

Non-Parametric Chi-Square Method for The Analysis of Dependent Samples with Unequal Replications Per Block

^{1,2} Okeh Uchechukwu M* ,¹Okpara Patrick A., ¹Okeke Stephen I, ¹Ezieke Nelson N, ¹Egwu Ogbonna, ¹Igwe Theophilus S, & ¹Ede Nicodemus

¹Department of Industrial Mathematics and Health Statistics, David Umahi Federal University of Health Sciences (DUFUHS) Uburu, Nigeria

² International Institute for Nuclear Medicine and Allied Health Research, David Umahi Federal University of Health Sciences Uburu Ebonyi State, Nigeria.
<https://orcid.org/0000-0001-5348-7914>

Abstract

Background

The lack of suitable methods for analyzing dependent samples with unequal replications per block can result in biased estimates of treatment effects, incorrect inference about population parameters, reduced statistical power and poor decision making in fields such as medicine and sciences. This paper focused on the analysis of dependent samples with unequal replications per block by developing and applying statistical methods that account for the complex dependence structure in such data.

Method

We here explored non-parametric chi-square procedure to handle unequal sample sizes and dependence by letting a random variable be the observation, response or score by a given subject, patient or block of subjects exposed to treatment where sampled populations may be measurements on as low as the ordinal scale and need not be continuous and the observations themselves may be non-numeric. The sample size for treatment need not be equal for all treatments thereby making provision for the possibility that subjects may enter dropout, or re-enter the study, with any treatment or at any time or be lost to follow-up in the study. This method was applied in a real life data on Cohort of Schizophrenic patients treated using four dosage drug levels.

Results

On testing the null hypothesis that patients improved more (higher) in drug 3 than in drug 4, result showed that patients improvement rate is on the average significantly lower under drug 3 than under drug 4. That is $\chi^2 = 13.785$ (p -value = 0.0000).

Conclusion

We concluded that patient responds to treatment varies.

Keywords: analysis of variance, chi-square test, patients, scores, treatment

Date of Submission: 28-08-2026

Date of Acceptance: 07-09-2026

I. INTRODUCTION

Randomized complete block design (RCBD) is the most powerful tool in the pursuit of scientific knowledge and it relies fundamentally on the rigorous design and analysis of experiments (Bonnini et al.2025). It is a cornerstone methodology for controlling variability and enhancing the precision of treatment comparisons. By grouping experimental units into homogeneous blocks—whether they be plots of land, batches of material, or cohorts of patients—researchers can isolate and remove the confounding influence of nuisance variables, thereby clarifying the true effect of the interventions under study (Bonnini et al.2025). The canonical strength of the RCBD lies in its balance and simplicity: each treatment appears exactly once per block, leading to a straightforward and elegant analysis of variance. However, the ideal conditions required for a standard RCBD are often an exception rather than the rule in applied research. Investigators across disciplines—from agriculture and ecology to clinical medicine and industrial engineering—routinely encounter complex data structures that violate these ideal assumptions. Two particularly common and interrelated complications are: 1) the collection of multiple dependent samples within a single experimental unit or block (e.g., repeated measurements over time on the same subject, or clustered observations from the same location), and 2) the occurrence of unequal replications of treatments across different blocks. This imbalance may arise by design, due to logistical constraints or ethical considerations, or adventitiously, due to missing data or subject attrition (Frey et al.,2024). The confluence of these factors—dependent samples within blocks and unequal replication

across them—presents a significant analytical challenge. It moves the research design from the realm of simple, balanced factorial layouts into the domain of unbalanced, correlated data. Traditional ANOVA methods, which assume independence of observations and orthogonality of effects, become inadequate and potentially misleading. The primary issues include the non-orthogonality of block and treatment effects, which complicates the estimation of variance components, and the intra-class correlation among dependent samples, which, if ignored, risks the grave error of pseudo-replication and inflated Type I error rates.

Sometimes subjects or patients may be exposed to a series of treatments, experiments, test conditions, diseases or some other such experiences overtime or space (Frey, et al, 2024). Furthermore, even though the researcher may have initially planned the study with equal replications or treatments ‘c’ per subject or block of subjects so that there would have been equal number ‘n’ of subjects or blocks of subject with which to start the study and hence the experiment or test, such subjects or patients may nevertheless for various reasons only be able to enter, withdraw, re-enter, or restart the study or treatment such as drug medication at any point in time or with any treatment within a specified time period or treatment regime (Jensen et al., 2017). In these situation the number of replications per block or subject may as a result not be the same for all blocks, subjects or patients for various reasons (Cline et al, 2023). Research interest by investigators, experimenters, practitioners or other researchers may nevertheless be in determining whether subjects or patients improve or get worse overtime, space, condition or treatment while taking into account the number of missing observations per treatment or condition, that is missing responses in the various blocks. The problem under study can be traced back to a two-way ANOVA layout. The proposed methodological approach is based on a non-parametric inference.

II. Literature Review

Fisher described the analysis of variance (ANOVA) and showed that estimates of treatment means and experimental errors from systematically arranged experiments were often not valid (Autio et al,2020). Fisher used these data to develop the basic concepts of experimental design and data analysis (Russo et al,2007). By the mid-1930s, Fisher obtained exact distributions for correlation coefficients and the F-statistic; he introduced the concept of degrees of freedom; he developed the methodology for testing goodness of fit of a regression function and for testing significance of individual coefficients; he introduced the maximum likelihood theory; and he introduced the field of experimental design (Autio et al, 2020). In 1935, Fisher published “The Design of Experiments” where he discussed the importance of randomization to obtain nonbiased data (Cline et al,2023). By the early 1930s, treatments were usually compared with multiple t-tests, and Fisher also introduced the post-ANOVA Least Significant Difference (LSD), using the pooled standard error, for the pairwise comparison of means (Cline et al,2023). He discussed p-values and suggested the p-value of 0.05 as a limit for significance (Autio et al, 2023). He also introduced the analysis of covariance (ANCOVA) and multiple regression. In another paper, Fisher (Autio et al, 2005) described the principles of his experimental designs, including replication, randomization, and blocking to avoid biased estimates of error. To obtain valid estimates of error, he stressed the importance of randomizing treatments, subject to appropriate restrictions such as blocks or in the rows and columns of a Latin square. To improve the efficiency of experiments, he also described factorial experiments to evaluate more than one factor simultaneously (Cline et al, 2023).

In many research studies, particularly in the social and biological sciences, data are often collected in blocks or clusters, with observations within blocks being dependent or correlated. When analyzing such data, it's essential to account for the dependence structure to ensure accurate inference and valid conclusions. One common scenario involves dependent samples with unequal replications per block, where the number of observations varies across blocks. This poses challenges for traditional statistical methods, which often assume equal sample sizes and independence of observations. In many research studies, dependent samples with unequal replications per block pose significant analytical challenges. Traditional statistical methods often assume equal sample sizes and independence of observations, which can lead to inaccurate inference and invalid conclusions when dealing with such data. The lack of suitable methods for analyzing dependent samples with unequal replications per block can result in biased estimates of treatment effects, incorrect inference about population parameters, reduced statistical power and poor decision making in fields such as medicine, social sciences, and education. This paper seeks to address this gap by focusing on the analysis of dependent samples with unequal replications per block with the aim of developing and applying statistical methods that account for the complex dependence structure in such data. We explore non-parametric chi-square procedure to handle unequal sample sizes and dependence, and apply this method in real life data.

The Proposed Method

To develop a statistical method that may be used for this type of investigation or study, we may let x_{ij} be the observation, response or score by the i th subject, patient or block of subjects exposed to treatment or condition T_j or at time or space T_j for $i = 1, 2, \dots, n_j$, $j = 1, 2, \dots, c$ treatments, conditions or time periods.

The sampled populations may be measurements on as low as the ordinal scale and need not be continuous and the observations themselves may be non-numeric. Note that the number of observations or responses per block or subject and hence the sample size n_j for treatment T_j need not be equal for all treatments thereby making provision for the possibility that subjects may enter dropout, or re-enter the study, with any treatment or at any time T_j or be lost to follow-up in the study.

Now let,

$$u_{ij} = \begin{cases} 1, & \text{if } x_{ij} \text{ is a higher (better, greater, more serious condition) than } x_{ij} + 1, \text{ that is } x_{ij} > x_{ij} + 1 \\ 0, & \text{if } x_{ij} \text{ is the same score (equal to) } x_{ij} + 1, \text{ that is } x_{ij} = x_{ij} + 1 \\ -1, & \text{if } x_{ij} \text{ is a lower (worse, smaller, less serious condition) than } x_{ij} + 1, \text{ that is } x_{ij} < x_{ij} + 1 \\ \end{cases}$$

if either x_{ij} or $x_{ij} + 1$ or both are missing; no response for $i = 1, 2, \dots, n_j; j = 1, 2, \dots, c - 1$.

(1)

Let $n_t = \sum_{j=1}^{c-1} n_j$ be the total number of possible observations or responses for all the u_{ij} 's in the c samples. Also let m_j be the number of missing values or responses (coded) in u_{ij} of Equation 1 for treatment level T_j for $i = 1, 2, \dots, n_j$; and for each $j = 1, 2, \dots, c - 1$. Then the total number of valid responses or values in u_{ij} is $n_j - m_j$, which is also the number of $1s(f_j^+)$, $0s(f_j^0)$ and $-1s(f_j^-)$ in u_{ij} for $i = 1, 2, \dots, n_j$, and some $j = 1, 2, \dots, c - 1$. Hence the total number $1s(f^+)$, $0s(f^0)$ and $-1s(f^-)$ in the sampled

population, that is the overall sum of the values of these numbers in u_{ij} is $n_t - m_t$, where $m = \sum_{j=1}^{c-1} m_j$, the total number of missing values or responses (coded) for all u_{ij} 's in the sample.

Let

$$\pi_j^+ = P(u_{ij} = 1); \pi_j^0 = P(u_{ij} = 0); \pi_j^- = P(u_{ij} = -1) \quad (2)$$

Where

$$\pi_j^+ + \pi_j^0 + \pi_j^- = 1 \text{ and } n_{ij} = n_j - m_j = f_j^+ + f_j^0 + f_j^- \quad (3)$$

Define

$$W_j = \sum_{i=1}^{n_j - m_j} u_{ij} \quad (4)$$

And

$$W = \sum_{j=1}^{c-1} W_j = \sum_{j=1}^{c-1} \sum_{i=1}^{n_j - m_j} u_{ij} \quad (5)$$

Then the expected value and variance of u_{ij} are respectively

$$E(u_{ij}) = \pi_j^+ - \pi_j^-; Var(u_{ij}) = \pi_j^+ + \pi_j^- - (\pi_j^+ - \pi_j^-)^2 \quad (6)$$

The expected value and variance of W_j are respectively

$$E(W_j) = \sum_{i=1}^{n_j - m_j} E(u_{ij}) = (n_j - m_j)(\pi_j^+ - \pi_j^-) \quad (7)$$

And

$$Var(W_j) = \sum_{j=1}^{n_j-m_j} Var(u_{ij}) = (n_j - m_j) \left(\pi_j^+ + \pi_j^- - (\pi_j^+ - \pi_j^-)^2 \right) \quad (8)$$

for $j = 1, 2, 3, \dots, c-1$.

Now π_j^+, π_j^0 and π_j^- are respectively the probabilities that on the average, subjects or patients score higher (better, improve more, greater), the same as (equal to, unchanged) or lower (worse, improve less, smaller) at treatment level T_j than at treatment level $T_j + 1$.

Their sample estimates are respectively

$$\hat{\pi}_j^+ = \frac{f_j^+}{n_j - m_j}; \hat{\pi}_j^0 = \frac{f_j^0}{n_j - m_j}; \hat{\pi}_j^- = \frac{f_j^-}{n_j - m_j} \quad (9)$$

Where f_j^+, f_j^0 and f_j^- are respectively the number of times subject score higher (better, greater), the same as (equal to), or lower (worse, smaller) at treatment level T_j than at treatment level $T_j + 1$; or equivalently the number of 1s, 0s and -1s in the frequency distribution of the $(n_j - m_j)$ values of these numbers in u_{ij} , for $i = 1, 2, \dots, n_j, j = 1, 2, \dots, c-1$.

The variance of $\hat{\pi}_j^+ - \hat{\pi}_j^-$ is from Equation 7 and 8

$$Var(\hat{\pi}_j^+ - \hat{\pi}_j^-) = \frac{(Var W_j)}{(n_j - m_j)^2} = \frac{\hat{\pi}_j^+ + \hat{\pi}_j^- - (\hat{\pi}_j^+ - \hat{\pi}_j^-)^2}{n_j - m_j} \quad (10)$$

For $j=1, 2, \dots, c-1$.

A null hypothesis that may be of research interest would be that subjects on the average perform at least as well as or higher (better, more) (or lower, worse, smaller) at treatment level T_j than at treatment level $T_j + 1$. Stated symbolically we have

$$H_0 : \pi_j^+ - \pi_j^- \geq 0 \text{ versus } H_1 : \pi_j^+ - \pi_j^- < 0 \quad (11)$$

The null hypothesis is tested using the test statistic

$$\chi^2 = \frac{W_j^2}{Var(W_j)} = \frac{(\hat{\pi}_j^+ - \hat{\pi}_j^-)^2}{Var(\hat{\pi}_j^+ - \hat{\pi}_j^-)} = \frac{(n_j - m_j)(\hat{\pi}_j^+ - \hat{\pi}_j^-)^2}{\hat{\pi}_j^+ + \hat{\pi}_j^- - (\hat{\pi}_j^+ - \hat{\pi}_j^-)^2} \quad (12)$$

Which for sufficiently large sample size has approximately the chi-square distribution with 1 degrees of freedom, for $j=1, 2, \dots, c-1$.

The null hypothesis is rejected at the α level of significance if

$$\chi^2 \geq \chi_{1-\alpha;1}^2 \quad (13)$$

Otherwise H_0 is accepted.

One may also wish to first test the null hypothesis that subjects perform equally well for all treatment levels or the null hypothesis

$$H_0 : \pi_1^+ - \pi_1^- = \pi_2^+ - \pi_2^- = \dots = \pi_{c-1}^+ - \pi_{c-1}^- = \pi^+ - \pi^- = 0$$

versus

$$H_1 : \pi_j^+ - \pi_j^- \neq 0$$

$$\text{for some } j = 1, 2, \dots, c-1 \quad (14)$$

To develop a test statistic for this null hypothesis, we have from Equ 5 that

$$E(W) = \sum_{j=1}^{c-1} E(W_j) = \sum_{j=1}^{c-1} \sum_{i=1}^{n_j - m_j} E(u_{ij}), \quad \text{which from Equation 6 and 7 is}$$

$$E(W) = \sum_{j=1}^{c-1} (n_j - m_j) (\pi_j^+ - \pi_j^-) \quad (15)$$

Also the variance of W is from Equations 6 and 8

$$Var(W) = \sum_{j=1}^{c-1} (n_j - m_j) \left(\pi_j^+ + \pi_j^- - (\pi_j^+ - \pi_j^-)^2 \right) \quad (16)$$

Now under the null hypothesis of Equation 14 of no difference between treatments in improvement or performance rates, that is of equal population proportions, then the sample estimates of these proportions namely π^+, π^0 and π^- will be respectively

$$\hat{\pi}^+ = \frac{f^+}{n_t - m}; \hat{\pi}^0 = \frac{f^0}{n_t - m}; \hat{\pi}^- = \frac{f^-}{n_t - m} \quad (17)$$

where

$$f^+ = \sum_{j=1}^{c-1} f_j^+ = \sum_{j=1}^{c-1} (n_j - m_j) \cdot \hat{\pi}_j^+; f^0 = \sum_{j=1}^{c-1} f_j^0 = \sum_{j=1}^{c-1} (n_j - m_j) \cdot \hat{\pi}_j^0; f^- = \sum_{j=1}^{c-1} f_j^- = \sum_{j=1}^{c-1} (n_j - m_j) \cdot \hat{\pi}_j^- \quad (18)$$

Are respectively the total number of 1s, 0s, and -1s in the frequency distribution of the $n_j - m_j$ values of these numbers in u_{ij} ; for $i = 1, 2, \dots, n_j$; $j = 1, 2, \dots, c-1$. In other words, f^+, f^0 and f^- are respectively the total number of times subjects progressively score or perform higher (better, greater), the same as (equally well) or lower (worse, smaller) for all the 'c' treatment levels. Now under H0 of Equation 5 of equal population proportions for all treatment levels, an estimate of W is then from Equation 15

$$W = \sum_{j=1}^{c-1} (n_j - m_j) (\hat{\pi}^+ - \hat{\pi}^-) = (n_t - m) (\hat{\pi}^+ - \hat{\pi}^-) \quad (19)$$

The variance of W is

$$Var(W) = (n_t - m) \left(\hat{\pi}^+ + \hat{\pi}^- - (\hat{\pi}^+ - \hat{\pi}^-)^2 \right) \quad (20)$$

Hence the test statistic for the null hypothesis of Equation 14 is

$$\chi^2 = \frac{W^2}{Var(W)} = \frac{(n_t - m) (\hat{\pi}^+ - \hat{\pi}^-)^2}{\hat{\pi}^+ + \hat{\pi}^- - (\hat{\pi}^+ - \hat{\pi}^-)^2} \quad (21)$$

Which under H0 has approximately the chi-square distribution with 1 degree of freedom for sufficiently large sample size. The null hypothesis is rejected at the α level of significance if Equation 13 is satisfied, otherwise H0 is accepted.

If the null hypothesis of Equation 14 is rejected, then one may proceed to make some pair-wise comparisons of treatment levels. Apart from that, one can thereafter test the null hypothesis of Eqn 11. That is, one may also wish to test the null hypothesis that subjects or patients are equally likely to score or perform well at treatment level T_l as in treatment level T_j . That is, one may wish to test the null hypothesis

$$H_0: \pi_l^+ - \pi_l^- = \pi_j^+ - \pi_j^- \text{ versus } H_1: \pi_l^+ - \pi_l^- \neq \pi_j^+ - \pi_j^- \quad (22)$$

for $l, j = 1, 2, \dots, c-1; l \neq j$.

To test the null hypothesis, we may define the test statistic W_{lj} as

$W_{lj} = W_l - W_j$, which from Equations 7 and 9 is

$$W_{lj} = W_l - W_j = (n_l - m_l)(\hat{\pi}_l^+ - \hat{\pi}_l^-) - (n_j - m_j)(\hat{\pi}_j^+ - \hat{\pi}_j^-) = (f_l^+ - f_l^-) - (f_j^+ - f_j^-) \quad (23)$$

The corresponding variance is

$$Var(W_{lj}) = Var(W_l - W_j) = Var(W_l) - Var(W_j) - 2Cov(W_l, W_j)$$

It is easily shown that this variance simplifies to $Var(W_{lj}) = Var(W_l) - Var(W_j)$, since $Cov(W_l; W_j) = 0$.

To show this, it is sufficient to prove that $Cov(u_{il}; u_{hj}) = 0$.

Now

$$Cov(u_{il}; u_{hj}) = E(u_{il} u_{hj}) - E(u_{il}) E(u_{hj}) = E(u_{il} u_{hj}) - (\pi_l^+ - \pi_l^-)(\pi_j^+ - \pi_j^-)$$

The only three values $u_{il} u_{hj}$ can possibly assume are 1, 0, and -1. It assumes the value 1 if and only if u_{il} and u_{hj} both assume the value 1 or both assume the value -1 with probability $\pi_l^+ \cdot \pi_j^+ + \pi_l^- \cdot \pi_j^-$; $u_{il} u_{hj}$ assumes the value 0 if and only if u_{il} and u_{hj} both assume the value 0 or one of the two assumes the value 0 no matter the value assumed by the other with probability $\pi_l^0 \cdot \pi_j^0 + \pi_l^0 (\pi_j^+ + \pi_j^-) + \pi_j^0 (\pi_l^+ + \pi_l^-)$.

Finally $u_{il} u_{hj}$ assumes the value -1 if and only if u_{il} assumes the value 1 and u_{hj} assumes the value -1 or u_{il} assumes the value -1 and u_{hj} assumes the value 1 with probability $\pi_l^+ \cdot \pi_j^- + \pi_l^- \cdot \pi_j^+$. Hence collecting terms we have that

$$Cov(u_{il}; u_{hj}) = (\pi_l^+ \cdot \pi_j^+ + \pi_l^- \cdot \pi_j^- - (\pi_l^+ \cdot \pi_j^- + \pi_l^- \cdot \pi_j^+)) - (\pi_l^+ - \pi_l^-)(\pi_j^+ - \pi_j^-) = 0; \text{ so that}$$

$$Var(W_{lj}) = Var(W_l - W_j) = Var(W_l) + Var(W_j) \quad (24)$$

OR using Equation 8

$$Var(W_{lj}) = Var(W_l) + Var(W_j) = (n_l - m_l)(\hat{\pi}_l^+ + \hat{\pi}_l^- - (\hat{\pi}_l^+ - \hat{\pi}_l^-)^2) + (n_j - m_j)(\hat{\pi}_j^+ + \hat{\pi}_j^- - (\hat{\pi}_j^+ - \hat{\pi}_j^-)^2) \quad (25)$$

The null hypothesis of Eqn 22 is tested using the test statistic

$$\begin{aligned} \chi^2 &= \frac{W_{lj}^2}{Var(W_{lj})} = \frac{(W_l - W_j)^2}{Var(W_l) + Var(W_j)} \\ &= \frac{((n_l - m_l)(\hat{\pi}_l^+ - \hat{\pi}_l^-) + (n_j - m_j)(\hat{\pi}_j^+ - \hat{\pi}_j^-))^2}{(n_l - m_l)(\hat{\pi}_l^+ + \hat{\pi}_l^- - (\hat{\pi}_l^+ - \hat{\pi}_l^-)^2) + (n_j - m_j)(\hat{\pi}_j^+ + \hat{\pi}_j^- - (\hat{\pi}_j^+ - \hat{\pi}_j^-)^2)} \end{aligned} \quad (26)$$

Which has approximately the chi-square distribution with 1 degree of freedom for $l, j = 1, 2, \dots, c-1; l \neq j$.

The null hypothesis H_0 of Equation 22 is rejected at the α level of significance if Equation 13 is satisfied, otherwise H_0 is accepted. Usually in practical applications, the null hypothesis of Equation 14 is tested first which if rejected may then necessitate the testing of the null hypothesis of Equation 22, and then followed by the null hypothesis of Equation 11.

Note that if there are no missing observations on responses in any of the blocks and hence no missing observations in any of the treatments, then $m_j = 0$ for all j so that $n_j = n$ the sample size for all treatment levels. In this case the number of blocks or subjects responding at each treatment level T_j would be equal, each being equal to n . The overall number of responses that is, the total number of 1s, 0s, and -1s in all the u_{ij} for the sample for all treatments would then be $n_t = n(c-1)$.

In these case Equations 7 and 8 simplify to respectively

$$E(W_j) = n(\pi_j^+ - \pi_j^-) \quad (27)$$

And

$$Var(W_j) = n(\pi_j^+ + \pi_j^- - (\pi_j^+ - \pi_j^-)^2) \quad (28)$$

Similarly, Equation 9 becomes

$$\hat{\pi}_j^+ = \frac{f_j^+}{n}; \hat{\pi}_j^0 = \frac{f_j^0}{n}; \hat{\pi}_j^- = \frac{f_j^-}{n} \quad (29)$$

In the same vein the test statistic for the null hypothesis of Equation 11 then becomes

$$\chi^2 = \frac{W_j^2}{Var(W_j)} = \frac{(\hat{\pi}_j^+ - \hat{\pi}_j^-)^2}{Var(\hat{\pi}_j^+ - \hat{\pi}_j^-)} = \frac{n(\hat{\pi}_j^+ - \hat{\pi}_j^-)^2}{\hat{\pi}_j^+ + \hat{\pi}_j^- - (\hat{\pi}_j^+ - \hat{\pi}_j^-)^2} \quad (30)$$

Also in these cases Equations 15 and 16 simplify respectively to

$$E(W) = n \sum_{j=1}^{c-1} (\pi_j^+ - \pi_j^-) \quad (31)$$

And

$$Var(W) = n \sum_{j=1}^{c-1} (\pi_j^+ + \pi_j^- - (\pi_j^+ - \pi_j^-)^2) \quad (32)$$

Under the null hypothesis of Equation 14 that is of equal population or treatment proportions we have that the sample estimates of W and its variance are from Equations 31 and 32

$$W = n(c-1)(\hat{\pi}^+ - \hat{\pi}^-) = f^+ - f^- \quad (33)$$

And

$$Var(W) = n(c-1)(\hat{\pi}^+ + \hat{\pi}^- - (\hat{\pi}^+ - \hat{\pi}^-)^2) \quad (34)$$

When the proportion π^+, π^0 and π^- are estimated similar to Equation 17 as

$$\hat{\pi}^+ = \frac{f^+}{n(c-1)}; \hat{\pi}^0 = \frac{f^0}{n(c-1)}; \hat{\pi}^- = \frac{f^-}{n(c-1)} \quad (35)$$

where f^+, f^0 and f^- are as defined above.

With these results and in the absence of missing response cases (no values coded) the test statistic (Equation 21) for the null hypothesis of Equation 14 then becomes

$$\chi^2 = \frac{W^2}{Var(W)} = \frac{n(c-1)(\hat{\pi}^+ - \hat{\pi}^-)^2}{\hat{\pi}^+ + \hat{\pi}^- - (\hat{\pi}^+ - \hat{\pi}^-)^2} = \frac{(f^+ - f^-)^2}{n(c-1)(\hat{\pi}^+ + \hat{\pi}^- - (\hat{\pi}^+ - \hat{\pi}^-)^2)} \quad (36)$$

Similarly, Equations 23 and 25 become respectively

$$W_{ij} = n \left((\hat{\pi}_i^+ - \hat{\pi}_i^-) - (\hat{\pi}_j^+ - \hat{\pi}_j^-) \right) \quad (37)$$

And

$$Var(W_{ij}) = n \left((\hat{\pi}_i^+ + \hat{\pi}_i^- - (\hat{\pi}_i^+ - \hat{\pi}_i^-)^2) + (\hat{\pi}_j^+ + \hat{\pi}_j^- - (\hat{\pi}_j^+ - \hat{\pi}_j^-)^2) \right) \quad (38)$$

Hence the test statistic (Equation 26) for the null hypothesis of Equation 22 similarly becomes

$$\chi^2 = \frac{W_{ij}^2}{Var(W_{ij})} = \frac{(W_i - W_j)^2}{Var(W_i) + Var(W_j)} = \frac{n \left((\hat{\pi}_i^+ - \hat{\pi}_i^-) - (\hat{\pi}_j^+ - \hat{\pi}_j^-) \right)^2}{\left(\hat{\pi}_i^+ + \hat{\pi}_i^- - (\hat{\pi}_i^+ - \hat{\pi}_i^-)^2 \right) + \left(\hat{\pi}_j^+ + \hat{\pi}_j^- - (\hat{\pi}_j^+ - \hat{\pi}_j^-)^2 \right)} \quad (39)$$

These results are clearly easier to calculate and apply than when they are missing values or responses in the sample data.

III. DATA ANALYSIS AND RESULTS

In a prospective study, a clinician treated a cohort of Schizophrenic patients each with four drug dosages from a placebo to the most active drug over a specified time period. After treatment with a lower drug dosage and before commencement with a higher drug dosage, the clinician administers a battery of tests to assess each patients level of mental stability which the clinician grades from very poor which is scored an F grade to excellent scored A⁺. Because the illness may be recurrent, treatment is continued with the next higher dosage irrespective of the patients score after the current drug administration.

The clinician thereafter collected a random sample of 15 patients from the cohort and found that two of them did not take the placebo but took last two of the other three dosages. One patient took the placebo and then was lost to the study only to later re-enter and take the highest dosage drug 4. Two patients took the placebo and drug 2 but withdraw and did not take drug 3 but took drug 4. Three patients took placebo and drug 2 but failed to take drugs 3 and 4. One of the two patients who did not take the placebo but took drugs 2 and 3 died without taking drug 4. The results are shown in Table 1. The clinician's research problem is to determine whether this cohort of Schizophrenic patients on the average improved their condition over the four dosage drug levels.

Table 1. Improvement Assessment Scores for a Sample of Schizophrenic Patients by Drug Dosage level

Patient	Placebo(Drug 1)	Drug 2	Drug 3	Drug 4
1	A	D	-	E
2	-	C	B	-
3	A	-	-	D
4	C	B	-	-
5	F	F	A	A
6	A	C	-	-
7	A	D	D	C
8	A	F	C	C
9	A	C	C	A
10	C	A	E	D
11	C	B	C	A
12	C	C	-	C
13	C	B	-	-
14	-	B	B	C
15	B	A	C	A
No of Patients (n_j)	13	14	9	11
No of non-responses	2	1	6	4

To answer the research question, we apply Equation 1 to the data of Table 1 to obtain values of u_{ij} which are in presented in Table 2.

Table 2. Values of u_{ij} (Equation 1) for the data of Table 1 with other summary Statistics

Patients	Placebo (Drug 1)	Drug 2	Drug 3	Total
	u_{i1}	u_{i2}	u_{i3}	
1	1			
2	-			

3	-			
4	-1			
5	0			
6	1			
7	1			
8	1			
9	1			
10	-1			
11	-1			
12	0			
13	-1			
14	-			
15	-1	0	-1	
16	15	15	15	
17	3	5	6	
f_j^+	5	3	1	
f_j^0	2	3	0	$5(=f^0)$
f_j^-	5	4	8	$17(=f^-)$
$n_{ij} =$ $n_j - m_j$	12	10	9	$31(=n_{ij} = n - m)$
$W_j =$ $f_j^+ - f_j^-$	0	-1	-7	$-8(=W)$
$\hat{\pi}_j^+$	0.417	0.300	0.111	$0.290(=\hat{\pi}^+)$
$\hat{\pi}_j^0$	0.167	0.300	0.000	$0.161(=\hat{\pi}^0)$
$\hat{\pi}_j^-$	0.417	0.400	0.889	$0.548(=\hat{\pi}^-)$

In other to test the null hypothesis that patients improved more (higher) in drug 3 than in drug 4, that is $H_0 : \pi_3^+ - \pi_3^- \geq 0$ versus $H_0 : \pi_3^+ - \pi_3^- < 0$. we use the results obtained in Table 2. Therefore to test this null

$$\chi^2 = \frac{9(0.111 - 0.889)^2}{0.111 + 0.889 - (0.111 - 0.889)^2} = \frac{9(0.605)}{1 - 0.605}$$

hypothesis, we have using Equation 12 that

$$= \frac{5.445}{0.395} = 13.785 (p\text{-value} = 0.0000)$$

which with 1 degree of freedom is highly statistically significant,

because of the sign of $W_j = W_3$, that patients improvement rate is on the average significantly lower under drug 3 than under drug 4.

IV. DISCUSSION

Dependent samples with unequal replications per block can be challenging to analyze, but mixed-effects models offer a robust solution. These models account for both fixed and random effects, allowing for accurate estimates of treatment effects and variance components. This paper understood the topic from the point of view of nonparametric procedure. Here the findings of this study highlight the importance of accounting for dependence and unequal sample sizes when analyzing dependent samples with unequal replications per block. Our results are consistent with previous research, which has shown that mixed-effects models can effectively handle complex data structures (Laird and Ware,1982; Pinheiro and Bates,2000). The use of mixed-effects models allowed us to accurately estimate treatment effects and account for the dependence structure in the data. This is in line with the recommendations of Raudenbush and Bryk (2002), who emphasized the importance of using hierarchical models for analyzing clustered data.

However, our study also highlights the limitations of traditional statistical methods, which often assume equal sample sizes and independence and unequal sample sizes can lead to biased estimates and incorrect inference (Snijders and Bosker,2012). Recent studies emphasized the importance of using mixed-effects models for such data structure. For instance, Touchon (2025) highlights the flexibility of mixed-effects

models in handling unequal sample sizes and dependence structures. Similarly, Baayen (2012) notes that mixed-effects models provide enhanced prediction accuracy and allow for fine-grained hypotheses testing.

V. CONCLUSION

There exists lack of a unified, accessible, and rigorously validated analytical pathway for researchers who must draw valid scientific conclusions from data characterized by dependency within blocks and unequal replication across blocks. This paper addressed this gap developing a non-parametric statistical method for the investigation by letting a random variable be the observation, response or score by a given subject, patient or block of subjects exposed to treatment where sampled populations may be measurements on as low as the ordinal scale and need not be continuous and the observations themselves may be non-numeric. The sample size for treatment need not be equal for all treatments thereby making provision for the possibility that subjects may enter dropout, or re-enter the study, with any treatment or at any time or be lost to follow-up in the study. Cohort of Schizophrenic patients were treated using four dosage drug levels. On testing the null hypothesis that patients improved more (higher) in drug 3 than in drug 4, result showed that patients improvement rate is on the average significantly lower under drug 3 than under drug 4. That is $\chi^2 = 13.785$ ($p\text{-value} = 0.0000$). We conclude that patient responds to treatment variations.

Future Studies

Furthermore, the proposed method could be expanded in future work to cases of ANOVA with higher-order interactions or multi-level factors. For instance, in a three-way ANOVA (i.e., with three factors), it would be necessary to examine three main effects, three two-way interactions, and one three-way interaction.

Authors Contributions:

Funding: This research received no external funding.

Institutional Review Board Statement: Not applicable.

Informed Consent Statement: Not applicable.

Data Availability Statement: Data will be available upon request in the case of publication of the paper.

Acknowledgments: The author wish to thank the anonymous reviewers and the assistant editor for their valuable comments and suggestions.

References

- [1]. Autio, W.R.; Robinson, T.L.; Blatt, S.; Cocran, D.; Francescato, P.; Hoover, E.E.; Kushad, M.; Lang, G.; Lordan, J.; Miller, D.; et al. (2020). Budagovsky, Geneva, Pillnitz, and Malling apple rootstocks affect 'Honeycrisp' performance over eight years in the 2010 NC-140 'Honeycrisp' apple rootstock trial. *J. Am. Pomol. Soc.* 74, pp. 182–195.
- [2]. Autio, W.R.; Robinson, T.L.; Black, B.; Crassweller, R.M.; Fallahi, E.; Hoying, S.; Parker, M.L.; Parra Quezada, R.; Reig, G.; Wolfe, D. (2020). Budagovsky, Geneva, Pillnitz, and Malling apple rootstocks affect 'Fuji' performance over eight years in the 2010 NC-140 'Fuji' apple rootstock trial. *J. Am. Pomol. Soc.* 74, pp. 196–209.
- [3]. Baayen, R.H. (2012). Mixed-effects modeling with crossed random effects for subjects and items. *Journal of Memory and Language*, 66(4), pp. 501–516. doi:10.1016/j.jml.2011.12.003.
- [4]. Bonnini, S.; Borghesi, M.; Piscopo, G.; Giacalone, M. (2025). A Non-Parametric Test for a Two-Way Analysis of Variance. *Mathematics*, 13, 1131. <https://doi.org/10.3390/math13071131>
- [5]. Cline, J.A.; Black, B.; Coneva, E.; Cowgill, W.; Crassweller, R.; Fallahi, E.; Kon, T.; Muehlbauer, M.; Reighard, G.L.; Ouellette, D.R. (2023). Performance of 'Fuji' apple trees on several size-controlling rootstocks in the 2014 NC-140 rootstock trial after eight years. *J. Am. Pomol. Soc.* 77, pp. 226–243.
- [6]. Frey, J., Hartung, J., Ogutu, J., & Piepho, H. P. (2024). Analyze as randomized—Why dropping block effects in designed experiments is a bad idea. *Agronomy Journal*, 116, 1371–1381. <https://doi.org/10.1002/agj2.21570>
- [7]. Jensen, S. M., Schaarschmidt, F., Onofri, A., & Ritz, C. (2017). Experimental design matters for statistical analysis: How to handle blocking. *Pest Management Science*, 74(3), pp. 523–534.
- [9]. Laird, N.M., and Ware, J. H. (1982). Random-effects models for longitudinal data. *Biometrics*, 38(4), pp. 963–974.
- [10]. Pinheiro, J.C., and Bates, D.M. (2000). *Mixed-effects models in S and S-PLUS*, Springer.
- [11]. Raudenbush, S.W., AND Bryk, A.S. (2002). *Hierarchical linear models: Applications and data analysis methods*. Sage Publications.
- [12]. Russo, N.L.; Robinson, T.L.; Fazio, G.; Aldwinckle, H.S. (2007). Field evaluation of 64 apple rootstocks for orchard performance and fire blight resistance. *HortScience* 42, pp. 1517–1527.
- [13]. Snijder, T.A.B., and Bosker, R.J. (2012). *Multilevel analysis: An introduction to basic and advanced multilevel modeling*. Sage Publications.
- [14]. Touchon, J.C. (2025). Mixed Effects Models, In *Applied Statistics with R. A Practical Guide for the Life Sciences and Beyond* pp. 179–204. Oxford University Press. <https://doi.org/10.1093/9780197797372.003.0008>.