



OSUN STATE UNIVERSITY

Osogbo, Nigeria

Department of Statistics, Faculty of Basic and Applied Sciences

Prepared by:

Timothy A. OGUNLEYE,

E-mail: timothy.ogunleye@uniosun.edu.ng

ESTIMATION AND TEST OF HYPOTHESES

CHAPTER 1: STATISTICAL CONCEPTS AND DEFINITIONS

1.0 Introduction

Statistical estimation is a fundamental aspect of inferential statistics that involves making informed guesses or estimates about population parameters based on sample data. Two common approaches to statistical estimation are point estimation and interval estimation. Under this chapter, we shall discuss the following sub-headings:

- 1.1 Hypothesis:** Hypothesis is a statistical statement which can be true or false and whose validity can easily be tested. The word hypothesis consists of two words: 'hypo + thesis = hypothesis' where 'hypo' means tentative or subject to the verification and 'thesis' means statement about solution of a problem. Therefore, the term hypothesis is a tentative statement about the solution of a particular problem. Simply, hypothesis is an intelligence guess or assumption which you make about the likely solution to the problem.
- 1.2 Statistical Hypothesis:** This can be defined as an assertion or a conjecture about the distribution of one or more random variable. It can be thought of as a statement about the possible solution to the problem you are trying to solve. In attempting to reach decisions, it is important to make assumption/guess about the population involved. Such assumption, which may or may not be true, is called Statistical Hypothesis.
- 1.3 Purpose of Hypothesis:** Scientists have found that it is good practice to state a hope for result ahead of time as hypothesis. Hypothesis guides us in searching for the solution to the problem at hand. Hypothesis as a generalization of facts needs be subjected to standard tests to find out if its claim is in agreement with empirical evidence and of course with past experience. Hypothesis performs the following functions:

- It helps the investigator to cave the procedures and methods to be used in solving his problem;
- It helps in deciding the proper statistical treatment (method) and in examining the results of the research;
- It is a way of reducing a research problem to statement, which can tested. Most problems are mere questions and are not testable unless reduced to hypothesis, which are subjected to confirmation on refutation;
- It acts as a frame work for the overall conclusion of the research;
- It is very necessary in studies where cause and effect relationships are to be discovered;
- It should be noted, however, that hypothesis is not an end in itself but rather a means for the investigator to understand the problem at hand clearly and to use the existing data at his disposal very well.

1.4 Characteristics of a Good Hypothesis: Care is needed in the formulation of hypothesis as earlier mentioned. Some of the following points must be borne in mind when formulating hypothesis:

- Base the hypothesis on the data available to you;
- Let the statements be in simple terms;
- The statement must be put in a form that can easily be tested, so that it can be readily accepted or rejected;
- Its statement must easily permit the development of a research design that is capable of providing a data necessary for testing its validity;
- The hypothesis must embrace all objectives set at the initial stage;
- It must adequately take care of the bias or prejudice against any preconceived outcome;
- Hypothesis must be consistent with known facts in order to avoid wasting time and energy.

1.5 Types of Hypothesis: There are two types of hypothesis, namely:

- **Null Hypothesis:** The null hypothesis is stated first and it is usually represented by H_0 . It is a statement of independence, insignificance; it is stated to show a neutral position. A null hypothesis states that there is no significant difference between two or more parameters. It is concerned with a judgment as to whether apparent differences are real differences or whether they merely result from sampling error. It refers to a general statement or default position that there is no relationship between two measured phenomena, or no difference among groups. A null hypothesis is, thus, a hypothesis that says there is no statistical

significance between the two variables in the hypothesis; it is the hypothesis that the researcher is trying to disprove.

- **Alternative Hypothesis:** After writing the null hypothesis, the alternative hypothesis is stated next. It is usually represented by H_1 . It is usually stated in form of prejudice or bias. It describes the result we hope to expect. It is stated in a form that will indicate the direction of the expected result. It should be stated clearly in anticipation of the type of analysis to be required. Simply, it directly negates the null hypothesis. When we get to practical work, we shall discuss more on hypothesis and the way it is formulated in relation with the operations in NAPTIP.

1.6 Level of Significance: The rejection or acceptance of a null hypothesis is based on a particular level of significance as a criterion. On many occasions, 5% level of significance is often used as a standard for rejection. Thus, we can define 'level of significance' as the level of probability at which it is agreed that the null hypothesis will be rejected, and conventionally it is represented by alpha (α) and is usually set at 0.05 for many statistical computations. Simply, it is the probability of committing a Type I error in statistical analysis.

1.7 P-Value: This is a procedure introduced to deal with those situations where it is difficult or impossible to derive a significance test because of the presence of nuisance parameters. Hence, p-value is the probability of observed data (or data showing a more extreme departure from the null hypothesis) when the null hypothesis is true. Hence, p-value is defined as the level of marginal significance within a statistical hypothesis test, representing the probability of the occurrence of a given event. The p-value is used as an alternative to rejection points to provide the smallest level of significance at which the null hypothesis would be rejected. When statistical hypothesis is performed, a p-value helps to determine the significance of our result.

1.8 Decision Rule: This is a procedure that the researcher uses to decide whether to accept or reject the null hypothesis. It is a rule that uses information from a sample to make a choice between two hypotheses. In statistics, we reject the null hypothesis if **p-value** is less than (or sometimes equal to) **level of significance**. On the other hand, we reject the null hypothesis if calculated test statistic is greater than (or sometimes equal to) the corresponding tabulated value. Thus, a Decision Rule is a formal rule that states, based on the data obtained, when to reject the null hypothesis H_0 .

1.9 Type I and Type II Errors: If we reject hypothesis when it is supposed to be accepted, we say a Type I error has been made. Type I error is represented by α and it is also called **Producer's**

Risk. On the other hand, if we accept hypothesis when it is supposed to be rejected, we say a Type II error has been made. It is represented by β and it is also called **Consumer's Risk**.

1.10 One-tailed Test: In hypothesis testing, when we are interested **only** in extreme values to one side of the mean, the test is called one-tailed test. For instance:

$$H_0 : \mu_1 = \mu_2 \text{ Versus } H_1 : \mu_1 < \mu_2$$

or

$$H_0 : \mu_1 = \mu_2 \text{ Versus } H_1 : \mu_1 > \mu_2$$

1.11 Two-tailed Test: In hypothesis testing, when we are interested in extreme values on **both** sides of the mean, the test is called two-tailed test. For instance:

$$H_0 : \mu_1 = \mu_2 \text{ Versus } H_1 : \mu_1 \neq \mu_2$$

1.12 An Estimator: An estimator is a random variable which is used to estimate or obtain population parameters, for example, μ, σ^2 , etc. There are two types of estimator: Point and Interval Estimates.

1.13 An Estimate: An estimate is simply a particular value of the estimator based on a specified sample, e.g. $\hat{\mu} = 2.89$, $\hat{\sigma} = 0.654$, etc.

1.14 Point Estimate: A Point Estimate is a single observed numerical value used as an estimate of the unknown population parameter. Examples is $\hat{\mu}$, which is a point estimate of $\bar{x}$.

1.15 Interval Estimate: This is an estimate of population parameter given by two numerical values between which the parameter may be considered to lie. For instance, computing 95% confidence interval for the proportion (p) of a non-defective battery may give the following result: $0.83 \leq p \leq 1.08$. This result is a very good example of interval estimate.

CHAPTER 2: TEST OF HYPOTHESIS

2.0 Introduction

In Statistics, one of the most, if not the most important goal is hypothesis testing. Hypothesis testing is presented through two independent statements: the null hypothesis and the alternative hypothesis. It is the statement of the alternative hypothesis that categorizes the tail of the test, and this should be articulated first. Before going deeper into hypothesis testing, let's discuss some associated terminologies of Testing of Hypothesis in Statistics.

Consider all possible samples of size n which can be drawn from a given population (either with or without replacement). For each sample, we can compute a statistic such as mean, variance and standard deviation that will vary from sample to sample. In this manner, we obtain a distribution of the statistic – this is called *Sampling Distribution*.

If, for instance, the particular statistic used is the sample mean, then the distribution is called *Sampling Distribution of Mean*. Similarly, we could have *Sampling Distribution of Standard Deviation, Variance, Proportion*, etc. For each sampling distribution, we can compute the mean, standard deviation, variance, as the case may be.

2.1 Sampling Distribution of Means

2.1.1 Mean

The mean, or arithmetic average, plays a crucial role in everyday activities. In the scope of this course, we are interested in two classes of mean value – a population mean and a sample mean.

- $\mu = [\text{miu}]$ is the population mean. That is, it is the average of all the individual elements in an entire population – for example, the population mean number of university students in the world. Obviously, it is difficult, if not impossible to know the true population mean, so it is estimated by the sample mean, $\bar{x}$, instead. The computation of population mean is:

$$\mu = \frac{X_1 + X_2 + \dots + X_N}{N} = \frac{\sum_{i=1}^N X_i}{N}$$

- $\bar{x} = [x \text{ bar, or overline } x]$ is the sample mean, or the arithmetic average of a sample that represents the entire population. The calculation of the sample mean is:

$$\bar{x} = \frac{x_1 + x_2 + \dots + x_n}{n} = \frac{\sum_{i=1}^n x_i}{n}$$

Note: Examples on computation of mean in connection with summation signs should be done during the class discussion.

Class Practical I:

1. Compute the sample mean of 12, 9, 11, 8, 5, 17, 20, 24, 13, 10, 7.

The results are shown below:

To compute the sample mean manually, you sum up all the values in the sample and then divide by the number of observations (sample size).
Let's compute the sample mean for the given dataset: 12, 9, 11, 8, 5, 17, 20, 24, 13, 10, 7.

$$\text{Sample Mean} = \frac{12+9+11+8+5+17+20+24+13+10+7}{11}$$

$$\text{Sample Mean} = \frac{136}{11}$$

$$\text{Sample Mean} \approx 12.36$$

Therefore, the sample mean for the given dataset is approximately 12.36.

2. If the sample mean of 5, 6, 3, 2, 3, y^2 , 9, 4 is 6, what are the two possible values of y and which of the two values is appropriate for the gap?

The results are shown below:

Let's solve for the missing value of y in the dataset such that the sample mean is 6.

The sample mean ($\bar{x}$) is calculated by summing up all the values in the sample and then dividing by the number of observations (sample size).

$$\bar{x} = \frac{5+6+3+2+3+y^2+9+4}{8}$$

Given that $\bar{x} = 6$, we can set up the equation and solve for y :

$$6 = \frac{5+6+3+2+3+y^2+9+4}{8}$$

Multiply both sides by 8 to clear the denominator:

$$48 = 5 + 6 + 3 + 2 + 3 + y^2 + 9 + 4$$

Combine like terms:

$$48 = 32 + y^2$$

Subtract 32 from both sides:

$$16 = y^2$$

Now, take the square root of both sides:

$$y = \pm 4$$

So, the two possible values for y are $y = 4$ and $y = -4$.

The appropriate value for y is 4.

2.1.2 Variance and Standard Deviation

The variance and standard deviation are statistics representing the difference of the x_i values from the $\bar{x}$, mean. For example, the mean of a data set is computed as follows:

$$\bar{x} = (118 + 111 + 126 + 137 + 148) / 5 = 128$$

The individual data point variability around the mean of 128 is displayed as follows:

$(x_i - \bar{x})$
$118 - 128 = -10$
$111 - 128 = -17$
$126 - 128 = -2$
$137 - 128 = +9$
$148 - 128 = +20$
$\sum (x_i - \bar{x}) = 0$

We can observe that the sum of the variability point $\sum (x_i - \bar{x})$ is zero; this makes sense because the $\bar{x}$ is the central weighted value, the summation will never provide a value other than zero. So, we need to square each variability value $(x_i - \bar{x})^2$ and then add them to find

their average. This average, $\frac{\sum (x_i - \bar{x})^2}{n-1}$, is referred to as the Variance of the data. This is computed as follows:

$$\text{Variance} = \frac{(-10)^2 + (-17)^2 + (-2)^2 + (9)^2 + (20)^2}{5-1} = \frac{874}{4} = 218.5$$

This is the sample variance. In short-cut, we can equally use the formular given below:

$$s^2 = \frac{1}{n-1} \left[\sum x^2 - \frac{(\sum x)^2}{n} \right]$$

Whereas the standard deviation shall be computed as the square root of the sample variance as follows:

$$s = \sqrt{\frac{1}{n-1} \left[\sum x^2 - \frac{(\sum x)^2}{n} \right]}$$

2.1.3 Standard Error of the Mean

To this point, we have discussed the variability of the individual x_i data around the mean, $\bar{x}$. Now, we will discuss the variability of the sample mean, $\bar{x}$ itself, as it relates to the theoretical population mean, μ . The standard deviation of the mean, *not of the data point*, is also termed standard error of the mean. Error, in this case, is variability. The computation for the standard error of the mean is as follows:

Case I: Suppose that all possible samples of size n are drawn **without replacement** from a finite population of size $N > n$, that is, N is known, the formular is

$$\sigma_{\bar{x}} = \frac{\sigma}{\sqrt{n}} \sqrt{\frac{N-n}{N-1}} \dots\dots\dots \text{(Population is known)}$$

Case II: If the sampling is **with replacement**, the formular to be used is

$$s_{\bar{x}} = \frac{s}{\sqrt{n}} \dots\dots\dots \text{(Population is unknown)}$$

Examples:

Example: Standard Error of the Mean for a Finite Population

Scenario:

A town has a total population of 100 residents ($N = 100$). You want to estimate the average income of the residents by taking a sample of 10 individuals.

Data:

Income of 10 residents (in thousands): 45, 55, 60, 50, 48, 52, 58, 53, 55, 49

Calculation:

1. Calculate the sample mean ($\bar{x}$):

$$\bar{x} = \frac{\sum_{i=1}^{10} x_i}{10}$$

$$\bar{x} = \frac{525}{10} = 52.5$$

1. Calculate the sample standard deviation (s):

$$s = \sqrt{\frac{\sum_{i=1}^{10} (x_i - \bar{x})^2}{9}}$$

$$s \approx 4.07$$

1. Calculate the standard error of the mean (SEM):

$$SEM = \sqrt{\frac{100-10}{100-1}} \times \frac{4.07}{\sqrt{10}}$$

$$SEM \approx 1.36$$

In this example, we used the finite population correction factor ($\sqrt{\frac{N-n}{N-1}}$) to adjust the standard error of the mean for a finite population. This correction factor accounts for the fact that sampling is done without replacement, and it adjusts the standard error to reflect the finite nature of the population.

Remember that this approach is appropriate when sampling without replacement from a known finite population. If the population is sufficiently large compared to the sample size, the correction factor becomes negligible, and the formula approaches the standard error formula for an infinite population.

Example 2:

Scenario:

In a drug trial, researchers want to assess the effect of a new medication on reducing cholesterol levels. They measure the cholesterol levels of a random sample of 20 patients before and after the treatment.

Data:

Cholesterol levels before treatment: 180, 190, 175, 200, 185, 195, 180, 170, 210, 200, 185, 190, 175, 185, 195, 200, 180, 190, 175, 195

Calculation:

1. Calculate the sample mean ($\bar{x}$):

$$\bar{x} = \frac{\sum_{i=1}^{20} x_i}{20}$$

$$\bar{x} = \frac{3775}{20} = 188.75$$

1. Calculate the sample standard deviation (s):

$$s = \sqrt{\frac{\sum_{i=1}^{20} (x_i - \bar{x})^2}{19}}$$

$$s \approx 10.53$$

1. Calculate the standard error of the mean (SEM):

$$SEM = \frac{s}{\sqrt{n}}$$

$$SEM = \frac{10.53}{\sqrt{20}} \approx 2.35$$

Example 3:

Scenario:

A standardized test is administered to a sample of 30 students to estimate the average score for all students who take the test.

Data:

Test scores: 85, 90, 88, 92, 87, 89, 91, 86, 88, 90, 87, 92, 85, 88, 89, 90, 91, 88, 87, 86, 90, 92, 89, 88, 86, 87, 91, 88, 90, 92

Calculation:

1. Calculate the sample mean ($\bar{x}$):

$$\bar{x} = \frac{\sum_{i=1}^{30} x_i}{30}$$

$$\bar{x} = \frac{2675}{30} = 89.17$$

$$s = \sqrt{\frac{\sum_{i=1}^{30} (x_i - \bar{x})^2}{29}}$$

$$s \approx 1.97$$

1. Calculate the standard error of the mean (SEM):

$$SEM = \frac{s}{\sqrt{n}}$$

$$SEM = \frac{1.97}{\sqrt{30}} \approx 0.36$$

These examples demonstrate the step-by-step calculation of the standard error of the mean using sample data from drug trial results and educational testing. The SEM provides an estimate of how much the sample mean is likely to vary from the true population mean.

2.2 Hypotheses for a Single Parameter

Under this section, we first introduce the concept of p -value. After that, we study hypothesis testing concerning a single parameter.

2.2.1 The p -Value: Corresponding to an observed value of a test statistic, the p -value (or attained significance level) is the lowest level of significance at which the null hypothesis would have been rejected.

Practical I:

The management of a local health club claims that its members lose on the average 15 pounds or less within the first 3 months after joining the club. To check this claim, a consumer agency took a random sample of 45 members of this health club and found that they lose an average of 13.8 pounds within the first 3 months of membership, with a sample standard deviation of 4.2 pounds.

- (i) Find the p -value for this test;
- (ii) Based on the p -value in (i), would you reject the null hypothesis at $\alpha = 0.01$?

Solution:

- (i) Let μ be the true mean weight loss in pounds within the first 3 months of membership in the club. Then, we have to test the hypothesis

$$H_0 : \mu = 15 \text{ versus } H_1 : \mu < 15$$

Here, $n = 45$, $\bar{x} = 13.8$, and $s = 4.2$. Hence, the test statistic is

$$z = \frac{\bar{x} - \mu}{s/\sqrt{n}} = \frac{13.8 - 15}{4.2/\sqrt{45}} = -1.9166$$

$$p\text{-value} = P(z < -1.9166) = \phi(-1.9166) = 0.0276$$

NB: We can use an α as small as 0.0276 and still reject the null hypothesis.

- (ii) No. Because the p -value = 0.0276 is greater than $\alpha = 0.01$, so we can't reject H_0 .

2.2.2 Test of Mean – One Sample Situation

Case I: Test of Mean when $n < 30$ (Small Sample) and Standard Deviation is Unknown

The test statistic to be applied is:

$$T_{cal} = \frac{\bar{x} - \mu}{s/\sqrt{n}}$$

The corresponding Critical Region is:

$$T_{tab} = T_{(n-1);(1-\alpha)} \dots\dots\dots \text{(For one-tail test)}$$

$$T_{tab} = T_{(n-1);(1-\alpha/2)} \dots\dots\dots \text{(For two-tail test)}$$

Case II: Test of Mean when $n \geq 30$ (Large Sample) and Standard Deviation is Known

The test statistic to be applied is:

$$Z_{cal} = \frac{\bar{x} - \mu}{\sigma / \sqrt{n}}$$

The corresponding Critical Region is:

$$Z_{tab} = Z_{(1-\alpha)} \dots\dots\dots \text{(For one-tail test)}$$

$$Z_{tab} = Z_{(1-\alpha/2)} \dots\dots\dots \text{(For two-tail test)}$$

2.2.3 Test of One-Sample Proportion

Case I: When Standard Deviation is Unknown

The test statistic to be applied is:

$$T_{cal} = \frac{\hat{p} - p_0}{\sqrt{\frac{p_0(1-p_0)}{n}}}$$

The corresponding Critical Region is:

$$T_{tab} = T_{(n-1);(1-\alpha)} \dots\dots\dots \text{(For one-tail test)}$$

$$T_{tab} = T_{(n-1);(1-\alpha/2)} \dots\dots\dots \text{(For two-tail test)}$$

Case II: When Standard Deviation is Known

The test statistic to be applied is:

$$Z_{cal} = \frac{\hat{p} - p_0}{\sqrt{\frac{p_0(1-p_0)}{n}}}$$

The corresponding Critical Region is:

$$Z_{tab} = Z_{(1-\alpha)} \text{ and } Z_{tab} = Z_{(1-\alpha/2)} \text{ (For one-tail and two-tail tests)}$$

Example:

Scenario:

A company claims that 80% of its customers are satisfied with its product. To test this claim, a random sample of 150 customers is taken, and it is found that 110 of them are satisfied.

Hypotheses:

- Null Hypothesis (H_0): The proportion of satisfied customers is equal to the claimed proportion.
 $H_0 : p = 0.80$
- Alternative Hypothesis (H_1): The proportion of satisfied customers is not equal to the claimed proportion.
 $H_1 : p \neq 0.80$

Test Statistic (Z-Test for Proportion):

The formula for the test statistic is given by:

$$Z = \frac{\hat{p} - p_0}{\sqrt{\frac{p_0(1-p_0)}{n}}}$$

Where:

- $\hat{p}$ is the sample proportion.
- p_0 is the claimed population proportion.
- n is the sample size.

Calculation:

$$\hat{p} = \frac{\text{Number of satisfied customers}}{\text{Total sample size}} = \frac{110}{150}$$

$$Z = \frac{\frac{110}{150} - 0.80}{\sqrt{\frac{0.80(1-0.80)}{150}}}$$

$$Z \approx \frac{0.7333 - 0.80}{\sqrt{\frac{0.80(0.20)}{150}}}$$

$$Z \approx \frac{-0.0667}{\sqrt{\frac{0.16}{150}}}$$

$$Z \approx \frac{-0.0667}{\sqrt{0.0010667}}$$

$$Z \approx \frac{-0.0667}{0.03266}$$

$$Z \approx -2.04$$

2.3 Testing of Hypotheses for Two Samples

2.3.1 When Standard Deviation is Known

The test statistic to be applied is:

$$Z_{cal} = \frac{\bar{x}_1 - \bar{x}_2}{\sqrt{\frac{\sigma_1^2}{n_1} + \frac{\sigma_2^2}{n_2}}}$$

The corresponding Critical Region is:

$$Z_{tab} = Z_{(1-\alpha)} \dots\dots\dots \text{(For one-tail test)}$$

$$Z_{tab} = Z_{(1-\alpha/2)} \dots\dots\dots \text{(For two-tail test)}$$

2.3.2 When Standard Deviation is Unknown

The test statistic to be applied is:

$$T_{cal} = \frac{\bar{x}_1 - \bar{x}_2}{S_p \sqrt{\frac{1}{n_1} + \frac{1}{n_2}}}$$

where the formular to compute S_p varies from one situation to the other as follows:

$$S_p = \sqrt{\frac{(n_1 - 1)s_1^2 + (n_2 - 1)s_2^2}{n_1 + n_2 - 2}}, \text{ for } n_1 + n_2 < 30$$

And:

$$S_p = \sqrt{\frac{n_1 s_1^2 + n_2 s_2^2}{n_1 + n_2 - 2}}, \text{ for } n_1 + n_2 \geq 30$$

The corresponding Critical Region is:

$$T_{tab} = T_{(n_1 + n_2 - 2); (1 - \alpha)} \dots \dots \dots \text{(For one-tail test)}$$

$$T_{tab} = T_{(n_1 + n_2 - 2); (1 - \alpha/2)} \dots \dots \dots \text{(For two-tail test)}$$

2.4.3 Test of Two-Sample Proportion

Case I: When Standard Deviation is Unknown

The test statistic to be applied is:

$$T_{cal} = \frac{\hat{p}_1 - \hat{p}_2}{\sqrt{\bar{p}(1 - \bar{p}) \left[\frac{1}{n_1} + \frac{1}{n_2} \right]}} \text{ for } \bar{p} = \frac{x_1 + x_2}{n_1 + n_2}$$

The corresponding Critical Region is:

$$T_{tab} = T_{(n_1 + n_2 - 2); (1 - \alpha)} \dots \dots \dots \text{(For one-tail test)}$$

$$T_{tab} = T_{(n_1 + n_2 - 2); (1 - \alpha/2)} \dots \dots \dots \text{(For two-tail test)}$$

Case II: When Standard Deviation is Known

The test statistic to be applied is:

$$Z_{cal} = \frac{\hat{p}_1 - \hat{p}_2}{\sqrt{\bar{p}(1 - \bar{p}) \left[\frac{1}{n_1} + \frac{1}{n_2} \right]}} \text{ for } \bar{p} = \frac{x_1 + x_2}{n_1 + n_2}$$

The corresponding Critical Region is:

$$Z_{tab} = Z_{(1-\alpha)} \dots\dots\dots \text{(For one-tail test)}$$

$$Z_{tab} = Z_{(1-\alpha/2)} \dots\dots\dots \text{(For two-tail test)}$$

CHAPTER 3: TESTING HYPOTHESIS ABOUT THE VARIANCE AND ITS COMPUTATION OF CONFIDENCE INTERVAL USING CHI-SQUARED STATISTICAL TOOL

3.0 Testing and Estimating a Single Variance or Standard Deviation

3.1 Testing σ^2 :

Many problems arise that require us to make decisions about variability. In this section, we will study two kinds of problems:

- We will test hypothesis about the variance (or standard deviation) of the population; and
- We will find confidence intervals for the variance (or standard deviation) of a population.

It is customary to talk about variance instead of standard deviation because our techniques employ the sample variance rather than the standard deviation. Of course, the standard deviation is just the square root of the variance, and any discussion about variance is easily converted to a similar discussion about standard deviation.

If we have a normal population with variance σ^2 and a random sample of n measurements is taken from this population with sample variance s^2 , then:

$$\chi_{cal}^2 = \frac{(n-1)s^2}{\sigma^2}$$

has a chi-square distribution with $n - 1$ degree of freedom.

Example I:

A large discount hardware store in Ketu Market, Lagos has been using the independent-lines procedure to check out customers. After long observation, the manager knows that the standard deviation of waiting times is 7 minutes. The manager decided to introduce the single-line procedure on a trial basis to see if a reduction in waiting time variability would occur. A random sample of 25 customers was monitored, and their waiting times for check-out were determined. The sample standard deviation was $s = 4$ minutes. Using 5% level of significance, test the claim that the variability of the waiting times has been reduced.

Sol:

As a null hypothesis, we assume that the variance in waiting times is the same as that of the former independent-lines procedure. The alternative hypothesis is that the variance for the single-line procedure is less than that for the independent-lines. If we let σ be the standard deviation of waiting times for the single-line procedure, then σ^2 is the variance, and we have:

The Hypothesis:

$$H_0: \sigma^2 = 49 \quad (\text{Since } \sigma = 7)$$

$$H_1: \sigma^2 < 49$$

The given parameters are:

$$\sigma = 7, \text{ so } \sigma^2 = 49,$$

$$s = 4, \text{ so } s^2 = 16 \text{ with } n = 25$$

Using the statistic:

$$\chi_{cal}^2 = \frac{(n-1)s^2}{\sigma^2} = \frac{(25-1)16}{49} = 7.837$$

The Critical Value:

$$\chi_{tab}^2 = \chi_{(n-1); \alpha}^2$$

$$\chi_{tab}^2 = \chi_{24; (1-0.05)}^2 = 36.42$$

Decision Rule:

Reject the null hypothesis if the calculated value > the tabulated value.

Decision:

Since the $\chi_{cal}^2 (= 7.837) < \chi_{tab}^2 (= 36.42)$, we do not reject the null hypothesis.

Conclusion:

It is reasonable to conclude that the variance of the single-line procedure is 49.

3.2 Estimating Confidence Interval for σ^2 :

Let a random sample of size n be taken from a normal population with population standard deviation σ , and let c be a chosen confidence level ($0 < c < 1$). Then, the Confidence Interval is:

$$\frac{(n-1)s^2}{\chi_U^2} < \sigma^2 < \frac{(n-1)s^2}{\chi_L^2}$$

Example II:

Mr Thommy is a truck farmer in Ghana who makes his living on a large single-vegetable crop of green beans. Because of modern machinery being used, the entire crop must be harvested at once. This means that Mr Thommy wants a small standard deviation between maturing times of individual plants. A seed company is trying to develop a new variety of green beans with a small standard deviation of maturing times. To test their new variety, Mr Thommy planted 30 of the new seeds and carefully observed the number of days required for each plant to arrive at its peak of maturity. The maturing times for these plants had a sample standard deviation of $s = 3.4$ days. Find a 95% confidence interval for the population standard deviation of maturing times of this variety of green bean.

Sol:

A random sample of $n = 30$ plants has a standard deviation of $s = 3.4$ days for maturity. We are requested to obtain a 95% confidence interval for the population standard deviation σ .

To find the Confidence Interval, we use the following values:

$c = 0.95$	(confidence level)
$n = 30$	(sample size)
$s = 3.4$	(sample standard deviation)

To find χ_U^2 , use the formular below to find U :

$$U = \left(\frac{1-c}{2} \right) = \left(\frac{1-0.95}{2} \right) = 0.025$$

Thus, the *Upper Chi-square* tabulated is:

$$\chi_U^2 = \chi_{(n-1);(1-0.025)}^2$$

$$\chi_U^2 = \chi_{29;0.975}^2 = \mathbf{45.72}$$

Also, to find χ_L^2 , use the formular below to find U :

$$L = \left(\frac{1+c}{2} \right) = \left(\frac{1+0.95}{2} \right) = 0.975$$

Thus, the *Lower Chi-square* tabulated is:

$$\chi_L^2 = \chi_{(n-1);(1-0.975)}^2$$

$$\chi_L^2 = \chi_{29;0.025}^2 = \mathbf{16.05}$$

Now, we apply the formular as follows:

$$\frac{(n-1)s^2}{\chi_U^2} < \sigma^2 < \frac{(n-1)s^2}{\chi_L^2}$$

$$\frac{(30-1)(3.4)^2}{45.72} < \sigma^2 < \frac{(30-1)(3.4)^2}{16.05}$$

$$7.332 < \sigma^2 < 20.887$$

To obtain the Confidence Interval for the Population Standard Deviation, take a square root of both sides as follows:

$$\sqrt{7.332} < \sigma < \sqrt{20.887}$$

Therefore:

$$2.708 < \sigma < 4.570$$

Class Practical II:

- (1) The mean weight of all yams of a farm is θ . A random sample of 125 yams from the farm yields a mean of 4.25kg. if the variance of the weight of all yams on the farm is known to be 1.44kg. Set up an appropriate hypothesis to test (at .05 level of significance) that:
 - (a) the mean is greater than 4.25kg;
 - (b) the mean is not equal to 4.25kg.
- (2) The breaking strengths of cables produced by a manufacturer have a mean of 150 pounds (lb) and a standard deviation of 36 pounds (lb). by a new technique in the manufacturing process, it is claimed that the breaking strengths can be decreased. Test this claim, when a sample of 22 cables were tested and it is found that the mean breaking strength is 124 (lb) at:
 - (a) .01 level of significance;
 - (b) .05 level of significance;
 - (c) .10 level of significance.

- (3) The Intelligence Quotient (IQ) of 20 students from a particular area of Epe, Lagos State has been measured by a psychometrician and the result showed a mean of 48 and standard deviation of 5.28, while IQ of 16 students from another part of Lagos State showed a standard deviation of 4.01 with a mean of 30. Is there a significant difference between the IQ of the two groups at:
- (a) .01 significant level;
 - (b) .05 significant level.
- (4) The mean height of 55 male students who showed above average participation in college athletics was 126.9cm with standard deviation of 5.3cm, while 46 male students who showed no interest in such participation had a mean of 111.2cm with a standard deviation of 3.99cm. Test the hypothesis that male students who participated in the college athletics are taller than other male students at both .01 and .05 levels of significance.
- (5) Two independent samples of 13 and 15 students respectively have the following heights:
Sample I: 9.34, 11.44, 13.05, 16.62, 5.77, 12.89, 19.02, 23.0, 16.02, 12.27, 22.01, 14.28, 19.08
Sample II: 11.02, 9.22, 20.26, 17.21, 12.03, 18.24, 15.11, 8.92, 20.2, 11.8, 14.2, 17.9, 14, 13, 18.1
Is there any significant difference between the mean heights of these two students? Verify this at .05 and .10 levels.
- (6) A polling firm samples 600 likely voters and asks them whether they favour a proposal involving school bonds. A total of 330 of these voters indicate that they favour the proposal. Can u conclude that more than 50% of all likely voters favour the proposal?
- (7) Do patients value interpersonal skills more than technical ability when choosing a primary care physician? The article 'Patients' Preferences for Technical Versus Interpersonal Quality when selecting a Primary Care Physician' (C. Fung, M. Elliot, *et al.*, *Health Services Research*, 2005: 957 - 977) reports the results of a study in which 304 people were asked to choose a physician based on two hypothetical descriptions. One physician was described as having high technical skills and average interpersonal skills, and the other was described as having average technical skills and high interpersonal skills. Sixty-two percent of the people choose the physician with high technical skills. Can we conclude that more than half of patients prefer a physician with high technical skills?

- (8) Suppose we have purchased a filling machine for candy bags that is supposed to fill each bag with 16 oz of candy. Assume that the weights of filled bags are approximately normally distributed. A random sample of 10 bags yields the following data (in oz):

15.87, 16.02, 15.78, 15.83, 15.69, 15.81, 16.04, 15.81, 15.92, 16.10

On the basis of these data, can we conclude that the mean fill weight is actually less than 16 oz?

- (9) The automobile manufacturer wishes to compare the lifetimes of two brands of tyre. She obtains samples of six tyres of each brand. On each of six cars, she mounts one tyre of each front wheel. The cars are driven until only 20% of the original tread remains. The distance, in miles, for each tyre are presented in the following table:

Car	Brand I	Brand II
1	36,925	34,318
2	45,300	42,280
3	36,240	35,500
4	32,100	31,950
5	37,210	38,015
6	41,099	40,128
7	31,712	30,781

Can

two

- (10) The Area

we conclude that there is a difference between the mean lifetimes of the brands of tyre?

article 'Modeling of Urban Stop-and-Go Traffic Noise' (P. Pamanikabud and C.

Tharasawatipipat, *Journal of Transportation Engineering*, 1999: 152 – 159) presents measurements of traffic noise from ten locations in Bangkok, Thailand. Measurements, presented in the following table, were made at each location, in both the acceleration and deceleration lanes.

Location	Acceleration	Deceleration
1	78.1	78.6
2	78.1	80.0
3	79.6	79.3
4	81.0	79.1
5	78.7	78.2
6	78.1	78.0
7	78.6	78.6
8	78.5	78.8

9	78.4	78.0
10	79.6	78.4

Can we conclude that there is a difference in the mean noise levels between acceleration and deceleration lanes?

2.5 Construction of Confidence Intervals for One-Sample Situation

2.5.1 Confidence Interval for Mean

Case I: When Standard Deviation is Unknown

$$(1 - \alpha)\% \text{C.I.} = \bar{x} \pm T_{tab} \cdot (s/\sqrt{n})$$

where $\bar{x}$ is the sample mean and s is the sample standard deviation.

Case II: When Standard Deviation is Known

$$(1 - \alpha)\% \text{C.I.} = \bar{x} \pm Z_{tab} \cdot (\sigma/\sqrt{n})$$

where μ is the population mean and σ is the population standard deviation.

2.5.2 Confidence Interval for Proportion

Case I: When Standard Deviation is Unknown

$$(1 - \alpha)\% \text{C.I.} = \hat{p} \pm T_{tab} \cdot \left(\sqrt{\frac{p_0(1-p_0)}{n}} \right)$$

Case II: When Standard Deviation is Known

$$(1 - \alpha)\% \text{C.I.} = \hat{p} \pm Z_{tab} \cdot \left(\sqrt{\frac{p_0(1-p_0)}{n}} \right)$$

2.6 Construction of Confidence Intervals for Two-Sample Situation

2.6.1 Confidence Interval for Mean

Case I: When Standard Deviation is Unknown

$$(1 - \alpha)\% \text{C.I.} = (\bar{x}_1 - \bar{x}_2) \pm T_{tab} \cdot \left[S_p \sqrt{\frac{1}{n_1} + \frac{1}{n_2}} \right]$$

where $(\bar{x}_1 - \bar{x}_2)$ is the difference between the sample means and S_p is the pooled standard deviation.

Case II: When Standard Deviation is Known

$$(1 - \alpha)\% \text{C.I.} = (\bar{x}_1 - \bar{x}_2) \pm Z_{tab} \cdot \left[\sqrt{\frac{\sigma_1^2}{n_1} + \frac{\sigma_2^2}{n_2}} \right]$$

where $(\bar{x}_1 - \bar{x}_2)$ is the difference between the sample means and σ_1^2 and σ_2^2 are the standard deviations of first and second populations.

2.5.2 Confidence Interval for Proportion

Case I: When Standard Deviation is Unknown

$$(1 - \alpha)\% \text{C.I.} = (\hat{p}_1 - \hat{p}_2) \pm T_{tab} \cdot \left[\sqrt{\bar{p}(1 - \bar{p}) \left[\frac{1}{n_1} + \frac{1}{n_2} \right]} \right]$$

$$\text{for } \bar{p} = \frac{x_1 + x_2}{n_1 + n_2}$$

Case II: When Standard Deviation is Known

$$(1 - \alpha)\% \text{C.I.} = (\hat{p}_1 - \hat{p}_2) \pm Z_{tab} \cdot \left[\sqrt{\bar{p}(1 - \bar{p}) \left[\frac{1}{n_1} + \frac{1}{n_2} \right]} \right]$$

2.7 The F-Distribution

This is used to test for the differences in **variances** between two populations or samples selected for examination. The test statistic for the F-distribution is:

$$F_{cal} = \frac{n_1 s_1^2 / (n_1 - 1) \sigma_1^2}{n_2 s_2^2 / (n_2 - 1) \sigma_2^2} \sim F_{(n_1 - 1), (n_2 - 1); (1 - \alpha)}$$

where:

s_1^2 = sample variance for the first sample selected;

s_2^2 = sample variance for the second sample selected;

σ_1^2 = population variance for the first population considered;

σ_2^2 = population variance for the second population considered.

But in a situation where the population variances are not available (given), we shall apply the following test statistic:

$$F_{cal} = \frac{s_1^2}{s_2^2} \sim F_{(n_1 - 1), (n_2 - 1); (1 - \alpha)}$$

Note that the larger variance shall be the numerator.

Class Practical III:

- (1) Suppose that the heights of 121 students of Augustine University, Ilara-Epe, Lagos State represents a random sample with mean 3.6m and variance 1.44m^2 . Obtain 95% confidence interval for the mean height of the students.
- (2) If the mean weight of a random sample of 10 packets of a certain commodity is 4.25kg with a standard deviation of 0.83kg, obtain a 99% confidence interval for the mean weight of the commodity.
- (3) A sample poll of 100 voters chosen at random from all voters in a given district indicated that 55% of them were in favour of a particular candidate. Find the:
(i) 90% (ii) 95% (iii) 99%
Confidence limits for the proportion of all the voters in favour of the candidate.
- (4) A sample of 150 brand 'A' light bulbs showed a mean lifetime of 1400 hrs (h) and a standard deviation of 120 hours (h). A sample of 200 brand 'B' light bulbs showed a mean lifetime of 1200 hours and a standard deviation of 80 hours. Construct a 95% confidence interval for the difference of the mean lifetimes of the populations of brands A and B.
- (5) Two samples of sizes 9 and 12 are drawn from two normally distributed populations having variances 16 and 25 respectively. If the sample variances are 20 and 8, determine whether the first sample has a significantly larger variance than the second sample at .05 level.
- (6) A study is conducted to determine whether there is less variability in the gold plating done by Process I than in the gold plating done by Process II. If the independent random samples with sizes $n_1 = 10$ and $n_2 = 12$ yield $s_1 = 0.032$ (mm) and $s_2 = 0.055$ (mm) respectively. Test the null hypothesis $\sigma_1 = \sigma_2$ against the alternative hypothesis $\sigma_1 < \sigma_2$.
- (7) A sample poll of 300 voters from the Western Region and 200 voters from the Eastern Region showed that 56% and 48% respectively were in favour of a given candidate.
 - (a) Is there any difference in the proportions of voters from the region, who favour the candidate?
 - (b) In Western Region, is the population proportion who favour the candidate 0.5?
- (8) A vote is to be taken among the residents of town and the surrounding community to determine whether a proposed chemical plant should be constructed. The construction site is within the towns' limits and for this reason many voters in the community feel that the proposal will pass because of the large proportion of town voters who favour the construction.

This is to determine if there is a significant difference in the proportion of town voters and country voters. If 120 of 200 town voters favour the proposal and 240 of 500 country residents favour it, would you agree that the proportion of town voters favouring the proposal is higher than the proportion of country voters?

- (9) An engineer claims that a new type of power supply for home computers lasts longer than the old type. Independent random samples of 75 of each of the two types are chosen, and the sample means and standard deviations of the lifetimes are computed: New: $\bar{x}_1 = 43\text{h}$, $s_1 = 2.5\text{h}$ and Old: $\bar{x}_2 = 42\text{h}$, $s_2 = 2.3$. Can we conclude that the mean lifetime of new power supplies is greater than that of the old power supplies?
- (10) Construct 95% confidence interval for question 9 above.

References

- Cyprian A. Oyeka (1990): *An Introduction to Applied Statistical Method in the Sciences*. Enugu: Nobern Avocation publishing coy.
- Osuntogun, E. O. (1997): *"Introduction to Social and Economic Statistics"* Unpublished paper.
- S. P. Gupta (2008). *Statistical Methods*, 36th Revised Edition. Sultan Chand & Sons, New Delhi.
- Descriptive Statistics Note (Unpublished).
- Brookes, B.C. and Dick, W.F.L. (1969). *An Introduction to Statistical Methods*, 2nd Edition, H.E.B. Publishers.