

# **Block 3**

## **Statistical Methods**



ignou  
THE PEOPLE'S  
UNIVERSITY

---

## UNIT 10 GRAPHICAL AND DIAGRAMMATIC PRESENTATION OF DATA\*

---

### Structure

- 10.0 Objectives
- 10.1 Introduction
- 10.2 Arrangement of Data
  - 10.2.1 Simple Array
  - 10.2.2 Frequency Array or Discrete Frequency Distribution
  - 10.2.3 Continuous or Grouped Frequency Distribution
  - 10.2.4 Various forms of Frequency Distribution
- 10.3 Tabulation of Data
  - 10.3.1 Meaning and types of Tables
  - 10.3.2 Parts of a Table
  - 10.3.3 Importance of Tables
- 10.4 Graphical Presentation of Data
  - 10.4.1 Line Graphs
  - 10.4.2 Histogram, Frequency Polygon and Frequency Curves
  - 10.4.3 Cumulative Frequency Curves - Ogives
- 10.5 Diagrammatic Presentation of Data
  - 10.5.1 One Dimensional Diagrams
  - 10.5.2 Two Dimensional Diagrams or Area Diagrams
  - 10.5.3 Pie Diagram or Pie Chart
  - 10.5.4 Three Dimensional Diagrams
  - 10.5.5 Pictograms and Statistical Maps
- 10.6 Let Us Sum Up
- 10.7 Answers or Hints to Check Your Progress Exercises

---

### 10.0 OBJECTIVES

---

After going through this Unit, you will be able to familiarize yourself with

- stages of statistical inquiry after data have been collected;
- methods of organizing (classification and arrangement) and condensing statistical data;
- concepts of frequency distribution and its various types; and
- different methods of presentation of statistical data such as tables, graphs, diagrams, pictograms, etc.

---

\* Adapted from IGNOU study material of BECC 107: Statistical Methods for Economics, Unit 2 written by Avatar Singh with modifications by K. Barik.

---

## 10.1 INTRODUCTION

---

In the preceding Unit, we discussed the methods of collection of data either by a statistical survey (or inquiry) or from some secondary source. Data collected either from census or sample inquiry, that is from primary source, are always hotchpotch and in rudimentary form. To start with, they are contained in hundreds and thousands of questionnaires. To make a head and tail out of them, they must be organised (i.e., classified and arranged), and condensed or summarised. For this purpose we can use various methods like preparing master sheets in which various information are recorded directly from the questionnaires. From these sheets small summary tables can be prepared manually. Now-a-days computers can be used for organisation and condensation of data more swiftly, efficiently and in much less time. Some computer softwares are available which help us to construct various types of graphs and diagrams.

Data can be summarized numerically also. Here we use summary measures like measures of central tendency of first order or degree (like Arithmetic, Geometric and Harmonic Means, Mode and Median); measures of central tendency of second order or degree, also called dispersion measures (like Range, Quartile Deviation, Mean Deviation, and Standard Deviation); measures of association in bivariate analysis (like Correlation and Regression), Index Numbers, etc. In this Unit we plan to discuss how data can be summarized using tables and graphs. Numerical summarization will be discussed subsequently in Units 3 and 4). It must be born in mind that a good summarization and presentation of data is not undertaken for its own sake. It is not an end in itself. In fact it sets the stage for useful analysis and interpretation of data. Again, a good presentation helps us to highlight significant facts and their comparisons. Figures can be made to speak out thereby making possible their intelligent use.

In this Unit we plan to concentrate on organising and condensing data in the form of simple array (ascending and descending order), frequency array and continuous frequency distributions, etc.; and presentation of statistical data in the form of tables and graphs.

---

## 10.2 ARRANGEMENT OF DATA

---

The mass of collected data is often voluminous, unintelligible and boring. It seems totally uninteresting and is not easily interpretable. For example, if you are provided with monthly income figures of 1000 families in a village it is difficult for you to infer anything. But if you are told that the average monthly income of the village is Rs. 2540, it is quite interesting and you are in a position to compare it with other figures.

The first step in the analysis and interpretation of data is its classification and tabulation. The process of arranging data into groups according to their common characteristics is known as its classification.

On the other hand tabulation implies a systematic presentation of data in rows and columns according to some salient features or characteristics.

Let us assume that we collected data on number of children and monthly income on the basis of a questionnaire from 50 families of C-III block of XYZ Colony, New Delhi. Let us assume that it produced the following data as given in Tables 10.1 and 10.2. *Can we make any head or tail out of it?*

**Table 10.1**

**Number of Children per family in C-III block, XYZ Colony, New Delhi**

|   |   |   |   |   |   |   |   |   |   |
|---|---|---|---|---|---|---|---|---|---|
| 2 | 0 | 1 | 5 | 3 | 1 | 2 | 1 | 0 | 2 |
| 4 | 3 | 2 | 2 | 3 | 4 | 1 | 0 | 2 | 3 |
| 1 | 4 | 2 | 3 | 1 | 2 | 5 | 4 | 1 | 3 |
| 2 | 1 | 3 | 2 | 3 | 4 | 1 | 2 | 3 | 1 |
| 4 | 5 | 2 | 1 | 1 | 0 | 3 | 2 | 0 | 2 |

**Table 10.2**

**Monthly Income of 50 families of C-III block, XYZ Colony, New Delhi**

|     |     |     |     |     |     |     |     |     |     |
|-----|-----|-----|-----|-----|-----|-----|-----|-----|-----|
| 547 | 622 | 691 | 684 | 567 | 586 | 680 | 578 | 583 | 578 |
| 708 | 544 | 528 | 540 | 730 | 541 | 720 | 698 | 763 | 633 |
| 640 | 637 | 598 | 631 | 618 | 692 | 600 | 650 | 604 | 640 |
| 646 | 654 | 689 | 736 | 731 | 844 | 798 | 712 | 772 | 820 |
| 678 | 663 | 800 | 692 | 700 | 781 | 658 | 798 | 709 | 720 |

As pointed out earlier, to make any head or tail out of the mass of raw data, such as presented above, we have to classify and arrange it. This can be done either by forming a simple array or a frequency array (discrete frequency distribution) or a continuous frequency distribution. Sections 10.3.1, 10.3.2 and 10.3.3 attempt to explain this aspect.

### **10.2.1 Simple Array**

It is an arrangement of given raw (univariate) data in ascending or descending order. In the ascending order, the observations are arranged in increasing order of magnitude. For example, numbers 3,5,7,8,9,10 are arranged in ascending order. In descending order, it is the reverse. For example, the numbers 10,9,8,7,6,5,3 are in descending order.

We can prepare both types of simple arrays from Table 10.1. In the following table, the figures have been arranged in ascending order. From the arrangement, it is clear that the lowest value is 0 and the highest one is 5.

**Table 10.3**  
**Number of Children per Family in C-III Block of XYZ Colony, New Delhi**  
**Simple Array -- Ascending Order**

|   |   |   |   |   |   |   |   |   |   |
|---|---|---|---|---|---|---|---|---|---|
| 0 | 0 | 0 | 0 | 0 | 1 | 1 | 1 | 1 | 1 |
| 1 | 1 | 1 | 1 | 1 | 1 | 1 | 2 | 2 | 2 |
| 2 | 2 | 2 | 2 | 2 | 2 | 2 | 2 | 2 | 2 |
| 2 | 3 | 3 | 3 | 3 | 3 | 3 | 3 | 3 | 3 |
| 3 | 4 | 4 | 4 | 4 | 4 | 4 | 5 | 5 | 5 |

After this arrangement the same figures make some sense.

The possible conclusions that can be drawn from this arrangement of data are that five families are issueless, twelve families have one child each, fourteen have two children each, ten families possess three children each, six families four children each and three families have five children each.

### 10.2.2 Frequency Array or Discrete Frequency Distribution

Here different observations are not repeatedly written as in simple array like 0, 0, 0, 0, 1, 1, 1, 1, 1, 1, 1, 1, etc. We count the number of times (i.e., frequency) an observation repeats itself. For example, in Table 10.3 the observation 4 is repeated 6 times. Thus the frequency of 4 above is 6. The frequency array, for the simple array given in Table 10.3, will look like as given below in Table 10.4.

A frequency array is a statistical table in which various observations are arranged in order of their magnitude along with their respective frequencies.

**Table 10.4**

|                     |   |    |    |    |   |   |       |
|---------------------|---|----|----|----|---|---|-------|
| Number of Children: | 0 | 1  | 2  | 3  | 4 | 5 | Total |
| Number of Families: | 5 | 12 | 14 | 10 | 6 | 3 | 50    |

When the number of observations is large enough, the counting process is often undertaken by the use of *tally bars*. In this method, all possible values of the variable are written in a column. For every observation, a tally bar denoted by ( | ) is noted against its corresponding values. Every fifth repetition is marked by crossing the previous four bars as (  $\overline{||||}$  ). In this way, we get blocks of five which simplify counting at the end. Thus a number or an observation repeated fourteen times will be marked as (  $\overline{||||} \overline{||||} ||||$  ). Note that after representing each observation by a bar on the *tally sheet*, the same will be ticked (✓) or crossed (×) so that it is not duplicated. The data of Table 10.1 is rewritten in the form of frequency distribution as shown in Table 10.5 below.

**Table 10.5**  
**Frequency Distribution of Number of Children per Family**

| <u>No. of Children</u> | <u>Tally Sheet</u> | <u>Frequency</u> |
|------------------------|--------------------|------------------|
| 0                      |                    | 5                |
| 1                      |                    | 12               |
| 2                      |                    | 14               |
| 3                      |                    | 10               |
| 4                      |                    | 6                |
| 5                      |                    | 3                |
| <b><u>Total</u></b>    |                    | <b><u>50</u></b> |

### 10.2.3 Continuous or Grouped Frequency Distribution

Numbers like 1, 2, 3, 4, 5, 20, 40, etc. are discrete numbers and are used where no value between the two consecutive numbers is possible. As in the case of the number of children, it will be impossible as well as funny to say that a particular family has 2.083 or 2.1 or 2.75 number of children. The family can have either 2 or 3 children and not a fraction of a child. Out of the two examples of raw data mentioned in Section 10.3, the number of children (Table 10.1) is an example of discrete data while monthly income (Table 10.2) is an example of continuous variable giving continuous data.

In this Section we propose to illustrate the construction of continuous or grouped frequency distribution from the raw data of Table 10.2 on monthly income of the 50 families.

To construct a grouped frequency distribution, the range of the given data, i.e., the difference of the highest and the lowest observations, is divided into various mutually exclusive and exhaustive sub-intervals, also known as class-intervals. The frequency of each class interval is then counted and written against it.

**Table 10.6**  
**Frequency Distribution of Monthly Income of Families**

| <u>Monthly Income</u><br>(Rs.) | <u>Tally Sheet</u> | <u>No. of Families</u><br>(Frequency) |
|--------------------------------|--------------------|---------------------------------------|
| 500 - 550                      |                    | 5                                     |
| 550 - 600                      |                    | 6                                     |
| 600 - 650                      |                    | 10                                    |
| 650 - 700                      |                    | 12                                    |
| 700 - 750                      |                    | 9                                     |
| 750 - 800                      |                    | 5                                     |
| 800 - 850                      |                    | 3                                     |
| <b><u>Total</u></b>            |                    | <b><u>50</u></b>                      |

We have just completed an exercise where the variable “*income of the family*” has been grouped in order to reduce it to a manageable form called grouped data or *Continuous Frequency Distribution*. However, prior to the construction of any grouped frequency distribution, it is very important to find answers to the following questions:

1. What should be the number of class intervals?
2. What should be the width of each class interval?
3. How will the class limits be designated?

### 1. What should be the number of class intervals?

Though there is no hard and fast rule regarding the number of classes to be formed, yet their number should be neither too small nor too large. If the number of classes is too small, i.e., width of each class is large, there is likelihood of greater loss of information due to grouping. On the other hand, if the number of classes is very large, the distribution may appear to be too fragmented and may not reveal any pattern of behaviour of the variable. Based on experience, it has been observed that the minimum number of classes should not be less than 5 or 6 and in any case, there should not be more than 20 classes.

Mr. Sturges has given a formula to determine the number of classes, that is,

$$\text{Number of classes} = 1 + 3.322 \times \log_{10} N,$$

where  $N$  is the total number of observations.

In our example of raw data on incomes of 50 families, the number of classes can be calculated as under:

$$\begin{aligned} \text{Number of classes} &= 1 + 3.322 \times \log_{10} 50 = 1 + 3.322 \times 1.6990 \\ &= 1 + 5.644 = 6.644 \approx 7. \end{aligned}$$

### 2. What should be the width of each class interval?

As far as possible, all the class intervals should be of equal width. However, when a frequency distribution, based on equal class intervals, does not reveal a regular pattern of behaviour of observations, it might become necessary to re-group the observations into class intervals of unequal width. By a regular pattern of behaviour we mean that there are no classes, with possible exclusion of extreme classes, where there are nil or very few observations while there is concentration of observations in their adjoining classes.

$$\text{Width of a Class} = \frac{\text{Largest Observation} - \text{Smallest Observation}}{\text{Number of Class Intervals}}$$

The approximate width of a class can be determined by the following formula:

However, the final decision, regarding width of class intervals, should also take into account the following points.

- i) As far as possible, the width should be a multiple of 5, because it is easy to grasp numbers like 5, 10, 15, ..... etc.
- ii) It should be convenient to find the mid-value of a class.
- iii) The observations in a class should be uniformly distributed.

### 3. How will the class limits be designated?

The smallest and the largest observations of a class interval are known as class limits. These are also termed as the lower and upper limits of a class, respectively. Since the mid-value of a class, which is used to compute mean, standard deviation, etc., is obtained from the class limits, it is very necessary to define these limits in an unambiguous manner. The following points should be kept in mind while defining class limits:

- It is not necessary that the lower limit of the first class be exactly equal to the smallest observation of the data.  
In fact it can be less than or equal to the smallest observation. Similarly, the upper limit of the last class may be greater than or equal to the largest observation of the data.
- It is convenient to have the lower limit of a class either equal to zero or some multiple of 5 or 10.
- The chosen class limits should be such that the observations in a class are uniformly distributed.

The class limits can be defined in either of the following methods:

- Exclusive Method*, and
- Inclusive Method*.

**i) Exclusive Method:** In this method, the upper limit of a class is taken to be equal to the lower limit of the following class. In order to keep various class intervals as mutually exclusive, it is decided that the observations with magnitude greater than or equal to lower limit but less than the upper limit of a class are included in it. For example, the class 500 - 550 shall include all observations with magnitude greater than or equal to 500 but less than 550. An observation with magnitude equal to 550 will be included in the next class, i.e., the class 550 - 600.

The major benefit of exclusive class intervals is that it ensures continuity of data because the upper limit of one class is the lower limit of the next class. In our example on monthly income (Table 10.6), there are 5 families whose income lies between Rs. 500 to Rs. 550, i.e., Rs. 500 to 549 and 6 families whose income lies between Rs. 550 to Rs. 600, i.e., Rs. 550 to 599, and so on. Based on this presumption we can rewrite this frequency distribution in the form of Table 10.7 also.

**Table 10.7**  
**Exclusive Class Intervals**

| <i>Monthly Income (Rs.)</i> | <i>Number of Families<br/>(Frequency)</i> |
|-----------------------------|-------------------------------------------|
| 500 but less than 550       | 5                                         |
| 550 but less than 600       | 6                                         |
| 600 but less than 650       | 10                                        |
| 650 but less than 700       | 12                                        |
| 700 but less than 750       | 9                                         |
| 750 but less than 800       | 5                                         |
| 800 but less than 850       | 3                                         |
| <b><u>Total</u></b>         | <b><u>50</u></b>                          |

- ii) **Inclusive Method:** In this method, all the observations with magnitude greater than or equal to the lower limit but less than or equal to the upper limit of a class is included in it. Now observe Table 10.8. Income of Rs. 549 is included in the class 500 to 549 so that an income of Rs. 550 automatically goes to the next class of 550 to 599. Since the upper limit of one class is not equal to the lower limit of the following class, this saves us from the confusion whether Rs. 550 goes to (500 to 549) or (550 to 599) class.

**Table 10.8**

| <b>Inclusive Class Intervals</b> |                                   |
|----------------------------------|-----------------------------------|
| Monthly Income<br>(Rs.)          | Number of Families<br>(Frequency) |
| 500 - 549                        | 5                                 |
| 550 - 599                        | 6                                 |
| 600 - 649                        | 10                                |
| 650 - 699                        | 12                                |
| 700 - 749                        | 9                                 |
| 750 - 799                        | 5                                 |
| 800 - 849                        | 3                                 |
| <b>Total</b>                     | <b>50</b>                         |

The *choice between exclusive and inclusive methods* depends upon whether we are dealing with continuous variable like income, heights, weights, etc. or a discrete variable like number of children in a family. For a continuous variable it is desirable to construct frequency distribution by the exclusive method because, as we have seen earlier, it ensures continuity. For a discrete variable like number of children in a family or number of students getting first division, the frequency distributions should be constructed by using inclusive type of class intervals.

#### **Mid-Value of a Class**

In exclusive type of class intervals, the mid-value or class mark of a class is defined as the arithmetic mean of its lower and upper limits. However, in case of inclusive class intervals, there is a gap between the upper limit of a class and the lower limit of the following class. This gap is eliminated by adding half of the gap to the upper limit and subtracting half of the gap from the lower limit. The new class limits, thus obtained, are known as *class boundaries*. The class boundaries of the inclusive class intervals in Table 10.8 are given in Table 10.9.

**Table 10.9**

| Monthly Income (Rs.) | No. of Families<br>(Frequency) |
|----------------------|--------------------------------|
| 499.5 - 549.5        | 5                              |
| 549.5 - 599.5        | 6                              |
| 599.5 - 649.5        | 10                             |
| 649.5 - 699.5        | 12                             |
| 699.5 - 749.5        | 9                              |
| 749.5 - 799.5        | 5                              |
| 799.5 - 849.5        | 3                              |
| <b>Total</b>         | <b>50</b>                      |

### 10.2.4 Various Forms of Frequency Distribution

Here we propose to introduce the meaning of the following frequency distributions:

- Open Ended Frequency Distribution
- A Frequency Distribution with Unequal Class Width
- Cumulative Frequency Distribution
- Relative Frequency Distribution

#### a) Open End Frequency Distribution

Open-end frequency distribution is one which has at least one of its ends open. Either the lower limit of the first class or upper limit of the last class or both are not specified. The words “below” or “less than” and “above” or “more than” are used. In the former the value extends to  $-\infty$  and in the latter to  $+\infty$ . Example of such a frequency distribution is given in Table 10.10.

**Table 10.10**

#### Open-end Class Frequency

| Class        | Frequency |
|--------------|-----------|
| Below 25     | 1         |
| 25 - 30      | 3         |
| 30 - 40      | 5         |
| 40 - 50      | 2         |
| 50 and above | 1         |
| <b>Total</b> | <b>12</b> |

**Table 10.11**

#### Unequal Class Frequency

| Class        | Frequency |
|--------------|-----------|
| 20 - 25      | 1         |
| 25 - 30      | 3         |
| 30 - 40      | 5         |
| 40 - 55      | 2         |
| 55 - 60      | 1         |
| <b>Total</b> | <b>12</b> |

#### b) A Frequency Distribution with Unequal Class Width

The classes of a frequency distribution may or may not be of equal width. A frequency distribution with unequal class width is reproduced in Table 10.11. Here, the width of 1st, 2nd and 5th classes is 5, while that of 3rd is 10 and that of 4th is 15. As we will see in Unit 5, *mode* is not representative value in such types of series and hence not defined.

#### c) Cumulative Frequency Distribution

Suppose that, with reference to data given in Table 10.6, we ask the following questions:

- How many families have their monthly income less than or equal to Rs. 700?
- How many families have their monthly income greater than or equal to Rs. 600?

The answers to the above questions can be easily obtained by forming an appropriate cumulative frequency distribution. To answer the first question, we need to form a “less than type” cumulative frequency distribution while a “greater than type” cumulative frequency distribution is required for answering the second question. These distributions are given in Tables 10.12 and 10.13 respectively.

Table 10.12

## “Less-than type” Cumulative Frequency Distribution

| Monthly Income (Rs.) | Frequencies |                 |
|----------------------|-------------|-----------------|
|                      | Simple      | Cumulative      |
| Less than 550        | 5           | 5               |
| Less than 600        | 6           | 5+6             |
| Less than 650        | 10          | 5+6+10          |
| Less than 700        | 12          | 5+6+10+12       |
| Less than 750        | 9           | 5+6+10+12+9     |
| Less than 800        | 5           | 5+6+10+12+9+5   |
| Less than 850        | 3           | 5+6+10+12+9+5+3 |

Table 10.13

## “More-than type” Cumulative Frequency Distribution

| Monthly Income (Rs.) | Frequencies |                 |
|----------------------|-------------|-----------------|
|                      | Simple      | Cumulative      |
| More than 500        | 5           | 3+5+9+12+10+6+5 |
| More than 550        | 6           | 3+5+9+12+10+6   |
| More than 600        | 10          | 3+5+9+12+10     |
| More than 650        | 12          | 3+5+9+12        |
| More than 700        | 9           | 3+5+9           |
| More than 750        | 5           | 3+5             |
| More than 800        | 3           | 3               |

## d) Relative Frequency Distribution

So far we have expressed the frequency of a value or that of a class as the number of times an observation is repeated. We can also express these frequencies as a *fraction* or a *percentage* of the total number of observations. Such frequencies are known as *the relative frequencies*. Table 10.14 demonstrates the construction of relative frequency distribution.

Table 10.14

## Relative Frequency Distribution of Monthly Income of 50 Families

| Class        | Frequency | Relative Frequency  |                        |
|--------------|-----------|---------------------|------------------------|
|              |           | As a fraction       | As a percentage        |
| 500 - 549    | 5         | $5 \div 50 = 0.10$  | $0.10 \times 100 = 10$ |
| 550 - 599    | 6         | $6 \div 50 = 0.12$  | $0.12 \times 100 = 12$ |
| 600 - 649    | 10        | $10 \div 50 = 0.20$ | $0.20 \times 100 = 20$ |
| 650 - 699    | 12        | $12 \div 50 = 0.24$ | $0.24 \times 100 = 24$ |
| 700 - 749    | 9         | $9 \div 50 = 0.18$  | $0.18 \times 100 = 18$ |
| 750 - 799    | 5         | $5 \div 50 = 0.10$  | $0.10 \times 100 = 10$ |
| 800 - 849    | 3         | $3 \div 50 = 0.06$  | $0.06 \times 100 = 6$  |
| <b>Total</b> | <b>50</b> | <b>1.00</b>         | <b>100</b>             |

From Table 2014 it is clear that the sum of relative frequencies should be either 1 (in the case of fraction) or 100 (in the case of percentages).

### Check Your Progress 1

1. Distinguish between the following, giving at least two points of distinction.

- a) Discrete and continuous frequency distributions
- b) Simple and cumulative frequency distributions
- c) Exclusive and inclusive class intervals
- d) Simple and frequency array

.....

.....

.....

.....

.....

2. Explain the following terms giving examples:

- a) Ungrouped data
- b) Class mark
- c) Open end classes
- d) Class limits
- e) Class boundaries
- f) Class frequencies
- g) Tally bar
- h) Relative frequencies

.....

.....

.....

.....

.....

3. Build a hypothetical frequency distribution on monthly pocket money of 20 students belonging to the lower middle class of a college. Prepare a relative frequency distribution from it.

.....

.....

.....

.....

.....

4. What points are to be kept in mind while taking decisions for preparing a frequency distribution in respect of:

- a) The number of classes, and
- b) Width of the class interval?

.....

.....

.....

.....

5. Construct less than and more than type cumulative frequency distributions from the following data:

|            |         |         |         |         |         |         |
|------------|---------|---------|---------|---------|---------|---------|
| Class:     | 10 - 20 | 20 - 30 | 30 - 40 | 40 - 50 | 50 - 60 | 60 - 70 |
| Frequency: | 5       | 8       | 10      | 12      | 8       | 7       |

.....

.....

.....

.....

.....

6. Construct a relative frequency distribution for the data given in question 5.

.....

.....

.....

.....

### 10.3 TABULATION OF DATA

Good presentation of data is as important as their satisfactory collection and arrangement. In fact, satisfactory collection and arrangement of data must be followed by good presentation. However, good presentation is not an end in itself. It may be necessary for satisfactory analysis and interpretation. A satisfactory presentation helps us in more than one ways. *Firstly*, it helps to highlight significant facts contained in the data. *Secondly*, it facilitates the comparison of data. *Finally*, it facilitates the easy understanding and an intelligent use of statistical information.

We will discuss presentation of statistical data under three heads.

- i) Formal tables
- ii) Graphic methods which will include line graphs, histograms, frequency polygon and curves, and cumulative frequency curves.
- iii) Geometric forms, pictures and statistical maps, which will include pie diagrams, bar diagrams, area and volume diagrams, etc.

In this Section we concentrate on tabular forms of presentation.

### 10.3.1 Meaning and Types of Tables

A table or a statistical table is a systematic arrangement of related statistical data in columns and rows, with a given predetermined and a well decided objective. A row of a table represents a horizontal while a column represents a vertical arrangement of data. To explain the nature of information given in a table, its rows and columns are designated by appropriate stubs and captions (or headings or sub-headings) respectively. Presentation of data in a tabular form should be simple, planned, unambiguous and logical.

Table 10.15 is based on hypothetical figures of exports and imports of country X with countries A, B, C, and D for three years 1995, 1996 and 1997.

**Table 10.15**

**Imports and Exports of X with A, B, C and D during 1995 - 1997**  
(in Crore of Rupees)

| Country      | 1995       |            | 1996       |            | 1997 <sup>@</sup> |            |
|--------------|------------|------------|------------|------------|-------------------|------------|
|              | Imports    | Exports    | Imports    | Exports    | Imports           | Exports    |
| A            | 60         | 70         | 65         | 75         | 70                | 65         |
| B            | 50         | 60         | 60         | 65         | 65                | 60         |
| C            | 40         | 30         | 40         | 40         | 42                | 50         |
| D            | 45         | 42         | 60         | 55         | 63                | 55         |
| <b>Total</b> | <b>195</b> | <b>202</b> | <b>225</b> | <b>235</b> | <b>240</b>        | <b>230</b> |

Note : <sup>@</sup> Figures are quick estimates.

Source : Trade Bulletin, 1998, Ministry of Foreign Trade of X.

In this table it is clear that the purpose is to show the imports and exports of country X vis-à-vis the rest of the world. Note that a particular entry of the table refers to a column and a row. For example, an entry at the intersection of second row and third column indicates that in 1996 country X imported goods and services worth Rs. 60 crores from country B. This figure then can be compared with other import and export figures to seek important interpretations.

### Types of Tables

Basically, we have two types of tables:

- i) Reference tables or general purpose tables
  - ii) Text tables or special purpose tables.
- i) *Reference Tables* are a general purpose tables and are a store of information with the aim of presenting detailed statistical information. From these tables, we can derive our information (i.e., secondary source). Tables presented by different government departments, ministries, Reserve Bank of India, Economic Surveys, etc. are reference tables and are a routine work of these departments. Another important example is the Population Census tables prepared by the Registrar General of India giving detailed information on the demographic features of India.

- ii) Students are advised to consult the latest issue of “Economic Survey” which is issued every year along with the union budget of India. Prepare from it a table on exports and imports of India to USA, UK, Russia, Canada and Germany for three or four years.
- iii) *Text Tables* are the special type of tables. They are smaller in size and are prepared from the reference tables. Their aim is to analyse only a particular aspect to bring out a specific point or to answer a particular question. For example from the Population Census tables we may pick out information on the number of people in Bombay and Delhi who speak different languages (mother tongue), profess different religions and come from different states of India. Similarly from various publications of Reserve Bank of India, we may be able to extract information, in tabular form, on money supply, rate of interest and bank rate for the last ten years or so.

Tables can be simple and one way, like the tables given in Section 10.2, where we deal with only one variable, say, income. Alternatively, it is called a univariate frequency distribution. In addition to this, we can have two-way or multi-way tables where we deal with two or more related characteristics (for example, Table 10.15).

### 10.3.2 Parts of a Table

*Parts* or *elements* of a table vary from table to table depending upon the nature of data and purpose of tabulation. Yet some points are common. These are:

1. **Table number** is required for the identification of a table particularly when there are more than one tables in a particular analysis. Table number is always mentioned in the centre at the top.
2. **Title of the table** gives the indication of the type of information contained in the body of the table. It is said that the *title is to the table what heading is to an essay*. Next to the table number, we mention the title of the table. Its purpose is to answer the questions like:
  - a) *What* is in the table?
  - b) *Where* is it in the table?
  - c) *When* did a particular information occur?
  - d) *How* has a particular information been arranged?

In respect of a sample of a table on exports and imports, (Table 10.15), these questions will be answered as below:

- a) The table contains values of exports and imports of country X.
- b) Information contained in the body of the table shows exports (sales to) and imports (purchases from) four countries A, B, C and D.
- c) These exports and imports occurred in 1995, 1996 and 1997.
- d) Information on exports and imports has been arranged according to year and countries.

## Dos and Don'ts of the Title

Do not opt for long sentences. Title should be brief and to the point. Present the title in bold letters and/or in capital letters. Expressions used should not convey more than one meaning.

Avoid the expressions like 'Table Presents ....., ' or 'A Detailed Comparison of Data Relating to ....., ' etc. It should be like a telegraphic message.

3. **Head note**, also called prefatory note, is written just below the title. It shows contents and unit of measurement like (rupees crore) or (lakh tonnes) or (thousand bales). It should be written in brackets and should appear on right side top just below the title. However, every table does not need a head note, like number of students in each class.
4. **Stubs** are used to designate rows. They appear on the left hand column of the table. Stubs consist of two parts:
  - a) *Stub head* describes the nature of stub entry.
  - b) *Stub entry* is the description of row entries.
5. **Captions**, also called box heads, designate the data presented in the columns of the table. It may contain more than one column heads, and each column head may be sub-divided in more than one sub-head. For example, we can divide the students of a college into hostelers and non-hostlers and then again into males and females. This will help us to know the number of male hostelers in, say, first year, second year and third year.
6. **Main body of the table**, also called *field* of the table, is its most important and bulky part. It contains the relevant numerical information about which a hint is already contained in the title of the table. In our example of Table 10.15 the title amply suggests that the body of the table contains numerical information on exports and imports of country X for a period of three years.
7. **Foot Note** is a qualifying statement put just below the table (at the bottom). Its purpose is to caution about the limitations of the data or certain omissions. For example in Table 10.15, the foot note reads that "figures are quick estimates" implying that these figures are not final. Similarly in the latest population census data the foot note may be "Excluding the State of Jammu and Kashmir".
8. **Source of data** may be the last part of a table, yet it is important one. It speaks about the authenticity of the data quoted. It also offers opportunity to the reader to check the data if he so desires and get more of it.

Taking all these points into consideration, the format of a hypothetical table is presented below:

Table No. 10.16

( -----TITLE----- )

Head note:

(in rupees crore)

| Stub Head    | ←----- Caption -----→ |          |                |          |
|--------------|-----------------------|----------|----------------|----------|
| Stub Entries | Column Head I         |          | Column Head II |          |
|              | Sub-head              | Sub-head | Sub-head       | Sub-head |
|              | MAIN                  | BODY     | OF THE         | TABLE    |
| Totals       |                       |          |                |          |

Foot note(s):

Source:

### 10.3.3 Importance of Tables

Numerical information arranged in tabular form has distinct advantage over other forms of presentation. First, tabulated data are easy to understand and interpret. Secondly, one can make quick comparison between different characteristics, for example, 'Are imports greater than exports over all the three years?' or 'Are exports increasing?' Thirdly, it opens doors for further investigations. Fourthly, they have a more lasting impression on human mind than the textual statements. Needless to say, that the statistical tables are used extensively in almost all fields of human inquiry.

### Check Your Progress 2

1. Distinguish between

- Caption, stub head and stub entries
- One-way and two-way tables
- Reference tables and text tables
- Column entry and row entry
- Head note and foot note

.....

.....

.....

.....

.....

2. Comment on the statement: "Title is to the table what heading is to an essay".

.....

.....

.....

.....

.....

3. Enumerate the various parts of a Statistical table.

.....

.....

.....

.....

4. Make a sketch of a two-way table to show the following information:

- a) Division of college according to 1st Year, 2nd Year and 3rd Year students
- b) Hosteler and non-hostelers
- c) Male and female students

Take hypothetical information.

.....

.....

.....

.....

## 10.4 GRAPHICAL PRESENTATION OF DATA

Besides formal tables, statistical data can also be presented in the form of various types of graphs. Graphs are a useful way of conveying information very quickly and briefly. With the same ease and efficiency, they help in comparing data over time and space. They are visual aids and have a powerful impact on the people. It is often said, “*a picture is worth a thousand words*”. They attract a reader’s attention to what they are supposed to convey about the data. Further, they may help us to estimate some values at a glance, and serve as a pictorial check on the accuracy of our solutions.

However, graphical presentation of data, although useful in different ways mentioned above, is only one method of describing data. This cannot and is not a substitute for other forms of presentation as well as further statistical analysis. In the following, we discuss some of the graphical methods of presentation.

### 10.4.1 Line Graphs

Although there are four quadrants on a plane, in economics we usually draw our diagrams only in the first quadrant where both the quantities measured on X-axis and Y-axis are positive. Economic quantities like price, quantity demanded and supplied, national income, consumption, production and host of other such variable are non-negative ( $\geq 0$ ).

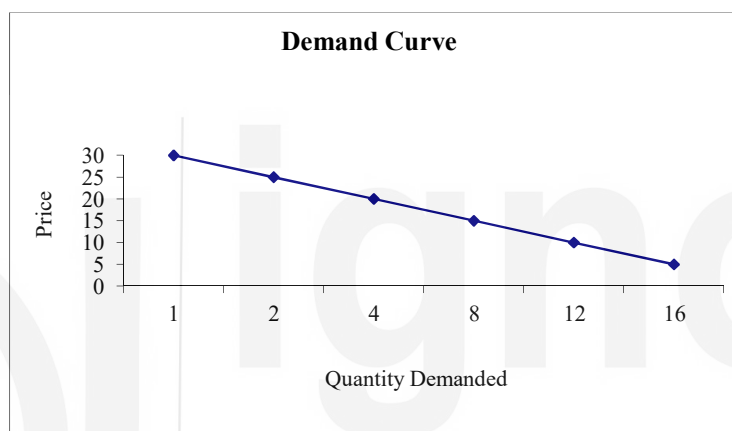
Let us take a demand schedule and plot it on the graph. The resultant curve on joining different points, assuming continuity, will give us line graph expressing relation between price and quantity demanded. Such a line graph in Economics is called a *demand curve*. Note that price is measured on Y-axis and quantity demanded on X-axis. The demand curve for data given in Table 10.17 is given in Fig. 10.1.

**Table 10.17**  
**Demand Schedule**

| Price of X (Rs.) | Quantity of X demanded |
|------------------|------------------------|
| 5                | 16                     |
| 10               | 12                     |
| 15               | 8                      |
| 20               | 4                      |
| 25               | 2                      |
| 30               | 1                      |

**Table 10.18**  
**Time Series Data**

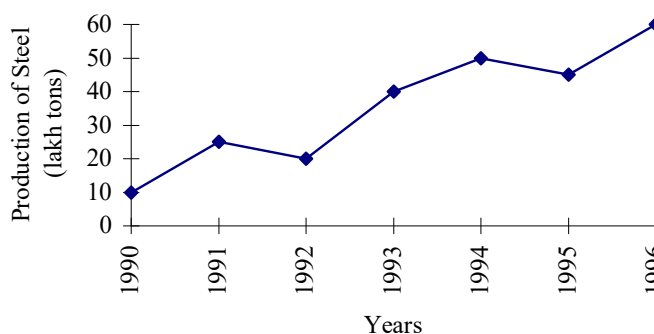
| Year | Production of Steel(tons) |
|------|---------------------------|
| 1990 | 10                        |
| 1991 | 25                        |
| 1992 | 20                        |
| 1993 | 40                        |
| 1994 | 50                        |
| 1995 | 45                        |
| 1996 | 60                        |



**Fig. 10.1**

A line graph may be used to show changes in some economic variable, say, steel production over time. In other words, if out of the two variables, one happens to be time (months, years, etc.), we get a line graph over time or simply *time series graph* or *historigram*. A time series expresses behaviour of an economic variable over time. An example of time series data is given in Table 10.18. Measuring years on X-axis and steel production on Y-axis, we can plot time series data on a graph, as shown in Fig.10.2.

**Historigram of Production of Steel**



**Fig. 10.2**

### 10.4.2 Histogram, Frequency Polygon and Frequency Curve

Histogram (do not confuse with historigram discussed earlier) is a very common type of graph for displaying classified data. It is a set of rectangles erected vertically. It has the following features:

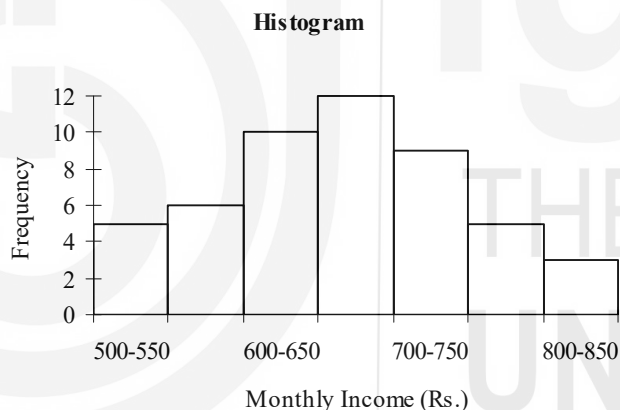
- It is a rectangular diagram.
- Since the rectangles are drawn with specified width and height, histogram is a two dimensional diagram. The width of a rectangle equals the class interval and height

$$= \frac{\text{Class frequency} \times \text{Width of the shortest class interval in the data}}{\text{Width of the class interval}}$$

- The area of each rectangle is proportional to the frequency of the respective class.

#### Construction of Histogram

To plot a histogram of the frequency distribution given in Table 10.6 on a graph paper, we mark off class intervals like 500 - 550, 550 - 600, etc. on the horizontal axis. Similarly, we mark off frequencies on the vertical axis. Since all the classes are of equal width, the height of each rectangle is taken to be equal to the frequency of the respective class. The histogram is shown in Fig. 10.3.



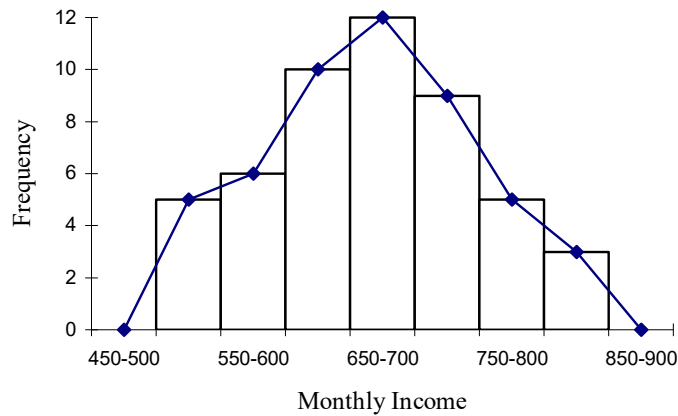
**Fig. 10.3**

Advantages of histogram are:

- The width of various rectangles show the nature of classes in the distribution, i.e., whether of equal width or not.
- Area of a rectangle shows the proportion of the class frequency in the total.

#### Frequency Polygon

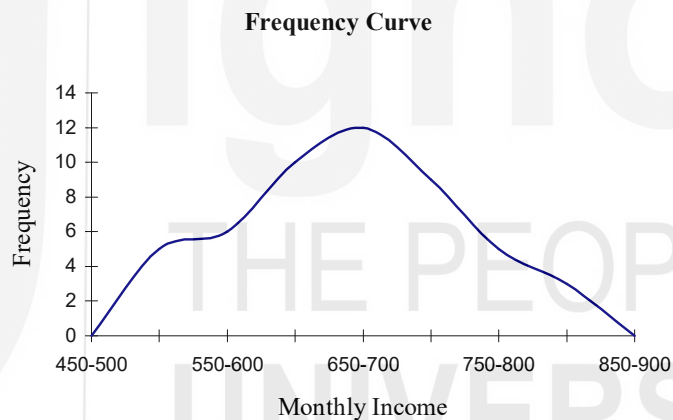
Frequency Polygon has been derived from the word “polygon” which means many sides. In statistics, it means a graph of a frequency distribution. A frequency polygon is obtained from a histogram by joining the mid-points of the top of various rectangles with the help of straight lines, as shown in Fig. 10.4. In order that total area under the polygon remains equal to the area under histogram, two arbitrary classes, each with zero frequency, are added on both ends, as shown below.



**Fig. 10.4: Frequency Polygon**

### Frequency Curve

If the points, obtained in case of frequency polygon are joined with the help of a smooth curve, we get a frequency curve as shown in Fig. 10.5.



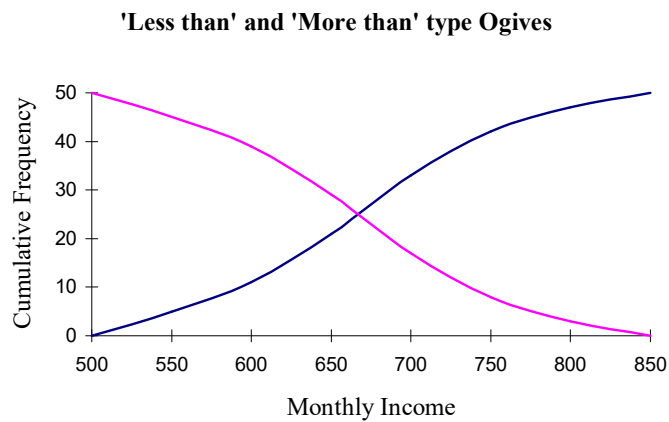
**Fig. 10.5**

### 10.4.3 Cumulative Frequency Curve – Ogive

The graph of a cumulative frequency distribution is known as cumulative frequency curve or ogive. Since a cumulative frequency distribution can be of 'less than' or 'greater than' type, accordingly, we can have 'less than' or 'greater than' type of ogives.

Ogives can be used to locate, graphically, certain partition values. We can also determine the percentage of observations lying between given limits. The ogives for the cumulative frequency distributions given in Tables 10.12 and 10.13 are drawn in Fig. 10.6.

Note that to draw a less than type ogive, we add a class interval of 'less than 500' with frequency equal to zero. Similarly, we add a class interval of 'more than 900' with frequency zero for the construction of a greater than type ogive.



**Fig 10.6**

## 10.5 DIAGRAMMATIC PRESENTATION OF DATA

A diagram is a visual form for the presentation of statistical data. Diagram refers to bars, squares, circles, maps, pictorials, cartograms, etc. Diagrams are different from graphs as the former are used only for presentation while the later can be used for analysis in addition to the presentation of data.

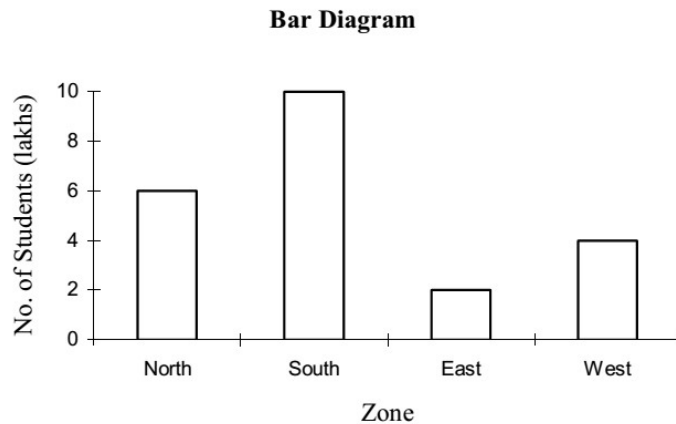
### 10.5.1 One Dimensional Diagrams

These are also known as *bar diagrams*. A *bar* is defined as a *thick line*, often made thicker to attract the attention of a reader. The height of the bar highlights the value of the variable with *width presenting nothing*. Therefore, it has nothing to do with the area of the bar. It is different from the histogram where both the width as well as the height of the bar are important. Further, the bars of the bar diagram are separated from one another so that the gap between the successive bars is same, whereas in histogram they are placed adjacent to one another with out gap. Finally, in histogram the bars are always vertically placed whereas in bar diagram they can be placed both vertically as well as horizontally. Let us take a simple example to demonstrate the construction of a bar diagram.

**Table 10.19**  
**Number of students in four zones of a country**

| <u>Zone</u> | <u>No. of students (lakhs)</u> |
|-------------|--------------------------------|
| North       | 6                              |
| South       | 10                             |
| East        | 2                              |
| West        | 4                              |

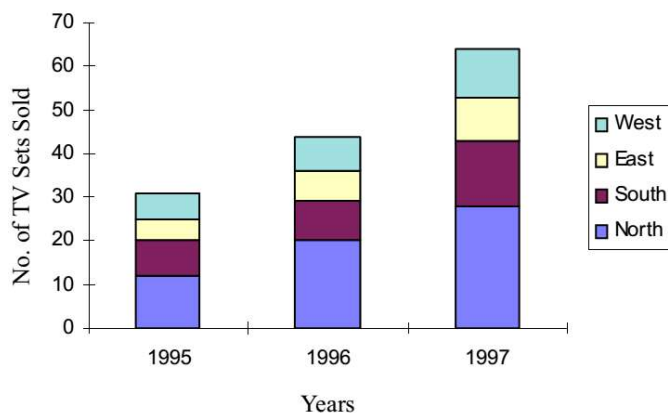
The bar diagram of the above data is drawn in Fig. 10.7. To make the bar diagram beautiful we can either colour the bars or shade them in different ways. This is left to the aesthetic taste of the investigator.

**Fig. 10.7****i) Sub-divided or Component Bar Diagram**

A sub-divided bar diagram is used when it is desired to represent the comparative values of different components of a phenomenon. In this diagram, the bars, corresponding to each phenomenon, is divided into various components. The portion of the bar occupied by each component denotes its share in the total. The sub-divisions of different bars must always be done in the same order and these should be distinguished from each other by using different colours or shades. A sub-divided bar diagram for the hypothetical data on sales of TV sets, given in Table 10.20 is drawn in Fig. 10.8.

**Table 10.20****Zone-wise sale of TV sets (1995-1997)**

| Zone  | Number of T.V. Sets sold (lakhs) |      |      |
|-------|----------------------------------|------|------|
|       | 1995                             | 1996 | 1997 |
| North | 12                               | 20   | 28   |
| South | 8                                | 9    | 15   |
| East  | 5                                | 7    | 10   |
| West  | 6                                | 8    | 11   |
| Total | 31                               | 44   | 64   |

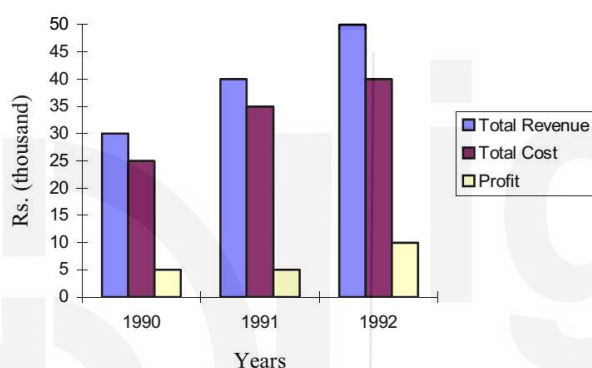
**Fig. 10.8: Sub-divided or Component Bar Diagram**

## ii) Multiple Bar Diagram

This diagram is used when comparisons are to be shown between two or more sets of data. A set of bars for a period, place or a related phenomenon are drawn side by side without gap. Different bars are distinguished by different shades or colours. A multiple bar diagram for the hypothetical data given in Table 10.21 is drawn in Fig. 10.9.

**Table 10.21**  
**Total revenue, total cost and profit of M/S XYZ (1990-92)**  
(Rupees thousand)

| Year | Total Revenue | Total cost | Profit |
|------|---------------|------------|--------|
| 1990 | 30            | 25         | 5      |
| 1991 | 40            | 35         | 5      |
| 1992 | 50            | 40         | 10     |



**Fig. 10.9**

### 10.5.2 Two Dimensional Diagrams or Area Diagrams

In the case of one dimensional diagrams only the height of the bar is important, and the width can be chosen according to convenience or aesthetic taste of the investigator. But in the case of two dimensional diagrams, area is more important. That is why they are also known as *Area diagrams*. There are three types of area diagrams.

- Rectangles*, where area equals width (or base) multiplied by the length ( or height) of the rectangle.
- Squares* where area equals square of side (or base).
- Circles* where area equals  $\pi r^2$ , with  $\pi=22/7$  and  $r$  = radius.

Let us consider data on, say, average salaries of three categories of University teachers, and prepare all the three types of area diagrams.

**Table 10.22**  
**Average Salaries of University Teachers as on 1/1/1998**

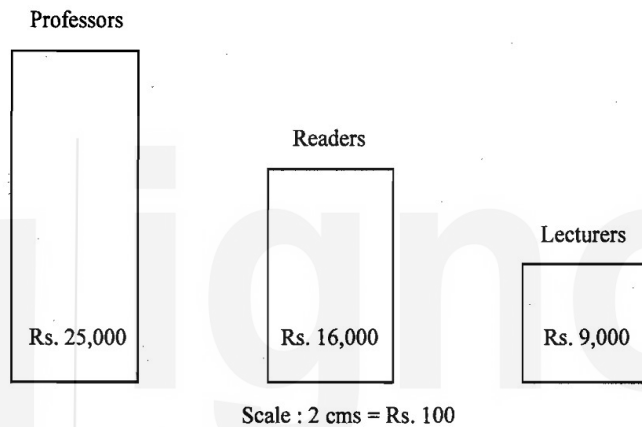
| <u>Class of Teachers</u> | <u>Average Salaries (Rs.)</u> |
|--------------------------|-------------------------------|
| Professors               | 25,000                        |
| Readers                  | 16,000                        |
| Lecturers                | 9,000                         |

- a) For drawing rectangles, a common base of, say, 100 is taken. Accordingly, the heights can be determined as:

1. Salary of Rs.25,000 = 100 (base)  $\times$  250 (height)
2. Salary of Rs.16,000 = 100 (base)  $\times$  160 (height)
3. Salary of Rs. 9,000 = 100 (base)  $\times$  90 (height)

Now take a scale of 2 cm = 100, so that the first rectangle has dimensions of 2 cm.  $\times$  5 cm, the second one has the dimensions of 2 cm  $\times$  3.2 cm and the third one has the dimensions of 2 cm  $\times$  1.8 cm. After this, we are in a position to draw the rectangles as area diagrams (Fig. 10.10).

#### Average Salaries of University Teachers (Rs.)



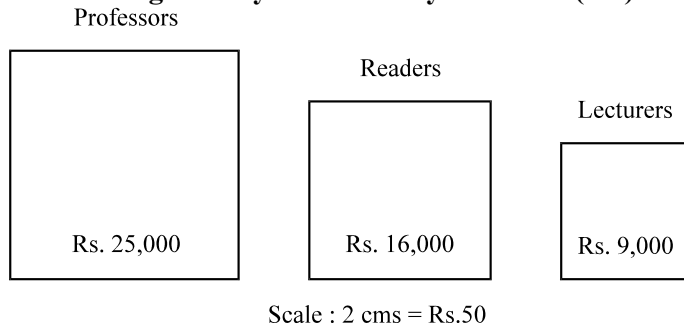
**Fig. 10.10**

- b) For drawing squares, we find the square root of various incomes. We have,

1.  $\sqrt{25,000} = 158.114$
2.  $\sqrt{16,000} = 126.491$
3.  $\sqrt{9000} = 94.868$

Choose a scale 1 cm = 50 so that the first square has each side approximately equal to 3.2 cm. (since  $158.114/50 \cong 3.2$ ), second has the side of 2.53 cm. and the third has the side of 1.9 cm. The relevant squares are drawn in Fig. 10.11.

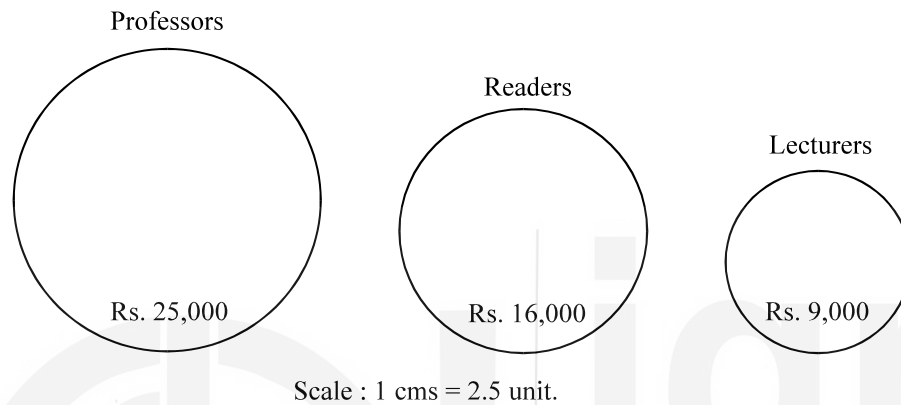
#### Average Salary of University Teachers (Rs.)



**Fig. 10.11**

- c) For drawing Circles we take the squares of their radii in the ratio of areas, i.e., 25000: 16000: 9000 or 25: 16: 9. This is based on the property of the circles that area of a circle is proportional to the square of its radius. Let  $r_1$ ,  $r_2$  and  $r_3$  denote the radii of the three circles, then we can write  $r_1^2 : r_2^2 : r_3^2 = 25 : 16 : 9$  or  $r_1 : r_2 : r_3 = 5 : 4 : 3$ . Taking 1 cm = 2.5 units, the radii of the three circles will be 2.0, 1.6 and 1.2 centimetres respectively. Let us draw the required circles.

**Average Salary of University Teachers (Rs.)**



**Fig. 10.12**

**10.5.3 Pie Diagram or Pie Chart**

It is also known as angular diagram. It is used to represent percentage break downs of the given data. For example the exports of a country to different countries and continents of the world can be expressed into ratios or percentages. These ratios or percentages can then be converted into angles by the formula

$$\frac{\text{Share of the sub - division}}{\text{Total}} \times 360^\circ$$

**Table 10.23**

**Exports of X to A, B, C and D in 1990**

| Country | Exports | Percentage Share                    | Degree                                            |
|---------|---------|-------------------------------------|---------------------------------------------------|
| A       | 300     | $(300 \times 100) \div 800 = 37.50$ | $(37.5 \times 360^\circ) \div 100 = 135^\circ$    |
| B       | 250     | $(250 \times 100) \div 800 = 31.25$ | $(31.25 \times 360^\circ) \div 100 = 112.5^\circ$ |
| C       | 150     | $(150 \times 100) \div 800 = 18.75$ | $(18.75 \times 360^\circ) \div 100 = 67.5^\circ$  |
| D       | 100     | $(100 \times 100) \div 800 = 12.50$ | $(12.5 \times 360^\circ) \div 100 = 45^\circ$     |
| Total   | 800     | 100                                 | $360^\circ$                                       |

## Pie Diagram Representing Exports of X

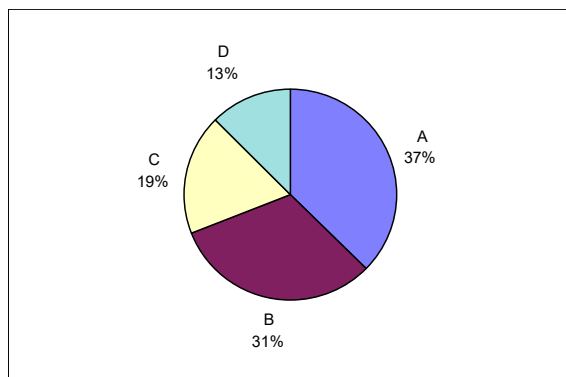


Fig. 10.13

## Steps in the construction of Pie diagram

1. Find the total of all components.
2. Find ratio or percentage of the share of sub-division to the total and multiply by  $360^\circ$  to get the angle corresponding to the each sub-division.
3. Draw a circle of a suitable size.
4. Use protractor to draw different angles at the centre. Preferably start with the largest one.
5. Shade the different segments with different colours or shades.
6. Write the components with percentage values in the marked, shaded or coloured areas.

## 10.5.4 Three Dimensional Diagrams

These diagrams are not very popular and are used very rarely. Since these diagrams are three dimensional (involving length, breadth and width), they denote volumes. They can take the form of boxes, cubes, blocks, spheres and cylinders. They are very useful when the variations in magnitudes of the observations are very marked. Here we will explain only the presentation of data by cubes for which we take the following steps:

- i) Find cube-root of each figure.
- ii) Take a convenient scale, preferably in centimeters.
- iii) Draw cubes, dimensions of which are calculated below for an example consisting of two classes of families: Poor and Very Rich.

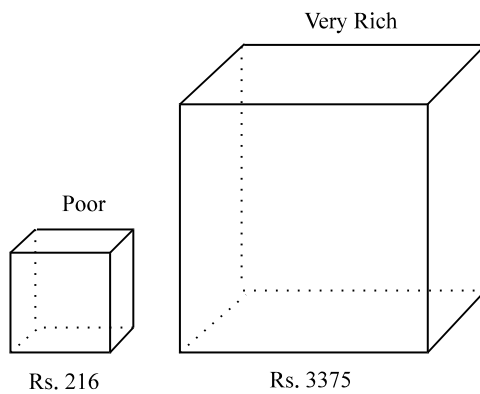
Table 10.24

|    | Income class | Income (Rs.) | Cube-root             | Side of cube |
|----|--------------|--------------|-----------------------|--------------|
| 1. | Poor         | 216          | $\sqrt[3]{216} = 6$   | 1.5 cms.     |
| 2. | Very Rich    | 3375         | $\sqrt[3]{3375} = 15$ | 3.75 cms.    |

Scale: 1 cm. = 4 units.

- iv) Now draw two cubes with sides equal to 1.5 cms. and 3.75 cms. respectively.

### Income levels of Poor and Very Rich people (Rs.)

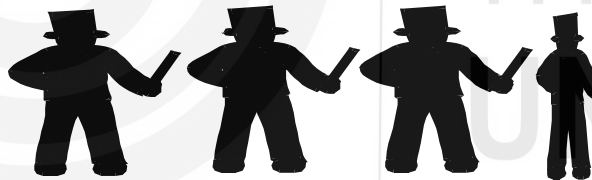


**Fig. 10.14**

#### 10.5.5 Pictograms and Statistical Maps

These are also known as catograms. Pictures are more attractive to laymen than other forms of graphic presentations.

But these are not suitable everywhere. It may suit cases involving population of people of a state or number of vehicles in a metropolitan city like Delhi or Bombay. For showing population of human beings, we draw human figures. Here also we have a scale. We may represent 1 lakh people by one human figure so that a population of three and half lakhs is shown by drawing 3 1/2 human figures, as given in Fig. 10.15.



**Fig. 10.15**

Pictograms suffer from a defect that they present only approximate values. For more accurate presentations bar diagrams are preferable.

#### Check Your Progress 3

- 1) Distinguish between the following giving at least two points of distinction.
  - a) Histogram and historigram.
  - b) Histogram and bar diagram.
  - c) Histogram and frequency polygon.
  - d) “Less-than” and “More-than” ogives.
  - e) Pie diagram and circle.

.....

.....

.....

2) Prepare a sub-divided bar chart and a pie diagram from the following data.

| Academic<br>Year | Expenditure on Books |          |       |           |       |
|------------------|----------------------|----------|-------|-----------|-------|
|                  | Economics            | Commerce | Maths | Languages | Total |
| 1996 - 97        | 5200                 | 10000    | 5000  | 4800      | 25000 |
| 1997 - 98        | 8000                 | 14000    | 7000  | 6000      | 35000 |

.....

.....

.....

.....

.....

3) Explain the following terms:

- Line graph
- Bar diagram
- Sub-divided or component bar diagram
- Multiple bar diagram
- Area diagram
- Volume diagram

.....

.....

.....

.....

.....

4) Fill in the blanks with a suitable word out of those given in brackets:

- A pie diagram is also called \_\_\_\_\_ diagram. (bar, angular, multiple bar).
- In the case of vertical bars, the variable is measured on the \_\_\_\_\_. (X-axis, plane, Y-axis).
- Bar diagrams, rectangles, squares, circles and pie charts are \_\_\_\_\_ forms of presenting data. (geometric, arithmetic, horizontal).
- By joining the mid-points of the top of each rectangle of a histogram, we get \_\_\_\_\_. (an Ogive, a frequency curve, a frequency polygon)
- Graph of “more-than” cumulative frequency distribution is also called “more - than” \_\_\_\_\_. (Ogive, frequency polygon, frequency curve)
- The caption of a table labels data presented in the \_\_\_\_\_ of a table. (rows, columns, foot-note)

5) Are the following statements true or false? If false, what should be the correct statement?

- A picture is worth a thousand words.
- Squares and circles are examples of area diagrams.

- c) We can have only vertical bar to present some data having one variable.
- d) The graph of an ordinary frequency distribution is called ogive.
- e) A time series graph is known as historigram.
- f) Histogram is same as bar diagram.

.....

.....

.....

.....

.....

---

## 10.6 LET US SUM UP

---

Collected data are unorganised and complex mass of figures. To draw some meaningful conclusions, they must be arranged in an orderly manner. This can be done in many ways, such as by forming simple and frequency array, discrete and continuous frequency distributions, etc.

Sometimes, it serves a useful purpose to form what is called “*less-than*” or “*more-than*” cumulative frequency distributions. Former is formed by successive totaling of frequencies from above and the latter by successive totaling from below.

After collection and condensation of data, good presentation of data is important. A good presentation helps to highlight important points of the data and makes possible useful comparisons and their intelligent use. This can be done through formal tables; line graphs; histograms, frequency polygon and frequency curves; “less-than” and “more-than” ogives; geometric forms – one, two and three dimensional diagrams such as bar diagrams, rectangles, squares, circles, cubes and pie diagrams; statistical maps. While using diagrams, their limitations must always be kept in mind. Diagrams give only a vague idea of the problem and can portray only a limited number of characteristics. Unlike a graphic presentation, the main limitation of a diagrammatic presentation is that it cannot be used as a tool of analysis. The level of accuracy of a graphic method is often lower than that of mathematical method.

---

## 10.7 ANSWER OR HINTS TO CHECK YOUR PROGRESS EXERCISES

---

### Check Your Progress 1

1. (a) See Sub-Section 10.2.2 and 10.2.3.
- (b) See Sub-Section 10.2.4.
- (c) See Sub-Section 10.2.3.
- (d) See Sub-Section 10.2.1 and 10.2.2.

2. You may give examples from your surrounding. For exact meaning of the terms refer to Section 10.2.
3. In the text we have converted the monthly income data in Table 10.2 to a frequency distribution in Table 10.6. From this you can take a clue.
4. Refer to Sub-Section 10.2.3.
5. Refer to Sub-Section 10.2.4(c).
6. Refer to Sub-Section 10.2.4(d).

### Check Your Progress 2

1. Refer to Table 10.16 and Sub-Section 10.3.2 for different parts of a table.
2. Refer to Sub-Section 10.3.2(2).
3. Refer to Table 10.16.
4. It can be presented in more than one ways. We have given one below. Try another.

**Division of Students of XY College**

| Year        | Hostelers |        | Non-Hostelers |        |
|-------------|-----------|--------|---------------|--------|
|             | Male      | Female | Male          | Female |
| First Year  |           |        |               |        |
| Second Year |           |        |               |        |
| Third Year  |           |        |               |        |

### Check Your Progress 3

1.
  - (a) See Sub-Sections 10.4.1 and 10.4.2
  - (b) See Sub-Sections 10.4.2 and 10.5.1
  - (c) See Sub-Section 10.4.2
  - (d) See Sub-Section 10.4.3
  - (e) See Sub-Section 10.5.2 and 10.5.3
2. Refer to Sub-Sections 10.5.1 and 10.5.3
3.
  - (a) See Sub-Section 10.4.1
  - (b) See Sub-Section 10.5.1
  - (c) See Sub-Section 10.5.1
  - (d) See Sub-Section 10.5.1
  - (e) See Sub-Section 10.5.2
  - (f) See Sub-Section 10.5.4
4.
  - (a) angular
  - (b) y-axis
  - (c) geometric
  - (d) a frequency polygon
  - (e) ogive
  - (f) columns
5. True: 1, 2, 5.  
False: 3, 4, 6.

---

# UNIT 11 MEASURES OF CENTRAL TENDENCY \*

---

## Structure

- 11.0 Objectives
- 11.1 Introduction
- 11.2 Measures of Central Tendency
  - 11.2.1 Arithmetic Mean
  - 11.2.2 Median
  - 11.2.3 Mode
- 11.3 Other Measures of Central Tendency
  - 11.3.1 Geometric Mean and Harmonic Mean
  - 11.3.2 Weighted Mean
  - 11.3.3 Pooled Mean
  - 11.3.4 Choosing a Measures of Central Tendency
- 11.4 Let Us Sum Up
- 11.5 Answers or Hints to Check Your Progress Exercises

---

## 11.0 OBJECTIVES

---

After going through this unit, you will be able to:

- compute numerical quantities that measure the central tendency of a set of data such as, mean, median, mode, geometric mean and harmonic mean.

---

## 11.1 INTRODUCTION

---

In the previous Unit we had discussed about condensation of raw data by grouping them into a few class intervals and presenting in the form of a table or diagram. Such tables or diagrams provide a rough idea of the distribution of observations. Often we need to compare between distributions. In such situations it is difficult to compare tables or diagrams simply by looking at them. It is much more convenient and useful for comparison if we could find out a single numerical value for describing the data.

Measures of Central Tendency (or Location) constitute one of the major statistics designed for this purpose. There are five main measures of central tendency. These are Arithmetic Mean, Geometric Mean, Median and Mode. You will learn about each one of these measures below.

---

## 11.2 MEASURES OF CENTRAL TENDENCY

---

In frequency distributions of observations, we notice that the observations tend to cluster around a central value. This phenomenon of clustering around a central value in a frequency distribution is called '*Central Tendency*'. Thus, it is of

---

\* Adopted from IGNOU course material of EEC 13: *Elementary Statistical Methods and Survey Techniques*, Unit 4 written by R S Bharadwaj with modifications by K. Barik

interest to locate such a value around which clustering of observations takes place. There are several measures of central tendency (or location) of a frequency distribution. These measures produce numbers that summarise a frequency distribution in terms of one its properties, namely, central tendency.

### 11.2.1 Arithmetic Mean

The *average* or the *arithmetic mean*, or simply the *mean* when there is no ambiguity, is the most common measure of central tendency. It is defined as the sum total of all values in the sample divided by the number of observations. It is denoted by a bar above the symbol of the variable being averaged. Thus  $\bar{X}$  stands for the mean of  $X$ -values in the sample. If in a sample a particular  $X$ -value, say  $X_i$  occurs with frequency  $f_i$  ( $i = 1, 2, \dots, n$ ), its contribution to the total of  $X$ -values is  $f_i X_i$ . Thus, we can compute the mean of  $X$ -values by

$$\bar{X} = \frac{1}{N} (f_1 X_1 + f_2 X_2 + \dots + f_n X_n) = \frac{\sum_{i=1}^n f_i X_i}{N}, \quad \text{where } N = \sum_{i=1}^n f_i.$$

When observations are classified into class intervals, as for continuous variables, individual observations falling into a class intervals are not separately identifiable and be contribution of the individual observations from a class intervals to the total cannot be calculated. To avoid this difficulty, it is assumed that every observation falling into a class interval has a value equal to the *mid-point* into which these observations fall. Such a procedure will not give the exact mean had we computed it from raw data and may require what is called corrections for grouping.

**Example 11.1:** Compute the mean for discrete frequency distribution of Table 11.1.

**Table 11.1**  
**Frequency distribution of 100 households by size**

| Household Size<br>( $X_i$ ) | Frequency ( $f_i$ ) |
|-----------------------------|---------------------|
| 1                           | 3                   |
| 2                           | 16                  |
| 3                           | 25                  |
| 4                           | 33                  |
| 5                           | 12                  |
| 6                           | 7                   |
| 7                           | 2                   |
| 8                           | 2                   |
| Total                       | 100                 |

Let us compute the arithmetic mean of the data given in the above table.

$$\bar{X} = \frac{\sum_{i=1}^n f_i X_i}{N} = \frac{1 \times 3 + 2 \times 16 + 3 \times 25 + 4 \times 33 + 5 \times 12 + 6 \times 7 + 7 \times 2 + 8 \times 2}{100} = \frac{374}{100} = 3.74$$

Thus, mean household size based on 100 households is 3.74.

**Example 11.2:** Compute the mean for grouped frequency distribution of Table 11.2.

**Table 11.2**

**Frequency distribution of 100 households by average monthly household expenditure on food**

| Expenditure class (Rs.) | Frequency  |
|-------------------------|------------|
| 262.5 – 286.5           | 1          |
| 286.5 – 310.5           | 14         |
| 310.5 – 334.5           | 16         |
| 334.5 – 358.5           | 28         |
| 358.5 – 382.5           | 26         |
| 382.5 – 406.5           | 15         |
| <b>Total</b>            | <b>100</b> |

For computation of the mean we have to construct table as given below.

| Class interval (Rs.) (1) | Mid-point (X <sub>i</sub> ) (2) | Frequency (f <sub>i</sub> ) (2) | f <sub>i</sub> X <sub>i</sub> (4) |
|--------------------------|---------------------------------|---------------------------------|-----------------------------------|
| 262.5 -286.5             | 274.5                           | 1                               | 274.5                             |
| 286.5 – 310.5            | 298.5                           | 14                              | 4179.0                            |
| 310.5 – 334.5            | 322.5                           | 16                              | 5160.0                            |
| 334.5 – 358.5            | 346.5                           | 28                              | 9702.0                            |
| 358.5 – 382.5            | 370.5                           | 26                              | 9633.0                            |
| 382.5 – 406.5            | 394.5                           | 15                              | 5917.5                            |
| <b>Total</b>             |                                 | <b>100</b>                      | <b>34866.0</b>                    |

Thus, mean of monthly average household expenditure on food is

$$\bar{X} = \frac{34866}{100} = \text{Rs.}348.66.$$

We should note from the above example that to find column (3) we need to multiply the corresponding values of column (1) and (2), and often hand computations are long for each multiplication. These computations can be

simplified, particularly when successive column (1) values are equidistant (but applicable otherwise also), by making the following simple transformation.

For  $i = 1, 2, \dots, n$

$$u_i = \frac{X_i - A}{h} \quad \text{i.e., } X_i = A + hu_i \quad \text{and so } \bar{X} = A + h\bar{u}.$$

Often  $A$  is called the ‘assumed mean’ and  $h\bar{u}$  as its correction to get  $\bar{X}$ . Choice of  $A$  and  $h$  are made so that computation of  $\bar{u}$  becomes simple. Usually  $A$  is taken as that  $X$  value for which the frequency is largest. For equidistant successive  $X$ -values is column (1),  $h$  may be taken as the difference between two successive  $X$ -values. For equal length class intervals, the difference between successive mid-points is the same as the length of each class interval.

We will explain this method by re-computing the mean of the monthly average household food expenditure data given in Table 11.2. We construct Table 11.3 by using  $A$  and  $h$  as explained below.

We define  $A = \text{Mid-point of the class with largest frequency} = 346.5$  and

$h = \text{Common length of each class interval} = 24$ .

Thus, 
$$u_i = \frac{X_i - 346.5}{24}$$

**Table 11.3**

**Computation of Mean of Frequency Distribution of Table 11.2**

| Class interval (Rs.) | Mid-point ( $X_i$ ) | $u_i = \frac{X_i - 346.5}{24}$ | frequency ( $f_i$ ) | $f_i u_i$ |
|----------------------|---------------------|--------------------------------|---------------------|-----------|
| 262.5 – 286.5        | 274.5               | –3                             | 1                   | –3        |
| 286.5 – 310.5        | 298.5               | –2                             | 14                  | –28       |
| 310.5 – 334.5        | 322.5               | –1                             | 16                  | –16       |
| 334.5 – 358.5        | 346.5               | 0                              | 28                  | 0         |
| 358.5 – 382.5        | 370.5               | 1                              | 26                  | 26        |
| 382.5 – 406.5        | 394.5               | 2                              | 15                  | 30        |
| <b>Total</b>         |                     |                                | <b>100</b>          | <b>9</b>  |

We find out that

$$\bar{u} = \frac{1}{N} \sum_{i=1}^n f_i u_i = \frac{1}{100} \times 9 = \frac{9}{100}$$

Thus,  $\bar{X} = A + h \times \bar{u} = 346.5 + 24 \times \frac{9}{100} = \text{Rs.} 348.66$  as was computed earlier.

### Properties of Arithmetic Mean

### Measures of Central Tendency

- 1) The algebraic sum of deviations of a given set of observations is zero when taken from the arithmetic mean.

Let  $X_1, X_2, \dots, X_n$  be  $n$  observations with respective frequencies as  $f_1, f_2, \dots, f_n$ .

Mathematically, this property implies that  $\sum_{i=1}^n f_i (X_i - \bar{X}) = 0$ , where  $X_i - \bar{X}$  is the deviation of  $i^{\text{th}}$  observation from mean. To prove the above property, we write

$$\sum_{i=1}^n f_i (X_i - \bar{X}) = \sum_{i=1}^n f_i X_i - \bar{X} \sum_{i=1}^n f_i = \sum_{i=1}^n f_i X_i - n \cdot \bar{X} = 0.$$

Hence, the result,

- 2) The sum of squares of deviations of a given set of observations is minimum when taken from the arithmetic mean.

Mathematically, this property implies that for any arbitrarily chosen origin,  $A$ ,

$$S = \sum_{i=1}^n f_i (X_i - A)^2 \text{ is minimum when } A = \bar{X}.$$

To prove this property, we note that the magnitude of  $S$  will depend upon the selected value of  $A$ . thus, we can say that  $S$  is a function of  $A$ .

We want to find that value of  $A$  for which  $S$  is minimum. Using calculus, this value is given by the equation  $\frac{dS}{dA} = 0$  such that  $\frac{d^2S}{dA^2} > 0$ .

(Remember that the value of a function is minimum when first derivative is zero and second derivative is positive.)

Differentiating  $S$  with respect to  $A$  and equating to zero, we get

$$\frac{dS}{dA} = -2 \sum_{i=1}^n f_i (X_i - A) = 0$$

This implies that

$$\sum_{i=1}^n f_i X_i - A \sum_{i=1}^n f_i = 0 \quad \text{or} \quad A = \frac{\sum_{i=1}^n f_i X_i}{\sum_{i=1}^n f_i} = \bar{X}.$$

Further, it can be shown that  $\frac{d^2S}{dA^2} > 0$  when  $A = \bar{X}$ .

### 11.2.2 Median

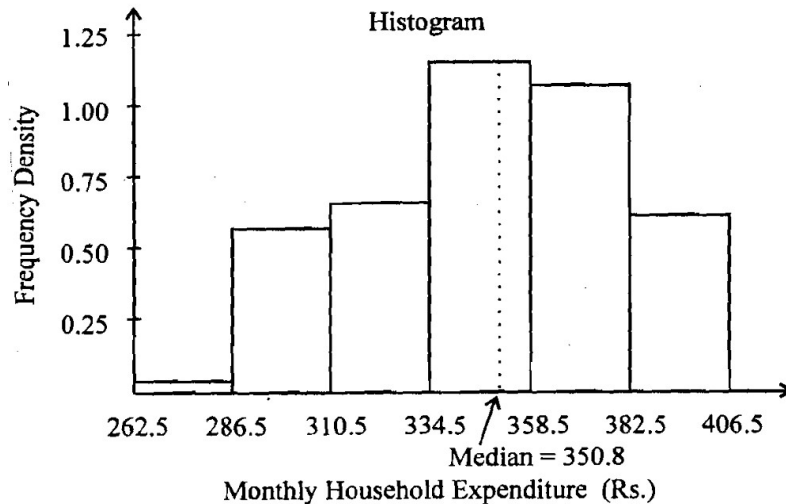
Median of a distribution locates a central point which divides a distribution into two equal halves, i.e., it is the middle most value among a set of observations. Let us start with examples in a discrete case. Consider a data set having 5 distinct observations: 2, 4, 9, 12, 19 (arranged in ascending order). Here 9 is the middle most value since an equal number of observations are to its left and to its right. Thus, 9 is the median of the above observations. Consider another data set having 6 distinct observations: 3, 8, 15, 25, 35, 43. Here any point between 15 and 25 has the property that equal number of observations are to its left and to its right. Any point in the interval 15 to 25 may be used as a median. Conventionally we take the middle point of such an interval to define median uniquely. Thus 20 is the median of 3, 8, 15, 25, 35, 43.

When a data set has non-distinct observations – a situation more common in practice – difficulties may arise. In such situations, it may not be always possible to locate the middle most value or the central point that divides the distribution into two equal halves. For example, in the case of the data set having 5 observations 2, 9, 9, 12, 19 the value 9 is repeated twice. Thus, a firm definition of median is needed to overcome such difficulties.

*A median of a distribution is a point or a central value such that **at least** 50% of the observations are less than or equal to it and **at least** 50% of the observations are greater than or equal to it.* With this definition of median and the convention of taking the middle point of a class in which each point is a median, median of a distribution can always be specified uniquely. Thus, median of observations 2, 9, 9, 12, 19 is 9 because 3 of the 5 observations (60%) are less than or equal to 9 and 4 of the 5 observations (80%) are greater than or equal to 9.

Let us find out the median household size from the frequency distribution in Table 11.1. We notice that 77 (out of 100) households have family size of less than or equal to 4 and 56 households have family size of more than or equal to 4. Thus median family size in this case is 4.

Median for a grouped frequency distribution of a continuous variable is easier to understand if we look at the associated histogram with height of a rectangle equal to the frequency density,  $\frac{f}{h}$ , of the class. In such a histogram, the area of a rectangle gives the frequency of the corresponding class. The median, in this case, is a point in one of the classes such that the areas to its left and to its right are 50% each. First step is to locate the class, up to the right boundary of which the total area is at least 50% (*called the median class*). Then the median is computed by adding, to the lower boundary value of this class, the length of a part of this class interval in proportion to the frequency needed to achieve 50%. A convenient method of finding out the median class is to compute the cumulative frequency (discussed in Unit 10, Sub-Section 10.2.4) and identifying the class interval in which the  $\frac{N}{2}$ th observation lies.



**Fig. 11.1**

Area up to the class boundary 334.5 is 131 and up to 358.5 is 59. Hence the median lies in the class 334.5 – 358.5. We now want to find a point in this classes so that the area from 334.5 to the point is  $(50 - 31) = 19$ , where are up to 334.5 is 31. Since the rectangle over the interval 334.5 – 358.5 has an area of 28, and is of length 24, to get an area of 19 we need  $\frac{19}{28}$ th part of 24. This works out to be

$\frac{19}{28} \times 24 = 16.3$ . Thus the median is  $334.5 + 16.3 = 350.8$ . Note also that the area in the class 350.8 to 358.5 is  $28 - 19 = 9$  and to the right of 350.8 is  $9 + 41 = 50$ , as it should be.

Based on the above procedure, we can write a formula for the computation of median.

$$M_d = l_m + \frac{\frac{N}{2} - C}{f_m} \times h, \text{ where}$$

$l_m$  is the lower limit of the median class, i.e., the class in which median lies,

$N$  is the total frequency,

$C$  is the cumulative frequency of classes preceding the median class (not that  $C = 31$  in the above example).

$f_m$  is the frequency of median class, and

$h$  is the width of median class.

### 11.2.3 Mode

As has been pointed out earlier, often observations tend to cluster around a central value. A simple measure of this phenomenon is called mode.

Mode or modal value of a discrete variable is defined as that value of the variable for which frequency is the maximum. Mode, however, is not the majority, i.e., it does not imply that most (50% or more) of the observations have the modal value.

From Table 11.1 we find that the mode or modal value of household size is 4 as this value occurs with largest frequency of 33 among 100 households.

There are, however, data sets when mode cannot be defined uniquely, i.e., the distribution has multiple mode. Raw data with 7 hypothetical observations with values 4, 3, 4, 1, 2, 5, 3 have two modes, 3 and 4. Distributions having two modes are called *bimodal distributions*, though the frequently encountered distributions have only one mode or are *unimodal*.

For observations on the continuous variable, like monthly household expenditure on food, no two observations are likely to have same value and so mode is not a meaningful measure of such raw data. However, central tendency comes out clearly when these raw data are grouped into various class intervals. For grouped data *modal class* is defined as the class having largest frequency. Since large class intervals are likely to include large number of observations and smaller class intervals are likely to have few observations, definition of modal class is meaningful only when class intervals have equal length.

For discrete data it is easier to find out the mode. But in the case of continuous data computation of the mode is done by the following formula:

$$M_0 = l_m + \frac{\Delta_1}{\Delta_1 + \Delta_2} \times h, \text{ where}$$

$l_m$  is the lower limit of the modal class, i.e., the class in which mode lies,

$\Delta_1 (= f_m - f_{m-1})$  is the difference of the frequencies of the modal class and its preceding class.

$\Delta_2 (= f_m - f_{m+1})$  is the difference of the frequencies of the modal class and its following class, and

$h$  is the width of the modal class.

Let us look back to Table 11.2. Here modal class is 334.5 – 358.5 as it has the highest frequency, 28.

Thus,  $l_m = 334.5$ ,  $\Delta_1 = 28 - 16 = 12$ ,  $\Delta_2 = 28 - 26 = 2$  and  $h = 24$ .

$$\text{Hence } M_0 = 334.5 + \frac{12}{12 + 2} \times 24 = 355.07$$

Mode is a useful measure of central tendency when a frequency distribution has a strong peak and it is particularly useless when a frequency distribution is almost flat.

### Check Your Progress 1

### Measures of Central Tendency

- 1) The frequency distribution of a family size for 250 families in a ward of an industrial town is given below:

Find the mean, median and mode.

| Family Size | Frequency | Family Size  | Frequency  |
|-------------|-----------|--------------|------------|
| 1           | 4         | 6            | 41         |
| 2           | 22        | 7            | 36         |
| 3           | 25        | 8            | 15         |
| 4           | 45        | 9            | 7          |
| 5           | 52        | 10           | 3          |
|             |           | <b>Total</b> | <b>250</b> |

.....

.....

.....

.....

- 2) Compute the mean, median and mode for the following frequency distribution.

| I.Q.    | Frequency | I.Q.         | Frequency  |
|---------|-----------|--------------|------------|
| 160-169 | 2         | 90-99        | 65         |
| 150-159 | 3         | 80-89        | 17         |
| 140-149 | 7         | 70-79        | 5          |
| 130-139 | 19        | 60-69        | 3          |
| 120-129 | 37        | 50-59        | 2          |
| 110-119 | 79        | 40-49        | 1          |
| 100-109 | 69        | <b>Total</b> | <b>309</b> |

.....

.....

.....

.....

.....

## 11.3 OTHER MEASURES OF CENTRAL TENDENCY

Besides the arithmetic mean, median and mode there are other averages which are relatively unimportant but may be appropriate in particular situations. These are Geometric Mean and Harmonic Mean.

Often we see that all the observations do not have equal importance. In such cases we need to give differential importance to different items. Here we use weighted means – arithmetic, geometric or harmonic – instead of simple means. This we will discuss in Section 11.3.2.

### 11.3.1 Geometric Mean and Harmonic Mean

Often we have to deal with that are time dependent, i.e., time series data which are unlike one-time data of Tables 11.1 and 11.2. For time dependent data, it is often of interest to find the pattern of change over time. Consider the following two data sets.

Set I: 1000 1100 1200 1300 1400 1500 1600

Set II: 1100 1210 1331 1464 1611 1772 1949

The first set looks like the basic salary (in Rs.) of an employee for 7 years with annual increment of Rs. 100 per year.

The second set looks more like his gross salary (in Rs.). Annual increase in the two sets is given below.

Set I: 100 100 100 100 100 100

Set II: 110 121 133 147 161 177

Arithmetic mean of the annual increase is 100 for Set I and 141.5 for Set II. On the basis of these average annual increases, if we find out the figures for the two sets, starting from the initial values, we would get the following:

Set I: 1000 1100 1200 1300 1400 1500 1600

Set II: 1100 1241.5 1383 1524.5 1666 1807.5 1949

We find that arithmetic mean has worked well for Set I. However, it has not worked well for Set II. It is because the progression of original numbers in the two sets is different. In set I, increment has been a fixed quantum whereas in Set II, figures have increased at a fixed rate. Fixed quantum of increase is called *arithmetic progression* and arithmetic mean is appropriate to describe the increase. Fixed rate of increase is called *geometric progression* and geometric mean is most appropriate to describe the increase.

For  $n$  numbers  $X_1, X_2, \dots, X_n$  the geometric mean (GM) is defined as the  $n$ th root of the product of these  $n$  numbers, i.e.,

$$GM = (X_1, X_2, \dots, X_n)^{\frac{1}{n}} = \left[ \prod_{i=1}^n X_i \right]^{\frac{1}{n}}$$

Clearly, GM is not defined unless all the  $n$  numbers are positive. If any number is negative or zero, we cannot calculate GM. By taking logarithm of GM, we have

$$\log GM = \left( \frac{1}{n} \right) (\log X_1 + \log X_2 + \dots + \log X_n) = \frac{1}{n} \sum_{i=1}^n \log X_i$$

which shows that now GM can be computed by using a log-table. Anti-logarithm of the arithmetic mean of  $\log X$  values is GM. For the second data set, gross salary increased at the rate of 11% every year. In practice, however, increase/decrease will not be at a fixed rate over the years; and it is meaningful to talk about average rate because fixed rate situation is rare.

In general, GM is more appropriate average for percentage (or proportionate) rates of change than arithmetic mean as in the case of rise in various price indices, cost of living indices, etc.

Finally, we discuss about another measure of location called the 'harmonic mean' (HM). This measure of central tendency comes naturally in many situations as in the following illustration. A stockist stocks Rs. 5000 worth of an item at the beginning of every month. Unit rate (in Rs.) of the item for five successive months had been 10.75, 11.80, 14.00, 11.45, and 12.00. The stockist wants to find average rate per unit of the item he has stocked for five months. Computation is presented below:

| Month        | Amount Spent (Rs.) | Unit Rate (Rs.) |
|--------------|--------------------|-----------------|
| 1            | 5000               | 10.75           |
| 2            | 5000               | 11.80           |
| 3            | 5000               | 14.00           |
| 4            | 5000               | 11.45           |
| 5            | 5000               | 12.00           |
| <b>Total</b> | <b>25000</b>       |                 |

$$\text{Average price (in Rs. of his entire stock)} = \frac{\text{Total Money Spent}}{\text{Total Quantity Purchased}}$$

$$= \frac{5 \times 5000}{\frac{5000}{10.75} + \frac{5000}{11.80} + \frac{5000}{14.00} + \frac{5000}{11.45} + \frac{5000}{12.00}}$$

$$= \frac{5}{\frac{1}{10.75} + \frac{1}{11.80} + \frac{1}{14.00} + \frac{1}{11.45} + \frac{1}{12.00}}$$

$$= \frac{1}{\frac{1}{5} \left( \frac{1}{10.75} + \frac{1}{11.80} + \frac{1}{14.00} + \frac{1}{11.45} + \frac{1}{12.00} \right)} = 11.91$$

The last expression is ‘the reciprocal of the arithmetic mean of the reciprocals’ and is called harmonic mean (HM). For a set of  $n$  values  $X_1, X_2, \dots, X_n$ , the HM is defined as

$$HM = \frac{n}{\frac{1}{X_1} + \frac{1}{X_2} + \dots + \frac{1}{X_n}} = \frac{n}{\sum_{i=1}^n \frac{1}{X_i}}$$

You should note that HM is not defined when any observation is zero.

If the stockiest, instead of stocking Rs. 5000 worth of items, stocks 3000 items at the beginning of every month at the given prices, the appropriate average would be arithmetic mean. To verify this, we can write

$$\begin{aligned} \text{Average Price} &= \frac{\text{Total Money Spent}}{\text{Total Quantity Purchased}} \\ &= \frac{3000 \times 10.75 + 3000 \times 11.80 + 3000 \times 14.00 + 3000 \times 11.45 + 3000 \times 12.00}{3000 \times 5} \\ &= \frac{10.75 + 11.80 + 14.00 + 11.45 + 12.00}{5} = \text{AM of the given prices.} \end{aligned}$$

### 11.3.2 Weighted Means

For many practical applications weighted means (arithmetic, geometric or harmonic) reflect phenomenon more clearly than unweighted or simple means that have been computed so far. For computation of, say, consumer price index, not all commodities are equally important. Increase in fuel cost may affect consumer price index more than an increase in agricultural prices. For stock market, stock of some key companies may be a trend setter. Weighted means are more appropriate in such situations. To find weighted mean, a weight  $w_i$  is attached to each  $X_i$  and the means are computed as if  $w_i$ 's are, symbolically, frequencies of the corresponding  $X_i$ 's. The computational formulae are as given below:

$$\text{Weighted AM} = \frac{\sum_{i=1}^n w_i X_i}{\sum_{i=1}^n w_i}$$

$$\text{Weighted GM} = \left( \prod_{i=1}^n X_i^{w_i} \right)^{\frac{1}{\sum w_i}} \quad \text{and}$$

$$\text{Weighted HM} = \frac{\sum_{i=1}^n w_i}{\sum_{i=1}^n \frac{w_i}{X_i}}$$

Weighted mean is equal to unweighted mean when each  $w_i$  is the same for each observation or equal to unity.

### 11.3.3 Pooled Mean

Often we come across situations when the means have been computed for different sources or samples. In such situations we become interested to find an overall mean if it is meaningful. This is done by computing what is called a *pooled mean*. The procedure of computing a pooled mean is given below.

Let  $m_1, m_2, \dots, m_r$  be  $r$  arithmetic (or geometric or harmonic) means, computed on the basis of  $n_1, n_2, \dots, n_r$  observations respectively. Then

$$\text{Pooled arithmetic mean} = \frac{1}{n} \sum_{i=1}^r m_i n_i, \text{ where } n = \sum_{i=1}^r n_i$$

$$\text{Pooled geometric mean} = \left( \prod_{i=1}^r m_i^{n_i} \right)^{\frac{1}{n}} \text{ and}$$

$$\text{Pooled harmonic mean} = \frac{n}{\sum_{i=1}^r \frac{n_i}{m_i}}$$

where  $n = n_1 + n_2 + \dots + n_r$

Note that the above expressions are similar to the expressions for weighted means.

### 11.3.4 Choosing a Measure of Central Tendency

It has already been discussed when a particular mean, AM or GM or HM, is more appropriate than the other two. However, when we have grouped data in which either of the end classes are open ended, i.e., of the type 'up to  $c_1$ ' and / or  $c_{k-1}$  and above', mid-points of such classes cannot be computed. Consequently, no mean can be computed. There is, however, no problem in computing median or mode in such cases. On the other hand, a pooled median or mode cannot be computed, like the case for mean, unless all the sets of data are made available in their entirety. These problems are related to computational difficulties and not to appropriateness of a measure.

Since graphical representation of data is more appealing, median or mode are more useful in such a situation because their crude values can be obtained easily without having to go through any computations. Also, median and mode are simple concepts for communication and comparison between graphs. It has, however, been observed that median is less stable than arithmetic mean in repeated sampling and we need to be careful when comparing graphs.

For data that has a distribution close in shape to what is called the normal distribution, with one peak and going down symmetrically on either side, we may use of mean, median or mode. It is because, for a normal distribution, these measures have the same value.

You should note that choosing an appropriate measure of central tendency is not an end to data analysis, and much still remains. For example, by saying that household average monthly expenditure on food is Rs. 348.66, it does not say whether a large number of households have very low monthly average expenditure on food or a few households have a very good menu. Next set of analysis aims at answering such questions.

### Check Your Progress 2

- 1) Given below are the prices is ratios for five commodities with the corresponding weights. Calculate the Weighted Arithmetic Mean and Geometric Mean.

| Commodity | Price Ratio | weight |
|-----------|-------------|--------|
| 1         | 2.20        | 30     |
| 2         | 1.85        | 25     |
| 3         | 1.80        | 22     |
| 4         | 2.05        | 13     |
| 5         | 1.75        | 10     |

.....

.....

.....

.....

- 2) The earnings of five nationalised banks, in crores of rupees, is given below.

217.40      330.50      682.55      1263.59      2249.63

Find the Geometric Mean of the earnings.

.....

.....

.....

.....

- 3) In a factory, a mechanic takes 15 days to fabricate a machine, the second mechanic takes 18 days, the third mechanic takes 30 days and the fourth mechanic takes 90 days. Find the average number of days taken the workers to fabricate the machine. Which average would you use, and why?

.....  
.....  
.....  
.....

- 4) The amount of interest paid on each of the three different sums of money yielding 10%, 12% and 15% simple interest per annum are equal. What is the average yield percent on the total sum invested?

.....  
.....  
.....  
.....  
.....

---

## 11.7 LET US SUM UP

---

In this unit you have learned to compute various measures of central tendency. These measures of central tendency can be divided into two broad categories, namely mathematical averages and positional averages. Positional averages are mode, median, quartile, percentiles, etc., while arithmetic mean, geometric mean and harmonic mean are mathematical averages. Geometric Mean is most suitable for averaging ratio and proportional rates of growth while Arithmetic mean or Harmonic mean can be used to find average rates like price, speed, etc. depending upon the nature of the given condition.

---

## 11.8 ANSWERS OR HINTS TO CHECK YOUR PROGRESS EXERCISES

---

### Check Your Progress 1

- 1) 5.1, 5, 5  
2) 108.48 ; 108.41 ; 111.42

### Check Your Progress 2

- 1) Rs. 1.96 ; Rs. 1.95  
2) Rs. 674.31 crore  
3) Arithmetic Mean, 38.25 days  
4) Harmonic Mean, 12%

---

## UNIT 12 MEASURES OF DISPERSION\*

---

### Structure

- 12.0 Objectives
- 12.1 Introduction
- 12.2 Range
- 12.3 Inter-quartile Range
- 12.4 Mean Deviation
- 12.5 Variance and Standard Deviation
- 12.6 Relationship between Dispersion and Standard Deviation
  - 12.6.1 Chebychev's Theorem
  - 12.6.2 Shape of Distribution
  - 12.6.3 Coefficient of Variation
  - 12.6.4 Concentration Ratio
- 12.7 Let Us Sum Up
- 12.8 Answers to Check Your Progress

---

### 12.0 OBJECTIVES

---

After going through this unit, you will be able to:

- explain the concept of dispersion;
- compute numerical quantities that measure the dispersion of a set of data;
- explain chebychev's inequality;
- compute the coefficient of variation; and
- find a measure for concentration of certain distribution of data.

---

### 12.1 INTRODUCTION

---

So far we have discussed various measures of central tendency, viz., arithmetic mean, median, mode geometric mean and harmonic mean. However, in many situations these measures do not represent the distribution of data. For example, look into the following three sets of data:

Set A: 2, 5, 17, 17, 44.

Set B: 17, 17, 17, 17, 17.

Set C: 13, 14, 17, 17, 24.

In all the sets the numerical value of the mean, median and mode are the same, that is, 17. Still all three sets are so different! While in Set B all the observations are equal, in Set A they are so dispersed. Definitely we need another measure which will account for such dispersion of data.

---

\* Adopted from IGNOU study material of EEC 13: *Elementary Statistical Methods and Survey Techniques*, Unit 5 written by R S Bharadwaj with modifications by K. Barik

The word dispersion is used to denote the degree of heterogeneity in the data. It is an important characteristic indicating the extent to which observations vary amongst themselves. The dispersion of a given set of observations will be zero when all of them are equal (as in Set B given above). The wider the discrepancy from one observation to another, the larger would be the dispersion. (Thus dispersion in Set A should be larger than that in Set C). A measure of dispersion is designed to state numerically the extent to which individual observations vary on the average. There are quite a few measures of dispersion. We discuss them below.

---

## 12.2 RANGE

---

Of all measures of dispersion, range is the smallest. It is defined as *the difference between the largest and the smallest observations*. Thus for the data given in Set A the range is  $44 - 2 = 42$ . Similarly, for Set B the range is  $17 - 17 = 0$  and for Set C it is 11. Now let us look into some grouped data. For Table 12.2 data (look back to the previous Unit), the range is Rs.  $406.5 - \text{Rs. } 262.5 = \text{Rs. } 144$ . Notice that, for grouped data, largest and the smallest observations are not identifiable. Hence we take *the difference between two extreme boundaries of the classes*.

It is intuitive that, because of central tendency, if one selects a small sample, observations are more likely to be around its mode than away from it. Less likely or extreme values will be included in the sample when its size is large. This, in other words, implies that range will increase with increase in sample size. Also, it is known that in repeated sampling with same sample size, range varies considerably making it a less suitable measure for comparisons. However, range is a measure which is easy to understand and can be computed quickly.

---

## 12.3 INTER-QUARTILE RANGE

---

Range as a measure of dispersion does not reflect a frequency distribution well, as it depends on the two extreme values. Even one very large or small observation, away from general pattern of other observations in the data set, makes the range very large. For example, in Set A, the range is found to be excessively large ( $44 - 2 = 42$ ) because of the presence of very large one observation, that is 44. To avoid such extreme observations, particularly when there is a strong central tendency, inter-quartile range is useful as a measure of dispersion. It is defined as

$$\text{Inter-quartile Range} = Q_3 - Q_1 = P_{75} - P_{25}.$$

Inter-quartile range is the range of the middle most 50% of the observations. If the observations are compact around median, i.e., a strong mode close to the median exists inter-quartile range will be smaller than half of the range. If the data are flat, having no central tendency, this measure will be large, and its value will be close to half of the range.

Let us look into the discrete data given in Table 12.1 of the previous Unit. Here,  $P_{75} = 4$  and  $P_{25} = 3$ . Hence, the inter-quartile range of household size is  $4 - 3 = 1$ . This shows that a strong central tendency exists in the distribution of household size range was observed to be 7 (since  $8 - 1 = 7$ ).

For Table 12.2 data,  $P_{25}$  of the average monthly expenditure on food was seen to be Rs. 325.50;  $P_{75}$  computed similarly works out to be Rs. 377.88 and inter-quartile range is Rs.  $377.88 - 325.50 = \text{Rs. } 52.38$ . Compared to Rs. 52.38, the range was observed to be Rs. 146.00 or 2.79 times larger. This shows not so strong central tendency for average monthly household expenditure on food.

## 12.4 MEAN DEVIATION

While the range depends on the two extreme observations, inter-quartile range depends on the two extreme observations among the middle most 50 percent of the observations. Thus, one talks only about the percentage of observations between minimum,  $P_{25}$  and maximum,  $P_{75}$ . Thus both range and inter-quartile range do not depend upon all the observations in the sample. Hence, while computing range or inter-quartile range we do not say anything about the distribution of observations within the group.

Among many possibilities to quantify spread or dispersion of observations, one possibility is to use the deviation of observations from some central value.

Since mean is the most commonly used measure of central tendency, it is often taken as the central value with reference to which the deviations are computed. These deviations are then suitably combined to get a measure of dispersion.

Mean deviation treats every single observation with equal weight, in the form of arithmetic mean of deviations based on each observation.

For observations  $X_1, X_2, \dots, X_n$ , if we take deviations as simple difference, then for the  $i^{\text{th}}$  observations the deviation is  $(X_i - \bar{X})$  where  $\bar{X}$  is the mean. Mean of these deviations is

$$\frac{1}{n} \sum_{i=1}^n (X_i - \bar{X}) = \frac{1}{n} \sum_{i=1}^n 1 = \bar{X} - \bar{X} = 0.$$

Since simple difference does not lead to any measure, absolute differences are used to define mean deviation.

$$\text{Mean Deviation} = \frac{1}{n} \sum_{i=1}^n |X_i - \bar{X}|, \text{ where}$$

the two vertical bars indicate that the sign of the difference with in two bars is to be taken as positive. For example,  $|2 - 4| = 4$  (and not  $-2$ ).

For frequency data, discrete or continuous type, the formula becomes

$$\text{Mean Deviation} = \frac{1}{N} \sum_{i=1}^n f_i |X_i - \bar{X}|,$$

where  $N = \sum_{i=1}^n f_i$  and  $X_i$ 's are distinct observations and  $f_i$  is the frequency of  $X_i$  in the discrete case and  $X_i$  is the mid-point of  $i^{\text{th}}$  class and  $f_i$  is its frequency for the continuous case. The need for such a measure is illustrated below.

Following summary values have been computed for two data sets.

|                        | Data Set I | Data Set II |
|------------------------|------------|-------------|
| Number of observations | 7          | 7           |
| $P_{25}$               | 12         | 12          |
| Median = $P_{75}$      | 17         | 17          |
| $P_{75}$               | 20         | 20          |
| Range                  | 10         | 10          |
| inter-quartile range   | 12         | 12          |
| Mean                   |            |             |

Thus, based on the above measures only, and not looking at the data sets I and II, it would appear that two persons separately may have worked out on the same data set. However, the two data sets may have been as given below.

**Data Set I:** 3      7      8      12      14      17      23

**Data Set II:** 2      7      11      12      13      17      22

One may construct much more different looking data sets having identical values for the above type of measures. This comparison indicates that more measures are needed and mean deviation is one such. This is not to imply that the above measures and mean deviation together completely describe a data set.

For data set I

Mean deviation =

$$\frac{1}{7} (|3-12| + |7-12| + |8-12| + |12-12| + |14-12| + |17-12| + |23-12|)$$

$$= \frac{9+5+4+0+2+5+11}{7} = \frac{36}{7} = 51.4$$

For data set II

Mean deviation =

$$\frac{1}{7} (|2-12| + |7-12| + |11-12| + |12-12| + |13-12| + |17-12| + |22-12|)$$

$$= \frac{10+5+1+0+1+5+10}{7} = \frac{32}{7} = 4.57.$$

Thus, observations in data set I are more dispersed from mean than that of data set II.

Let us now compute mean deviation of household size and household average monthly food expenditure.

For household size data of Table 12.1, mean =  $\bar{X} = 3.74$ . Mean deviation is now computed as

$$\begin{aligned}\text{Mean deviation} &= \frac{1}{N} \sum_{i=1}^N f_i |X_i - \bar{X}| \\ &= \frac{1}{100} (3|1-3.74| + 16|2-3.74| + \dots + 2|8-3.74|) = \frac{109.12}{100} = 1.0912. \text{ For}\end{aligned}$$

Table 12.2 distribution on average household expenditure on food, mean  $\bar{X} = \text{Rs.} 348.66$ .

The mean deviation =

$$= \frac{1}{100} (2|274.5 - 348.66| + \dots + 15|394.5 - 348.66|) = \frac{2510.88}{100} = 25.11$$

So far we have considered mean deviation from mean. The mean deviation from median or from mode can also be defined in a similar way.

## 12.5 VARIANCE AND STANDARD DEVIATION

The most frequently used measures of dispersion are variance and standard deviation. Variance is so commonly used that it is also called dispersion.

Variance is a measure which suitably combines individual deviations from the mean, treating each observation with equal weight as in mean deviation. For variance, however, measure of individual deviation is taken as the *squared difference from the mean*. Since it is more manageable to use the squared difference rather than absolute difference, particularly while doing format mathematics, use of variance has become more popular. Conventionally variance for a population is denoted by  $\sigma^2$  (pronounced *sigma-squared*) and variance for a sample is denoted by  $s^2$ . Variance is defined as the mean of the squared deviations of observations from their mean. Variance from raw data is computed by

$$\text{Variance} = \sigma^2 = \frac{1}{n} \sum_{i=1}^n (X_i - \bar{X})^2$$

For frequency data, discrete or continuous type, the formula becomes

$$\sigma^2 = \frac{1}{N} \sum_{i=1}^N f_i (X_i - \bar{X})^2, \text{ where } N = \sum_{i=1}^N f_i$$

In the same scale of measurement, for example, observations with a variance of 2 are less dispersed than observations with variance more than 2. To talk about a distribution in terms of a measure of central tendency and a measure dispersion, it is a practical need to use both measures in the same unit. Mean and mean deviation are in the same unit. Since each deviation has been squared for Based

on variance, an equally or more popular measure of dispersion in the same unit as that of observations is *standard deviation*, abbreviated as s.d. Standard deviation as the positive *square root of variance*, i.e., s.d. =  $\sigma$ . As it is the positive square root of variance, it cannot be negative.

Let us compute the s.d. for household size data of Table 12.1

$$\sigma^2 = \frac{1}{100} [3(1 - 3.74)^2 + 16(2 - 3.74)^2 + \dots + 2(8 - 3.74)^2] = \frac{199.24}{100} = 1.9924 \text{ and}$$

$$\sigma = 1.4115.$$

Similarly for Table 12.2 distribution of average monthly household expenditure on food, variance in Rs. Square is given by

$$\sigma^2 = \frac{1}{100} [2(274.50 - 348.66)^2 + \dots + 15(394.5 - 348.66)^2] = \frac{95725.437}{100} = 957.25,$$

and s.d. is

$$\sigma = \text{Rs. } 30.94.$$

For computational convenience, the formula for variance is written in alternative form as

$$\sigma^2 = \frac{1}{n} \sum_{i=1}^n f_i (X_i - \bar{X})^2 = \frac{1}{n} \sum_{i=1}^N X_i^2 - \bar{X}^2$$

or

$$\sigma^2 = \frac{1}{n} \sum_{i=1}^n f_i (X_i - \bar{X})^2 = \frac{1}{n} \sum_{i=1}^N f_i X_i^2 - \bar{X}^2$$

as the case may be. Thus, variance is viewed as

variance = Mean of Squares – Square of the Mean

Using the above formulae, you may compute the variance for the data given in Tables 12.1 and 12.2 and verify the earlier results.

The computation of variance may be greatly simplified by changing  $X_i$  to

$$u_i = \frac{X_i - A}{h}, \text{ as was done in the computation of mean.}$$

Note that, since

$$u_i - \bar{u} = \frac{X_i - A}{h} - \frac{\bar{X} - A}{h} = \frac{X_i - \bar{X}}{h}, \text{ we can write}$$

$$X_i - \bar{X} = h(u_i - \bar{u})$$

Hence,

$$\sigma_x^2 = \frac{1}{n} \sum_{i=1}^n (X_i - \bar{X})^2 = \frac{1}{n} \sum_{i=1}^n \{h(u_i - \bar{u})\}^2 = h^2 \sigma_u^2$$

where  $\sigma_x^2$  is the variance of  $X_i$  and  $\sigma_u^2$  is the variance of the  $u$  values.

Since the magnitude of  $u$  values is smaller, it is easier to compute variance of  $u$  values. Then the variance of  $X$  values can be easily computed by using the above formula.

Let us compute the variance by applying the above method for the data given in Table 12.2.

If we write  $u_i = \frac{X_i - 346.5}{24}$ , the  $u$  values are

$-3, -2, -1, 0, 1, 2$  and the respective frequencies are 1, 14, 16, 28, 26, 15.

The mean of  $u$  values =  $\frac{-3 \times 1 - 2 \times 14 - 1 \times 16 + 0 \times 28 + 1 \times 26 + 2 \times 15}{100} = 0.09$

The mean of squares of  $u$  values =

$$\frac{9 \times 1 + 4 \times 14 + 1 \times 16 + 0 \times 28 + 1 \times 26 + 4 \times 15}{100} = 1.67$$

Thus  $\sigma_u^2 = 1.67 - (0.09)^2 = 1.6619$  and

$$\sigma_x^2 = (24)^2 \cdot (1.6619) = 957.25.$$

Even though change from  $X$  to  $u$  is for computational ease, it brings up an important issue. Notice that  $\sigma_u^2 = 1.6619$  but  $\sigma_x^2 = 957.25$ , where  $u$  was obtained from  $X$  by a simple linear transformation, i.e., by change of origin and scale of  $X$  values. Typical such natural case are pounds and kilograms for weight, gallons and litres for liquid volume, etc. Since 1 kg. = 2.2046 lbs., s.d. of 5 kg. when measured in kilograms is same as 11.023 lbs. when measured in pounds; or since 1 litre = 0.22 gallon, s.d. of 5 litres when measured in litres is same as s.d. of 1.1 gallons when measured in gallons. Thus, whereas variance and standard deviation are supposed to measure spread of observations, not much can be made out of these measures due to their dependence on the unit of measurement.

In this context, the single most useful result about the spread of observations based on mean and standard deviation, irrespective of unit of measurement, is due to Chebychev.

### Check Your Progress 3

1) What is dispersion? What are the common measures of dispersion?

.....

.....

.....

.....

.....

.....

- 2) In a batch of 10 children the marks obtained by a dull boy are 25 marks below the average marks of other children. Show that the standard deviation of marks for all the children is at least 7.5. If this standard deviation is actually 12.0, find the standard deviation when the dull boy is left out.

.....

.....

.....

.....

.....

.....

- 3) The following data shows the daily profits (in Rs.) made by a shopkeeper on 15 successive days.

116, 87, 91, 81, 98, 102, 97, 100, 105, 101, 115, 98, 102, 98, 93

Determine the range, the mean deviation about mean and the standard deviation for the data.

.....

.....

.....

.....

.....

.....

- 4) Compute the arithmetic mean, standard deviation and the mean deviation of the following data.

| Scores   | 4 – 5 | 6 – 7 | 8 – 9 | 10 – 11 | 12 – 13 | 14 – 15 | Total |
|----------|-------|-------|-------|---------|---------|---------|-------|
| <b>f</b> | 4     | 10    | 20    | 15      | 8       | 3       | 60    |

.....

.....

.....

.....

.....

- 5) The mean and the s.d. of a sample of 100 observations were calculated as 40 and 5.1 respectively by a student who by mistake took one observation as 50 instead of 40. Calculate the correct s.d.

.....

.....

.....

.....

## 12.6 RELATIONSHIP BETWEEN DISPERSION AND STANDARD DEVIATION

You have earlier learnt that when all the values in a set of data are located near their mean, then exhibit a small amount of dispersion or variation and those set of data in which some values are located far from their mean have a large amount of dispersion. A useful rule that illustrates the relationship between dispersion and standard deviation is given by Chebychev's theorem.

### 12.6.1 Chebychev's Theorem

For any set of observation and positive constant  $k (>1)$ , the proportion of observations lying within  $k$  standard deviations of the mean is certain to be at least  $1 - \frac{1}{k^2}$ .

Note that the theorem is not useful for any positive  $k$  less than or equal to 1, since  $1 - \frac{1}{k^2}$  is at the most equal to zero. For other values of  $k$ , the minimum proportion can be computed easily. For example, proportion of observations within 1.5 s.d. of the mean is certain to be at least  $1 - \frac{1}{1.5^2} = 0.556$  or 55.6%. The following

figure indicates spread of data based on Chebychev's theorem. For the household size data of Table 12.1,  $\bar{X} = 3.74$  and  $s = 1.4115$ . If we take  $k=2$ , we can say that at least  $\left[ \left( 1 - \frac{1}{2^2} \right) \times 100 \right] = 75\%$  of the households are certain to have their size between  $3.74 \pm 2 \times 1.4115$ , i.e., between 0.917 and 6.563.

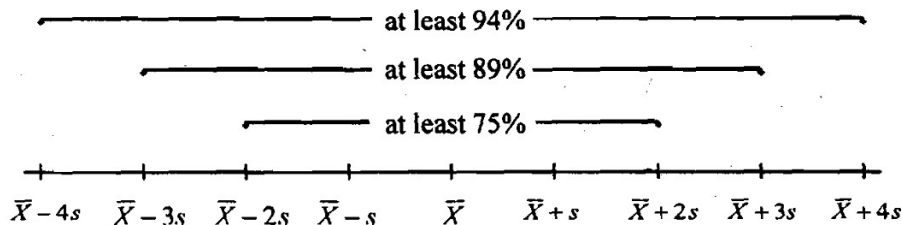


Fig. 12.2

For the Table 12.2 distribution of average monthly household expenditure on food  $\bar{X} = \text{Rs.}348.66$  and  $s = \text{Rs.} 30.94$ , at least 55.6% (for  $k = 1.5$ ) of households are certain to have monthly average food expenditure between Rs. 302.25 and Rs. 395.07. You can find the relevance of this theorem when we study normal distribution later in Unit 11.

### 12.6.2 Shape of Distribution

For methodological studies in many situations, a distribution is adequately described by measures of central tendency and dispersion. Yet other measures are also in use to describe distribution in practical situations, particularly for economic variables such as income, consumption, economic assets, etc., which are non-negative. Two such measures are *coefficient of variation* and *concentration ratio*. These measures will be viewed here essentially as measures of inequality in the distribution of economic variables.

### 12.6.3 Coefficient of Variation

Let us propose to economic status of households in two villages. The summary figures of monthly calories intake of households are given below for the two villages.

|                                     | Villages |      |
|-------------------------------------|----------|------|
|                                     | A        | B    |
| Number of Households (n)            | 817      | 561  |
| Mean calorie intake ( $\bar{X}$ )   | 2417     | 2235 |
| s.d. of calorie intake ( $\sigma$ ) | 418      | 232  |

The problem is to identify the village that has more inequality as far as calorie intake is concerned. Village A has higher mean calorie intake but has larger s.d. and larger number of households compared to village B. Village A may actually have more number of poorer households in than in village B. Therefore, in village A, inequality between households may be more than that in village B. One index which measures the quantum of such disparity is called the coefficient of variation, abbreviated as c.v. It is defined as percentage standard deviation per unit of mean, i.e.,

$$\text{c.v} = \frac{\sigma}{\bar{X}} \times 100$$

Since  $\sigma$  and  $\bar{X}$  have the same unit of measurement, c.v. is unit free and is not affected by the choice of unit of measurement.

For village A,  $\text{c.v.} = \frac{418}{2417} \times 100 = 17.29$  and for village B,

$$\text{c.v.} = \frac{232}{2235} \times 100 = 10.38.$$

Since the coefficient of variation in village A is greater than the coefficient of variation in village B, the inequalities are given in village A compared to village B.

To compare the extent of inequalities, we compute

$$\frac{17.29 - 10.38}{10.38} \times 100 = 66.57 \text{ which implies that compared to village B, } 66.57\%$$

more inequality exists in village A.

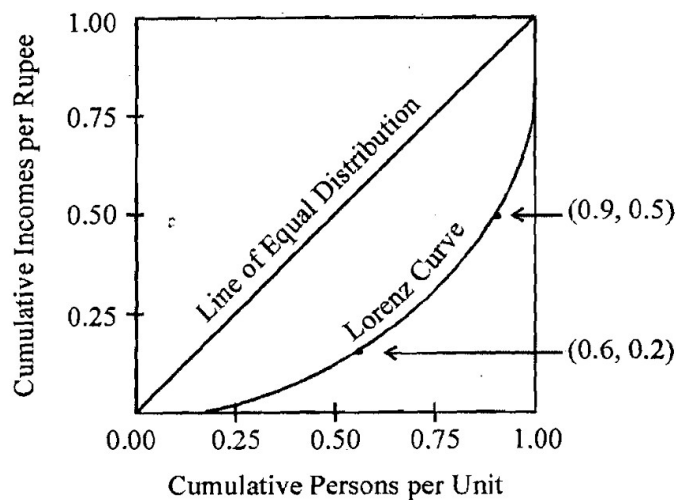
#### 12.6.4 Concentration Ratio

Above was a comparison of inequality between two villages, without quantifying the level of inequality within each village. If a distribution has a long right tail, it shows that a few have a large share. In other words, a majority of population has a very small share. Let us consider the distribution of income of a hypothetical economy.

Suppose there are three classes of people in the economy – the upper class, the middle class and the lower class. Let 10%, 30% and 60% be the share of population in these three classes respectively. Suppose the lower class receives only 20% of the national income, the middle class 30% and the upper class the rest, i.e., the remaining 50%. We can now present the data in a percentage cumulative frequency distribution form. Thus, the lowest 60% of the population receives only 20% of the income, the lowest 90% receive 50% (= 20 + 30) of the income and obviously, 100% of the population receive 100% of the income. If we take a graph paper where the percent cumulative total income is plotted on the vertical axis and we plot the point (0, 0), (60, 20), (90, 50) and (100, 100), then the curve joining these points is what we call the **curve of concentration** or **Lorenz curve**. The straight line joining the points (0,0) and (100, 100) give the line of *equal distribution* or the *equitable line*. The equitable line is that one which shows that the proportion of share is exactly the same as the proportion of population who are supposed to share. The area between the line of equal distribution and the curve of concentration, called the *area of concentration* is an indicator of the degree of concentration; the larger the area the greater is the concentration.

#### Coefficient of Inequality

Let us take coordinates of the above points in per unit terms instead of percentage terms. Thus, the coordinates of the points, in the above example, can be written as (0,0), (0.60, 0.20), (0.90, 0.50) and (1.00, 1.00). The coefficient of inequality of income distribution is then defined as the ratio of the area of concentration to total area of the triangle. Since the area of the triangle is 0.5 (since  $\frac{1}{2} \times 1 \times 1 = 0.5$ ), the coefficient of inequality is equal to twice the area of concentration when coordinates of various points are taken per unit rather than in percentage.



**Fig. 12.3: Lorenz Curve**

**Check Your Progress 4**

- 1) The following figures give the crude birth rate per 1000 people in Switzerland from 1968 to 1980.

Crude birth rate ( $X$ ): 17.1, 16.5, 15.8, 15.2, 14.3, 13.6, 12.9, 12.3, 11.7, 11.5, 11.3, 11.3, 11.6.

Calculate the Variance, Standard Deviation and Coefficient of Variation.

.....

.....

.....

.....

.....

- 2) The following table gives the distribution of age of lady teachers of a school as revealed by records.

| Age Group (years) | No. of lady teachers |
|-------------------|----------------------|
| 15 – 19           | 3                    |
| 20 – 24           | 13                   |
| 25 – 29           | 21                   |
| 30 – 34           | 15                   |
| 35 – 39           | 5                    |
| 40 – 44           | 4                    |
| 45 – 49           | 2                    |

Calculate coefficient of variation, and (ii) number of teachers between the age 26 and 33 years.

.....

.....

.....

.....

.....

---

## 12.7 LET US SUM UP

---

In this unit you learned about the measures of dispersion. The most important measures of dispersion you learned about in the unit are the variance, standard deviation and the concentration ratio. You have also learned to compute variance, standard deviation and coefficient of variation using both ungrouped and grouped data. The coefficient of variation is used to compare the dispersion of two distributions having either different means (even when their variables are measured in same units) or different units of measurement of their variables.

---

## 12.8 ANSWERS OR HINTS TO CHECK YOUR PROGRESS EXERCISES

---

### Check Your Progress 3

- 1) Do it yourself.
- 2) 9.9
- 3) 35, 6.46, 8.85
- 4) 9.23, 2.49, 2.03
- 5) 5.0

### Check Your Progress 4

- 1) 3.085, 2.021, 15.004%
- 2) 23.47%, 25 (rounded figure)