

Appropriate Statistical Analysis in Sustainable Procurement, Environmental and Social Standard Practices

G.O. Nwafor^{1*}, H. C. Iwu¹, D. C. Bartholomew¹, T.W Owolabi¹, C. H. Izunobi¹, E.N. Egbo²

¹CE-sPESS Federal University of Technology, Owerri /Social Standard/Department of Statistics

²Department of Statistics/Federal University of Technology, Owerri

³Department of Statistics, Chukwuemeka Odumegwu Ojukwu, University

DOI: <https://doi.org/10.36347/sjpms.2025.v12i03.004>

| Received: 11.02.2025 | Accepted: 15.03.2025 | Published: 21.03.2025

*Corresponding author: G.O. Nwafor

CE-sPESS Federal University of Technology, Owerri /Social Standard/Department of Statistics

Abstract

Original Research Article

This paper terms to address the many challenges faced by procurement managers, environmental analysts and social standard practitioners on the appropriate use of statistical tools and analysis on their respective filed of study. The selection of statistical techniques is based on the characteristics of the collected data and questionnaire design as well as the goal of the scientific inquiry. The current work is methodologically complete and helpful to a wide range of readers because all possible scenarios are almost exhaustive. The chosen statistical approach must appropriately represent the features of the data and the study's goals since statistical analysis techniques and research hypotheses interact. This study highlights how crucial it is to choose statistical analysis techniques using a logical and scientific process. Offering thorough instructions, an appreciate information on the null and alternative hypothesis. With regard to sustainable procurement, environmental, and social standard practices, this study seeks to empower researchers and improve the overall quality and dependability of scientific studies, as well as the selection of suitable statistical methods and analysis.

Keywords: Statistical Analysis, Questionnaire Design, Procurement, Environment, Social Standards.

Copyright © 2025 The Author(s): This is an open-access article distributed under the terms of the Creative Commons Attribution **4.0 International License (CC BY-NC 4.0)** which permits unrestricted use, distribution, and reproduction in any medium for non-commercial use provided the original author and source are credited.

1. INTRODUCTION

Efficient procurement management is crucial for cost control, improving organizational efficiency, and supporting overall strategic goals. It directly impacts an organization's ability to deliver products or services on time, within budget, and with the expected level of quality. Procurement, environmental, and social standards increasingly intersect, especially as organizations aim to promote sustainable and ethical sourcing practices. Statistical practices can be used to assess, monitor, and improve compliance with these standards, creating more data-driven, transparent, and accountable procurement systems.

Environmental literature underscores the importance of data analysis for evaluating ecological impacts and monitoring compliance with environmental regulations. Statistics allow researchers to identify patterns in environmental data, such as air or water quality, and model future scenarios to predict potential environmental outcomes. For example, *multivariate*

analysis helps in examining the interaction of multiple environmental factors to understand their collective impact on ecosystems.

Choosing statistical analysis techniques for research is a crucial and complex process that calls for a methodical, scientific approach. Since the chosen method might have a substantial impact on the dependability and calibre of the research findings, it is crucial to match it with the particulars of the research design and hypothesis.

Types of variables

A variable is a "changeable number," as the name implies. Because of their intrinsic variability, variables can be measured or observed and can take on a variety of values depending on the subject of the investigation. Anthropometric measurements like height, demographic characteristics like age, and health indicators like body mass index (BMI) are a few examples of variables. These varied entities enable

researchers to measure and record key features for their research. In general, variables can be divided into two groups: quantitative and categorical (qualitative). Nominal and ordinal variables are two further subcategories of categorical variables, which capture traits that are difficult to quantify. Nominal variables have no intrinsic order and function as descriptors for names, labels, or categories. Sex is a prime example, where the classifications of male and female is obtained.

For researchers and practitioners to get significant conclusions from statistical studies, it is essential to comprehend the subtleties of null and alternative hypotheses as well as the results of their testing. This section clarifies these ideas and provides information on the nuances of testing hypotheses using a variety of statistical techniques. Test of normalcy a statistical technique for determining if the gathered data fulfills normality or has a normal distribution is the normality test [1].

Compliance with environmental standards requires continuous data collection and appropriate analysis. Statistical tools help agencies and researchers monitor whether environmental conditions meet legal requirements and sustainability goals. Methods like trend analysis and hypothesis testing are used in literature to highlight changes over time and assess their significance.

2. RELATED LITERATURE

One of the most frequent mistakes in statistical analysis is misinterpreting p-values. Although it just shows the likelihood of finding data at least as extreme under the null hypothesis, many people mistakenly believe that a p-value of less than 0.05 is "proof" of a hypothesis [2]. There is a possibility that the Shapiro-Wilks test, one of the most effective normality tests, might be conducted with three samples. The data does not necessarily follow a normal distribution, though, even if the P value is higher than the significance level of 0.05. All hypothesis tests have Type I and Type II errors, which are identified by the power and significance levels [3].

In finding out if the averages of three or more independent groups are the same, a statistical hypothesis testing technique called an ANOVA is utilized [4]. Inaccurate results may result from choosing a statistical test without taking the distribution, sample size, or data type into account. For instance, it is improper to use a t-test for non-normally distributed data without making the necessary adjustments [5]. If confounding variables are not taken into account, the measured effect may appear stronger or less than it actually is. To take these factors into consideration, regression analysis or other multivariate methods are frequently required [5].

Poor generalization on new data results from overfitting, which happens when a model is overly complicated and catches noise instead of the underlying

data pattern. To avoid overfitting, cross-validation is essential [6]. The likelihood of a Type I error rises while performing several hypothesis tests. To maintain the overall significance level, adjustments like the Bonferroni correction are frequently required [7]. Analysis of linear regression a statistical analysis technique called linear regression is used to estimate a regression model that establishes a linear relationship between a quantitative response variable and one or more explanatory variables.

The following are included in linear regression analyses: First, each explanatory variable's regression coefficient is estimated; next, the regression model is tested to see if all estimated regression coefficients are equal to zero; finally, the final regression model is constructed; and finally, the coefficient of determination (R^2) is computed to demonstrate how well the regression model explains the data used to construct the model [8].

Analyses of linear regression include the following: 1) Each explanatory variable's regression coefficient is estimated; 2) the regression model is tested to see if all estimated regression coefficients are equal to zero; 3) the regression model is tested to see if the regression coefficients of each explanatory variable are equal to zero; 4) the final regression model is constructed; and 5) the coefficient of determination (R^2) is computed to demonstrate how well the regression model explains the data used to construct the model [8].

However, compared to the parametric test, the nonparametric test has a little less power. Furthermore, the difference between the values of groups can only be detected; the size of these differences cannot be compared. Consequently, if at all possible, statistical analysis should be conducted using the parametric test, and the parametric test should be used to check the data's normality first [9].

3. METHODOLOGY

With an emphasis on the statistical hypothesis testing phase and variable classification, this paper examines a thorough guideline for methodically selecting suitable statistical analysis techniques. This study attempts to give researchers a strong basis for well-informed methodological decision making by offering a thorough analysis of these factors. By investigating the null and alternative hypotheses that are specific to particular statistical methods of analysis, this study goes beyond theoretical considerations and into the real world. A well-designed flowchart for choosing the statistical analysis method is suggested, and the dynamic link between these hypotheses and statistical methods is completely examined.

In sampling surveys, the sample size is determined before the actual selection of sample units from the population units. The determination of sample size is done in consideration of the population parameter to be

estimated from the sample, the required margin of error permissible in estimating the parameter and the desired precision of the survey. The level of error tolerance or the margin of error is usually denoted by e . The size of e is the difference between the value of a population parameter and the value of a sample statistic that estimates the parameter. The sample size n is determined by evaluating a $100(1 - \alpha)\%$ confidence that e is less than or equal to two times the square root of the variance of the sample statistic that estimates the population parameter.

Ordinal variables, on the other hand, add a sense of order by defining values according to a system of ranking between several categories. Using a Likert-type scale to rate satisfaction (e.g., "very dissatisfied," "dissatisfied," "neither dissatisfied nor satisfied," "satisfied," and "very satisfied") is the classic example with balancing principle of positive and negative equal statement in the design of questionnaire with respect to all the variables of interest. Quantitative variables, on the other hand, indicate traits that are accurately measurable and represent numerical values. Continuous and discrete variables are further segmented from a quantitative variable. A nuanced representation of the qualities is provided by continuous variables, which can take on an endless number of real values within a specified interval. Examples are height, which covers a continuous range of measures, and age, which captures a spectrum of genuine values. Discrete variables, on the other hand, are limited to a finite range of real values. For instance, the number of children can only have values that are zero or positive integers (such as 0, 1, 2, etc.). As their names suggest, variables are essentially a concept of change that can be quantified. In research and analysis, a sophisticated understanding of data is made possible by this complex tapestry of diversity.

Statistical analysis is ideally conducted in a rigorous manner with considerations of not only the study design but also the underlying model assumptions governing sample data. Descriptive statistics help to characterize the basic features of data by performing statistical calculations and a summary that yields insight into the nature of its distribution. More sophisticated models underlie modern statistical tools; overall, records can be classified into a few categories through contingency tables modeling or logistic regression, and fit more complex models on data arising from cross-sectional or longitudinal studies. In cross-sectional studies, a well-established standard multiple linear regression model can accommodate both continuous and discrete predictors; the variables that are to be tested through another procedure following model estimation can be compiled in a hypothesis-testing data set. Longitudinal data are needed to model the correlations in measures of the dependent variables, and panel analysis of variance or multilevel mixed-effects models can be applied under this structured model assumption.

The reference can be selected by the researcher based on the circumstances of the study. On the other hand, naming a reference is not required if the explanatory variable is quantitative. In this instance, if the other explanatory variables are held constant, a one-unit increase in the explanatory variable raises the odds of the response variable by the odds ratio (for example, if the odds ratio of a quantitative explanatory variable is 1.2, a one-unit increase in the explanatory variable raises the odds ratio by 20%). Depending on the number of explanatory variables, logistic regression analyses can be classified as simple or multiple logistic regression analyses, just like linear regression analyses. A basic logistic regression analysis is performed when there is only one explanatory variable, while a multiple logistic regression analysis is performed when there are two or more explanatory factors. Binary and multinomial logistic regression analysis can be further separated based on how many levels the response variable has. A binary logistic regression analysis is performed when the response variable has two levels, while a multinomial logistic regression analysis is performed when the answer variable has three or more levels.

The null hypothesis asserts that "the odds ratio is one," while the alternative hypothesis asserts that "the odds ratio is not one" in order to assess the significance of the odds ratio for each explanatory variable. "The odds ratio cannot be said to statistically be one under the significance level" if the null hypothesis is rejected. The link between the explanatory and response variables can be described as follows since the computed odds ratio is not 1. Depending on whether the explanatory variable in the logistic regression analysis is a quantitative or categorical variable, the odds ratio is interpreted differently. In the event that the explanatory variable is categorical, the odds ratios to the other levels are computed, with one level of the explanatory variable serving as the reference. For instance, the odds ratio of females to males can be computed if the explanatory variable is sex, with male (reference) and female levels. Since male is the reference, the chances ratio for male is 1. On the other hand, if a woman were to serve as the reference, her chances ratio would be 1.

As a result, four divisions of logistic regression analysis can be applied based on the number of levels for the response variable and the number of explanatory variables: 1) A simple binary logistic regression analysis is used for classification when there are one explanatory variable and two levels for the response variable; 2) A simple multinomial logistic regression analysis is used for classification when there are one explanatory variable and three or more levels for the response variable; 3) the classification is a multiple binary logistic regression analysis if there are two or more explanatory variables and two levels for the response variable; and 4) a multiple multinomial logistic regression analysis if there are two or more explanatory variables and three or more

levels for the response variable. The term "multinomial" is substituted with "ordinal" if the answer variable is an ordinal variable.

4. RESULTS AND DISCUSSION

There are five essential steps in the systematic process of statistical hypothesis testing [10]. The hypothesis is developed first, followed by the determination of the significance level, test statistic, rejection area, or significant probability (P value), and conclusions. At the conclusion step, "the null hypothesis cannot be rejected at the predetermined significance level" if the test statistic is outside the rejection area or the P value exceeds the predefined significance threshold. On the other hand, "the null hypothesis is rejected at the significance level" if the test statistic is inside the rejection area or if the P value is below the predefined significance threshold. In this instance, the alternative hypothesis—rather than the null hypothesis—is supported by the conclusions that are reached and understood. For instance, the null hypothesis is rejected in a statistical hypothesis test when the significance level is set at 0.05 and the significance probability is computed to be 0.002. Likewise, the null hypothesis is rejected in a statistical hypothesis test with a significance level of 0.1 and a computed significance probability of 0.07.

The alternative hypothesis's content serves as the foundation for the conclusion. This procedure gives researchers a methodical framework for thoroughly assessing hypotheses and deriving significant findings supported by statistical data. The pivotal phases of hypothesis testing contribute to the integrity and dependability of the study findings and provide a strong basis for gaining understanding of the underlying dynamics of null and alternative hypotheses. Calculations are made for Pearson's correlation coefficient and Spearman's correlation coefficient if one of the two variables is a rank scale. The degree of the linear link between two variables is shown by the two correlation coefficients. Furthermore, the stronger the positive linear relationship between the two variables, the closer the correlation coefficient is to +1; conversely, the stronger the negative linear relationship, the closer the correlation coefficient is to -1. There is no indication of a linear relationship if the correlation value is 0. To ascertain whether the computed correlation coefficient is zero, a significance test must be run. "The correlation coefficient is zero," according to the null hypothesis, and "the correlation coefficient is not zero," according to the alternative hypothesis. A significant positive or negative linear correlation is shown, depending on the sign of the correlation coefficient, if the null hypothesis is rejected. This leads to the conclusion that "the correlation coefficient cannot be said to be statistically zero under the significance level." All regression coefficients are zero" is the null hypothesis for the regression model's significance test, whereas at least one regression coefficient is not zero is the alternative hypothesis. The conclusion is that "it cannot be stated that all the

regression coefficients are zero under the significance level" if the null hypothesis is rejected. The computed regression model indicates a meaningful linear relationship between the response and explanatory variables since at least one regression coefficient is not zero. The alternative hypothesis asserts that "the regression coefficient is not zero," while the null hypothesis asserts that "the regression coefficient is zero" in order to assess the significance of the regression coefficient. Since the computed regression coefficient is not zero, if the other explanatory variables are held constant, a one-unit change in the explanatory variable causes the response variable to change by the regression coefficient's value. To be used, linear regression analyses need to meet a number of requirements. First, the residuals' distribution needs to be normal. Otherwise, a linear regression analysis should be replaced with a generalized linear model (GLM). Second, homoscedasticity must be met by the residuals. It is necessary to alter the regression model by changing the response variable in order to achieve homoscedasticity if the residuals do not. Third, independence must be satisfied by the residuals. A time-series regression analysis should be carried out instead of a linear regression analysis if the residuals do not fulfill independence, which suggests a dependent relationship between the residuals. Fourth, the regression model's linearity needs to be met. The explanatory or response variable must be changed, or the regression model must be reset, in order to achieve linearity if it is not satisfied. Lastly, the explanatory variables in the linear regression model shouldn't be multicollinear [11]. The regression model should be modified using variable selection techniques such stepwise selection, forward selection, and backward elimination if multicollinearity is present. Yamane's formula for determining sample size from a finite population is summed up in the document [12].

5. CONCLUSIONS AND RECOMENDATIONS

The application of appropriate statistics analysis in procurement, environmental, and social standards studies is fundamental for evidence-based decision-making. Through quantitative analysis, organizations can evaluate their practices, comply with regulatory requirements, and design programs that balance economic, environmental, and social objectives. The normality test, one-group mean and independent two-group mean difference test, dependent or before-and-after group mean difference test, one-way analysis of variance (ANOVA), repeated-measures ANOVA, chi-square test, correlation analysis, linear regression analysis, and logistic regression analysis are among the statistical analysis techniques covered in this study. The study thoroughly investigates the null and alternative hypotheses and demonstrates how to appropriately interpret the findings in the event that the null hypothesis is disproved. The outcomes of the statistical analysis are compared with predefined significance thresholds to

decide if the null hypothesis should be rejected. The observed data offers enough evidence to refute the idea that there is no effect or difference when the null hypothesis is rejected. The null hypothesis's rejection region in a two-tailed test indicates that the observed result deviates statistically significantly from the predicted result by falling on either the extreme left or right tail of the distribution.

REFERENCES

1. Kim, T. K., & Park, J. H. (2019). "More about the basic assumptions of t-test: normality and sample size." *Korean J Anesthesiol*, 72, 331-5.
2. Goodman, S. N. (2008). "A Dirty Dozen: Twelve P-Value Misinterpretations." *Seminars in Hematology*, 45(3), 135–140.
3. Khan, R. A., & Ahmad, F. (2015). "Power Comparison of Various Normality Tests." *Pak J Stat Oper Res*, 11, 331–45.
4. Kim, H. Y. (2015). "Statistical notes for clinical researchers: Assessing normal distribution (2) using skewness and kurtosis." *Korean Journal of Anesthesiology*, 68(4), 276–277. This paper emphasizes understanding the assumptions of statistical tests.
5. Schisterman, E. F., Cole, S. R., & Platt, R. W. (2009). "Overadjustment Bias and Unnecessary Adjustment in Epidemiologic Studies." *Epidemiology*, 20(4), 488–495. This article explains the importance of proper adjustments in statistical models.
6. Babyak, M. A. (2004). "What You See May Not Be What You Get: A Brief, Nontechnical Introduction to Overfitting in Regression-Type Models." *Psychosomatic Medicine*, 66(3), 411–421. This provides a practical overview of overfitting.
7. Bennett, C. M., Wolford, G. L., & Miller, M. B. (2009). "The Principled Control of False Positives in Neuroimaging." *Social Cognitive and Affective Neuroscience*, 4(4), 417–422.
8. Lewis-Beck, C., & Lewis-Beck, M. (2015). *Applied regression: An introduction*. 2nd ed. Thousand Oaks, Sage publications, Inc, 55-60.
9. Nahm, F. S. (2016). "Nonparametric statistical tests for the continuous data: the basic concept and the practical use." *Korean J Anesthesiol*, 69, 8–14. doi: 10.4097/kjae.2016.69.1.8.
10. Kwak, S. (2023). "Are only p-values less than 0.05 significant?" A p-value greater than 0.05 is also significant! *J Lipid Atheroscler*, 12, 89-95.
11. Kim, J. H. (2019). "Multicollinearity and misleading statistical results." *Korean J Anesthesiol*, 72, 558-69.
12. Yamane, T. (1967) *Statistics: An Introductory Analysis*. 2nd Edition, Harper and Row, New York.