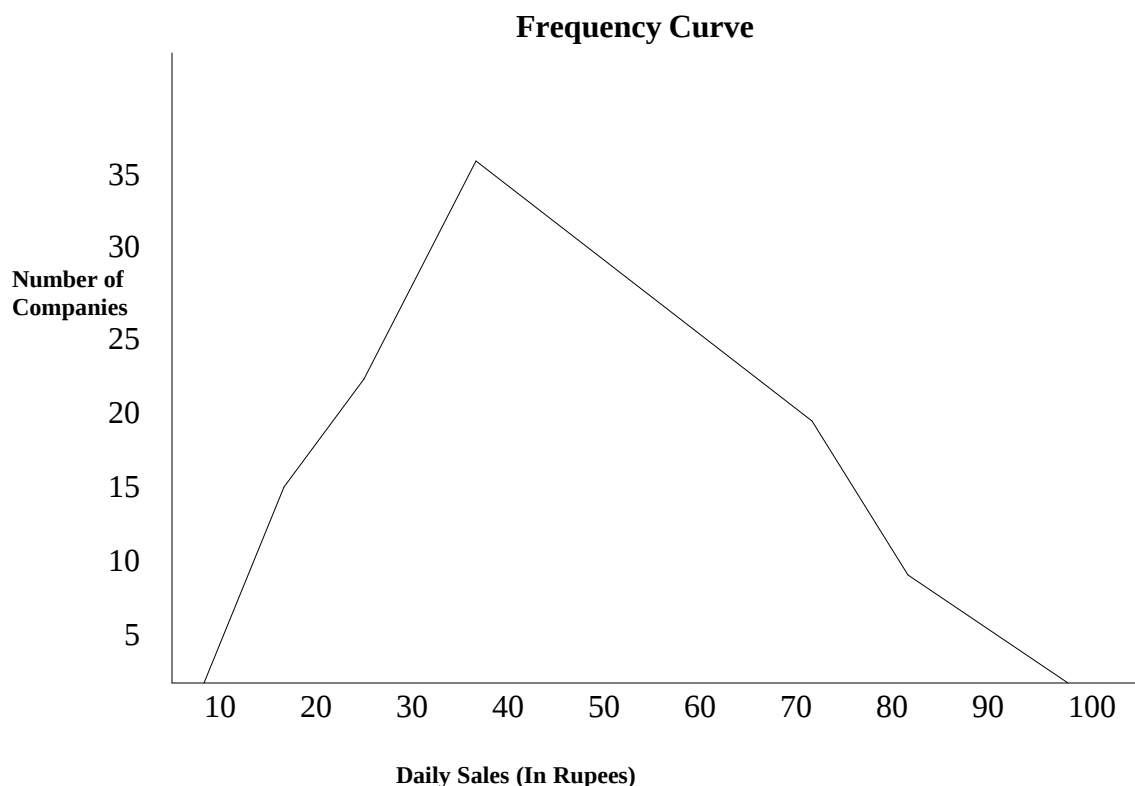
Frequency Polygon: The frequency polygon is a graphical presentation of frequency distribution. A polygon is a many sided figure. A frequency polygon is



Daily Sales (In Rupees)

Constructed by taking the mid-points of the upper horizontal side of each rectangle on the histogram and connecting these mid-points by straight lines. In order to close the polygon, an additional class is assumed at each end, having a zero frequency. To illustrate the frequency polygon of this distribution is shown above.

If we draw a smooth curve over these points in such a way that the area included under the curve is approximately the same as that of the polygon, then such a curve is known as frequency curve. The following figure shows the same data smoothed out to form a frequency curve, which is another form of presenting the same data.



Remark: The histogram is usually associated with discrete data and a frequency polygon is appropriate for continuous data. But this distinction is not always followed in practice and many factors may influence the choice of graph.

The frequency polygon and frequency curve have a special advantage over the histogram particularly when we want to compare two or more frequency distributions.

Ogives or Cumulative frequency Curve: An ogive is the graphical presentation of a cumulative frequency distribution and therefore when the graph of such a distribution is drawn, it is called cumulative frequency curve or ogive. There are two methods of constructing ogive, viz.,

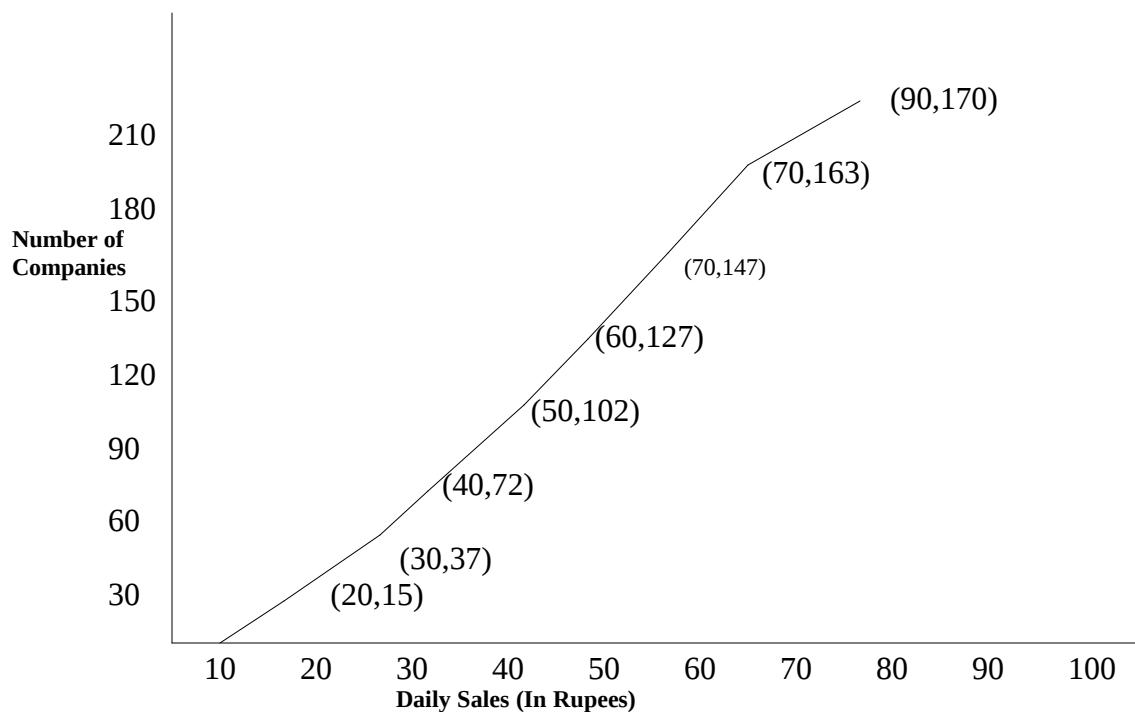
- i) Less than ogive
- ii) More than ogive

Less than Ogive: In this method, the upper limit of the various classes are taken on the X-axis and the frequencies obtained by the process of cumulating the preceding frequencies on the Y-axis. By joining these points we get less than ogive. Consider the example relating to daily sales discussed earlier.

Daily sales (Rs. Thousand)	No. of companies	Daily sales (Rs. Thousand)	No. of Companies
10-20	15	Less than 20	15

20-30	22	Less than 30	37
30-40	35	Less than 40	72
40-50	30	Less than 50	102
50-60	25	Less than 60	127
60-70	20	Less than 70	147
70-80	16	Less than 80	163
80-90	7	Less than 90	170

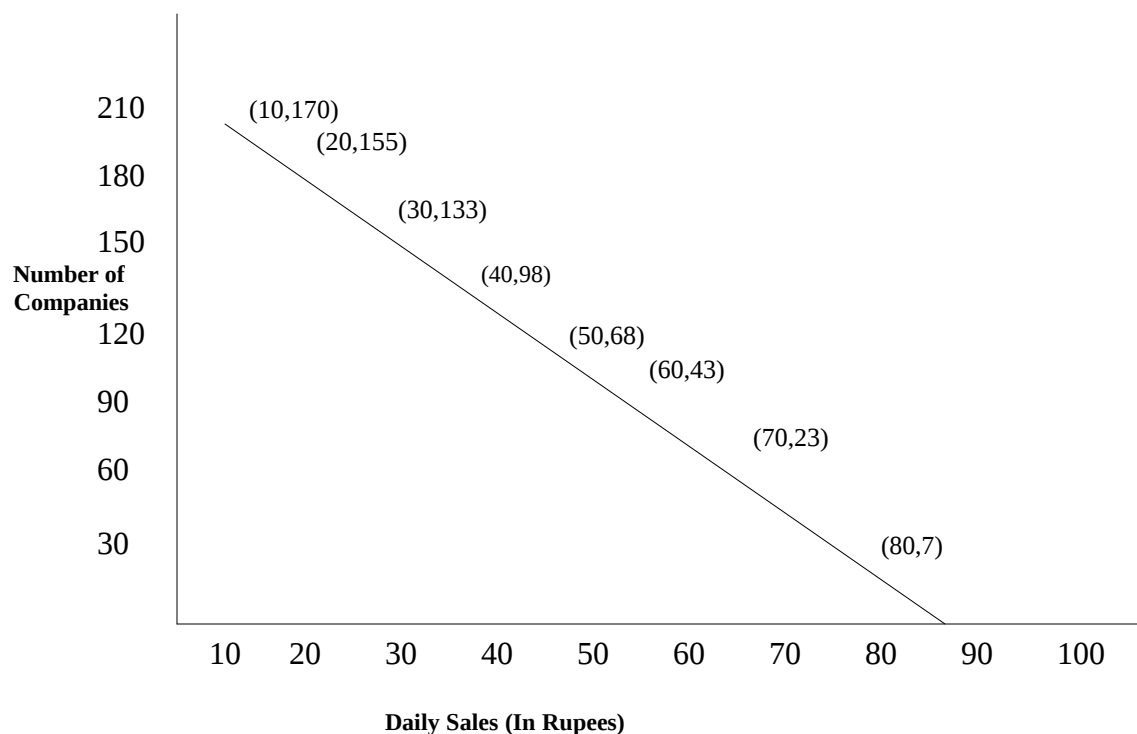
The less than **Ogive Curve** is shown below:



More than Ogive: Similarly more than ogive or cumulative frequency curve can be drawn by taking the lower limits on X-axis and cumulative frequencies on the Y-axis. By joining these points, we get more than ogive. The table and the curve for this case is shown below:

Daily sales (Rs. Thousand)	No. of companies	Daily sales (Rs. Thousand)	No. of Companies
10-20	15	More than 10	170
20-30	22	More than 20	155
30-40	35	More than 30	133
40-50	30	More than 40	98
50-60	25	More than 50	68
60-70	20	More than 60	43
70-80	16	More than 70	23
80-90	7	More than 80	7

The more than ogive curve is shown below:



The shape of less than ogive curve would be a rising one whereas the shape of more than ogive curve should be falling one.

The concept of ogive is useful in answering questions such as : How many companies are having sales less than Rs. 52,000 per day or more than Rs. 24,000 per day or between Rs. 24,000 and Rs. 52,000?

Exercises

1. Explain the purpose and methods of classification of data giving suitable examples.
2. What are the general guidelines of forming a frequency distribution with particular reference to the choice of class intervals and number of classes?
3. Explain the various diagrams and graphs that can be used for charting a frequency distribution.
4. What are ogives? Point out the role. Discuss the method of constructing ogives with the help of an example.
5. The following data relate to the number of family members in 30 families of a village.

4 3 2 3 4 5 5 7 3 2
 3 4 2 1 1 6 3 4 5 4
 2 7 3 4 5 6 2 1 5 3

Classify the above data in the form of a discrete frequency distribution.

6. The profits (Rs. Lakhs) of 50 companies are given below:

20 12 15 27 28 40 42 35 37 43
 55 65 53 62 29 64 69 36 25 18
 56 55 43 35 26 21 48 43 50 67
 14 23 34 59 68 22 41 42 43 52
 60 26 26 37 49 53 40 20 18 17

Classify the above data taking first class as 10-20 and form a frequency distribution.

7. The income(Rs.) of 24 employees of a company are given below:

1800 1250 1760 3500 6000 2500
 2700 3600 3850 6600 3000 1500
 4500 4400 3700 1900 1850 3750
 6500 6800 5300 2700 4370 3300

Form a continuous frequency distribution after selecting a suitable class interval.

8. Draw a histogram and a frequency polygon from the following data:

Marks	No. of students	Marks	No. of students
0-20	8	60-80	12
20-40	12	80-100	3
40-60	15		

9. Go through the following data carefully and then construct a histogram.

Income (Rs.)	No. of Persons	Income (Rs.)	No. of persons
500-1000	18	3000-4500	22
1000-1500	20	4500-5000	12
1500-2500	30	5000-7000	5
2500-3000	25		

10. The following data relating to sales of 100 companies is given below:

Sales (Rs. Lakhs)	No. of companies	Sales (Rs. Lakhs)	No. of companies
5-10	5	25-30	18
10-15	12	30-35	15
15-20	13	35-40	10
20-25	20	40-45	7

Draw less than and more than ogives. Determine the number of companies whose sales are (i) less than Rs. 13 lakhs (ii) more than 36 lakhs and (iii) between Rs. 13 lakhs and Rs. 36 lakhs.

9.10 Let us Sum Up

This Lesson illustrated the Presentation of data through tables and charts which is essential for a management student to understand. A frequency distribution is the principal tabular Let us Sum Up of either discrete or continuous data. The frequency distribution may show actual, relative or cumulative frequencies. Actual and relative frequencies may be charted as either histogram (a bar chart) or a frequency polygon. Two graphs of cumulative frequencies are: less than ogive or more than ogive. These aspects discussed in this Lesson find major applications while presenting any data with a managerial perspective.

9.11 Lesson – End Activities

1. How the data is classified?
2. What are the guidelines for choosing the classes?

9.12 References

1. Statistics – R.S.N. Pillai, Mrs. Bhavathi.
2. Statistical Methods – Gupta G.S.S.

UNIT – IV

Lesson 10 - Descriptive Statistics – Measures of Central Tendency

Contents

- 10.1 Aims and Objectives
- 10.2 Significance of Measures of Central Tendency
- 10.3 Properties of a Good Measure of Central Tendency
- 10.4 Arithmetic Mean
- 10.5 Combined Mean of Two Groups
- 10.6 Weighted AM
- 10.7 Median
- 10.8 Median for a grouped frequency distribution
- 10.9 Mode
- 10.10 Mode of a grouped frequency distribution
- 10.11 Let us Sum Up
- 10.12 Lesson – End Activities
- 10.13 References

10.1 Aims and Objectives

This Lesson deals with the statistical methods for summarizing and describing numerical methods for summarizing and describing numerical data. The objective here is to find one representative value, which can be used to locate and summarise the entire set of varying values. This one value can be used to make many decisions concerning the entire set. We can define measures of central tendency (or location) to find some central value around which the data tend to cluster. Needless to say the content of this Lesson is important for a manager in taking decisions and also while communicating the decisions.

10.2 Significance of Measures of Central Tendency

Measures of central tendency i.e condensing the mass of data in one single value, enable us to get an idea of the entire data. For example, it is impossible to remember the individual incomes of millions of earning people of India. But if the average income is obtained, we get one single value that represents the entire population.

Measures of central tendency also enable us to compare two or more sets of data to facilitate comparison. For example, the average sales figures of April may be compared with the sales figures of previous months.

10.3 Properties of a Good Measure of Central Tendency

A good measure of central tendency should possess, as far as possible, the following properties.

- i) It should be easy to understand.

- ii) It should be simple to compute.
- iii) It should be based on all observations.
- iv) It should be uniquely defined.
- v) It should be capable of further algebraic treatment.
- vi) It should not be unduly affected by extreme values.

Following are some of the important measures of central tendency which are commonly used in business and industry.

Arithmetic Mean

Weighted Arithmetic Mean

Median

Quantiles

Mode

Geometric Mean

Harmonic Mean

10.4 Arithmetic Mean

The arithmetic mean (or mean or average) is the most commonly used and readily understood measure of central tendency. In statistics, the term average refers to any of the measures of central tendency. The arithmetic mean is defined as being equal to the sum of the numerical values of each and every observation divided by the total number of observations. Symbolically, it can be represented as:

$$\overline{X} = \frac{\sum X}{N}$$

Where $\sum X$ indicates the sum of the values of all the observations, and N is the total number of observations. For example, let us consider the monthly salary (Rs.) of 10 employees of a firm :

2500, 2700, 2400, 2300, 2550, 2650, 2750, 2450, 2600, 2400

If we compute the arithmetic mean, then

$$\overline{X} = \frac{2500 + 2700 + 2400 + 2300 + 2550 + 2650 + 2750 + 2450 + 2600 + 2400}{10}$$

$$= \frac{25300}{10} = \text{Rs. } 2530$$

Therefore, the average monthly salary is Rs. 2530.

We have seen how to compute the arithmetic mean for ungrouped data. Now let us consider what modifications are necessary for grouped data. When the observations are classified into a frequency distribution, the midpoint of the class interval would be treated as the representative average value of that class. Therefore, for grouped data, the arithmetic mean is defined as

$$\overline{X} = \frac{\sum fX}{N}$$

Where X is midpoint of various classes, f is the frequency for corresponding class and N is the total frequency. i.e. $N = \sum f$.

This method is illustrated for the following data which relate to the monthly sales of 200 firms.

Monthly sales (Rs. Thousand)	No. of Firms	Monthly Sales (Rs. Thousand)	No. of Firms
300-350	5	550-600	25
350-400	14	600-650	22
400-450	23	650-700	7
500-550	52	700-750	2

For computation of arithmetic mean, we need the following table:

Monthly Sales (Rs. Thousand)	Mid point X	No. of firms f	fX
300-350	325	5	1625
350-400	375	14	5250
400-450	425	23	9775
450-500	475	50	23750
500-550	525	52	27300
550-600	575	25	14375
600-650	625	22	13750
650-700	675	7	4725
700-750	725	2	1450
		N=200	fx=102000

$$\bar{X} = \frac{\sum fX}{N} = \frac{102000}{200} = 510$$

Hence the average monthly sales are Rs. 510.

To simplify calculations, the following formula for arithmetic mean may be more convenient to use.

$$\bar{X} = A + \sum \frac{fd}{N} \times i$$

Where A is an arbitrary point, $d = \frac{X-A}{i}$, and i=size of the equal class interval.

REMARK: A justification of this formula is as follows. When $d = \frac{X-A}{i}$, then $X = A + i d$. Taking summation on both sides and dividing by N, we get

$$\bar{X} = A + \sum \frac{fd}{N} \times i$$

This formula makes the computations very simple and takes less time. To apply this formula, let us consider the same example discussed earlier and shown again in the following table.

Monthly Sales (Rs. Thousand)	Mid point X	No. of firms f	(x-525)/50	fd
300-350	325	5	-4	-20
350-400	375	14	-3	-42
400-450	425	23	-2	-46
450-500	475	50	-1	-50
500-550	525	52	0	0
550-600	575	25	+1	+25
600-650	625	22	+2	+44
650-700	675	7	+3	+21
700-750	725	2	+4	+8
N = 200			fd = -60	

$$\overline{X} = A + \sum_N fd \cdot x_i = 525 - \frac{60}{200} \times 50$$

$$= 525 - 15 = 510 \text{ or Rs. 510}$$

It may be observed that this formula is much faster than the previous one and the value of arithmetic mean remains the same.

Properties of AM

1. The algebraic sum of deviations of a set of values from their AM is zero.
2. Sum of squares of deviations of a set of values is minimum when deviations taken about AM.

10.5 Combined Mean of Two Groups

Let $\overline{x}_1$ and $\overline{x}_2$ be the means of two groups. Let there be n_1 observations in the first group and n_2 observations in the second group. Then $\overline{x}$, the mean of the combined group can be obtained as

$$\overline{x} = \frac{n_1 \overline{x}_1 + n_2 \overline{x}_2}{n_1 + n_2}$$

Example : Average daily wage of 60 male workers in a firm is Rs. 120 and that of 40 females is Rs.100. Find the mean wage of all the workers.

Solution: Here $n_1 = 60$, $\overline{x}_1 = 120$ and $n_2 = 40$, $\overline{x}_2 = 100$

$$\begin{aligned}\text{Combined Mean} &= \frac{60 \times 120 + 40 \times 100}{60 + 40} \\ &= 112\end{aligned}$$

10.6 Weighted AM

When calculating AM we assume that all the observations have equal importance. If some items are more important than others, proper weightage should be given in accordance with their importance. Let $w_1, w_2, \dots, w_n$ be the weights attached to the items $x_1, x_2, \dots, x_n$, then the weighted AM is defined as

$$\text{Weighted mean} = \frac{w_1 x_1 + w_2 x_2 + \dots + w_n x_n}{w_1 + w_2 + \dots + w_n}$$

Example: A teacher has decided to use a weighted average in figuring final grades for his students. The midterm examination will count 40%, the final examination will count 50% and quizzes 10%. Compute the average mark obtained for a student who got 90 marks for midterm examination, 80 marks for final and 70 for quizzes.

Solution: Here $w_1 = 40, \quad x_1 = 90$
 $w_2 = 50, \quad x_2 = 80$
 $w_3 = 10, \quad x_3 = 70$

$$\begin{aligned}\text{Weighted mean} &= \frac{40 \times 90 + 50 \times 80 + 10 \times 70}{40 + 50 + 10} \\ &= \frac{8300}{100} \\ &= 83\end{aligned}$$

10.7 Median

The median of a set of observations is a value that divides the set of observations in half, so that the observations in one half are less than or equal to the median and the observations in the other half are greater than or equal to the median value.

In finding the median of a set of data it is often convenient to put the observations in ascending or descending order. If the number of observations is odd, the median is the middle observation. For example, if the values are 52, 55, 61, 67, and 72, the median is 61. If there were 4 values instead of 5, say 52, 55, 61, and 67, there would not be a middle value. Here any number between 55 and 61 could serve as a median; but it is desirable to use a specific number for the median and we usually take the AM of two middle values, i.e, $(55+61)/2 = 58$.

Median is the primary measure of location for variables measured on ordinal scale because it indicates which observation is central without attention to how far above or below the median the other observations fall.

Example: Find the median of 10, 2, 4, 8, 5, 1, 7

Solution: Observations in ascending order of magnitude are 1, 2, 4, 5, 7, 8, 10

Here there are 7 observations, so median is the 4th observation.

That is, median = 5

10.8 Median for a grouped frequency distribution

In a grouped frequency distribution, we do not know the exact values falling in each class. So, the median can be approximated by interpolation. Let the total number of observations be N . for calculating median we assume that the observations in the median class are uniformly distributed. Median class is the class in which the $(N/2)^{\text{th}}$ observation belongs. Also assume that median is the $(N/2)^{\text{th}}$ observation.

Here the frequency table must be continuous. If it is not, convert it into continuous table. Prepare a less than cumulative frequency table and find the median class. Let ' l ' be the lower limit of the median class, ' f ' the frequency of the median class, and ' c ' is the class width of the median class. By the assumption of uniform distribution, the ' f ' observations in the median class are $l + \frac{c}{f}, l + \frac{2c}{f}, \dots, l + \frac{fc}{f}$. Let ' m ' be the cumulative

frequency of the class above the median class. Then the median will be the $(\frac{N}{2} - m)^{\text{th}}$ observation in the median class.

That is, median = $l + (\frac{N}{2} - m) \frac{c}{f}$

Example : Calculate the median of the following data:

class	frequency
0 - 10	4
10 - 20	12
20 - 30	24
31 - 40	36
40 - 50	20
50 - 60	16
60 - 70	8
71 - 80	5

Solution: Since the frequency table is of inclusive, convert it into exclusive by subtracting 0.5 from the lower limits and adding 0.5 to the upper limits.

Class	Frequency	Cumulative frequency
0.5 - 10.5	4	4
10.5 - 20.5	12	16
20.5 - 30.5	24	40
30.5 - 40.5	36	76
40.5 - 50.5	20	96
50.5 - 60.5	16	112
60.5 - 70.5	8	120
70.5 - 80.5	5	125

Here $\frac{N}{2} = \frac{125}{2} = 62.5$, which lies in the 30.5 - 40.5 class (median class)

So, $l = 30.5$, $f = 36$, $m = 40$ and $c = 10$

$$\begin{aligned}
 \text{Median} &= l + \left(\frac{N}{2} - m \right) \frac{c}{f} \\
 &= 30.5 + (62.5 - 40) \frac{10}{36} \\
 &= 36.75
 \end{aligned}$$

Property of Median: The sum of absolute deviations of a set values is minimum when the deviations are taken from *median*.

10.9 Mode

The mode of a categorical or a discrete numerical variable is that category or value which occurs with the greatest frequency.

Example : The mode of the data 2, 5, 4, 4, 7, 8, 3, 4, 6, 4, 3 is 4 because 4 repeated the greatest number of times.

10.10 Mode of a grouped frequency distribution

In a grouped frequency distribution, to find the mode, first locate the modal class. Modal class is that class with maximum frequency. Let l be the lower limit of the modal class, 'c' be the class interval, f_1 be the frequency of the modal class, f_0 be the frequency of the class preceding and f_2 be the frequency of the class succeeding the modal class.

$$\text{Then, Mode} = l + \frac{c(f_1 - f_2)}{2f_1 - f_0 - f_2}$$

Example : Find the mode of the distribution given below

class	frequency
-------	-----------

10 – 15	3
15 – 20	9
20 – 25	16
25 – 30	12
30 – 35	7
35 – 40	5
40 – 45	2

Solution:

Here the modal class is the class 20 – 25.

That is, $l = 20$, $c = 5$, $f_0 = 9$, $f_1 = 16$ and $f_2 = 12$

$$\begin{aligned}\text{Mode} &= l + \frac{c(f_1 - f_2)}{2f_1 - f_0 - f_2} \\ &= 20 + \frac{5(16 - 12)}{32 - 9 - 12} = 21.8\end{aligned}$$

Exercises

- Find the arithmetic mean, median, and mode of the following data: 38, 28, 12, 18, 28, 44, 28, 19, 21.
- Calculate the mean, median and mode of the following data:
 Class: 10 – 20 20 – 30 30 – 40 40 – 50 50 – 60
 Frequency: 25 52 73 40 10
- From the following data of income distribution, calculate the AM. It is given that
 i) the total income of persons in the highest group is Rs. 435, and ii) none is earning less than Rs. 20.

Income (Rs)	No. of persons
Below 30	16
“ 40	36
“ 50	61
“ 60	76
“ 70	87
“ 80	95
80 and above	5

- Mean of 20 values is 45. If one of these values is to be taken 64 instead of 46. Find the correct mean.
- The mean yearly salary of employees of a company was Rs. 20,000. The mean yearly salaries of male and female employees were Rs. 20,800 and Rs. 16,800 respectively. Find out the percentage of males employed.

6. The average wage of 100 male workers is Rs. 80 and that 50 female workers is 75. Find the mean wage of workers in the company.

10.11 Let us Sum Up

The importance of measures of central tendency is described in this Lesson followed with different terms like mean, median, mode, etc. Measures of central tendency give one of the very important characteristics of data. Any one of the various measures of central tendency may be chosen as the most representative or typical measure. The AM is widely used and understood as a measure of central tendency. The concepts of weighted arithmetic mean, geometric mean and harmonic mean, are useful for specific types of applications. The median is a more representative measure for open-end distribution and highly skewed distribution. The mode should be used when the most demanded or customary value is needed. The examples shown in the Lesson clearly brings out the probable applications and the solution for specific problems.

10.12 Lesson – End Activities

1. Define Arithmetic mean, Genetic Mean.
2. Mention the properties of a good measure of central tendency.

10.13 References

Sundaresan and Jayaselan – An Introduction to Business Mathematics and Statistical Methods.

Lesson 11 - Quartiles, Deciles and Percentiles

Contents

- 11.1 Aims and Objectives
- 11.2 Measures of Dispersion
- 11.3 Quartile Deviation
- 11.4 Relative Measures
- 11.5 Skewness and Kurtosis
- 11.6 Let us Sum Up
- 11.7 Lesson – End Activities
- 11.8 References

11.1 Aims and Objectives

In the previous Lesson, we have discussed about the common measures of central tendency which are widely used in statistics. Median, as has been indicated, is a locational average, which divides the frequency distribution into two equal parts. Quartiles, deciles and percentiles are not averages. They are the partition values, which divides the distribution into certain equal parts.

Quartiles

Quartiles are the values, which divides a frequency distribution into four equal parts so that 25% of the data fall below the first quartile (Q_1), 50% below the second quartile (Q_2), and 75% below the third quartile (Q_3). The values of Q_1 and Q_3 can be found out as in the case of Q_2 (Median). For a raw data, Q_1 is the $(n/4)^{\text{th}}$ observation and Q_3 is the $(3n/4)^{\text{th}}$ observation.

For a grouped table, $Q_1 = l_1 + \left(\frac{N}{4} - m_1\right) \frac{c_1}{f_1}$

Where N is the total frequency, l_1 is the lower limit of the first quartile class (class in which $(N/4)^{\text{th}}$ observation belongs), m_1 is the cumulative frequency of the class above the first quartile class, f_1 is the frequency of the first quartile class and c_1 is the width of the first quartile class.

$$Q_3 = l_3 + \left(\frac{3N}{4} - m_3\right) \frac{C_3}{f_3}$$

Where l_3 is the lower limit of the third quartile class (class in which $(3N/4)^{\text{th}}$ observation belongs), m_3 is the cumulative frequency of the class above the third quartile class, f_3 is the frequency of the third quartile class and C_3 is the width of the third quartile class.

Deciles and Percentiles

Deciles are nine in number and divide the frequency distribution into 10 equal parts. Percentiles are 99 in number and divide the frequency distribution into 100 equal parts.

Selecting the Most Appropriate Measure of Central Tendency

Generally speaking, in analyzing the distribution of a variable only one of the possible measures of central tendency would be used. Its selection is largely a matter of judgment based upon the kind of data, the aspect of the data to be examined, and the research question. Some of the points that must be considered are following.

Central tendency for interval data is generally represented by the A.M., which takes into account the available information about distances between scores. For ranked (ordinal) data, the median is generally most appropriate, and for nominal data, the mode.

If the distribution is badly skewed, one may prefer the median to the mean, because the example, the median income of people is usually reported rather than the A.M.

If one is interested in prediction, the mode is the best value to predict if an *exact* score in a group has to be picked.

11.2 Measures of Dispersion

So far we have discussed averages as sample values used to represent data. But the average cannot describe the data completely.

Consider two sets of data : 5, 10, 15, 20, 25
 15, 15, 15, 15, 15

Here we observe that both the sets are with the same mean 15. But in the set I, the observations are more scattered about the mean. This shows that, even though they have the same mean, the two sets differ. This reveals the necessity to introduce measures of dispersion.

A measure of dispersion is defined as a mean of the scatter of observations from an average.

Commonly used measures of dispersion are *Range, Mean deviation, Standard deviation, and quartile deviation.*

11.2.1 Range

Range of a set of observations is the difference between the largest and the smallest observations. In the case of grouped frequency table, range is the difference between the upper bound of last class and the lower bound of the first class.

Example : The range of the set of data 9, 12, 25, 42, 45, 62, 65 is $65 - 9 = 56$

Range is the simplest measure of dispersion but its demerit is that it depends only on the extreme values.

11.2.2 Mean Deviation about the Mean

You have seen that range is a measure of dispersion, which does not depend on all observations. Let us think about another measure of dispersion, which will depend on all observations.

One measure of dispersion that you may suggest now is the sum of the deviations of observations from mean. But we know that the sum of deviations of observations from the A.M is always zero. So we cannot take the sum of deviations of observations from the mean as a measure.

One method to overcome this is to take the sum of absolute values of these deviations. But if we have two sets with different numbers of observations this cannot be justified. To make it meaningful we will take the average of the absolute deviations. Thus mean deviation (MD) about the mean is the mean of the absolute deviations of observations from arithmetic mean.

If $x_1, x_2, \dots, x_n$ are n observations, then, $MD = \frac{1}{n} \sum_{i=1}^n |x_i - \bar{x}|$

Example : Find the MD for the following data 12, 15, 21, 24, 28

Solution:

$$\bar{X} = \frac{12 + 15 + 21 + 24 + 28}{5} = 20$$

x	$ x_i - \bar{x} $
12	8
15	5
21	1
24	4
28	8
Total	26

$$MD = \frac{26}{5} = 5.2$$

Mean deviation about mean for a frequency table

Let $x_1, x_2, \dots, x_n$ be the values and $f_1, f_2, \dots, f_n$ are the corresponding frequencies. Let N be the sum of the frequencies. Then, $MD = \frac{1}{N} \sum_{i=1}^n |x_i - \bar{x}| f_i$

In the case of a grouped frequency table, take the mid-values as x -values and use the same method given above.

Example : Find the mean deviation of the heights of 100 students given below:

Height in cm	frequency
160 – 162	5
163 – 165	18
166 – 168	42
169 – 171	27
172 - 174	8

Solution:

Height in cm	Mid-value (x)	Frequency (f)	fx	$ x_i - \bar{x} $	$f_i x_i - \bar{x} $
160 – 162	161	5	805	6.45	32.25
163 – 165	164	18	2952	3.45	62.10
166 – 168	167	42	7014	0.45	18.90
169 – 171	170	27	4590	2.55	68.85
172 - 174	173	8	1384	5.55	44.40
Total		100	16745		226.50

$$\bar{X} = \frac{16745}{100} = 167.45$$

$$\begin{aligned} \text{MD} &= \frac{1}{N} \sum_{i=1}^n |x_i - \bar{x}| f_i \\ &= \frac{226.5}{100} = 2.265 \end{aligned}$$

11.2.3 Variance and Standard Deviation

When we take the deviations of the observations from their A.M both positive and negative values occurs. For defining mean deviation we took absolute values of the deviations. Another method to avoid this problem is to take the square of the deviations. So, variance is the mean of squares of deviations from A.M. Positive square root of variance is called **standard deviation**.

If $x_1, x_2, \dots, x_n$ are n observations, then, the variance = $\frac{1}{n} \sum_{i=1}^n (x_i - \bar{x})^2$ and standard

deviation(SD) is defined as, $SD = \sqrt{\frac{1}{n} \sum_{i=1}^n (x_i - \bar{x})^2}$

Example : Find the variance and standard deviation of the following data:

42, 39, 44, 40, 36, 39, 30, 46, 48, 36

Solution: Arithmetic mean $\bar{X} = \frac{400}{10} = 40$

$$\begin{aligned}
 \frac{1}{n} \sum_{i=1}^n (x_i - \bar{x})^2 &= \frac{1}{10} [(42 - 40)^2 + (39 - 40)^2 + \dots + (36 - 40)^2] \\
 &= \frac{254}{10} = 25.4 \\
 \text{Variance} &= 25.4 \\
 \text{S.D} &= \sqrt{25.4} = 5.04
 \end{aligned}$$

Variance and Standard deviation for a frequency table

Let $x_1, x_2, \dots, x_n$ be the values and $f_1, f_2, \dots, f_n$ are the corresponding frequencies. Let N be the sum of the frequencies. Then, $\text{Variance} = \frac{1}{N} \sum_{i=1}^n (x_i - \bar{x})^2 f_i$ and

$$\text{Standard deviation} = \sqrt{\frac{1}{N} \sum_{i=1}^n (x_i - \bar{x})^2 f_i}$$

The above formulae for variance can be expressed as, $\text{variance} = \frac{1}{N} \sum f_i x_i^2 - \bar{X}^2$

In the case of a grouped frequency table, take the mid-values as x -values and use the same method given above.

Example : Find the variance and standard deviation of the following data:

class	frequency
0 – 10	3
10 – 20	4
20 – 30	6
30 – 40	10
40 – 50	7

Solution:

class	mid-value (x)	frequency (f)	fx	fx ²
0 – 10	5	3	15	75
10 – 20	15	4	60	900
20 – 30	25	6	150	3750
30 – 40	35	10	350	12250
40 – 50	45	7	315	14175
Total		30	890	31150

$$\text{Variance} = \frac{1}{N} \sum f_i x_i^2 - \bar{X}^2$$

$$N = 30, \quad \bar{X} = \frac{890}{30} = 29.67, \quad \sum f_i x_i^2 = 31150$$

$$\begin{aligned}
 \text{Variance} &= \frac{31150}{30} - (29.67)^2 \\
 &= 1038.33 - 880.31 \\
 &= 158.02 \\
 \text{Standard deviation} &= \sqrt{158.02} = 12.57
 \end{aligned}$$

Short-cut method to find standard deviation

If the values of x are very large, the calculation of SD becomes time consuming.

Let the mid-values of k classes be $x_1, x_2, \dots, x_k$ and $f_1, f_2, \dots, f_k$ be the corresponding frequencies. We use the transformation of the form $u_i = \frac{x_i - A}{C}$ for $i = 1, 2, \dots, k$.

Here A and C can be any two numbers. But it is better to take A as a number among the middle part of the mid-values. If all the classes are of equal width, C can be taken as the class width.

$$\text{Variance of } u_i\text{'s, } \text{Var}(u) = \frac{1}{N} \sum f_i u_i^2 - \bar{u}^2$$

$$\text{Then variance of } x_i\text{'s, } \text{Var}(x) = C^2 \times \text{Var}(u)$$

$$\text{That is, } \text{SD}(x) = C \times \text{SD}(u)$$

Example : Consider the problem in example 5, let us find out the SD using short-cut method.

Solution:

class	mid-value (x)	$u_i = \frac{x_i - 25}{10}$	frequency (f)	fu	fu^2
0 – 10	5	-2	3	-6	12
10 – 20	15	-1	4	-4	4
20 – 30	25	0	6	0	0
30 – 40	35	1	10	10	10
40 – 50	45	2	7	14	28
Total			30	14	54

$$\bar{u} = \frac{\sum fu}{N} = \frac{14}{30} = 0.467, \quad \sum f_i u_i^2 = 54, \quad N = 30$$

$$\begin{aligned}
 \text{Variance}(u) &= \frac{54}{30} - (0.467)^2 \\
 &= 1.8 - 0.21809
 \end{aligned}$$

$$\begin{aligned} &= 1.5819 \\ \text{Variance}(x) &= 10^2 \times 1.5819 = 158.19 \end{aligned}$$

$$\text{SD}(x) = \sqrt{158.19} = 12.57$$

Combined Variance

If there are two sets of data consisting of n_1 and n_2 observations with s_1^2 and s_2^2 as their respective variances, then the variance of the combined set consisting of n_1+n_2 observations is :

$$S^2 = [n_1(s_1^2 + d_1^2) + n_2(s_2^2 + d_2^2)] / (n_1 + n_2)$$

Where d_1 and d_2 are the differences of the means, $\bar{x}_1$ and $\bar{x}_2$, from the combined mean $\bar{x}$ respectively.

Example : Find the combined standard deviation of two series A and B

	Series A	Series B
Mean	50	40
Standard deviation	5	6
No. of items	100	150

Solution:

Given $\bar{x}_1 = 50$ and $\bar{x}_2 = 40$, $s_1^2 = 25$ and $s_2^2 = 36$, $n_1 = 100$ and $n_2 = 150$

$$\text{Combined mean } \bar{x} = \frac{100 \times 50 + 150 \times 40}{100 + 150} = 44,$$

$$d_1 = \bar{x}_1 - \bar{x} = 50 - 44 = 6, \text{ and } d_2 = \bar{x}_2 - \bar{x} = 40 - 44 = -4$$

$$\begin{aligned} \text{Combined variance} &= \frac{100(25 + 36) + 150(36 + 16)}{100 + 150} \\ &= 55.6 \end{aligned}$$

$$\text{Therefore, combined SD} = \sqrt{55.6} = 7.46$$

11.3 Quartile Deviation

Quartile deviation (Semi inter-quartile range) is one-half of the difference between the third quartile and first quartile.

$$\text{That is, Quartile deviation, Q.D} = \frac{Q_3 - Q_1}{2}$$

Example : Estimate an appropriate measure of dispersion for the following data:

Income (Rs.)	No. of persons
Less than 50	54
50 – 70	100
70 – 90	140
90 – 110	300
110 – 130	230

130 – 150	125
Above 150	51
	1000

Solution:

Since the data has open ends, Q.D would be a suitable measure

Income (Rs.) x	No. of persons f	Cumulative frequency
Less than 50	54	54
50 – 70	100	154
70 – 90	140	294
90 – 110	300	594
110 – 130	230	824
130 – 150	125	949
Above 150	51	1000
	1000	

$$Q_1 = l_1 + \left(\frac{N}{4} - m_1 \right) \frac{c_1}{f_1}$$

$$Q_3 = l_3 + \left(\frac{3N}{4} - m_3 \right) \frac{c_3}{f_3}$$

$$\text{Here } N = 1000, \frac{N}{4} = 250, \frac{3N}{4} = 750$$

The class 70 – 90 is the first quartile class and 110 – 130 is the third quartile class

$$l_1 = 70, \quad m_1 = 154, \quad c_1 = 20, \quad f_1 = 140$$

$$l_3 = 110, \quad m_3 = 594, \quad c_3 = 20, \quad f_3 = 230$$

$$Q_1 = 70 + (250 - 154) \frac{20}{140}$$

$$= 83.7$$

$$Q_3 = 110 + (750 - 594) \frac{20}{230}$$

$$= 123.5$$

$$Q.D = \frac{123.5 - 83.7}{2} = 19.9 \text{ Rs.}$$

11.4 Relative Measures

The absolute measures of dispersion discussed above do not facilitate comparison of two or more data sets in terms of their variability. If the units of measurement of two or more sets of data are same, comparison between such sets of data is possible directly in terms of absolute measures. But conditions of direct comparison are not met, the desired comparison can be made in terms of the *relative measures*.

Coefficient of Variation is a relative measure of dispersion which express standard deviation(σ) as percent of the mean. That is Coefficient of variation, $C.V = (\sigma / \bar{x})100$.

Another relative measure in terms of quartile deviations is **Coefficient of quartile**

deviation and is defined as $Q_r = \frac{Q_3 - Q_1}{Q_3 + Q_1} \times 100$.

Example: An analysis of the monthly wages paid to workers in two firms A and B, belonging to the same industry, gives the following results:

	Firm A	Firm B
Number of workers	586	648
Average monthly wage	52.5	47.5
Standard deviation	10	11

In which firm, A or B, is there greater variability in individual wages?

Solution: Coefficient of variation for firm A = $\frac{10}{52.5} \times 100$
 $= 19\%$

Coefficient of variation for firm B = $\frac{11}{47.5} \times 100$
 $= 23\%$

There is greater variability in wages in firm B.

11.5 Skewness and Kurtosis

Skewness

Very often it becomes necessary to have a measure that reveals the direction of dispersion about the center of the distribution. Measures of dispersion indicate only the extent to which individual values are scattered about an average. These do not give information about the direction of scatter. *Skewness* refers to the direction of dispersion leading departures from symmetry, or lack of symmetry in a direction.

If the frequency curve of a distribution has longer tail to the right of the center of the distribution, then the distribution is said to be positively skewed. On the other hand, if the

distribution has a longer tail to the left of the center of the distribution, then distribution is said to be negatively skewed. Measures of skewness indicate the magnitude as well as the direction of skewness in a distribution.

Empirical Relationship between Mean, Median and Mode

The relationship between these three measures depends on the shape of the frequency distribution. In a symmetrical distribution the value of the mean, median and the mode is the same. But as the distribution deviates from symmetry and tends to become skewed, the extreme values in the data start affecting the mean.

In a positively skewed distribution, the presence of exceptionally high values affects the mean more than those of the median and the mode. Consequently the mean is highest, followed, in a descending order, by the median and the mode. That is, for a *positively skewed distribution*, $\text{Mean} > \text{Median} > \text{Mode}$. In a negatively skewed distribution, on the other hand, the presence of exceptionally low values makes the values of the mean the least, followed, in an ascending order, by the median and the mode. That is, for a negatively skewed distribution, $\text{Mean} < \text{Median} < \text{Mode}$.

Empirically, if the number of observations in any set of data is large enough to make its frequency distribution smooth and moderately skewed, then, $\text{Mean} - \text{Mode} = 3(\text{Mean} - \text{Median})$

Measures of Skewness

3. **Karl Pearson's measure of skewness:** Prof. Karl Pearson has been developed this measure from the fact that when a distribution drifts away from symmetry, its mean, median and mode tend to deviate from each other.

Karl Pearson's measure of skewness is defined as, $S_{kP} = \frac{\text{Mean} - \text{Mode}}{\text{SD}}$

4. **Bowley's measure of skewness:** developed by Prof. Bowley, this measure of skewness is derived from quartile values.

It is defined as $S_{kB} = \frac{Q_3 + Q_1 - 2Q_2}{Q_3 - Q_1}$

5. **Moment measure of skewness:**

If $x_1, x_2, \dots, x_n$ are n observations, then the r^{th} moment about mean is defined as

$$m_r = \frac{1}{n} \sum_{i=1}^n (x_i - \bar{x})^r$$

The moment measure of skewness is defined as $\beta_1 = m_3 / (\text{SD})^3$

In a perfectly symmetrical distribution $\beta_1 = 0$, and a greater or smaller value of β_1 results in a greater or smaller degree of skewness.

Kurtosis

Kurtosis refers to the degree of peakedness, or flatness of the frequency Curve. If the curve is more peaked than the normal curve, the curve is said to be *lepto* kurtic. If the curve is more flat than the normal curve, the curve is said to be *platy* kurtic. The normal curve is also called *meso* kurtic. The moment measure of kurtosis is $\beta_2 = \frac{m_4}{m_2^2}$. The value of $\beta_2=3$, if the distribution is normal; more than 3, if the distribution is *lepto* kurtic; and less than 3, if the distribution is *platy* kurtic.

Example : Given m_2 (variance) = 40, $m_3 = -100$. Find a measure of skewness.

Solution:

$$\begin{aligned}\text{Moment measure of skewness, } \beta_1 &= m_3 / (SD)^3 \\ &= \frac{-100}{(\sqrt{40})^3} = -0.4\end{aligned}$$

Hence, there is negative skewness

Example : The first four moments of a distribution about mean are 0, 2.5, 0.7, and 18.75. Comment on the Kurtosis of the distribution

$$\begin{aligned}\text{Moment measure of kurtosis is, } \beta_2 &= \frac{m_4}{m_2^2} \\ &= \frac{18.75}{2.5^2} = 3\end{aligned}$$

So, the curve is normal.

Exercises

- Find the standard deviation of the values: 11, 18, 9, 17, 7, 6, 15, 6, 4, 1
- Daily sales of a retail shop are given below:

Daily sales(Rs):	102	106	110	114	118	122	126
No. of days:	3	9	25	35	17	10	1

Calculate the mean and standard deviation of the above data and explain what they indicate about the distribution of daily sales?

- Goals scored by two teams A and B in a foot ball season were as follows:

No. of goals scored:	0	1	2	3	4
No. of matches	A: 2	9	8	5	4
	B: 1	7	6	5	3

Find which team may be considered more consistent?

- The mean of two samples of sizes 50 and 100 respectively are 54.1 and 50.3 and the standard deviations are 19 and 8. Find the mean and the standard deviation of the combined sample.

5. Find the quartile deviation of the following data:

Class	Frequency
< 15	5
15 – 20	12
20 – 25	22
25 – 30	31
30 – 35	19
35 – 40	9
>40	2

6. Find the skewness of the data 2, 3, 5, 8, 7, 6, 8, 7, 6, 5
7. Find the kurtosis of the data 7, 6, 9, 1, 0, 5, 5, 6, 5, 4
8. Find the Karl Pearson's measure of skewness of the following data:

Class	Frequency
< 15	5
15 – 20	12
20 – 25	22
25 – 30	31
30 – 35	19
35 – 40	9
>40	2

11.6 Let us Sum Up

In this Lesson we have discussed about how the concepts of measures of variation and skewness are important. Measures of variation considered were the range, average deviation, quartile deviation and standard deviation. The concept of coefficient of variation was used to compare relative variations of different data. The skewness was used in relation to lack of symmetry. Some example problems were also shown solved for a better understanding.

11.7 Lesson – End Activities

1. Define Quartile deviation.
2. Give the necessity for finding the skewness of the data.

11.8 References

R.S.N. Pillai and Mrs. Bhagavathi – Statistics.

Lesson 12 - Statistical Inference

Contents

- 12.1 Aims and Objectives
- 12.2 Sampling Distributions
- 12.3 Sampling Distribution of the Sample Mean
- 12.4 Distribution of Sample mean
- 12.5 Some Uses of Sampling distribution of Mean\
- 12.6 The Chi- Square Distribution
- 12.7 The Student – t Distribution
- 12.8 Student ‘t’ table
- 12.9 The F- Distribution
- 12.10 Estimation of Parameters
- 12.11 Testing Hypotheses
- 12.12 Let us Sum Up
- 12.13 Lesson – End Activities
- 12.14 References

12.1 Introduction

Sample statistics form the basis of all inferences drawn about populations. If we know the probability distribution of the sample statistic, then we can calculate the probability of that the sample statistic assumes a particular value or has a value in a given interval. This ability to calculate the probability that the simple statistic lies in a particular interval is the most important factor in all statistical inferences. Such aspects are covered in this Lesson. Examples are shown for better understanding of the subject.

12.2 Sampling Distributions

Suppose we wish to draw conclusions about a characteristic of a population. We draw a random sample of size n and take measurements about the characteristic, which we interested to study. Let the sample values be $x_1, x_2, x_3, \dots, x_n$. Then any quantity which can be determined as a function of the sample values $x_1, x_2, x_3, \dots, x_n$ is called a **statistic**. Since the sample values are the results of random selections, a statistic is a random variable. Therefore, a statistic has a probability distribution. It is known as **sampling distribution**. The standard deviation of the sampling distribution is called **standard error**.

The process of inferring certain facts about a population based on a sample is known as *statistical inference*. Sample statistics and their distributions are the basis of all inferences drawn about the population.

12.3 Sampling Distribution of the Sample Mean

Suppose we have a sample of size n from a population. Let $x_1, x_2, x_3, \dots, x_n$ be the values of the characteristic under study corresponding to the selected units. Then the sample mean $\bar{X}$ is defined as $\bar{X} = \frac{x_1 + x_2 + x_3 + \dots + x_n}{n}$.

If we draw another sample of size n from the same population, we may end up with a different set of sample values and so a different sample mean. Thus the value of the sample mean is determined by chance causes. The distribution of the sample mean is called sampling distribution of the sample mean.

12.4 Distribution of Sample mean

12.4.1 Distribution of sample mean of sample taken from any infinite population

If $x_1, x_2, x_3, \dots, x_n$ constitute a random sample from an infinite population having the mean μ and variance σ^2 , then the distribution of sample mean will be normal with mean μ and variance σ^2/n , when n is large.

12.4.2 Distribution of sample mean of sample taken from the normal population

If $\bar{X}$ is the mean of a random sample of size n from a normal population with the mean μ and variance σ^2 , its sampling distribution is a normal distribution with the mean μ and variance σ^2/n .

Example 1: a random sample of size 100 is taken from a normal population with $\sigma = 25$. What is the probability that the mean of the sample will greater from the mean of the population by atleast 3.

Solution: Let μ be the population mean and $\bar{x}$ be the sample mean. Given that $n = 100$, $\sigma = 25$.

Required probability = $P(\bar{x} - \mu > 3)$

$$\begin{aligned} &= P\left(\frac{\bar{x} - \mu}{\sigma} \sqrt{n} > \frac{3}{\sigma} \sqrt{n}\right) \\ &= P(z > 1.2) \\ &= 0.1151 \text{ (from } N(0,1) \text{ table, since } z \sim N(0,1)) \end{aligned}$$

Example 2: A random sample of size 64 is taken from an infinite population with the mean 22 and variance 196. What is the probability that the mean of the sample will greater than 23.

Solution: Given $n = 64$, $\mu = 22$, $\sigma = 14$. Let $\bar{x}$ be the sample mean.

We have to find out $P(\bar{x} > 23)$

$$\begin{aligned} P(\bar{x} > 23) &= P\left(\frac{\bar{x} - 22}{14} \sqrt{64} > \frac{23 - 22}{14} \sqrt{64}\right) \\ &= P\left(z > \frac{8}{14}\right) = P(z > 0.57) = 0.2843 \end{aligned}$$

12.5 Some Uses of Sampling distribution of Mean

1. To test the mean of a normal population when population standard deviation is known
2. To test the mean of any population when sample size is large (usually $n > 30$)
3. To test the equality of means of two populations when sample sizes large.
4. To test the equality of means of two normal populations when population standard deviations are known.
5. To find out the confidence interval for population mean; difference of population means of two populations. (both cases sample sizes are large).

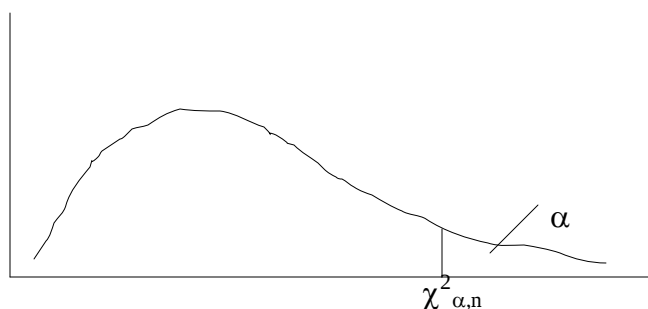
12.6 The Chi- Square Distribution

If a random variable X has the standard normal distribution, then the distribution X^2 is called chi-square (χ^2) distribution with one degree of freedom. This distribution would be quite different from a normal distribution because X^2 , being a square term, can assume only non-negative values. The probability curve of χ^2 will be higher near 0, because most of the x -values are close to 0 in a standard normal distribution.

If $X_1, X_2, \dots, X_n$ are independent standard normal variables, then $X_1^2 + X_2^2 + \dots + X_n^2$ has the χ^2 distribution with n degrees of freedom. Here 'n' is the only one parameter.

χ^2 – table

Since χ^2 -distribution arises in many important applications, especially in statistical inference, integrals of its density has been tabulated. The table gives the value of $\chi^2_{\alpha, n}$ such that probability that χ^2 is greater than $\chi^2_{\alpha, n}$ is equal to α for $\alpha = 0.005, 0.01, 0.025, 0.05$ etc. and $n = 1, 2, 3, \dots$. That is, the table gives $P(\chi^2 > \chi^2_{\alpha, n}) = \alpha$



Some Uses of Chi – Square Distribution

1. To test the variance of a normal population.
2. To test the independence of two attributes.
3. To test the homogeneity of two attributes.
4. To find the confidence interval for the variance of a normal population.

12.7 The Student – t Distribution

If X and Y are two independent random variables, X has the standard normal distribution and Y has a chi-square distribution with ‘n’ degrees of freedom, then the distribution of

the statistic $t = \frac{X}{\sqrt{\frac{Y}{n}}}$ is called Student ‘t’ distribution. The t-distribution was first obtained

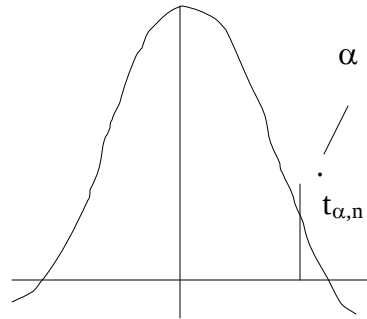
by W.S. Gosset, who is known under the pen name ‘Student’.

An example of a t-statistic is $t = \frac{\bar{x} - \mu}{s} \sqrt{n}$, which follows t-distribution with (n-1) degrees of

freedom, where $\bar{x}$ and s are mean and standard deviation of a random sample of size n from a normal population with mean μ and variance σ^2 .

12.8 Student ‘t’ table

Student ‘t’ table has many applications in statistical inference. The t-table gives the values $t_{\alpha, n}$ for $\alpha = 0.25, 0.125, 0.10, 0.05$ etc. and $n = 1, 2, 3, \dots$, where $t_{\alpha, n}$ is such that the area to its right under the curve of the t-distribution with ‘n’ degrees of freedom is equal to α . That is, $t_{\alpha, n}$ is such that $P(t > t_{\alpha, n}) = \alpha$. Also note that the t-distribution is a symmetric distribution.



Some Uses of t-distribution

1. To test the mean of a normal population when the sample size is small and population variance is unknown.
2. To test the equality of means of two normal populations when the sample sizes are small and population variances are unknown but same.
3. To test the correlation coefficient is zero.
4. To find the confidence interval of mean of normal population when sample size is small and population variance is unknown.

12.9 The F- Distribution

If U and V are independent random variables having chi-square distribution with m and n degrees of freedom, then the distribution of $\frac{U/m}{V/n}$ is called the F-distribution with m and n degrees of freedom.

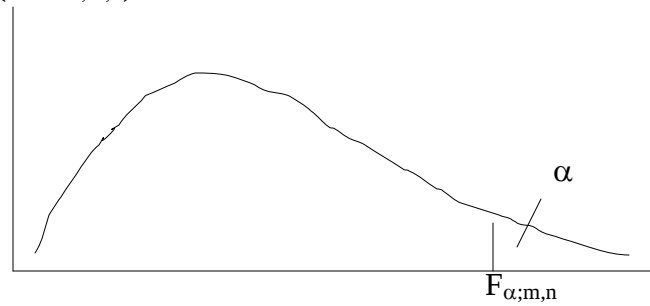
For example, if S_1^2 and S_2^2 are the variances of independent random samples of sizes m and n from normal populations with variances σ_1^2 and σ_2^2 , then,

$$F = \frac{S_1^2 \sigma_2^2}{S_2^2 \sigma_1^2} \text{ has an F-distribution with } m-1 \text{ and } n-1 \text{ degrees of freedom.}$$

Table of F-distribution

The table of F-distribution gives the values $F_{\alpha;m,n}$ for $\alpha=0.05$ and 0.01 for various values of m and n where $F_{\alpha;m,n}$ is such that the area to the right under the curve of F-distribution with m , n degrees of freedom is equal to α .

That is $F_{\alpha;m,n}$ is such that $P(F > F_{\alpha;m,n}) = \alpha$



Some Uses of F-distribution

1. To test the equality of variances of two normal populations.
2. F-distribution is used in analysis of variance.

12.10 Estimation of Parameters

The problem of estimation is of finding out a value for unknown population parameters, which we cannot directly observe, as precisely as possible. Managers deal this problem most frequently. They make quick estimates too. Since our estimates are based only on a sample, the estimates are not likely to be exactly equal to the value we are looking for. Still we will be able to obtain estimates whose possible values are around the true, but unknown value. The difference between the true value and the estimate is the error in estimation.

There are two types of estimates 1. Point Estimate 2. Interval Estimate

If an estimate of a population parameter is given by a single value, then the estimate is called *point estimate* of the parameter. But if two distinct numbers give an estimate of a population parameter between which the parameter may be considered to lie, then the estimate is called an *interval estimate* of the parameter.

A function, T , used for estimating a parameter θ , is called an *estimator* and its value given a sample is known as *estimate*.

Required Properties of an Estimator

1. **Unbiasedness:** An estimator must be an *unbiased* estimator of the parameter. That is an estimator T is said to be unbiased for a parameter θ if $E(T) = \theta$.
2. **Efficiency:** Efficiency refers to the size of the standard error of the estimator. That is, an estimator T_1 is said to be more efficient than another estimator T_2 if standard error of T_1 is less than the standard error of T_2 .
3. **Consistency:** As the sample size increases the value of the estimator must get close to the parameter.
4. **Sufficiency:** An estimator T is said to be sufficient for a parameter θ if T contains all information which the sample contains and furnishes about θ .

Some Point Estimators

1. The sample mean $\bar{X}$ is a point estimator of the population mean μ
2. The sample proportion is a point estimate of the population proportion.
3. The sample variance is a point estimator of population variance.

12.11 Testing Hypotheses

Statistical testing or testing hypotheses, is one of the most important aspects of the theory of decision-making. Testing hypotheses consists of decision rules required for drawing probabilistic inferences about the population parameters.

Definition: A **Statistical Hypothesis** is a statement concerning a probability distribution or population parameters and a process by which a decision is arrived at, whether or not a hypothesis is true is **Testing Hypothesis**.

For example, the statement, mean of a normal population is 30, the variance of a population is greater than 12 are statistical hypotheses.

Null Hypothesis and Alternate Hypothesis

The hypothesis under test is known as the *null hypothesis* and the hypothesis that will be accepted when the null hypothesis is rejected is known as the *alternate hypothesis*.

The null hypothesis is usually denoted by H_0 and the alternate hypothesis by H_1 . For example, if the population mean is represented by μ , we can set up our hypothesis as follows: $H_0: \mu \leq 30$; $H_1: \mu > 30$.

The following are the steps in testing a statistical hypothesis. We draw a sample from the concerned population. Then choose the appropriate test statistic. A **test statistic** is a statistic, based on the value of it we decide either to reject or accept a hypothesis. Divide the sample space of the test statistic into two regions, namely, rejection region and acceptance region. (The set of sample points, which lead to the rejection of the null hypothesis, is called the **Critical Region or Rejection Region**). Calculate the value of the test statistic for our sampled data. If this value falls in the rejection region, reject the hypothesis; otherwise accept it.

Type I Error and Type II Error

Since we have to depend on the sample there is no way to know, which of the two hypotheses is actually true. The test procedure is to fix the rejection region, in which the value of test statistic observed, the null hypothesis would be rejected. The null hypothesis may be true, but the test procedure may reject the null hypothesis. This error is known as the first kind of error. It is also possible that the null hypothesis is actually false but the test accepts it. This error is known as the second kind of error. Thus, *the error committed in rejecting a true null hypothesis is called type I error and the error in accepting a false null hypothesis is called the type II error*.

Significance Level

The probabilities of two errors cannot be simultaneously reduced, since if we increase the rejection region the probability of type I error will increase whereas the reduction in rejection region will increase type II error. The procedure usually adopted is to keep the probability of type I error below a pre-assigned number and subject to this condition minimize the type II error. A pre-assigned number α between 0 and 1 chosen as an upper bound of type I error is called **the level of significance**.

Two-tailed and One-tailed Tests

A test where the critical region is found to lie under one tail of the distribution of the test statistic is called *One-tailed test*. In *two-tailed tests* the critical region lies under both the tails of the distribution of the test statistic.

Example: Let μ be the mean of a population. Then,

1. $H_0: \mu = 30$; $H_1: \mu \neq 30$ is a two tailed test
2. $H_0: \mu = 30$; $H_1: \mu > 30$ is a single tailed test.

Exercise

3. A population is normally distributed with mean 90. A sample of size 10 is taken at random from the population. Find the probability that the population mean is greater than 85.
4. In the above problem, suppose we have to test whether the population mean is equal to 85. Formulate the null hypothesis and alternate hypothesis.

12.12 Let us Sum Up

The concept of sampling distribution is introduced in this Lesson. Some of the commonly used sampling distributions used in statistics and some of the applications are also shown. The sampling distribution is very important in statistical calculations and inferences.

12.13 Lesson – End Activities

1. Define students – t distribution and F – distribution.
2. List the uses of sampling distribution of mean.

12.14 References

1. Gupta. S.P. – Statistical Methods.
2. R.S.N. Pillai and Mrs. Bhagavathi – Statistics.

Lesson 13 - Correlation and Regression

Contents

- 13.1 Aims and Objectives
- 13.2 Correlation
- 13.3 The Scatter Diagram
- 13.4 The Correlation Coefficient
- 13.5 Karl Pearson's Correlation Coefficient
- 13.6 Relation between Regression Coefficients and Correlation Coefficient
- 13.7 Coefficient of Determination
- 13.8 Spearman's Rank Correlation Coefficient
- 13.9 Tied Ranks
- 13.10 Regression
- 13.11 Linear Regression
- 13.12 Let us Sum Up
- 13.13 Lesson – End Activities
- 13.14 References

13.1 Introduction

There are situations where data appears as pairs of figures relating to two variables. A correlation problem considers the joint variation of two measurements neither of which is restricted by the experimenter. The regression problem discussed in this Lesson considers the frequency distribution of one variable (called the dependent variable) when another (independent variable) is held fixed at each of several levels.

Examples of correlation problems are found in the study of the relationship between IQ and aggregate percentage of marks obtained by a person in the SSC examination, blood pressure and metabolism or the relation between height and weight of individuals. In these examples both variables are observed as they naturally occur, since neither variable is fixed at predetermined levels.

Examples of regression problems can be found in the study of the yields of crops grown with different amount of fertilizer, the length of life of certain animals exposed to different levels of radiation, and so on. In these problems the variation in one measurement is studied for particular levels of the other variable selected by the experimenter.

13.2 Correlation

Correlation measures the degree of linear relation between the variables. *The existence of correlation between variables does not necessarily mean that one is the cause of the change in the other. It should be noted that the correlation analysis merely helps in*

determining the degree of association between two variables, but it does not tell any thing about the cause and effect relationship. While interpreting the correlation coefficient, it is necessary to see whether there is any cause and effect relationship between variables under study. If there is no such relationship, the observed is meaningless.

In correlation analysis, all variables are assumed to be random variables.

13.3 The Scatter Diagram

The first step in correlation and regression analysis is to visualize the relationship between the variables. A scatter diagram is obtained by plotting the points (x_1, y_1) , (x_2, y_2) , ..., (x_n, y_n) on a two-dimensional plane. If the points are scattered around a straight line, we may infer that there exist a linear relationship between the variables. If the points are clustered around a straight line with negative slope, then there exist negative correlation or the variables are inversely related (i.e, when x increases y decreases and vice versa.). If the points are clustered around a straight line with positive slope, then there exist positive correlation or the variables are directly related (i.e, when x increases y also increases and vice versa.).

For example, we may have figures on advertisement expenditure (X) and Sales (Y) of a firm for the last ten years, as shown in Table 1. When this data is plotted on a graph as in Figure 1 we obtain a scatter diagram. A scatter diagram gives two very useful types of information. First, we can observe patterns between variables that indicate whether the variables are related. Secondly, if the variables are related we can get an idea of what kind of relationship (linear or non-linear) would describe the relationship.

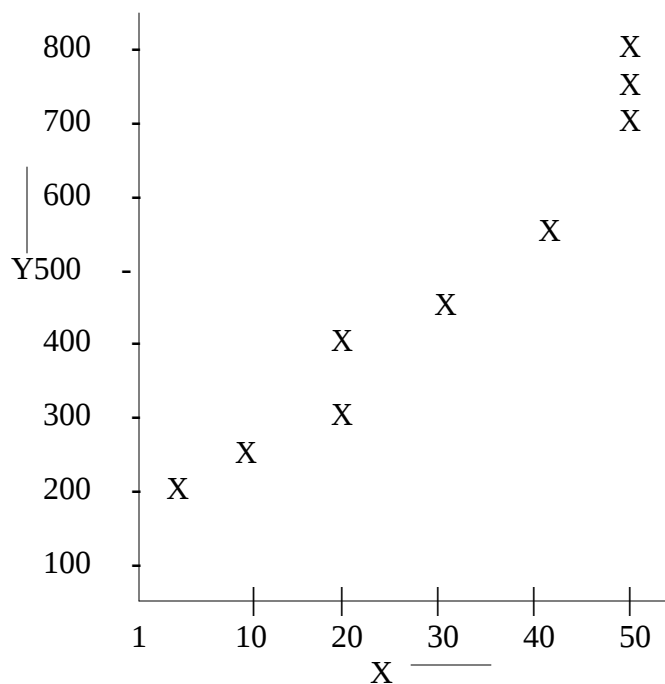
Table 1

Year-wise data on Advertisement Expenditure and Sales

Year	Advertisement Expenditure In thousand Rs. (X)	Sales in Thousand Rs. (Y)
1988	50	700
1987	50	650
1986	50	600
1985	40	500
1984	30	450
1983	20	400
1982	20	300
1981	15	250
1980	10	210
1979	5	200

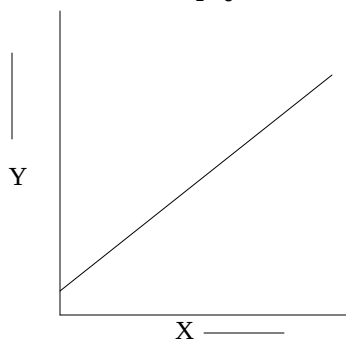
Correlation examines the first Question of determining whether an association exists between the two variables, and if it does, to what extent. Regression examines the second question of establishing an appropriate relation between the variables.

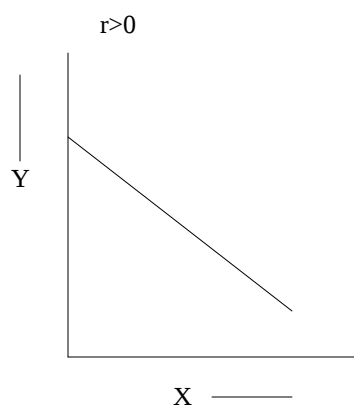
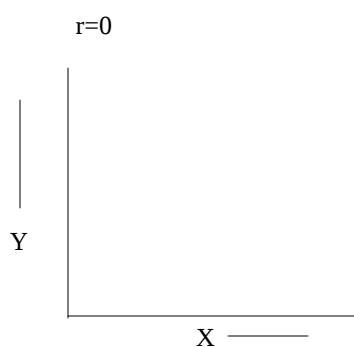
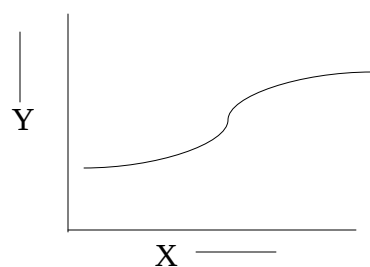
Figure 1 : Scatter Diagram



The scatter diagram may exhibit different kinds of patterns. Some typical patterns indicating different correlations between two variables are shown in Figure 2.

Figure 2: Different Types of Association Between Variables
 $r > 0$



(a) **Positive Correlation**(b) **Negative Correlation**(c) **No Correlation**(d) **Non-linear Association**

13.4 The Correlation Coefficient

Definition and Interpretation

The correlation coefficient measure the degree of association between two variables X and Y. Pearson's formula for correlation coefficient is given as

$$r = \frac{\frac{1}{n} \sum (X - \bar{X})(Y - \bar{Y})}{\sigma_x \sigma_y}$$

Where r is the correlation coefficient between X and Y, σ_x and σ_y are the standard deviation of X and Y respectively and n is the number of values of the pair of variables X

and Y in the given data. The expression $\frac{1}{n} \sum (X - \bar{X})(Y - \bar{Y})$ is known as the covariance between X and Y. Here r is also called the Pearson's product moment correlation coefficient. You should note that r is a dimensionless number whose numerical value lies between +1 and -1. Positive values of r indicate positive (or direct) correlation between the two variables X and Y i.e. as X increase Y will also increase or as X decreases Y will also decrease. Negative values of r indicate negative (or inverse) correlation, thereby meaning that an increase in one variable results in a decrease in the value of the other variable. A zero correlation means that there is no association between the two variables. Figure II shown a number of scatter plots with corresponding values for the correlation coefficient r.

The following form for carrying out computations of the correlation coefficient is perhaps more convenient :

$$r = \frac{\sum xy}{\sqrt{\sum X^2} \sqrt{\sum y^2}} \quad \text{.....(18.2)}$$

where

$x = X - \bar{X}$ = deviation of a particular X value from the mean- $\bar{X}$

$y = Y - \bar{Y}$ = deviation of a particular Y value from the mean $\bar{Y}$

Equation (18.2) can be derived from equation (18.1) by substituting for σ_x and σ_y as follows:

$$\sigma_x = \sqrt{\frac{1}{n} \sum (X - \bar{X})^2} \quad \text{and} \quad \sigma_y = \sqrt{\frac{1}{n} \sum (Y - \bar{Y})^2} \quad \text{.....(18.3)}$$

13.5 Karl Pearson's Correlation Coefficient

If $(x_1, y_1), (x_2, y_2), \dots, (x_n, y_n)$ be n given observations, then the Karl Pearson's correlation coefficient is defined as, $r = \frac{S_{xy}}{S_x S_y}$, where S_{xy} is the covariance and S_x, S_y are the standard deviations of X and Y respectively.

$$\text{That is, } r = \frac{\frac{1}{n} \sum xy - \bar{x} \bar{y}}{\sqrt{\frac{1}{n} \sum x^2 - \bar{x}^2} \sqrt{\frac{1}{n} \sum y^2 - \bar{y}^2}}$$

The value of r is in between -1 and 1. That is, $-1 \leq r \leq 1$. When $r = 1$, there exist a perfect positive linear relation between x and y. when $r = -1$, there exist perfect negative linear relationship between x and y. when $r = 0$, there is no linear relationship between x and y.

13.6 Relation between Regression Coefficients and Correlation Coefficient

Correlation coefficient is the geometric mean of the regression coefficients.

We know that $b_{yx} = \frac{S_{xy}}{S_x^2}$ and $b_{xy} = \frac{S_{xy}}{S_y^2}$

The geometric mean of b_{yx} and b_{xy} is $\sqrt{b_{yx}b_{xy}} = \sqrt{\frac{S_{xy}S_{xy}}{S_y^2 S_x^2}}$

$$= \frac{S_{xy}}{S_x S_y}$$

= r , the correlation coefficient.

Also note that the sign of both the regression coefficients will be same, so the sign of correlation coefficient is same as the sign of regression coefficient.

13.7 Coefficient of Determination

Coefficient of determination is the square of correlation coefficient and which gives the proportion of variation in y explained by x . That is, coefficient of determination is the ratio of explained variance to the total variance. For example, $r^2 = 0.879$ means that 87.9% of the total variances in y are explained by x . When $r^2 = 1$, it means that all the points on the scatter diagram fall on the regression line and the entire variations are explained by the straight line. On the other hand, if $r^2 = 0$ it means that none of the points on scatter diagram falls on the regression line, meaning thereby that there is no linear relationship between the variables.

Example: Consider the following data:

X:	15	16	17	18	19	20
Y:	80	75	60	40	30	20

1. Fit both regression lines
2. Find the correlation coefficient
3. Verify the correlation coefficient is the geometric mean of the regression coefficients
4. Find the value of y when $x = 17.5$

Solution:

X	Y	XY	X ²	Y ²
15	80	1200	225	6400
16	75	1200	256	5625
17	60	1020	289	3600
18	40	720	324	1600
19	30	570	361	900
20	20	400	400	400
105	305	5110	1855	18525

$$\bar{x} = \frac{\sum x}{n} = \frac{105}{6} = 17.5, \quad \bar{y} = \frac{\sum y}{n} = \frac{305}{6} = 50.83$$

$$S_{xy} = \frac{1}{n} \sum x_i y_i - \bar{x} \bar{y} = \frac{5110}{6} - 17.5 \times 50.83 = -37.86$$

$$S_x^2 = \frac{1}{n} \sum x_i^2 - (\bar{x})^2 = \frac{1855}{6} - 17.5^2 = 2.92$$

$$S_y^2 = \frac{1}{n} \sum y_i^2 - (\bar{y})^2 = \frac{18525}{6} - 50.83^2 = 503.81$$

$$b_{yx} = \frac{S_{xy}}{S_x^2} = \frac{-37.86}{2.92} = -12.96 \text{ and } b_{xy} = \frac{S_{xy}}{S_y^2} = \frac{-37.86}{503.81} = -0.075$$

1. Regression line of y on x is $y - \bar{y} = \frac{S_{xy}}{S_x^2} (x - \bar{x})$

$$\text{i.e., } y - 50.83 = -12.96(x - 17.5)$$

$$\mathbf{y = -12.96x + 277.63}$$

Regression line of x on y is $x - \bar{x} = \frac{S_{xy}}{S_y^2} (y - \bar{y})$

$$\text{i.e., } x - 17.5 = -0.075(y - 50.83)$$

$$\mathbf{x = -0.075y + 21.31}$$

2. Correlation coefficient, $r = \frac{S_{xy}}{S_x S_y}$

$$= \frac{-37.86}{1.71 \times 22.45} = \mathbf{0.986}$$

3. $b_{yx} \times b_{xy} = -12.96 \times -0.075 = 0.972$

$$\text{Then, } \sqrt{0.972} = \mathbf{0.986}$$

$$\text{So, } \mathbf{r = -0.986}$$

4. To predict the value of y, use regression line of y on x.

$$\text{When } x = 17.5, \quad y = -12.96 \times 17.5 + 277.63 = \mathbf{50.83}$$

Short-Cut Method: The correlation coefficient is invariant under linear transformations.

Let us take the transformations, $u = \frac{x-18}{1}$ and $v = \frac{y-40}{10}$

X	Y	u	v	uv	u ²	v ²
15	80	-3	4	-12	9	16
16	75	-2	3.5	-7	4	12.25
17	60	-1	2	-2	1	4
18	40	0	0	0	0	0
19	30	1	-1	-1	1	1
20	20	2	-2	-4	4	4
85	305	-3	6.5	-26	19	37.25

$$\bar{u} = \frac{\sum u}{n} = \frac{-3}{6} = -0.5, \quad \bar{v} = \frac{\sum v}{n} = \frac{6.5}{6} = 1.083$$

$$S_{uv} = \frac{1}{n} \sum u_i v_i - \bar{u} \bar{v} = \frac{-26}{6} - (-0.5 \times 1.083) = -3.79$$

$$S_u^2 = \frac{1}{n} \sum u_i^2 - (\bar{u})^2 = \frac{19}{6} - (-0.5)^2 = 2.92$$

$$S_v^2 = \frac{1}{n} \sum v_i^2 - (\bar{v})^2 = \frac{37.25}{6} - 1.083^2 = 5.077$$

$$b_{vu} = \frac{S_{uv}}{S_u^2} = \frac{-3.79}{2.92} = -1.297 \text{ and } b_{uv} = \frac{S_{uv}}{S_v^2} = \frac{-3.79}{5.077} = -0.75$$

1. Regression line of v on u is $v - \bar{v} = b_{vu}(u - \bar{u})$
 i.e., $v - 1.083 = -1.297(u - -0.5)$
 $v = -1.297u + 0.4345$

Therefore, the regression line of y on x is $\frac{y-40}{10} = -1.297 \frac{x-18}{1} + 0.4345$
 i.e., $y = -12.97x + 277.8$

Regression line of u on v is $u - \bar{u} = b_{uv}(v - \bar{v})$
 i.e., $u - -0.5 = -0.75(v - 1.083)$
 $u = -0.75v + 0.31225$

Therefore, the regression line of x on y is $\frac{x-18}{1} = -0.75 \frac{y-40}{10} + 0.31225$
 i.e., $x = -0.075y + 21.31$

2. Correlation coefficient, $r = \frac{S_{uv}}{S_u S_v}$
 $= \frac{-3.79}{1.71 \times 2.253} = -0.986$

3. $b_{vu} \times b_{uv} = -1.297 \times -0.75 = 0.97275$
 Then, $\sqrt{0.972} = 0.986$
 So, $r = -0.986$

13.8 Spearman's Rank Correlation Coefficient

Sometimes the characteristics whose possible correlation is being investigated, cannot be measured but individuals can only be ranked on the basis of the characteristics to be measured. We then have two sets of ranks available for working out the correlation coefficient. Sometimes the data on one variable may be in the form of ranks while the data on the other variable are in the form of measurements which can be converted into ranks. Thus, when both the variables are ordinal or when the data are available in the ordinal form irrespective of the type variable, we use the rank correlation coefficient.

The Spearman's rank correlation coefficient is defined as , $r = 1 - \frac{6\sum d_i^2}{n(n^2 - 1)}$

Example: Ten competitors in a beauty contest were ranked by two judges in the following orders:

First judge: 1 6 5 10 3 2 4 9 7 8

Second judge: 3 5 8 4 7 10 2 1 6 9

Find the correlation between the rankings.

Solution:

x_i	y_i	$d_i = x_i - y_i$	d_i^2
1	3	-2	4
6	5	1	1
5	8	-3	9
10	4	6	36
3	7	-4	16
2	10	-8	64
4	2	2	4
9	1	8	64
7	6	1	1
8	9	-1	1

The Spearman's rank correlation coefficient is defined as , $r = 1 - \frac{6\sum d_i^2}{n(n^2 - 1)}$

$$= 1 - \frac{6 \times 200}{10(10^2 - 1)} = -0.212$$

That is, their opinions regarding beauty test are apposite of each other.

13.9 Tied Ranks

Sometimes where there is more than one item with the same value a common rank is given to such items. This rank is the average of the ranks which these items would have got had they differed slightly from each other. When this is done, the coefficient of rank correlation needs some correction, because the above formula is based on the supposition that the ranks of various items are different. If in a series, ' m_i ' be the frequency of i^{th} tied ranks,

$$\text{Then, } r = 1 - \frac{6[\sum d_i^2 + \sum \frac{1}{12}(m^3 - m)]}{n(n^2 - 1)}$$

Example: Calculate the rank correlation coefficient from the sales and expenses of 10 firms are below:

Sales(X): 50 50 55 60 65 65 65 60 60 50
 Expenses(Y): 11 13 14 16 16 15 15 14 13 13

Solution:

x	R ₁	y	R ₂	d= R ₁ - R ₂	d ²
50	9	11	10	-1	1
50	9	13	8	1	1
55	7	14	5.5	1.5	2.25
60	5	16	1.5	3.5	12.25
65	2	16	1.5	0.5	0.25
65	2	16	3.5	-1.5	2.25
65	2	15	3.5	-1.5	2.25
60	5	14	5.5	-0.5	0.25
60	5	13	8	-3	9
50	9	13	8	1	1
					31.5

Here there are 7 tied ranks, $m_1 = 3, m_2 = 3, m_3 = 3, m_4 = 2, m_5 = 2, m_6 = 2, m_7 = 3$.

$$\begin{aligned}
 r &= 1 - \frac{6[\sum d_i^2 + \sum \frac{1}{12}(m^3 - m)]}{n(n^2 - 1)} \\
 &= 1 - \frac{6[31.5 + \frac{1}{12}[(3^3 - 3) + (3^3 - 3) + (3^3 - 3) + (2^3 - 2) + (2^3 - 2) + (2^3 - 2) + (3^3 - 3)]]}{10(10^2 - 1)} \\
 &= 0.75
 \end{aligned}$$

Exercises

1. A company selling household appliances wants to determine if there is any relationship between advertising expenditures and sales. The following data was compiled for 6 major sales regions. The expenditure is in thousands of rupees and the sales are in millions of rupees.

Region :	1	2	3	4	5	6
Expenditure(X):	40	45	80	20	15	50
Sales (Y):	25	30	45	20	20	40

- a) Compute the line of regression to predict sales
 - b) Compute the expected sales for a region where Rs.72000 is being spent on advertising
2. The following data represents the scores in the final exam., of 10 students, in the subjects of Economics and Finance.

Economics:	61	78	77	97	65	95	30	74	55
Finance:	84	70	93	93	77	99	43	80	67

- a) Compute the correlation coefficient?
3. Calculate the rank correlation coefficient from the sales and expenses of 9

firms are below:

Sales(X):	42	40	54	62	55	65	65	66	62
Expenses(Y):	10	18	18	17	17	14	13	10	13

13.10 Regression

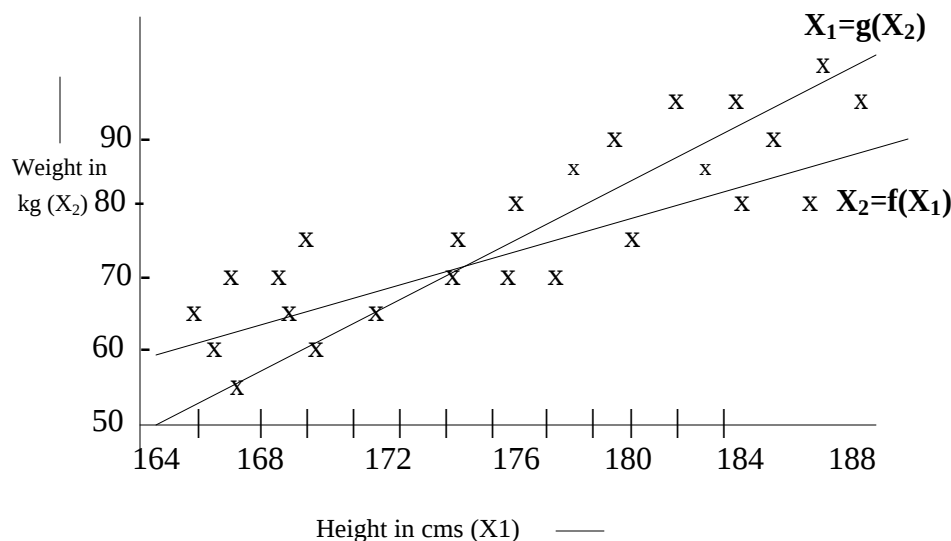
In industry and business today, large amounts of data are continuously being generated. This may be data pertaining, for instance, to a company's annual production, annual sales, capacity utilisation, turnover, profits, manpower levels, absenteeism or some other variable of direct interest to management. Or there might be technical data regarding a process such as temperature or pressure at certain crucial points, concentration of a certain chemical in the product or the braking strength of the sample produced or one of a large number of quality attributes.

The accumulated data may be used to gain information about the system (as for instance what happens to the output of the plant when temperature is reduced by half) or to visually depict the past pattern of behaviours (as often happens in company's annual meetings where records of company progress are projected) or simply used for control purposes to check if the process or system is operating as designed (as for instance in quality control). Our interest in regression is primarily for the first purpose, mainly to extract the main features of the relationships hidden in or implied by the mass of data.

What is Regression?

Suppose we consider the height and weight of adult males for some given population. If we plot the pair $(X_1, X_2) = (\text{height}, \text{weight})$, a diagram like figure 1 will result. Such a diagram, you would recall from the previous Lesson, is conventionally called a scatter diagram.

Note that for any given height there is a range of observed weights and vice-versa. This variation will be partially due to measurement errors but primarily due to variations between individuals. Thus no unique relationship between actual height and weight can be expected. But we can note that average observed weight for a given observed height increases as height increases. The locus of average observed weight for given observed height (as height varies) is called the **regression curve** of weight on height. Let us denote it by $X_2 = f(X_1)$. There also exists a regression curve of height on weight similarly defined which we can denote by $X_1 = g(X_2)$. Let us assume that these two "curves" are both straight lines (which in general they may not be). In general these two curves are not the same as indicated by the two lines in Figure 3.

Figure 3: Height and Weight of thirty Adult Males

A pair of random variables such as (height, weight) follows some sort of bivariate probability distribution. When we are concerned with the dependence of a random variable Y on quantity X , which is variable but not a random variable, an equation that relates Y to X is usually called a **regression equation**. Simply when more than one independent variable is involved, we may wish to examine the way in which a response Y depends on variables $X_1 X_2 \dots X_k$. We determine a regression equation from data which cover certain areas of the X -space as $Y = f(X_1, X_2, \dots, X_k)$

13.11 Linear Regression

Regression analysis is a set of statistical techniques for analyzing the relationship between two numerical variables. One variable is viewed as the *dependent variable* and the other as the *independent variable*. The purpose of regression analysis is to understand the direction and extent to which values of dependent variable can be predicted by the corresponding values of the independent variable. The regression gives the nature of relationship between the variables.

Often the relationship between two variable x and y is not an exact mathematical relationship, but rather several y values corresponding to a given x value scatter about a value that depends on the x value. For example, although not all persons of the same height have exactly the same weight, their weights bear some relation to that height. On the average, people who are 6 feet tall are heavier than those who are 5 feet tall; the mean weight in the population of 6-footers exceeds the mean weight in the population of 5-footers.

This relationship is modeled statistically as follows: For every value of x there is a corresponding population of y values. The population mean of y for a particular value of x is denoted by $f(x)$. As a function of x it is called the *regression function*. If this regression function is linear it may be written as $f(x) = a + bx$. The quantities a and b are parameters that define the relationship between x and $f(x)$.

In conducting a regression analysis, we use a sample of data to estimate the values of these parameters. The population of y values at a particular x value also has a variance; the usual assumption is that the variance is the same for all values of x .

Principle of Least Squares

Principle of least squares is used to estimate the parameters of a linear regression. The principle states that the best estimates of the parameters are those values of the parameters, which minimize the sum of squares of residual errors. The residual error is the difference between the actual value of the dependent variable and the estimated value of the dependent variable.

Fitting of Regression Line $y = a + bx$

By the principle of least squares, the best estimates of a and b are

$$b = \frac{S_{xy}}{S_x^2} \quad \text{and} \quad a = \bar{y} - b\bar{x}$$

Where S_{xy} is the covariance between x and y and is defined as $S_{xy} = \frac{1}{n} \sum x_i y_i - \bar{x} \bar{y}$

And S_x^2 is the variance of x , that is, $S_x^2 = \frac{1}{n} \sum x_i^2 - (\bar{x})^2$

Example: Fit a straight line $y = a + bx$ for the following data.

Y	3.5	4.3	5.2	5.8	6.4	7.3	7.2	7.5	7.8	8.3
X	6	8	9	12	10	15	17	20	18	24

Solution:

Y	X	XY	X ²
3.5	6	21	36
4.3	8	34.4	64
5.2	9	46.8	81
5.8	12	69.6	144
6.4	10	64	100
7.3	15	109.5	225
7.2	17	122.4	289
7.5	20	150	400
7.8	18	140.4	324

8.3	24	199.2	576
63.3	139	957.3	2239

$$\bar{x} = \frac{139}{10} = 13.9 \quad \bar{y} = \frac{63.3}{10} = 6.33$$

$$S_{xy} = \frac{1}{n} \sum x_i y_i - \bar{x} \bar{y} = \frac{957.3}{10} - 13.9 \times 6.33 = 7.743$$

$$S_x^2 = \frac{1}{n} \sum x_i^2 - (\bar{x})^2 = \frac{2239}{10} - 13.9^2 = 30.69$$

$$\text{So, } b = \frac{S_{xy}}{S_x^2} = \frac{7.743}{30.69} = 0.252$$

$$\text{and } a = \bar{y} - b \bar{x} = 6.33 - 0.252 \times 13.9 = 2.8272$$

Therefore, the straight line is $y = 2.8272 + 0.252 x$

Two Regression Lines

There are two regression lines; regression line of y on x and regression line of x on y . In the regression line of y on x , y is the dependent variable and x is the independent variable and it is used to predict the value of y for a given value of x . But in the regression line of x on y , x is the dependent variable and y is the independent variable and it is used to predict the value of x for a given value of y .

The regression line of y on x is given by

$$y - \bar{y} = \frac{S_{xy}}{S_x^2} (x - \bar{x})$$

and the regression line of x on y is given by

$$x - \bar{x} = \frac{S_{xy}}{S_y^2} (y - \bar{y})$$

Regression Coefficients

The quantity $\frac{S_{xy}}{S_x^2}$ is the regression coefficient of y on x and is denoted by b_{yx} , which gives the slope of the line. That is, $b_{yx} = \frac{S_{xy}}{S_x^2}$ is the rate of change in y for the unit change in x .

The quantity $\frac{S_{xy}}{S_y^2}$ is the regression coefficient of x on y and is denoted by b_{xy} , which gives the slope of the line. That is, $b_{xy} = \frac{S_{xy}}{S_y^2}$ is the rate of change in x for the unit change in y .

13.12 Let us Sum Up

In this Lesson the concept of correlation and regression are discussed. The correlation is the association between two variables. A scatter plot of the variables may suggest that the two variables are related but the value of the Pearson's correlation coefficient r quantifies this association. The correlation coefficient r may assume values from -1 and $+1$. The sign indicates whether the association is direct (+ve) or inverse (-ve). A numerical value of 1 indicates perfect association while a value of zero indicates no association. Regression is a device for establishing relationships between variables from the given data. The discovered relationship can be used for predictive purposes. Some simple examples are shown to understand the concepts.

13.13 Lesson – End Activities

1. Define correlation, Regression.
2. Give the purpose of drawing scatter diagram.

13.14 References

1. P.R. Vital – Business Mathematics and Statistics.
2. Gupta S.P. – Statistical Methods.

UNIT - V

Lesson 14 - Time Series

Contents

- 14.1 Aims and Objectives
- 14.2 Definition of a time series
- 14.3 Time series cycle
- 14.4 Time series models
- 14.5 Time series analysis
- 14.6 Standard time series models
- 14.7 Description of time series components
- 14.8 Graphing a time series
- 14.9 Let us Sum Up
- 14.10 Lesson – End Activities
- 14.11 References

14.1 Aims and Objectives

This Lesson defines a time series and describes the structure (called the time series model) within which time series' movements can be explained and understood. The various components that go to make up each time series value are then discussed and, finally, brief mention is made of graphical techniques.

14.2 Definition of a time series

A time series is the name given to the value of some statistical variables measured over a uniform set of time points. Any business, large or small, will need to keep records of such things as sales, purchases, value of stock held and VAT and these could be recorded daily, weekly, monthly, quarterly or yearly. These are examples of time series.

A time series is a name given to numerical data that is described over a uniform set of time points.

Time series occur naturally in all spheres of business activity as demonstrated in the following example.

Example 1 (Situations in which time series occur naturally)

- a) Annual turnover of a firm for ten successive years.
- b) Numbers unemployed (in thousands) for each quarter of four successive years.
- c) Total monthly sales for a small business for three successive years.

- d) Daily takings for a supermarket over a two month period.
 e) Number of registered journeys for a Home Removals firm (see table below)

	Qtr 1	Qtr 2	Qtr3	Qtr 4
Year 1	73	90	121	98
Year 2	69	92	145	107
Year 3	86	111	157	122
Year 4	88	109	159	131

14.3 Time series cycle

Normally, time series data exhibits a general pattern which broadly repeats, called a cycle. Sales of domestic electricity always have a distinct four-quarterly cycle; monthly sales for a business will exhibit some natural 12-monthly cycle; daily takings for a supermarket will display a definite 6-daily cycle. The cycle for the Home Removals data in above can be seen to be 4-quarterly.

14.4 Time series models

Business records, and in particular certain time series of sales and purchases, need to be kept by law. Of course they are also used to help control current (and plan future) business activities. To use time series effectively for such purposes, the data have to be organized and analysed. In order to explain the movements of time series data, models can be constructed which describe how various components combine to form individual data values.

As an example, a Sales Manger could set up the following model to explain the expense claims of his sales force each week:

$$y = f + t$$

Where, y = total expenses for week,

f = fixed expenses (meals, insurance etc), and

t = travelling expenses (petrol, car maintenance, incidentals, etc.)

14.5 Time series analysis

It is the evaluation and extraction of components of a model that 'break down' a particular series into understandable and explainable portions and enables :

- Trends to be identified.
- Extraneous factors to be eliminated and
- Forecasts to be made

The understanding, description and use of these processes is known as time series analysis.

14.6 Standard time series models

Depending on the nature, complexity and extent of the analysis required, there are various types of model that can be used to describe time series data. However, for the purposes of this manual, two main models will be referred to. They are known as the simple additive and multiplicative models.

The components that go to make up each value of a time series are described in the following definitions.

The time series additive model

$$y = t + s + r$$

where, y is a given time series value

t is the trend component

s is the seasonal component

r is the residual component.

The time series multiplicative model

$$y = t \times S \times R$$

where, y is a given time series value

t is the trend component

S is the seasonal component

R is the residual component.

Put another way, given a set of time series data, every single given (y) value can be expressed as the sum or product of three components. It is the evaluation and interpretation of these components that is the main aim of the overall analysis.

Note that although the trend component will be constant no matter which of the two models are used, the values of the seasonal and residual components will depend on which model is being used. In other words, given a set of data to which both models are being applied, both trend values would be identical whereas the respective seasonal and residual components would be quite different.

14.7 Description of time series components

a) **Trend.** The underlying, long-term tendency of the data.

b) **Seasonal variation.** These are short-term cyclic fluctuations in the data about the trend which take their name from the standard business quarters of the year. Note however that the word 'season' in this context can have many different meanings. For example:

- i. daily 'seasons' over a weekly cycle for sales in a supermarket,
- ii. monthly 'seasons' over a yearly cycle for purchases of a company,
- iii. quarterly 'seasons' over a yearly cycle for sales of electricity in the domestic sector.

c) **Residual variation.** These include other factors not explained by a) and b) above. This variation normally consists of two components:

- i. **Random factors.** These are disturbances due to ‘everyday’ unpredictable influences, such as weather conditions, illness, transport breakdowns, and so on.
- ii. **Long-term cyclic factor.** This can be thought of (if it exists) as due to underlying economic causes outside the scope of the immediate environment. Examples are standard trade cycles or minor recessions.

Example 2 (general comments on a given time series)

Comment on the following data, which relates to visitors (in hundreds) to a hotel over a period of three years. Do not use any quantitative techniques or analyses.

	Qtr 1	Qtr 2	Qtr 3	Qtr 4
Year 1	57	85	97	73
Year 2	64	96	107	89
Year 3	76	102	115	95

Answer

The data displays a distinct 4-quarterly cycle over the three year period, with the underlying trend showing a steady increase overall, as well as in each particular quarter. It shows a significant seasonal effect with (not unexpectedly) the cycle peak in the summer quarter and a trough in the winter quarter. Increases are significantly less in the second and third quarters from year 2 to year 3, which may be due to an upper capacity limit in accommodation for those periods or some other random factor. There is not enough data to identify and possible long-term cyclic factors.

14.8 Graphing a time series

- a) The standard graph for a time series is a line diagram, known technically as a historigram. It is obtained by plotting the time series values (on the vertical axis) against time (on the horizontal axis) as single points which are joined by straight line segments.
- b) **Historigrams** can be shown on their own but it is quite common to see both a historigram and the graph of associated derived data, such as a trend, plotted together on the same chart.

Exercise

1. What is a time series ?
2. What are the aims of time series analysis ?
3. Describe the simple additive time series model and name its components.
4. Describe what a ‘season’ is in the context of a time series and give some examples.
5. For an additive time series model, what does the term ‘residual variation’ mean? Describe briefly its two main constituents.
6. What might contribute towards random variation for data pertaining to daily sales in a supermarket over a period of four weeks. Try to list at least six factors.
7. Graph the following data and comment on significant features.

Sales of a company (Rs.000)

	Qtr 1	Qtr 2	Qtr 3	Qtr 4
1982	19	31	62	9
1983	20	32	65	17

1984	24	36	78	14
1985	24	39	83	20
1986	25	42	85	24

14.9 Let us Sum Up

In this Lesson, we have discussed about a time series which is a set of data that is described over a uniform set of time points. Cycles are general patterns that repeat and occur in most types of time series. Time series models are used to gain an understanding of the factors that effect time series. The time series additive model describes the way that the trend, seasonal and residual components independently make up each time series value. A historigram is the standard way of displaying a time series diagrammatically. The applications of time series is obviously occurring while analyzing sales data, marketing related data, advertisement pattern and costs, inventory analyis, etc.

14.10 Lesson – End Activities

1. Define time series.
2. How to graph a time series?

14.11 References

1. Gupta S.P. – Statistical Methods.

Lesson 15 - Time Series Trend

Contents

- 15.1 Aims and Objectives
- 15.2 The significance of trend values
- 15.3 Techniques for extracting the trend
- 15.4 The method of semi-averages
- 15.5 Working data (for rest of the Lesson)
- 15.6 The method of least squares regression
- 15.7 The method of moving averages
- 15.8 Moving average centering
- 15.9 Comparison of techniques for trend
- 15.10 Let us Sum Up
- 15.11 Lesson – End Activities
- 15.12 References

15.1 Aims and Objectives

This Lesson describes the significance of trend values and the three most common methods of extracting a trend from a given time series. Each method is demonstrated using a common time series and the results compared graphically. Significant features of the three techniques are listed, including their advantages and disadvantages.

15.2 The significance of trend values

It will be recalled from the previous Lesson that the object of finding the time series trend is to enable the underlying tendency of the data to be highlighted. Thus, a business sales trend will normally show whether sales are moving up or down (or remaining static) in the long term.

The trend can also be thought of as the core component of the additive time series model about which the two other components, seasonal (s) and residual (r) variation, fluctuate. This component is found by identifying separate trend (f) values, each corresponding to a time point. In other words, at each time point of the series, a value of t can be obtained which forms one of the components that go to make up the observed value of y. The following section summarizes three different ways of obtaining trend values for a given time series.

15.3 Techniques for extracting the trend

There are three techniques that can be used to extract a trend from a set of time series values.

- a) **Semi-averages.** This is the simplest technique, involving the calculation of two (x,y) averages which, when plotted on a chart as two separate points and joined up, form a straight line. A similar method was introduced in Lesson 15, to find a regression line.
- b) **Least squares regression.** This method, also introduced in Lesson 15 similarly results in a straight line.

- c) **Moving averages.** This is the most commonly used method for identifying a trend and involves the calculation of a set of averages. The trend, when obtained and charted, consists of straight line segments.

15.4 The method of semi-averages

The method of semi-averages for obtaining a trend for a time series is now demonstrated with a simple example.

Suppose the following sales (Rs. in 1000) were recorded for a firm and it is required to obtain a semi-average trend.

	Week 1					Week 2				
	Mon	Tue	Wed	Thu	Fri	Mon	Tue	Wed	Thu	Fri
Sales(y)	250	320	340	520	410	260	380	410	670	420

Note that the data is time-ordered, which is normal and natural for a time series. The procedure for obtaining a trend using the method of semi-averages is:

- STEP 1 Split the data into a lower and an upper group.
For the data given:
the lower group is 250,320,340,520 and 410;
the upper group is 260,380,410,670 and 420.
- STEP 2 Find the mean value of each group.
The mean of the lower group (L) is $1840/5 = 368$.
The mean of the upper group (U) is $2140/5 = 428$.
- STEP 3 Plot, on a graph, each mean against an appropriate time point.
'An appropriate time point' can always be taken as the median time point of the respective group. Thus L would be plotted against Wednesday of week 1 and U against Wednesday of week 2.
- STEP 4 The line joining the two plotted points is the required trend.
Note that it is important that the two groups in question have an equal number of data values. If the given data, however, contains an odd number of data values, the middle value can be ignored (for the purposes of obtaining the trend line).

Once a trend line has been obtained, the trend values corresponding to each time point can be read off from the graph.

A fully worked example follows.

15.5 Working data (for rest of the Lesson)

The following set of data will be referred to throughout the Lesson in order to demonstrate the calculations involved in using each of the three methods for obtaining a time series trend.

	UK outward passenger movements by sea											
	Year 1				Year 2				Year 3			
Quarter	1	2	3	4	1	2	3	4	1	2	3	4

Number of 2.2 5.0 7.9 3.2 2.9 5.2 8.2 3.8 3.2 5.8 9.1 4.1
Passengers (millions)

Example 1 (calculating a time series trend using semi-averages)

Question

Using the working data, given above:

- a) Use the method of semi-averages to obtain and plot a trend line.
- b) Draw up a table showing the original data (y) values against the trend (t) values (obtained from the graph).

Answer

- a) The data has been split up into lower and upper groups, each one being totaled and then averaged.

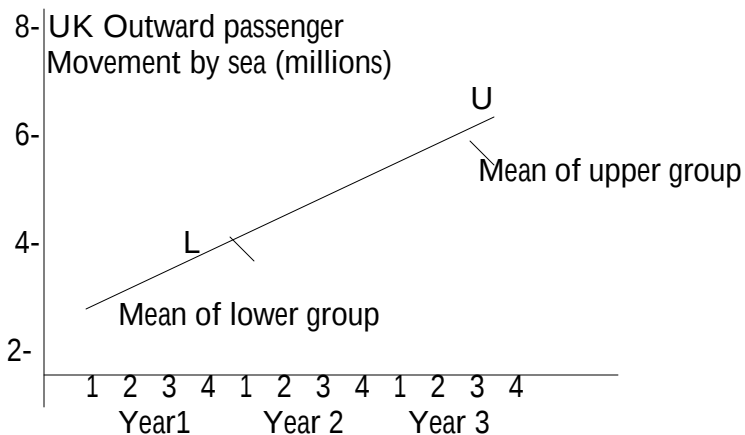
Year 1	Q1	2.2	Year 2	Q3	8.2
	Q2	5.0		Q4	3.8
	Q3	7.9	Year 3	Q1	3.2
	Q4	3.2		Q2	5.8
Year 2	Q1	2.9		Q3	9.1
	Q2	2.2		Q4	4.1
	Total	26.4		Total	34.2
	Mean(L)	4.4		Mean(U)	5.7

In this situation, both L and U must be plotted against a hypothetical point between the middle two time points in their respective sets. That is, L is plotted at a time point between Year 1 Q3 and Year 1 Q4 and L is plotted corresponding to a point between Year 3 Q1 and Year 3 Q2.

In Figure 1, the two means have been plotted and joined by a straight line to form the trend line.

- b) The trend values have been read from the graph and are tabulated below, together with the original data values.

	Year 1				Year 2				Year 3			
Quarter	1	2	3	4	1	2	3	4	1	2	3	4
Data(y)	2.2	5.0	7.9	3.2	2.9	5.2	8.2	3.8	3.2	5.8	9.1	4.1
Trend(t)	3.9	4.1	4.3	4.5	4.7	4.9	5.2	5.4	5.6	5.8	6.0	6.2



15.6 The method of least squares regression

The technique of least squares regression was explained and demonstrated in earlier Lesson. In order to use this method to obtain a trend line for a time series, it is necessary to consider the time series data as bivariate. The procedure is given as follows.

- STEP 1 Take the physical time points as values (coded as 1,2,3 etc if necessary) of the independent variable x .
- STEP 2 Take the data values themselves as values of the dependent variable y .
- STEP 3 Calculate the least squares regression line of y on x , $y=a+bx$.
- STEP 4 Translate the regression line as $t=a+bx$, where any given value of time point x will yield a corresponding value of the trend, t .

An example of the use of this technique follows.

Example 2 (calculating a time series trend using least squares)

Question

For the working data of section 5, calculate, using least squares regression, a trend component for each time point given.

Answer

Put y =number of passengers and x = time point, coded from 1 to 10. i.e. 1=Year 1(Qtr 1) and 10=Year 3(Qtr 2)

x	y	xy	x^2	trend (t)
1	2.2	2.2	1	4.11
2	5.0	10.0	4	4.28
3	7.9	23.7	9	4.45
4	3.2	12.8	16	4.62
5	2.9	14.5	25	4.79
6	5.2	31.2	36	4.96
7	8.2	57.4	49	5.13
8	3.8	30.4	64	5.30

9	3.2	28.8	81	5.47
10	5.8	58.0	100	5.64
11	9.1	100.1	121	5.81
12	4.1	49.2	144	5.98
78	60.6	418.3	650	

From the table : $x=78$; $y=60.6$; $xy=418.3$; $x^2=650$; $n=12$.

Putting the regression line as $y=a+bx$, a and b are now calculated.

$$\text{Thus: } b = \frac{n \sum xy - \sum x \sum y}{n \sum x^2 - (\sum x)^2} = \frac{12 \times 418.3 - 78 \times 60.6}{12 \times 650 - 78^2}$$

$$= \frac{292.8}{1716}$$

$$\text{i.e. } b = 0.17 \text{ (2D)}$$

$$\text{and: } a = \frac{\sum y}{n} - b \frac{\sum x}{n} = \frac{60.6}{12} - 0.17 \times \frac{78}{12}$$

$$\text{i.e. } a = 3.94 \text{ (2D)}$$

thus, the regression line for the trend is $t = 3.94 + (0.17)(x)$ (2D)

(Remember that once the regression line is determined, it will be used for calculating trend values. So the normal 'y' has been replaced by 't')

The time point values ($x=1,2,3$ etc) can now be substituted into the above regression line to give the trend values required.

When $x=1$ (Year 1 Qtr1), $t=3.94+0.17(1)$ i.e., $t=4.11$ (2D)

When $x=2$ (Year 2 Qtr2), $t=3.94+0.17(2)$ i.e., $t=4.28$ (2D) ...etc.

These and other values of t are tabulated in the previous table.

15.7 The method of moving averages

This method of obtaining a time series trend involves calculating a set of averages, each one corresponding to a trend (t) value for a time point of the series. These are known as moving averages, since each average is calculated by moving from one overlapping set of values to the next. The number of values in each set is always the same and is known as the period of the moving average.

To demonstrate the technique, a set of moving averages of period 5 has been calculated below for a set of values.

Original values:	12	10	11	11	9	11	10	10	11	10
Moving totals:			53	52	52	51	51	52		
Moving averages:			10.6	10.4	10.4	10.2	10.2	10.4		

The first total, 53, is formed from adding the first 5 items; i.e. $53=12+10+11+11+9$. Similarly, the second total, $52=10+11+11+9+11$, and so on. The averages are then obtained by dividing each total by 5.

Notice that the totals and averages are written down in line with the middle value of the set being worked on. These averages are the trend (t) values required.

It should also be noticed that there are no trend values corresponding to the first and last two original values. This is always the case with moving averages and is a disadvantage of this particular method of obtaining a trend.

15.7.1 Let us Sum Up of the moving average technique

Moving averages (of period n) for the values of a time series are arithmetic means of successive and overlapping values, taken n at a time.

The (moving) average values calculated form the required trend components (t) for the given series.

The following points should be noted when considering a moving average trend.

- a) The period of the moving average must coincide with the length of the natural cycle of the series. Some examples follows.
 - i. Moving averages for the trend of numbers unemployed for the quarters of the year must have a period of 4.
 - ii. Total monthly sales of a business for a number of years would be described by a moving average trend of period 12.
 - iii. A moving average trend of period 6 would be appropriate to describe the daily takings for a supermarket (open six days per week) over a number of months.
- b) Each moving average trend value calculated must correspond with an appropriate time point. This can always be determined as the median of the time points for the values being averaged. For moving averages with an odd-numbered period, 3,5,7, etc, the relevant time point is that corresponding to the 2nd, 3rd, 4th, etc value. See the example in the previous section, where the moving averages had a period of 5 and thus each average obtained was set against the 3rd value of the respective set being averaged.

However, when the moving averages have an even-numbered period (2,4,6,8,etc). There is no obvious and natural time point corresponding to each calculated average. The following section describes the technique known as ‘centering’, which is used in these circumstances.

15.8 Moving average centering

When calculating moving averages with an even period (i.e. 4,6 or 8), the resulting moving average would seem to have to be placed in between two corresponding time points. As an example, the following data has a 4-period moving average calculated and shows its placing

Time point	1	2	3	4	5	6	7	8	9	10
Data value	9	14	17	12	10	14	19	15	10	16
Totals(of 4)			52	53	53	55	58	58	60	
Averages (of 4)			13.00	13.25	13.25	13.75	14.50	15.00		

The placing of these averages as described above would not be satisfactory when the averages are being used to represent a trend, since the trend values need to coincide with particular time points. A method known as centering is used in this type of situation, where the calculated averages are themselves averaged in successive overlapping pairs.

This ensures that each calculated (trend) value 'lines up' with a time point.

This technique is now shown for the previous data.

Time point	2	3	4	5	6	7	8	9
Averages(of 4)		13.00	13.25	13.25	13.75	14.50	14.50	15.00
Averages (of 2)		13.125	13.250	13.500	14.125	14.500	14.750	

A worked example follows which uses this technique.

Example 3 (calculating trend values using moving average centering)

Question

Calculate trend values for the working data of section 5, using moving averages with an appropriate period. Plot a graph of the original data with the trend superimposed.

Answer : The cycle of the data is clearly 4-quarterly and we thus need a (centered) 4-quarterly moving average trend, using the technique described in section 11 above. Table 1 demonstrates the standard columnar layout of the calculations.

Qtr		Original Data(y)	Moving totals of 4	Moving average	Centered moving average(t)
Year 1	1	2.2			
	2	5.0			
			18.3	4.575	4.66
	3	7.9	19.0	4.750	4.78
Year 2	4	3.2	19.2	4.800	4.84
	1	2.9	19.5	4.875	4.95
	2	5.2	20.1	5.025	5.06
	3	8.2	20.4	5.100	5.18
Year 3	4	3.8	21.0	5.250	5.18
					5.36
	1	3.2	21.9	5.475	
	2	5.8	22.2	5.550	5.51
	4	4.1			

Table 1

Notice that the two starting and ending time points do not have a trend value. As mentioned previously, this type of omission will always occur with a moving average trend.

Figure 2 shows a graph of the original data with the trend values superimposed.

15.9 Comparison of techniques for trend

For the working data given in section 4, all three methods of obtaining a trend have now been demonstrated. The method of semi-averages (Example 1), least squares (Example 2) and moving averages (Example 3). Figure 3 shows the graphs for comparison.

The fact that the three sets of trend values are quite distinct underlines the fact that there is no unique set of trend values for a time series. Each method will yield a different trend, as has been evidenced.

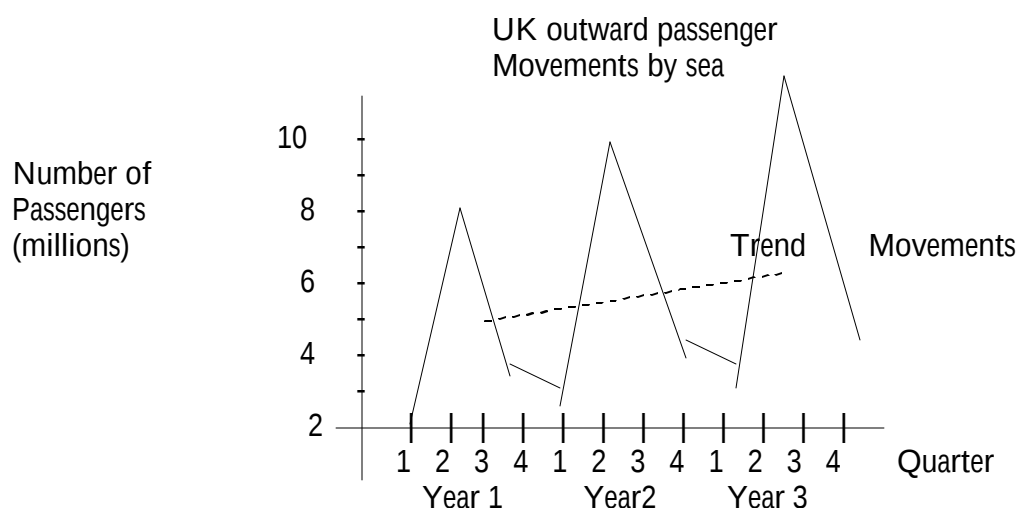


Figure 2

Significant features of each method are now summarized.

- Semi-averages.** Although simple to apply, the fact that only two plotted points are used in its construction leads to the general feeling that it is unrepresentative. It also assumes that a strictly linear trend is appropriate to the data.
- Least squares.** Although mathematically representative of the data, it assumes that a linear trend is appropriate. It is generally though unsuitable for highly 'seasonal' data.
- Moving averages.** The most widely used technique for obtaining a trend. If the period of the averages is chosen appropriately, it will show the true nature of the trend, whether linear or non-linear. One disadvantage is the fact that no trend values are obtained for the beginning and end time points of a series.

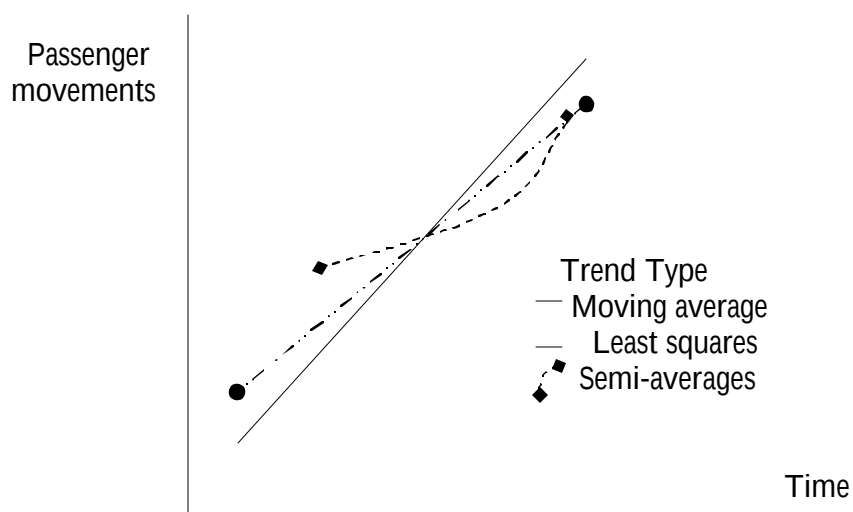


Figure 3

Exercises

1. Calculate a set of trend values (to ID) using the method of semi-averages, for the following data:

16, 12, 15, 14, 18, 12, 14, 13, 18, 13.

2. Calculate a set of moving averages of period: (a) 3 (b) 5 for the following time series data:
8, 11, 10, 21, 4, 9, 12, 10, 23, 5, 10, 13, 11, 26, 6.

Which set of moving averages is the correct one to use for obtaining a trend for the series?

3. Draw a histogram for the data described in question 2 above, superimposing the correct trend values.

4. The number of houses (in thousands) built each year between 1953 and 1969 (inclusive) are given as:

Year	1	2	3	4	5	6	7	8	9
Number of houses	319	348	317	308	308	329	332	354	378
Year	10	11	12	13	14	15	16	17	
Number of houses	364	358	383	391	396	415	426	378	

Assuming a seven-year cycle, eliminate the cyclical movement by producing a moving average trend and plot this, together with the original data on the same chart.

5. The following figures relate to Rate receipts (in £m) for a Local Authority.

	Year1	Year2	Year3
Qtr1	2.8	3.0	3.0
Qtr2	4.2	4.2	4.7
Qtr3	3.0	3.5	3.6
Qtr4	4.6	5.0	5.3

Plot a histogram for the data, together with a least squares regression trend

15.10 Let us Sum Up

In this Lesson the time series trend is discussed and three common techniques for identifying trend components are discussed. They are : (i) semi-averages (ii) least squares regression and (iii) moving averages. For time series that have a significant seasonal effect, the moving average technique is generally preferred. When moving averages are used for identifying trend components, the period of the average must coincide with the cycle of the data being analysed. This is done in order to remove possible cyclical fluctuations. Even-period moving averages must be centered in order that their values coincide with actual time points. It is also to be noted that there is no unique set of trend values for a given time series. The particular method chosen needs to take into account the nature of the data and the use to which trend values will be put.

15.11 Lesson – End Activities

1. What is meant by moving average?

15.12 References

R.S.N. Pillai and Mrs. Bhagavathi – Statistics.

Lesson 16 - Seasonal Variation and Forecasting

Contents

- 16.1 Aims and Objectives**
- 16.2 The nature of seasonal variation**
- 16.3 Technique for calculating seasonal variation**
- 16.4 Seasonally adjusted time series**
- 16.5 Notes on Example 3**
- 16.6 Forecasting**
- 16.7 Technique for forecasting**
- 16.8 Projecting the trend**
- 16.9 Forecasting and residual variation**
- 16.10 Let us Sum Up**
- 16.11 Lesson – End Activities**
- 16.12 References**

16.1 Aims and Objectives

The Lesson described the nature of seasonal variation in a time series and how it can be calculated. Forecasting, or the ability to estimate future values of a given time series using seasonal variation, is dealt with in this Lesson.

16.2 The nature of seasonal variation

Seasonal (or short-term cyclic) variation is present in many time series. Winter sports-wear will sell well in autumn and winter, and badly in spring and summer; supermarket sales are higher at the end of the week than at the beginning; sales of umbrellas are at the peak during the end of the summer and just at the beginning of the rainy season, etc.

When values are obtained to describe seasonal variation, they are sometimes known as seasonal values or factors and are expressed as deviations (i.e. '+' or '-') from the underlying trend. They show, on average, by how much a particular season will tend to increase or decrease the underlying trend. Thus we would expect the seasonal variation for winter sportswear to be positive in autumn and winter and negative in spring and summer.

Seasonal variation components give an average effect on the trend which is solely attributable to the 'season' itself. They are expressed in terms of deviations from (additive model) or percentages of (multiplicative model) the trend.

The use of seasonal variation figures are of great importance to organizations operating in environments where a seasonal factor is significant. For example, a regional Electricity Board needs to know the average increase in demand expected in the winter months in order to be able to meet this demand. The following two sections describe and demonstrate the technique for calculating seasonal variation.

16.3 Technique for calculating seasonal variation

a) Additive model

Given the original time series (y) values, together with the trend (t) values, the procedure for calculating the seasonal variation is given as follows.

STEP 1 Calculate, for each time point, the value of $y-t$ (the difference between the original value and the trend).

STEP 2 For each season in turn, find the average (arithmetic mean) of the $y-t$ values.

STEP 3 If the total of the averages differs from zero, adjust one or more of them so that their total is zero.

The values so obtained are the appropriate seasonal variation values; i.e. the 's' figures in the additive model $y = t + s + r$.

b) Multiplicative model

Given the original time series (y) values, together with the trend (t) values, the procedure for calculating the seasonal variation is given as follows.

STEP 1 Calculate, for each time point, the value of $(y-t)/t$ (the difference between the original value and the trend expressed as a proportion of the trend).

STEP 2 For each season in turn, find the arithmetic mean of the above proportional changes.

[Note that this should strictly involve calculating the geometric mean of $1 +$ proportional change values. In practice however this is felt to be too complex!]

STEP 3 If the total of the averages differs from zero, adjust one or more of them so that their total is zero.

The values so obtained are the appropriate seasonal variation values; i.e. the 'S' figures in the multiplicative model $y = t + S + R$.

Example 1 (Calculating seasonal variation figures using the additive model)

The sales of a company (y, in Rs. 000) are given below, together with a previously calculated trend (t). The subsequent calculations to find the seasonal variation are shown, laid out in a standardized way.

STEP 1				STEP 2					
	y	t	y-t	Deviations (y-t)					
				Q1	Q2	Q3	Q4	Sum	
Year 1 Qtr 1	20	23	-3	Year1	-3	-14	26	-9	
2	15	29	-14		-10	-25	45	-11	
3	60	34	26						
4	30	39	-9						
Year 2 Qtr 1	35	45	-10	Totals	-13	-39	71	-20	
2	25	50	-25	Averages	-6.5	-19.	35.5	-10.0	-0.5
3	100	55	45						
4	50	61	-11						

STEP 3

Since the averages sum to -0.5 (and not zero), it is necessary to adjust one or more of them accordingly. In this case, since the difference is so small, only one will be adjusted. In order to make the smallest percentage error, the largest value (35.5) is changed to 36.0. this adjustment is shown in the following table:

	Q1	Q2	Q3	Q4
Initial s values	-6.5	-19.5	35.5	-10.0
Adjustment	0	0	+0.5	0
Adjusted s values	-6.5	-19.5	36.0	-10.0 (Sum = 0)

The interpretation of the figures is that the average seasonal effect for quarter 1, for instance, is to deflate the trend by 6.5 (Rs. 000) and that for quarter 3 is to inflate the trend by 36 (Rs. 000).

Example 2 (Calculating seasonal variation figures using the multiplicative model)

The sales of a company (y, in Rs. 000) are given below, together with a previously calculated trend (t). The subsequent calculations to find the seasonal variation are shown, laid out in a standardized way.

Step 1					Step 2					
	y	t	$\frac{y-t}{t}$	$s=1+\frac{y-t}{t}$	Deviations $\frac{(1+y-t)}{t}$					
					Q1	Q2	Q3	Q4	Sum	
Year 1 Qtr1	20	23	-0.13	0.87	Year1	0.87	0.52	1.76	0.77	
2	15	29	-0.48	0.52						
3	60	34	0.76	1.76						
4	30	39	-0.23	0.77						
Year2 Qtr 1	35	45	-0.22	0.78	Year 2	0.78	0.50	1.82	0.82	
2	25	50	-0.50	0.50	G. Means	0.82	0.51	1.79	0.79	3.91
3	100	55	0.82	1.82						
4	50	61	-0.18	0.82						

STEP 3

Since the averages sum to 3.91 (and not 4), it is necessary to add 0.09 to one or more of them accordingly. In this case, as in the previous Example, since the difference is so small, only one will be adjusted. In order to make the smallest percentage error, the largest value (1.79) is changed to 1.88. This adjustment is shown in the following table.

	Q1	Q2	Q3	Q4
Initial S values	0.82	0.51	1.79	0.79
Adjustment	0	0	+0.9	0
Adjusted S values	0.82	0.51	1.88	0.79 (Sum = 4.00)

The interpretation of the figures is that the average seasonal effect for quarter 1, for instance, is to deflate the trend by 18% (since 0.82 is 0.18 less than 1) and that for quarter 3 is to inflate the trend by 88%.

16.4 Seasonally adjusted time series

One particular and important use of seasonal values is to seasonally adjust the original data. The effect of seasonal adjustment is to smooth away seasonal fluctuations, leaving a clear view of what might be expected 'had seasons not existed'.

The techniques is similar for both models but is shown separately for clarity.

Additive model:

The adjustment is performed by subtracting the appropriate seasonal figure from each of the original time series values and represented algebraically by $y-s$.

As an example, the data of Examples 1 and 2 are seasonally adjusted below.

	Y	s	y-s	
Year1 Qtr 1	20	-6.5	$20-(-6.5)=26.5$	
2	15	-19.5	$15-(-19.5)=34.5$	
3	60	36.0	$60-36.0=24.0$	
4	30	-10.0	$30-(-10.0)=40.0$	Seasonal
Year2 Qtr 1	35	-6.5	$35-(-6.5)=41.5$	adjusted values
2	25	-19.5	$25-(-19.5)=44.5$	
3	100	36.0	$100-36.0=64.0$	
4	50	-10.0	$50-(-10.0)=60.0$	

Multiplicative model:

The adjustment is performed by dividing each of the original time series values by S and is represented algebraically by y/S .

As an example, the data of Example 1 are again seasonally adjusted below.

	Y	S	y/S	
Year1 Qtr 1	20	0.82	$20/0.82=24.3$	
2	15	0.51	$15/0.51=29.5$	
3	60	1.88	$60/1.88=31.9$	
4	30	0.79	$30/0.79=37.8$	Seasonally
Year2 Qtr 1	35	0.82	$35/0.82=42.6$	adjusted values
2	25	0.51	$25/0.51=49.2$	
3	100	1.88	$100/1.88=53.2$	
4	50	0.79	$50/0.79=63.0$	

Seasonally adjusted time series data are obtained by subtraction (additive model) or division (multiplicative model) as follows:

Multiplicative model: seasonally adjusted value = y/s .

Example 3 (Seasonal adjustment of a time series)

The following data gives UK outward passenger movements (in millions) by sea, together with a 4-quarterly moving average trend (calculated previously in the earlier Lesson). Find the values of the seasonal variation for each of the four quarters (using an additive model) and hence obtain seasonally adjusted outward passenger movements. Plot the result on a graph.

Answer :

The deviations are calculated and displayed in column 5, and the calculations for the seasonal variation are shown in the lower table and the results, together with the seasonally adjusted data, have been added at column 6 and 7.

	Original data (y)	Centered moving average (t)	Deviations (y-t)	Seasonal variation (s)	Seasonally adjusted data (y-s)
Year1 Qtr 1	2.2			-2.03	4.23
2	5.0			0.28	4.72
3	7.9	4.66	3.24	3.21	4.69
4	3.2	4.78	-1.58	-1.46	4.66
Year2 Qtr 1	2.9	4.84	-1.94	-2.03	4.93
2	5.2	4.95	0.25	0.28	4.92
3	8.2	5.06	3.14	3.21	4.92
4	3.8	5.18	-1.38	-1.46	5.26
Year3 Qtr 1	3.2	5.36	-2.16	-2.03	5.23
2	5.8	5.51	0.29	0.28	5.52
3	9.1			3.21	5.89
4	4.1			-1.46	5.56

	Q1	Q2	Q3	Q4	Sum
Year 1			3.24	-1.58	
Year 2	-1.94	0.25	3.14	-1.38	
Year 3	-2.16	0.29			
Totals	-4.10	0.54	6.38	-2.96	
Averages	-2.05	0.27	3.19	-1.48	-0.07
Adjustments	+0.02	+0.01	+0.02	+0.02	
Adjusted averages	-2.03	0.28	3.21	-1.46	0.00

2

The required graph is plotted in Figure 1.

UK outward passenger movements by sea

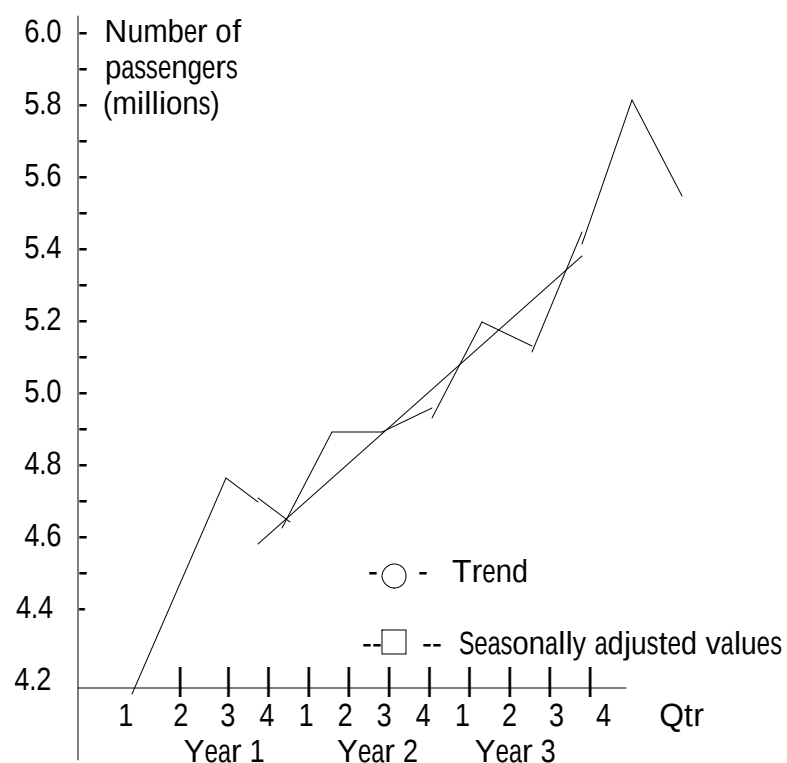


Figure 1

16.5 Notes on Example 3

1. It is usual to show the calculation of the seasonal values in rectangular form as demonstrated above.
2. Notice that the adjustment needed above was +0.07. However, rather than adding all of this to just one of the averages, it was divided up into the four parts +0.02, +0.02,

+0.02 and +0.01, each being added to a separate average. This is generally regarded as a fairer way to adjust.

3. Even though the moving average trend values are missing at the beginning and end time points, the seasonal values calculated can still be used at these points and thus seasonal adjustment can be performed for all original data items.

16.6 Forecasting

a) A particular use of time series analysis is in forecasting, sometimes called projecting the time series. Clearly, business life would be much easier if monthly sales for the next year were known or the number of transport breakdowns next month could be determined. However, no-one can predict the future; the best that can be done is to estimate the most likely future values, given the analysis of previous years' sales or last month's breakdowns.

b) Forecasting can be performed at different levels depending on the use to which it will be put. Simple guessing, based on previous figures, is occasionally adequate. However where there is a large investment at stake (in plant, stock and manpower for example), structured forecasting is essential.

c) any forecasts made, however technical or structured, should be treated with caution, since the analysis is based on past data and there could be unknown factors present in the future. However, it is often reasonable to assume that patterns that have been identified in the analysis of past data will be broadly continued, at least into the short-term future.

16.7 Technique for forecasting

Forecasting a value for a future time point involves the following steps.

- STEP 1 Estimate a trend value for the time point. There are a number of ways of estimating future trend values and some of these are described in section 12.
- STEP 2 Identify the seasonal variation value appropriate to the time point. Seasonal variation values are calculated in the manner already described in section 5.
- STEP 3 Add (or multiply, depending on the model) these two values together, giving the required forecast.

Time series forecasting can be attempted using the simple additive or multiplicative model in the following adapted form:

$$\text{Additive:} \quad y_{\text{est}} = t_{\text{est}} + S$$

$$\text{Multiplicative:} \quad y_{\text{est}} = t_{\text{est}} \times S$$

Where: y_{est} = estimated data value
 t_{est} = projected trend value
 S = appropriate seasonal variation value.

Notice that there is no provision for residual variation in the above forecasting models.

Example 4 (Time series forecasting)

Forecast the values for the four quarters of year 4, given the following information which has been calculated from a time series. Assume that the trend in year 4 will follow the same pattern as in year 1 to 3 and an additive model is appropriate.

	Year1				Year2				Year3			
Quarter	1	2	3	4	1	2	3	4	1	2	3	4
Trend (t)	42	44	46	48	50	52	54	56	58	60	62	64

S_1 =seasonal factor for quarter 1=-15; s_2 =-8; s_3 =+6; s_4 =+17

STEP 1 Estimate trend values for the relevant time points. Note that, in this case, the trend values increase by exactly 2 per quarter.

Trend for year 4, quarter 1= $t_{4,1}$ =66.

Similarly, $t_{4,2}$ =68, $t_{4,3}$ =70 and $t_{4,4}$ =72

STEP 2 Identify the appropriate seasonal factors. The seasonal factors for year 4 are taken as the given seasonal factors. That is, seasonal factor for year 4, quarter 1= s_1 =-15 etc.

STEP 3 Add the trend estimates to the seasonal factors, giving the required forecasts.

Forecast for year 4, quarter 1= $t_{4,1}+s_1=66-15=51$;

Forecast for year 4, quarter 2= $t_{4,2}+s_2=68-8=60$;

Forecast for year 4, quarter 3= $t_{4,3}+s_3=70+6=76$;

Forecast for year 4, quarter 4= $t_{4,4}+s_4=72+17=89$.

16.8 Projecting the trend

Projecting the trend for the data of Example 3 was straightforward since the given trend values increased uniformly, thereby displaying a distinct linear pattern. In general, trend values will not conform conveniently in this way.

There are a number of techniques available for projecting the trend, depending on the method used in obtaining the trend values themselves. The most common are now listed.

- a) Linear trend. Whether the method of least squares or semi-averages has been used, the projection involves simply extending the trend line already calculated.
- b) Moving average trend. There is no one universal method. Three common means of projecting are listed below.
 - i. 'By eye' (or inspection) from the graph. This involves adding a projection freehand in a manner that seems most appropriate. This might seem fairly arbitrary, but remember that any form of projection (no matter how technical) is still only an estimate. This particular method can be employed when the calculated trend values are distinctly 'non-linear'.
 - ii. Using the method of semi-average on the calculated trend values to obtain a linear projection of the trend. This method can be employed with 'fluctuating linear' trend values.
 - iii. Using the first and last of the calculated trend values to obtain a linear projection of the trend. This method can be employed with fairly 'steady linear' trend values.

Example 5 (Time series forecasting)**Question**

Forecast the four quarterly values for year 4 for the following data, which relates to UK outward passenger movements by sea (in millions). The trend (calculated previously) and the seasonal variation components (using the multiplicative model) are given below.

	Year1				Year2				Year3			
Quarter	1	2	3	4	1	2	3	4	1	2	3	4
Number of Passengers(y)	2.2	5.0	7.9	3.2	2.9	5.2	8.2	3.8	3.2	5.8	9.1	4.1
Trend(t)			4.66	4.78	4.84	4.95	5.06	5.18	5.36	5.51		
Seasonal variation(S):	Qtr1=0.60; Qtr2=1.05; Qtr3=1.65; Qtr4=0.70											

Plot the original values, trend and forecast on a single chart.

Answer

STEP1

Estimate trend values for the relevant time points. Since there is a fairly steady increase in the trend values, demonstrating an approximate linear relationship, method iii (from section 11 (b)) is appropriate for projecting the trend.

Range of trend values = $5.51 - 4.66 = 0.85$

Therefore, average change per time period = $0.85/7 = 0.12$ (approx).

[Note that since there are 8 trend values, there are correspondingly only 7 'jumps' from the first to the last. Hence the divisor of 7 in the above calculation.]

The last trend value given is 5.51 for Year 3 Quarter 2 and this must be used as the base value to which is added the appropriate number of multiples of 0.12.

Thus, the trend estimates are:

$t(\text{Year4Qtr1}) = 5.51 + 3(0.12) = 5.87$; $t(\text{Year4Qtr2}) = 5.51 + 4(0.12) = 5.99$;

$t(\text{Year4Qtr3}) = 5.51 + 5(0.12) = 6.11$; $t(\text{Year4Qtr4}) = 5.51 + 6(0.12) = 6.23$.

STEP 2

Identify the appropriate seasonal factors.

These values are given in the question as:

$S_1 = 0.60$; $S_2 = 1.05$; $S_3 = 1.65$; $S_4 = 0.70$.

STEP 3

Multiply the trend estimates by the respective S values, giving the required forecasts.

$y(\text{Year4Qtr1}) = 5.87 \times 0.60 = 3.51$; $y(\text{Year4Qtr2}) = 5.99 \times 1.05 = 6.30$;

$y(\text{Year4Qtr3}) = 6.11 \times 1.65 = 10.10$; $y(\text{Year4Qtr4}) = 6.23 \times 0.70 = 4.37$.

These values are plotted in Figure 2, along with the original data and trend.

UK outward passenger movements by sea

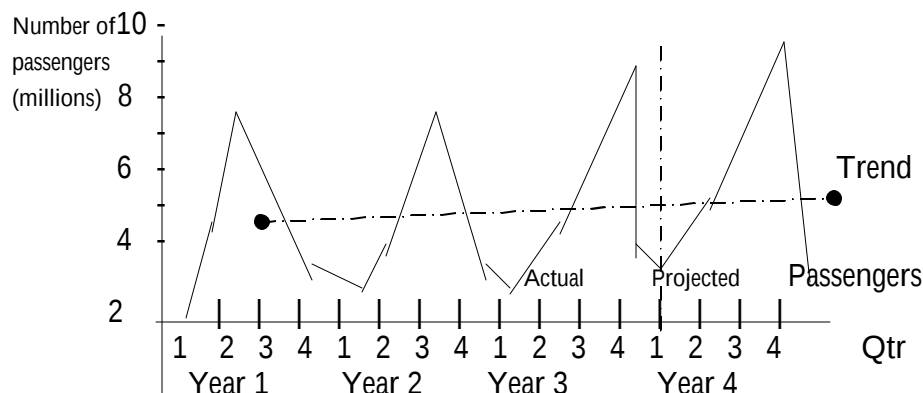


Figure 2

16.9 Forecasting and residual variation

Residual variation is the variation which takes into account everything else other than trend and seasonal factors. In the main it consists of small random fluctuations which, although not controllable, have little effect. If the residual variation is relatively large however, it will make forecasts less dependable, since they effectively ignore residual elements.

Thus, being able to identify a residual element in a time series will normally be a pointer to how reliable any projection will be.

Exercises

1. The following data describes the sales of components for a particular firm:

	Quarters			
	1	2	3	4
Year1				130
Year2	140	160	90	140
Year3	160	170	120	170
Year4	180	200	130	

Seasonally adjust these sales, using:

- an additive model
- a multiplicative model

2. The data below relates to Rate receipts (in Rs. in Lakh) for a Local Authority with a corresponding trend value in brackets.

	1982	1983	1984
Quarter1	2.8(3.3)	3.0(3.7)	3.0(4.2)
Quarter2	4.2(3.4)	4.2(3.9)	4.7(4.3)
Quarter3	3.0(3.5)	3.5(4.0)	3.6(4.4)
Quarter4	4.6(3.6)	5.0(4.1)	5.3(4.5)

Assuming an additive model:

- a) calculate the seasonal variation
- b) estimate the receipts for 1985.

3. The following data describes personal savings as a percentage of earned income for a particular region of the country.

	1980	1981	1982
Quarter1	0.1	12.6	11.9
Quarter2	8.6	7.6	8.7
Quarter3	8.0	7.6	8.3
Quarter4	5.8	6.2	7.2

Use both additive and multiplicative models to seasonally adjust the above percentages and forecast the percentage saving for quarter 1 or 1983. Comment on the results.

16.10 Let us Sum Up

Seasonal factors are of importance of management as a control factor wherever seasonal effects are significant. Seasonal factors: (i) are individually expressed as deviations from (additive model) or percentages of (multiplicative model) the trend; (ii) should collectively sum to either 0 (additive model) or 4 (multiplicative model). Seasonally adjusted values are calculated by subtracting seasonal factors from trend values (additive model); or dividing trend value by seasonal factor (multiplicative model). Seasonally adjusted values are used to eliminate the effect of seasonal variation. Time series forecasting involves adding the appropriate seasonal factors to calculated trend projections (additive model); multiplying the calculated trend projections by the appropriate seasonal factor (multiplicative model).

16.11. Lesson – End Activities

1. Give the importance of forecasting.
2. Describe the techniques for forecasting.

16.12. References

1. Statistical Methods – Gupta S.P.

Lesson 17 -Index relatives

Contents

- 17.1 Aims and Objectives
- 17.2 Definition of an Index Number
- 17.3 Simple index number construction
- 17.4 Some notation
- 17.5 Index relatives
- 17.6 Time series of relatives
- 17.7 Changing the base of fixed-base relatives
- 17.8 Comparing sets of fixed base relatives
- 17.9 Actual and real values of a commodity
- 17.10 Time series deflation
- 17.11 Let us Sum Up
- 17.12 Lesson – End Activities
- 17.13 References

17.1 Aims and Objectives

Index numbers provide a standardized way of comparing the values, over time, of commodities such as prices, volume of output and wages. They are used extensively, in various forms, in Business, commerce and Government. This Lesson introduces index numbers and describes the most simple form; the index relative. Relatives are defined, calculated as time series and compared (using a base-changing techniques). Finally, time series deflation is described, which is a method of calculating an index of the real values of time series.

This Lesson also describes index relatives, the simplest form of index number, and some of the ways that they can be presented and manipulated and composite index numbers, which describe the change over time of groups or classes of commodities that have something in common. The two forms of composite index covered are the weighted average of relatives and the weighted aggregate.

17.2 Definition of an Index Number

An index number measures the percentage change in the value of some economic commodity over a period of time.

It is always expressed in term of a base of 100.

‘Economic commodity’ is a term of convenience, used to describe anything measurable which has some economic relevance. For example: price, quantity, wage, productivity, expenditure, and so on.

Examples of typical index number values are:

125 (an increase of 25%), 90 (a decrease of 10%), 300 (an increase of 200%).

17.3 Simple index number construction

a) Suppose that the price of standard boxes of ball-point pens was Rs. 60 in January and rose to Rs. 63 in April. We can calculate as follows:

$$\text{percentage increase} = \frac{63-60}{60} = \frac{100}{20} = 5$$

In other words, the price of ball-point pens rose by 5% from January to April. To put this into index number form, the 5% increase is added to the base of 100, giving 105. This is then described as follows:

the price index of ball-point pens in April was 105(January = 100).

Note that any increase must always be related to some time period, otherwise it is meaningless. Index numbers are no exception, hence the (January=100) in the above statement, which:

- i. gives the starting point (January) over which the increases in price is being measured;
 - ii. emphasizes the base value (100) of the index number.
- b) If the productivity of a firm (measured in units of production per man per day) decreased by 3% over the period from 1983 to 1985, this percentage would be subtracted from 100 to give an index number of 97. Thus we would say:

'the productivity index for the company in 1985 was 97 (1983=100)'.

17.4 Some notation

a) It is convenient, particularly when giving formulae for certain types of index numbers, to be able to refer to an economic commodity at some general time point. Prices and quantities (since they are commonly quoted indices) have their own special letters, p and q respectively. In order to bring in the idea of time, the following standard convention is used.

Index number notation

p_0 = price at base time point
 p_n = price at some other time point
 q_0 = quantity at base time point
 q_n = quantity at some other time point.

In the example in 5.4.3(a) above, time point 0 was January and time point n was April, with $p_0=60$ and $p_n=63$.

b) It is also convenient on occasions to label index numbers themselves in a compact way. There is no standard form for this but, for example (from section 5.4.3 b), the following is sometimes used:

$$I_{1985}(1983=100)97$$

or

$$I_{1985/1983}=97$$

Which is translated as:

‘the index for 1985, based on 1983 (as 100), is 97’.

17.5 Index relatives

An index relative (sometimes just called a relative) is the name given to an index number which measures the change in a single distinct commodity. A price relative was calculated in section 5.4.3 (a) and a productivity relative was found in section 5.4.3 (b). However, there is a more direct way of calculating relatives than that demonstrated in section 3. the following shows the method of calculating a price and quantity relative.

Price and quantity relatives	
Price relative:	$I_p = \frac{P_n}{P_0} \times 100$
Quantity relative:	$I_Q = \frac{q_n}{q_0} \times 100$

Expenditure and productivity relatives can be calculated in a similar fashion.

Example 1 (Calculation of price and quantity relatives)

The following table gives details of prices and quantities sold of two particular items in a department store over two years.

Item	1984	Number sold	1985	Number sold
	Price		Price	
	P_0	q_0	P_n	q_n
Product I	Rs. 438	37	Rs. 462	18
Product II	Rs. 322	26	Rs. 384	45

We wish to find price and quantity relatives for 1985 (1984=100) for both items.

Year 0=1984 and year n=1985.

For the Product I:

$$\text{Price relative} = I_{85/84} = \frac{p_n}{p_0} \times 100 = \frac{462}{438} \times 100$$

$$\text{Quantity relative} = I_{85/84} = \frac{q_n}{q_0} \times 100 = \frac{18}{37} \times 100$$

$$= 48.6$$

For the Product II :

$$\begin{aligned}\text{Price relative} &= \Pi_{85/84} = \frac{pn}{p0} \times 100 = \frac{384}{322} \times 100 \\ &= 119.3\end{aligned}$$

$$\begin{aligned}\text{Quantity relative} &= \Pi_{85/84} = \frac{qn}{q0} \times 100 = \frac{45}{26} \times 100 \\ &= 173.1\end{aligned}$$

The above calculations and presentation demonstrates typical index number notation.

Thus it can be seen that an index number is a compact way of describing percentage changes over time.

17.6 Time series of relatives

It is often necessary to see how the values of an index relative change over time. Given the values of some commodity over time (i.e a time series), there are two distinct ways in which relatives can be calculated.

a) **Fixed base relatives.** Here, each relative is calculated based on the same fixed time point. This approach can only be used when the basic nature of the commodity is comparing 'like with like'. For example, the price of rice in a supermarket over six monthly periods or weekly family expenditure on entertainment.

b) **Chain base relatives.** In this case, each relative is calculated with respect to the immediately preceding time point. This approach can be used with any set of commodity values, but must be used when the basic nature of the commodity changes over the whole time period. For example, a company might wish to construct a monthly index of total petrol costs of the standard model of car that its salesmen use. However, the model is likely to change yearly with, for instance, different tyres or 'lean-burn' engines being fitted as standard. Both of these features would affect petrol consumption and thus, also, the petrol cost index. Therefore, in this case, a chain base relative should be used.

Example 2 demonstrates the use of the two techniques for the values of a commodity over time.

Example 2 (Fixed and chain base set of relatives for a given time series)

The data in Table 1 relate to the production of beer (thousands of hectoliters) in the United Kingdom for the first six months of a year.

Table 2 shown the calculation of both fixed and chain base relatives, together with some descriptive calculations.

Year	Jan	Feb	Mar	Apr	May	Jun
Production	4,563	4,245	4,841	4,644	5,290	5,166
Table 1						
Fixed base relative (Mar=100)	94.3	87.7	100	95.9	109.3	106.7
chain base relative	-	93.0	114.0	95.9	113.9	97.7
Table 2						
	$\frac{4563}{4841} \times 100$	$\frac{4245}{4563} \times 100$	$\frac{4841}{4245} \times 100$		$\frac{5290}{4841} \times 100$	

In Table 2, the fixed base relative have been calculated by dividing each month's production by the March production (4841) and multiplying by 100. they enable each month's production to be compared with the March production. Thus, for example, May's production (relative=109.3) was 9.3% up on March.

The chain base relatives in Table 2 have been calculated by dividing each month's production by the previous month's production and multiplying by 100. they enable changes from month to month to be highlighted. Thus, for example, February's production (chain relative=93.0) was 7% down on January, March's production (chain relative=114.0) was 14% up on February, and so on.

17.7 Changing the base of fixed-base relatives

Given a time series of relatives, it is sometimes necessary to change the base. One of the reasons for doing this might be that the original base time point is too far in the past to be relevant today and amore recent one is needed. For example, the following set has a base of 1965, which would probably now be considered out-of-date.

	1987	1988	1989	1990	1991	1992	1993
Index(1965=100)	324	351	377	384	391	404	428

The procedure for changing the base of a time series of relatives is essentially the same as that for calculating a set of relatives for a given time series of values. However, the procedure is given below and demonstrated, using the above set of relatives:

STEP 1 Choose the required new base time point and thus identify the corresponding relative.

We will choose 1987 as the base year, with a corresponding relative of 324.

STEP 2 Divide each relation in the set by the value of the relative identified above and multiply the result by 100.

Thus, each index relative given needs to be divided by 324 and multiplied by 100. Table 3 shows the new index numbers.

	1987	1988	1989	1990	1991	1992	1993
OLD Index (1965=100)	324	351	377	384	391	404	428
NEW Index(1987=100)	100	108	116	119	121	125	132
	$\frac{324}{324} \times 100$		$\frac{377}{324} \times 100$			$\frac{404}{324} \times 100$	

Table 3

17.8 Comparing sets of fixed base relatives

Sometimes it is necessary to compare two given sets of time series relatives. For example, the annual index for the number of televisions sold might be compared with the annual index for television licenses taken out, or the monthly consumer prices index compared with the monthly index for wages. In cases such as these, it is usually found that the bases on which the two sets of indices are calculated are different. For example, the consumer index might have a base of 1974, while the wage index has a base of 1983. This can make comparisons difficult because the two sets of index relatives will be of different magnitudes. As an illustration, consider the data of Table 4.

Year	1986	1987	1988	1989	1990	1991	1992
Number of TV sets sold (1988 = 100)	61	88	100	135	165	192	210
Number of TV licences taken out (1970 =100)	210	230	250	300	360	410	500

Comparing the indices given above is not easy. Many percentage increases will have to be calculated before any worthwhile comparisons can be made. This type of problem can be overcome by changing the base of one set of indices to match the base of the other. The following example shows the calculations necessary.

Example 3 (Time series comparison by changing the base of one of the sets)

Question

Compare the figures given in Table 4 by changing the base of one of the sets and comment on the results.

Answer

The base of the television licence relatives will be changed to coincide with the base of the televisions sold relatives. The following table shows the new figures.

Year	1986	1987	1988	1989	1990	1991	1992
Number of TV sets sold (1988 = 100)	61	88	100	135	165	192	210
Number of TV licences taken out (1970 =100)	210	230	250	300	360	410	500
Number of television licences Taken out (1988=100)	84	92	100	120	144	164	200
		$\frac{230}{250} \times 100$			$\frac{360}{250} \times 100$		

The two sets of relatives are now much easier to compare. Before 1988 and up to 1991, sales of television sets increased at a much faster rate. However, over the last year, the number of television licenses taken out increased dramatically, showing the same percentage increase (over 1988) as the sales of television sets (possibly due to detector van publicity).

17.9 Actual and real values of a commodity

In times of significant inflation, the actual value of some commodity is not the best guide of its 'real' value (or worth). The worth of any commodity can only be measured relative to the value of some associated commodity. In other words, some relevant 'indicator' is necessary against which to judge value.

For example, suppose that the annual rent of some business premises last year was Rs. 2200. Clearly the actual cost is higher. However, if we are now given the information (as an indicator) that the cost of business premises in the region as a whole has risen by 15% over the past year, we can rightly argue that the real cost of the given premises has decreased. On the other hand, if business turnover for the premises (as an alternative indicator) has only increased by 5%, we might consider that the real cost of the premises has increased. Thus, depending on the particular indicator chosen, the real value of a commodity can change.

Two standard national indicators are the rate of inflation (normally represented by the Retail Price Index) and the Index of Output of the Production Industries

The following section describes a method of constructing a series of relatives to measure the real value of some commodity over time. This is known as time series deflation.

17.10 Time series deflation

Time series deflation is a technique used to obtain a set of index relatives that measure the changes in the real value of some commodity with respect to some given indicator.

Month	1	2	3	4	5	6	7	8
Average daily wage(Rs.)	17.60	18.10	18.90	19.60	20.25	20.30	20.60	21.40
Retail price index	106.1	107.9	112.0	113.1	116.0	117.4	119.5	119.7

Table 5

The procedure for calculating each index relative is given below, using the data of Table 5 to demonstrate calculating the real wage index for month 7 (month 1 = 100) as an example.

STEP1 Choose a base for the index of real values of the series.

In this case, month 1 has been chosen.

Then, for each time point of the series:

STEP2 Find the ratio of the current value to the base value.

For month 7, this gives: $\frac{20.60}{17.60} = 1.17$

This step expresses the increase in the actual value as a multiple.

STEP3 Multiply by the ratio of the base indicator to the current indicator (notice that these two values are in reverse order compared with the two in the previous step).

For month 7, this gives: $1.17 \times \frac{106.1}{119.5} = 1.039$, 'deflating' the above

wage multiple.

STEP4 Multiply by 100

For month 7, this gives: $1.039 \times 100 = 103.9$.

This step changes the multiple of the previous step into an index (based on 100).

The above steps can be summed up both in symbols and words as follows.

Real Value Index (RVI)

Given a time series (x-values) and some indicator index series (I – values) for comparison, the real value index for period n is given by:

$$\begin{aligned} RVI &= \frac{\text{currentvalue}}{\text{basevalue}} \times \frac{\text{baseindicator}}{\text{currentindicator}} \times 100 \\ &= \\ &= \frac{X_n}{X_0} \times \frac{I_0}{I_n} \times 100 \end{aligned}$$

The following example duplicates the data of Table 5 and shows the real wage index relatives, the calculations (using the above steps) being demonstrated for selected values.

Example 4 (Index relatives of real values)

Table 6 below shows the values of the real wage index relative for the data of Table 5.

Month	1	2	3	4	5	6	7	8
Average daily wage (£)	17.60	18.10	18.90	19.60	20.25	20.30	20.60	21.40
Retail price index	106.1	107.9	112.0	113.1	116.0	117.4	119.5	119.7
Real Wage Index	100	101.1	101.7	104.5	105.2	104.2	103.9	107.8

Table 6

$$\frac{18.90}{17.60} \times \frac{106.1}{112.0} \times 100 \quad \frac{20.25}{17.60} \times \frac{106.1}{116.0} \times 100 \quad \frac{21.40}{17.60} \times \frac{106.1}{119.7} \times 100$$

The real wage index shows that the real value of the average weekly wage has increased by 7.8% over the nine-month period. In real terms, wages increased steadily with larger than usual increases in months 4 and 8 and small decreases in months 6 and 7.

Exercises

1. The average price of a product this year was Rs. 33.3, which represented a decrease of 10% over last year's average price. The number bought (at these prices) last year was 2500, but increased by 750 this year. Calculate price, quantity and expenditure relatives for these cassettes for this year (based on last year).

2. The following data relate to the production of cars from a particular assembly line over a number of months.

	Mar	Apr	May	Jun	Jul	Aug	Sep	Oct
Production	142	126	128	104	108	146	158	137

Calculate sets of productivity relatives (to ID) with: a) Mar = 100 b) May = 100 c) Aug = 100.

3. Butter stocks (thousand tones) in a particular year

Mar	Apr	May	Jun	Jul	Aug	Sep	Oct	Nov
216.9	225.1	234.6	237.2	235.2	230.1	224.4	226.1	220.2

Calculate (to ID) a set of:

- fixed base relatives (Mar = 100);
- chain base relatives.

Comment on the results.

4. The yearly index for the production of an important product for a firm is contrasted with a national production index for the same type of product.

	19X0	19X1	19X2	19X3	19X4	19X5	19X6	19X7
Production index for firm(19X2=100)	101	96	100	107	98	98	103	107
National production index(19X0=100)	384	382	427	445	416	410	427	444

Compare the firm's production record with national production by changing the base of the national index.

5. Compare the following series, using the same fixed base, and comment on the results.

	<i>Average earnings index numbers</i>									
	Feb	Mar	Apr	May	Jun	Jul	Aug	Sep	Oct	Nov
Whole economy	164.6	168.1	169.4	169.4	171.9	173.7	173.4	176.1	173.9	176.8
Coal and Coke	78.2	122.5	137.9	139.5	148.0	149.5	150.7	152.9	153.6	159.3

6. The figures below compare the fuel costs of a small garage with a national price index.

Time point	1	2	3	4	5	6	7
Cost of fuel (in Rs.000)	34.1	34.8	33.6	33.6	33.4	33.1	33.4
Producer(fuel)price index	169.8	173.9	163.8	151.1	148.9	147.4	147.4

Produce an index (time point 1 = 100), to ID, of the real cost of fuel to the garage by deflating the given fuel costs by the Producer (Fuel) Price Index.

7. The data below show the gross income of a particular category of family compared with the Retail price Index over a seven year period.

	19X5	19X6	19X7	19X8	19X9	19Y0	19Y1
Family income(Rs.000)	6,989	8,105	8,416	10,037	11,475	13,443	16,140
Retail price index	134.8	157.1	182.0	197.1	223.5	263.7	295.0

Calculate:

- an index of real gross income (19X5 = 100)
- a chain base index of real gross income, using the Retail Price Index as an indicator.

17.11 Let us Sum Up

This Lesson discussed about indices and its common applications. An index number measures the percentage change in the value of some economic commodity over a period of time. It is always expressed in terms of a base of 100. An index relative is the name given to an index number which measures the change in a single distinct commodity. A price relative can be calculated as the ratio of the current price to the base price multiplied by one hundred. Quantity, expenditure and productivity relatives are calculated in a similar manner. Fixed base relatives are found by calculating relatives for each value of a time series based on the same fixed time point. Chain base relatives are found by calculating relatives for each value of time series based on the immediately preceding time point. In order to compare two time series of relatives, each series should have the same base time point. The real value of some commodity can only be measured in terms of some 'indicator'. Standard indicators are the Retail Price Index or the Index of Output of the Production Industries. Time series deflation is also discussed which is a technique used to obtain a set of index relatives that measure the changes in the real value of some commodity with respect to some given indicator.

17.12. Lesson – End Activities

1. Define Index Number.

17.13. References

1. Gupta S.P. – Statistical Methods

Lesson 18 - Special Published Indices

Contents

- 18.1 Aims and Objectives**
- 18.2 The Retail Prices Index**
- 18.3 Main RPI groups and their weights**
- 18.4 The family Expenditure Survey**
- 18.5 Price collection and calculation of the RPI**
- 18.6 The Purchasing Power (index)**
- 18.7 The Tax and Price index**
- 18.8 Index numbers of producer Prices**
- 18.9 Indices of average earnings**
- 18.10 Index of output of the production industries**
- 18.11 Other index numbers**
- 18.12 Let us Sum Up**
- 18.13 Lesson – End Activities**
- 18.14 References**

18.1 Aims and Objectives

This Lesson describes some of the most important official index numbers. The price indices described are the Retail Price Index (which includes the important Family Expenditure Survey), Purchasing Power, the Tax and Price Index and Index numbers of Producer Prices. Indices of Average Earnings are also covered. Volume (or quantity) indices described are the Index of Output of the Production Industries and the Index of Retail Sales. Some indices described cover more than one section.

18.2 The Retail Prices Index

The Retail Prices Index (or RPI), is probably the best known of all the published indices.

- a) It is published monthly by the Department of Employment and Displayed (to different levels of complexity) in the following publications: Monthly Digest of Statistics, the Annual abstract of Statistics, the Department of Employment Gazette and Economic Trends.
- b) The RPI measures the percentage changes, month by month, in the average level of prices of the commodities and services purchased by the great majority of households in the Country. It takes account of practically all wage earners and most small and medium salary earners.
- c) The items covered by the RPI are classified into several groups. For example, Food, Housing, Transport and Vehicles etc). Each group is sub-divided into sections. For

example, Transport is sub-divided into Motoring/cycling and Fares). These sections may be further split up into separate items. For example, Fares are split up into Rail and Road.

- d) Each month, an overall index is published, together with separate indices for each group, section and individual item (of which there are approximately 350).
- e) Each group (and further sections and specific items) is weighted according to expenditure by a 'typical family' and the weights are updated annually.
- f) The weights are obtained from a continuous investigation known as the Family Expenditure Survey.

18.3 Main RPI groups and their weights

Table 1 shows the main groups of the RPI, their separate price indices (as at January 1986) and their weights for three different dates.

Main groups	Price index	Weights		
	January 1986 (1974 = 100)	1962	1973	1985
Food	341.1	350	248	190
Alcoholic drink	423.8	71	73	75
Tobacco	545.7	80	49	37
Housing	463.7	87	126	153
Fuel and light	507.0	55	58	65
Durable household goods	265.2	66	58	65
Clothing and footwear	225.2	106	89	75
Transport and vehicles	393.1	68	135	156
Miscellaneous goods	402.9	59	65	77
Meals bought out	426.7	-	46	45
Overall	379.7	1000	1000	1000

Notes on Table 1:

- a) Weights are always calculated to add to 1000.
- b) 'Meals bought out' was not included in the 1962 weightings
- c) Certain items of expenditure are not included in the RPI. These include:
 - i. Income tax and National Insurance payments;
 - ii. Insurance and pension payments;
 - iii. Mortgage payments for house purchase (except for interest payments which are included);
 - iv. Gambling, gifts, charity, etc.

Example 1 (Comments on the data in Table 1)

- a) The Retail Prices Index for January 1986 (1974 = 100) was 379.7. This represents an overall increase in prices of approximately 280% since 1974.
- b) Food has been subject to below average price increase (341.1 index = 241% increase) and expenditure has continued to decrease significantly. Since food is a basic necessity of life, this is a good indication of our increasing affluence.
- c) Tobacco has seen the highest increase in price (index = 545.7) with a definite downward trend in expenditure. The latter trend is obviously due to both high price and health warnings.
- d) Clothing and Footwear has had the lowest increase in price (index = 225.2), representing only a doubling in price over the previous 10 years, but this group has still seen a downward trend in expenditure. Since there is no reason to suppose that we now buy fewer clothes, it probably means that clothes are much cheaper in real terms.
- e) Housing and Transport and Vehicles both show a similar upward trend in expenditure. However, where Transport is only showing an average price increase, Housing shows the third highest (index = 463.7). Upward expenditure on transport clearly signifies our increasing mobility (in both work and recreation). Extra expenditure on housing probably reflects social and ecological factors as much as increase in price.

18.4 The family Expenditure Survey

The Family Expenditure Survey (FES) is a continuous major investigation which, among other things, measures average consumption levels. These are used to obtain the (annually revised) weights for items included in the RPI. The FES involves a stratified random sample, spread over the course of a year, of about 10000 households. Each household is visited by an interviewer. Each member of the household over the age of sixteen years is required to keep a detailed diary of all expenditure for a continuous 14-day period, which is checked and retained by the interviewer. The interviewer also completes a Household Schedule, which contains information on longer term spending such as rent, rates, carpets, cars, and so on. (An Income Schedule is also filled out for the members of the household.). The published weights are calculated, not from a single year's FES data, but as an average of the previous three year's data. This ensures that large items of expenditure do not unfairly influence average patterns of spending. The pattern of FES varies from country to country.

18.5 Price collection and calculation of the RPI

Prices are collected by Department of Employment staff. Different types of retail outlets, from village shops to large supermarkets, are visited. To ensure uniformity, the same ones are used each month and these will be the type of retail outlet used by households examined by the FES. Price relatives are calculated (for each item covered by the RPI) for each retail outlet and averaged for a local area. Average relatives for all local areas are in turn averaged to obtain a national average of relatives (for each of the 350 items covered by the index). Weights are then used to calculate composite indices using the average of relatives method for items within sections, sections within groups and, finally, groups. Thus the RPI is a weighted average of relatives of each group.

18.6 The Purchasing Power (index)

The Purchasing Power is an index which has been based solely on the annual average of the RPI. The philosophy behind the index is: when prices go up, the amount which can be purchased with a given sum of money goes down. The index is described in terms of two particular years. If the purchasing power of the Rupee is taken to be 100 in the first year, the comparable purchasing power in a later year is calculated as:

For example, the PP index for 1984 (1980 = 100) is given as 75. This can be interpreted as:

$$100 \times \frac{\text{Average Price Index for First Year}}{\text{Average Price Index for Later Year}}$$

- i. the Rupee (in 1984) is worth only 75% of its 1980 value, of
- ii. 100 rupees buys (in 1984) what would only have cost 75 rupees in 1980.

18.7 The Tax and Price index

The Tax and Price Index (TPI), published monthly, is another index which is linked to the Retail Prices Index. The TPI measures the increase in gross taxable income needed to compensate taxpayers for any increase in retail prices (as measured by the RPI). It is considered as a more comprehensive index than the RPI since, while the RPI measures changes in retail prices, the TPI additionally takes account of the changes in liability to direct taxes (including employees' national insurance contributions) facing a representative cross-section of taxpayers. Some people would argue that the TPI is a better measure of the cost of living than the RPI since it takes direct taxes into account. However, whether or not this is acceptable depends on the meaning of the phrase 'cost of living' – it has different meanings to different people and circumstances. Another complicating factor is that the TPI (a relatively new index) is regarded suspiciously by some political opponents of the Government in office at the time of its introduction.

Example 2 (Comparison of the TPI and RPI)

The TPI for June 1985 (January 1978 = 100) was 191.7 [INDEX 1]

The RPI for June 1985 (January 1974 = 100) was 376.4 [INDEX 2]

The RPI for 1978 (1974 = 100) was 197.1 [INDEX 3]

Note that it is difficult to compare the first two indices, since their base dates are different. However, the information contained in INDEX 3 allows the RPI (INDEX 2) to be base-changed to coincide with the base of the TPI (INDEX 1), for a direct comparison.

$$\begin{aligned} \text{Thus: RPI}_{85/78} &= \frac{\text{RPI}_{85/74}}{\text{RPI}_{78/74}} \times 100 \\ &= \frac{376.4}{197.1} \times 100 \\ &= 191.0 \text{ (ID)} \end{aligned}$$

Therefore the TPI for June 1985 shows a slightly higher increase (91.7%) than the RPI for June 1985 (91.0%)

Note, however, that INDEX 3 is based on annual averages whereas the other two indices are based on actual months of the year. Hence the above base change will cause the resultant figure to be slightly in error.

18.8 Index numbers of producer Prices

The Producer Price Indices (PPI) measure manufacturers' prices and were formerly known as the wholesale Price Indices. The data for the indices are collected by the Business Statistics Office. Indices are produced for a wide range of prices including output (home sales), materials and fuel purchased, commodities produced and imported. They are quoted for main industrial groupings, such as Motor Vehicles and Parts, Food Manufacturing industries, Textile industry, and son on. In some publications, the groupings are sub-divided into items. The various index numbers produced are calculated from the price movements of about 10,000 closely defined materials and products representative of goods purchased and manufactured. All the indices express the current prices as a percentage of their annual average price in 1980, the base year.

18.9 Indices of average earnings

The Indices of Average Earnings measure the changes in average gross income. They are published for manual workers and all workers and given for industry groups. Actual and seasonally adjusted indices are given for certain tables. The series as at June 1986 are all based on 1980 = 100.

18.10 Index of output of the production industries

The Index of Output of the Production Industries was formerly known as the Index of Industrial Production. It provides a general measure of monthly changes in the volume of output of the production industries. Energy, water supply and manufacturing are included in the index. However, agriculture, construction, distribution, transport, communications, fiancé and all other public and private services are excluded. The index covers the production of intermediate, investment and consumer goods, both for home and export. Many of the series presented are seasonally adjusted. This excludes any changes in production resulting from public and other holidays and from other seasonal factors. The adjustments are designed to eliminate normal month to month fluctuations and thus to show the trend more clearly.

18.11 Other index numbers

Some other index numbers that are given in main publications are:

- a) Index numbers of Output (at constant factor cost);
- b) Index of retail sales;
- c) Index numbers of Expenditure (at 1980 prices, currently);
- d) Volume Index of Sales of Manufactured Goods;

- e) Indices of Labour costs;
- f) (external Trade) Volume and Unit Value Index numbers; and an important non-official publication:
- g) The Financial Times Ordinary Share Index.

Exercises

1. What is the Retail Prices Index (RPI)?
2. Name at least five of the eleven main groups into which the RPI is divided.
3. Name some of the items of expenditure that are not included in the calculation of the RPI.
4. How are prices collected for the RPI?
5. Explain what the 'Purchasing Power' and how it is calculated.
6. What does the Tax and Price Index (TPI) measure?
7. Compare the RPI and TPI.
8. Describe some aspects of the Index Numbers of Producer Prices.
9. What is the Index of Retail Sales and how are the data in its construction collected?

18.12 Let us Sum Up

This Lesson discussed special published indices which finds applications in economics and in financial management. The Retail Prices Index (RPI) is published monthly and measures the percentage changes in the average level of prices of the commodities and services purchased by most households. Purchasing Power (index) gives the percentage worth of a current pound compared with a pound in a previous period. The Tax and Price Index measures the increase in gross taxable income needed to compensate taxpayers for any increases in retail prices (as measured by the RPI). It takes account of direct taxation. The Indices of Average Earning measure the changes in average gross income for manual and other workers. The Index of Output of the Production Industries provides a general measure of monthly changes in the volume of output of the production industries. Index numbers of Retail Sales give both volume and value indices and are compiled on the type of business rather than on a commodity basis.

18.13 Lesson – End Activities

Give the importance of the Retail Price Index.

18.14 References

1. Gupta S.P. – Statistical Methods.
2. P.R. Vital – Business Mathematics and Statistics.