

## Article

# Rank-Poisson Transformation for Use with Count Data in Poisson Regression

Daniel B. Wright \*  and Sage N. Stafford 

Department of Educational Psychology, Leadership, and Higher Education, University of Nevada,  
Las Vegas, NV 89154, USA; sage.stafford@unlv.edu

\* Correspondence: daniel.wright@unlv.edu

## Abstract

Count outcomes are commonly analyzed using Poisson regression, but empirical data often exhibit overdispersion, excess ties, heaping, or other departures from the Poisson distribution. This paper evaluates a rank-Poisson transformation, denoted *poisrank*, designed to map observed counts onto Poisson quantiles before fitting a Poisson regression model. Our goal is to test whether a rank-Poisson transformation offers a useful general-purpose strategy when count data do not satisfy Poisson assumptions. Using an empirical example and a Monte Carlo simulation study with Poisson, overdispersed, rounded, and gapped count distributions, we compared Poisson regression on raw counts, Poisson regression after the *poisrank* transformation, quasi-Poisson regression, and additional comparison approaches. Although the transformation made the marginal distribution more similar to a Poisson distribution, it generally did not outperform standard alternatives for inference. In particular, quasi-Poisson regression more consistently maintained appropriate rejection rates with overdispersion whereas *poisrank* tended to be conservative and often reduced power. These findings suggest that the rank-Poisson transformation is better understood as an exploratory robustness device than as a preferred replacement for established count-data methods.

**Keywords:** robust statistics; Poisson regression; overdispersion; transformations; statistical power

## 1. Introduction

Count data occur in the medical, educational, social, and physical sciences, and Poisson regression is one of the most common tools for modeling them. However, the Poisson model assumes that the conditional variance is approximately equal to the conditional mean, termed equidispersion [1,2], an assumption that is often violated in practice. Real count outcomes may be overdispersed, rounded, heaped at preferred values, censored, zero-inflated, etc. When these departures are ignored, Poisson regression can yield misleading standard errors and overly liberal significance tests [1]. The practical problem is how to model count data when assumptions are sufficiently violated to affect the validity of the conclusions.

One possible response to distributional misspecification is to transform the outcome. In settings involving non-normal continuous data, rank-based inverse transformations (INT) can be used to make variables more compatible with methods that assume normality [3]. In van der Waerden's article and a series of letters [4,5], he showed that ranking values for the response variable and then applying the inverse normal function improved performance



Academic Editors: Tommi Sottinen  
and John D. Clayton

Received: 23 February 2026

Revised: 6 May 2026

Accepted: 14 May 2026

Published: 20 May 2026

**Copyright:** © 2026 by the authors.  
Licensee MDPI, Basel, Switzerland.  
This article is an open access article  
distributed under the terms and  
conditions of the [Creative Commons  
Attribution \(CC BY\)](https://creativecommons.org/licenses/by/4.0/) license.

for ANOVA and  $t$ -tests when the data were not normally distributed. Beasley et al. note that this approach has “gained in popularity among genetic researchers... INTs have been applied in a variety of genetic research designs” (p. 580) [6]. Transforming both variables also improves performance of correlations, structural equation models, etc., [7]. A simple R function follows (several R packages also have functions for this).

```
normrank <- function(x,delta=.7)
  qnorm((rank(x) - delta)/(length(x) - 2 * delta + 1))
```

The present paper explores an analogous idea for count data. Specifically, we consider whether ranking observed counts and mapping those ranks to Poisson quantiles make the data more suitable for Poisson regression when the original counts are not well described by a Poisson distribution. Thus, the goal of this paper is to define and evaluate a rank-Poisson transformation, which we call *poisrank*. Our contribution is not simply to introduce the transformation as an algorithm, but to assess whether it yields useful statistical behavior relative to standard alternatives. In particular, we ask whether applying the transformation before Poisson regression provides more accurate inference than fitting Poisson regression directly to non-Poisson count data, and whether it offers any practical advantage over established remedies such as quasi-Poisson or negative binomial regression approaches.

For both an empirical example and simulation conditions, the rank-Poisson transformation often reduced the discrepancy between the observed marginal distribution and a Poisson distribution. However, this distributional improvement did not consistently translate into better inference. In the settings we test, especially with overdispersion, quasi-Poisson regression more closely tracked the intended rejection behavior, whereas the *poisrank* approach tended to be conservative and often had lower power. Accordingly, the paper concludes that the transformation is of methodological interest but does not currently justify replacing standard count-data models as a default analytic strategy.

### 1.1. Rank-Poisson Transformation

Let  $x_1, \dots, x_n$  denote the observed counts. The rank-Poisson transformation proceeds in three steps. First, the counts are ranked to obtain  $r_i = \text{rank}(x_i)$ . Second, each rank  $r_i$  is converted to a pseudo-percentile  $p_i$  using

$$p_i = \frac{r_i - \delta}{n - 2\delta + 1} \quad (1)$$

where  $\delta$  is a small offset chosen to avoid endpoint problems. Third, the pseudo-percentile is mapped to the corresponding quantile of a Poisson distribution with parameter  $\lambda$ :

$$x_i^* = Q_{\text{Pois}(\lambda)}(p_i), \quad (2)$$

where  $Q_{\text{Pois}(\lambda)}$  denotes the Poisson quantile function. The transformed values  $x_i^*$  are then analyzed using Poisson regression. In the implementation examined here,  $\lambda$  is estimated from the sample mean or, optionally, a trimmed mean of the observed counts.

As the *normrank* transformation seeks to transform the response variable so the data satisfy the some assumptions of OLS (normal) regression, a *poisrank* transformation seeks to transform count variables to better approximate the Poisson distribution. An analog to the *normrank* function follows (the output is turned into type integer because *qpois* does not return this type). This function has two options: the *delta* serves a similar role as used in the *normrank* transform in transforming the ranked values into percentiles, and the second option allows the mean to be trimmed [8]. See the Supplementary Materials for

details of how defaults were selected for the function's arguments (accessed 15 May 2026 <https://github.com/dbrookswr/poisrank>).

```
poisrank <- function(x,delta=.0001,tr=0){
  x <- qpois((rank(x)-delta)/(length(x) -2 * delta + 1),mean(x,tr=tr))
  return(as.integer(x))
}
```

A central difficulty for any rank-based approach to count data is that counts are discrete and frequently contain many ties. This matters because tied observations receive the same or nearly the same rank information, and distinct observed values may map to the same transformed value after application of the Poisson quantile function. For example,

```
x1 <- c(0,0,0,0,0,0,0,0,1,2,2,3)
poisrank(x1)

## [1] 0 0 0 0 0 0 0 0 1 1 1 2
```

$1 \rightarrow 1$  and  $2 \rightarrow 1$ . The ranking algorithm assigns tied values the mid-rank, as is the norm in statistics.

Different original counts can collapse to the same transformed count, so the mapping is not invertible in the usual sense. This is not a minor technicality but a substantive limitation of the approach because ties are common in count variables.

The transformation also alters interpretability. In an ordinary Poisson regression, coefficients can often be discussed in terms of multiplicative changes in expected counts. After a rank-based Poisson transformation, that interpretation is no longer available on the original outcome scale. The transformed counts preserve some order information (not always with ties) but not metric spacing, which means the resulting coefficients are most naturally interpreted as operating on the transformed scale rather than as original-count incidence-rate effects. For that reason, the method is better suited to inferential comparison than to direct substantive interpretation.

From a computational standpoint, the transformation is simple. Ranking  $n$  observations generally requires  $O(n \log n)$  time, and the subsequent percentile conversion and Poisson quantile lookup are linear in  $n$ . Thus, the preprocessing cost of the method is dominated by ranking, after which the usual Poisson model fitting proceeds. This low computational burden is one practical advantage of the approach, although computational simplicity does not imply good inferential performance.

### 1.2. Relation to Existing Approaches

When count data depart from Poisson assumptions, researchers already have several model-based options. Quasi-Poisson regression adjusts standard errors to account for dispersion while leaving the mean structure unchanged. Negative binomial regression accommodates stronger overdispersion by modeling the variance as a nonlinear function of the mean [9]. Zero-inflated and hurdle models are appropriate when excess zeros arise from distinct data-generating processes, and censored-count settings may call for other specialized models. These methods are grounded in explicit assumptions about how and why the data depart from the Poisson distribution. Therefore, it is valuable for researchers to consider the idiosyncrasies that give rise to these distributions. While there are procedures for censored and zero-inflated data, statisticians have not developed procedures for each of the many ways that give rise to these distributions.

The rank-Poisson transformation occupies a different conceptual niche. Rather than specifying a distributional mechanism for overdispersion, excess zeros, or heaping, it

attempts to reshape the observed outcome so that Poisson regression may become more tenable. Its potential appeal lies in its generality: one need not commit to a specific non-Poisson data-generating model. Simultaneously, that generality may come at a cost because a broadly defined transformation may fail to target the reasons for the extra-Poisson variation as model-based alternatives do. Accordingly, the relevant question is not whether `poisrank` can make a distribution look more Poisson-like, but whether it can improve inference relative to approaches already recommended for non-Poisson counts. The transformation is evaluated against methods that are already widely used when Poisson assumptions are suspect.

## 2. Does Getting Tweeted Increase Citations?

We evaluate the rank-Poisson transformation from two complementary perspectives. The first is empirical: we re-analyze a published count dataset to illustrate how inference changes when the outcome is modeled using methods that differ in how they handle non-Poisson variation. The second is statistical: we conduct a Monte Carlo simulation study to compare rejection rates across methods under known data-generating processes. The simulation includes a true Poisson condition as well as several non-Poisson conditions chosen to reflect substantively important departures from Poisson assumptions. These analyses allow us to assess both the practical consequences and the repeated-sampling properties of the proposed transformation.

To illustrate the practical issue, we re-examined citation counts from a study of whether articles promoted on Twitter subsequently received more Google Scholar citations. In these data, the mean citation count was 12.32, whereas the variance was 214.61, indicating substantial overdispersion relative to the Poisson model. Thus, these data provide a useful test case because the raw Poisson model yields a statistically significant result despite an outcome distribution that clearly violates equidispersion.

Branch and colleagues examined whether tweeting about an article increased the article's citation rate [10]. They had eleven academic social media influencers choose five articles for each of ten months, then randomly choose one of the five to tweet about. In total there were 550 articles and 110 were tweeted about. The data can be analyzed taking into account the multilevel structure [11], but here we do not consider who the influencer was as a factor for simplicity. The outcome measure was the number of citations on Google Scholar during the next three years and the predictor variable was whether the article was tweeted about. Their data can be found at <https://doi.org/10.1371/journal.pone.0292201.s001>.

Table 1 summarizes the results. Poisson regression results are statistically significant, suggesting being tweeted about increases the citation rate. However, because the variance is much greater than the mean, the assumptions for the Poisson regression are not met. The quasi-Poisson model was estimated with the base R function `glm` [12]. The negative binomial regression was estimated using the `glm.nb` function from **MASS** [13]. Both of these regressions produced non-significant findings, which coincide with the descriptive statistics and the small Cohen's  $d$ . If the `poisrank` function is used on these data before a Poisson regression conducted, it similarly produces a non-significant finding. In addition, the `normrank` function [7] was used with a linear regression, and it also produced a non-significant difference. Thus, when the outcome is overdispersed, raw Poisson regression can produce a conclusion that is not supported by other analytic approaches, including quasi-Poisson, negative binomial, and the rank-Poisson transformation analysis.

The preceding example provides a practical demonstration of the proposed method and highlights why a Poisson regression on the untransformed data can be misleading. Although the `poisrank` method did not produce the misleading result that the Poisson regression on the raw data did, a single example cannot show whether it performs well in

general. To determine whether that pattern generalizes beyond this single dataset, the next section reports a simulation study comparing rejection rates across methods under known Poisson- and non-Poisson-generating conditions.

**Table 1.** Testing whether being tweeted significantly affected the number of citations for Branch et al. [10].

| Model                                | Coefficient | se   | t     | p      |
|--------------------------------------|-------------|------|-------|--------|
| Poisson regression no transformation | 0.11        | 0.03 | 3.71  | <0.001 |
| Quasi-Poisson regression             | 0.11        | 0.12 | 0.89  | 0.373  |
| Negative binomial regression         | 0.11        | 0.10 | 1.10  | 0.272  |
| Poisson regression with poisrank     | −0.02       | 0.03 | −0.52 | 0.601  |
| Linear regression with normrank      | −0.07       | 0.11 | −0.68 | 0.497  |

### 3. Testing a Two Group Comparison

The simulation study was designed to evaluate whether the rank-Poisson transformation improves inferential performance across several classes of count distributions. Of particular interest was whether the method could maintain appropriate rejection rates when the data were not Poisson-distributed while still retaining reasonable power when group differences were present. To answer this question, we compared the proposed method with raw Poisson regression, quasi-Poisson regression, and normrank-based linear regression across data-generating conditions that differed in practical (e.g., the data generating based on how some survey responses arise) and important ways from the Poisson model (e.g., both over- and underdispersion). The simulation was designed not merely to assess whether the transformation made the data look more Poisson-like, but whether it made the rejection rate more appropriate when the data-generating process was known.

#### 3.1. Methods

The Monte Carlo simulation examined a simple two-group comparison in which the outcome was a count variable. For each condition, data were generated for a control group and an experimental group under three effect sizes, approximately 0.0, 0.2, and 0.5 standard deviations of the control group. Sample size was fixed at  $n = 100$ , and 10,000 replications were generated for each condition. This design allowed us to assess both false-positive behavior under the null condition and rejection rates when a true group difference was present.

Six types of count distributions were examined. The first was a true Poisson condition, included as a benchmark in which ordinary Poisson regression should perform well. The second was an overdispersed condition, representing the common case in which the variance exceeds the mean. The third introduced rounding or heaping, so that some values were disproportionately concentrated at preferred “prototypical” response options. The fourth introduced gaps in the outcome distribution while preserving approximate equidispersion. The fifth is a zero-inflated Poisson distribution, and the final is an underdispersed distribution, where values from a Poisson distribution are truncated above a certain value. These conditions collectively represent several substantively different ways that count data may depart from a Poisson model.

The Poisson condition was generated with  $\lambda = 4$  in the control group, and the experimental group mean was varied to produce the target effect sizes. The overdispersed condition was generated using a negative binomial distribution parameterized to produce a control-group mean near 4 and a substantially larger variance. The rounded condition began with Poisson data and then applied rounding to create heaping at selected values. The gap condition also began with Poisson data, after which observations in the middle range were probabilistically displaced to create gaps in the distribution. The zero-inflated

began with a Poisson distribution and then a random 40% were assigned zero. The truncated distribution began with a Poisson ( $\lambda = 4$ ), with all values above 4 (about 37%) then changed to 4. Additional generation details are provided in Table 2 and the Supplementary Materials.

**Table 2.** The means, variances, and overdispersion ratios for the distributions examined.

| Distribution           | Statistics   | Effect = 0 |       | Effect = $0.2\sigma$ |       | Effect = $0.5\sigma$ |       |
|------------------------|--------------|------------|-------|----------------------|-------|----------------------|-------|
|                        |              | Control    | Exp   | Control              | Exp   | Control              | Exp   |
| Poisson                | mean         | 4.00       | 4.00  | 4.00                 | 4.40  | 4.00                 | 5.01  |
|                        | variance     | 4.00       | 4.01  | 3.99                 | 4.38  | 3.99                 | 5.00  |
|                        | overd. ratio | 1.00       | 1.00  | 1.00                 | 1.00  | 1.00                 | 1.00  |
| $\Delta\bar{x}/sd_c =$ |              | 0.00       |       | 0.20                 |       | 0.50                 |       |
| overdispersed          | mean         | 4.00       | 4.00  | 4.00                 | 4.69  | 4.00                 | 5.73  |
|                        | variance     | 12.05      | 12.06 | 11.99                | 15.78 | 12.03                | 22.09 |
|                        | ratio        | 3.01       | 3.01  | 3.00                 | 3.36  | 3.01                 | 3.86  |
| $\Delta\bar{x}/sd_c =$ |              | 0.00       |       | 0.20                 |       | 0.50                 |       |
| Rounded                | mean         | 4.02       | 4.02  | 4.02                 | 4.45  | 4.02                 | 5.10  |
|                        | variance     | 4.73       | 4.71  | 4.71                 | 5.39  | 4.70                 | 8.62  |
|                        | ratio        | 1.18       | 1.17  | 1.17                 | 1.21  | 1.17                 | 1.69  |
| $\Delta\bar{x}/sd_c =$ |              | −0.00      |       | 0.20                 |       | 0.50                 |       |
| Gaps                   | mean         | 4.00       | 4.00  | 4.00                 | 4.40  | 4.00                 | 5.01  |
|                        | variance     | 4.01       | 4.01  | 4.02                 | 4.41  | 4.03                 | 5.03  |
|                        | ratio        | 1.00       | 1.00  | 1.00                 | 1.00  | 1.01                 | 1.00  |
| $\Delta\bar{x}/sd_c =$ |              | 0.00       |       | 0.20                 |       | 0.50                 |       |
| Zero-Inflated          | mean         | 2.40       | 2.40  | 2.40                 | 2.91  | 2.40                 | 3.65  |
|                        | variance     | 6.33       | 6.32  | 6.32                 | 8.65  | 6.30                 | 12.70 |
|                        | ratio        | 2.64       | 2.63  | 2.63                 | 2.98  | 2.63                 | 3.48  |
| $\Delta\bar{x}/sd_c =$ |              | −0.00      |       | 0.20                 |       | 0.50                 |       |
| Truncated              | mean         | 3.22       | 3.22  | 3.21                 | 3.43  | 3.22                 | 3.75  |
|                        | variance     | 1.12       | 1.12  | 1.13                 | 0.88  | 1.12                 | 0.42  |
|                        | ratio        | 0.35       | 0.35  | 0.35                 | 0.26  | 0.35                 | 0.11  |
| $\Delta\bar{x}/sd_c =$ |              | 0.00       |       | 0.20                 |       | 0.50                 |       |

Each simulated dataset was analyzed using four methods. First, a standard Poisson regression was fitted to the raw counts. Second, the proposed `poisrank` transformation was applied, and the transformed counts were analyzed with Poisson regression. Third, a quasi-Poisson regression was fitted to the raw counts in order to adjust for dispersion while retaining the same mean structure. Fourth, a `normrank` transformation was applied, and the transformed values were analyzed using linear regression. These comparison methods were chosen to distinguish the performance of the proposed transformation from both the uncorrected Poisson model, an established alternative for overdispersed data, and an approach that has been shown useful with regression assuming normality.

Performance was assessed using the proportion of replications rejecting the null hypothesis at the 5% level. When the true effect size is zero, this proportion estimates the Type I error rate and indicates whether a method is too liberal or too conservative. When the true effect size is greater than zero, the same quantity reflects empirical rejection performance and can be interpreted as statistical power. A useful method should therefore maintain rejection rates near the nominal level under the null while retaining strong rejection rates when true effects are present. The supplementary materials detail procedures used to choose default transformation settings, and the Anderson-Darling distance from

a Poisson distribution was also examined as a measure of how closely the transformed marginal distribution resembled a Poisson distribution.

### 3.2. Findings

Table 3 summarizes the results. When the data were genuinely Poisson distributed, the ordinary Poisson model performed as expected. Under the null condition, its rejection rate was close to the nominal 5% level, and rejection rates increased as the simulated group difference increased. Quasi-Poisson and normrank-based analyses behaved similarly in this setting, indicating that these alternatives did not substantially distort inference when the Poisson assumptions were satisfied. The poisrank approach, however, was somewhat more conservative, with slightly lower rejection rates across effect sizes. Thus, even in the best-case scenario for Poisson regression, the transformation did not provide a performance advantage and instead reduced the power of the comparison.

**Table 3.** Proportion of rejected null hypothesis using 5%.  $\lambda = 4$  for control,  $n = 100$ . There were 10,000 replications for each. Dispersion estimate is the mean for the three control conditions and the one experimental group with an effect size of zero (so is the same as the control conditions).

| Distribution     | Dispersion | Effect Size | Poisraw | Poisrank | Quasipois | Normrank |
|------------------|------------|-------------|---------|----------|-----------|----------|
| Poisson          | 1.000      | 0.0         | 0.049   | 0.039    | 0.048     | 0.048    |
|                  |            | 0.2         | 0.168   | 0.147    | 0.166     | 0.164    |
|                  |            | 0.5         | 0.657   | 0.616    | 0.650     | 0.639    |
| Extra-Poisson    | 3.007      | 0.0         | 0.254   | 0.038    | 0.047     | 0.047    |
|                  |            | 0.2         | 0.453   | 0.123    | 0.151     | 0.136    |
|                  |            | 0.5         | 0.863   | 0.472    | 0.561     | 0.510    |
| Rounded Poisson  | 1.173      | 0.0         | 0.071   | 0.036    | 0.051     | 0.052    |
|                  |            | 0.2         | 0.202   | 0.097    | 0.151     | 0.134    |
|                  |            | 0.5         | 0.705   | 0.382    | 0.569     | 0.433    |
| Gaps Poisson     | 1.005      | 0.0         | 0.047   | 0.038    | 0.045     | 0.045    |
|                  |            | 0.2         | 0.165   | 0.140    | 0.163     | 0.158    |
|                  |            | 0.5         | 0.661   | 0.617    | 0.653     | 0.643    |
| Zero-inf Poisson | 2.632      | 0.0         | 0.052   | 0.001    | 0.001     | 0.000    |
|                  |            | 0.2         | 0.336   | 0.021    | 0.039     | 0.017    |
|                  |            | 0.5         | 0.950   | 0.312    | 0.563     | 0.433    |
| Truncated        | 0.349      | 0.0         | 0.001   | 0.007    | 0.049     | 0.052    |
|                  |            | 0.2         | 0.006   | 0.051    | 0.180     | 0.189    |
|                  |            | 0.5         | 0.137   | 0.556    | 0.854     | 0.876    |

A different pattern emerged with the overdispersion distribution. In this condition, Poisson regression applied to the raw counts rejected the null far too often, including under the zero-effect condition, showing the expected inflation in false-positive findings when overdispersion is ignored. Quasi-Poisson regression substantially corrected this problem and yielded rejection rates much closer to those observed in the true Poisson benchmark. The poisrank transformation also moved results in the correct direction, but it did so more aggressively, producing noticeably lower rejection rates than quasi-Poisson when true effects were present. In practical terms, the transformation reduced the liberal behavior of the misspecified Poisson model but often at the cost of lower power. The normrank transformation with OLS regression produced similar rejection rates to the quasi-Poisson and poisrank procedures.

The rounded-count condition produced a similar, though less extreme, pattern. Because rounding created heaping and modest overdispersion, the raw Poisson model again rejected somewhat more often than would be desirable. Quasi-Poisson regression remained

closer to the expected rejection behavior, whereas the *poisrank* approach again yielded lower rejection rates than the comparison methods. This suggests that although the transformation can partially accommodate non-Poisson features created by heaping, it does not recover the same level of inferential efficiency as the better known alternatives. The quasi-Poisson and *normrank* procedures produced similar rejection rates, with the *normrank* approach producing lower values.

The condition with gaps in the count distribution differed from the overdispersed and rounded conditions. Despite its irregular shape, it had close to equidispersion. In this setting, Poisson regression, quasi-Poisson regression, and *normrank*-based analysis performed similarly to their behavior under the true Poisson condition. The *poisrank* approach again showed lower rejection rates. This result suggests that a nonstandard distributional shape by itself does not necessarily create the same inferential problems as overdispersion, and that the proposed transformation offers little advantage when the mean and variance remain approximately equal.

The results for the zero-inflated distribution were interesting. Under the null effect size, the Poisson regression on the raw scores produced approximately a 5% rejection rate while the remaining methods were extremely conservative. When a small or medium effect was present, the Poisson regression on the raw scores displayed the greatest power. The *poisrank*, quasi-Poisson, and *normrank* performed similarly and were substantially underpowered when a small or medium effect was present. Therefore, when dispersion departed from unity due to overdispersion, a Poisson regression on the raw scores tended to be most sensitive but poorly calibrated depending on the type of misspecification. By contrast, *poisrank*, quasi-Poisson, and *normrank* were more conservative, with quasi-Poisson maintaining the highest power.

Finally, for the truncated distribution, the only underdispersed distribution considered, the Poisson regression on the raw counts failed to reject the null hypothesis at a higher rate than the other approaches (0.1% when the null hypothesis was true). The *poisrank* approach also produced a low rejection rate. The quasi-Poisson and *normrank* approaches produced similar rejection rates for the three effect sizes. The rejection rate when the null hypothesis was true was approximately the nominal 5%.

#### 4. Discussion

This study examined whether a rank-Poisson transformation can serve as a useful general-purpose strategy when count outcomes depart from the Poisson distribution. The results support a cautious conclusion. Although the transformation often made the observed counts more similar to a Poisson distribution, it did not generally improve inferential performance relative to a standard alternative when data are over- and underdispersed. In both the empirical illustration and the simulation study, the proposed method moved results away from the overly liberal (i.e., standard errors too small) results from Poisson regression under misspecification, but it typically did not outperform quasi-Poisson regression and often yielded lower rejection rates when true effects were present. Thus, the main contribution of the paper is not to recommend *poisrank* as a new default method but to clarify both its potential and its limits. It is important to stress that the quasi-Poisson regression performed well both with underdispersed and overdispersed data. None of the analytic approaches tested performed well when data were drawn from a zero-inflated distribution, but there are alternatives available that are suited for data like these [1,14].

The empirical citation example illustrated the practical motivation for the study. In those data, raw Poisson regression suggested a significant effect of tweeting on later citations, whereas quasi-Poisson, negative binomial, *poisrank*-transformed Poisson, and *normrank*-based analyses all yielded non-significant results. Because the citation counts

showed substantial overdispersion, this pattern strongly suggests that the raw Poisson result reflected model misspecification rather than strong evidence for an effect.

The simulation study showed that this was not an isolated case. With under and overdispersion, raw Poisson regression rejected the null hypothesis too often, whereas quasi-Poisson more closely maintained the intended rejection rates. For the overdispersed data, the poisrank transformation corrected in the same direction as the quasi-Poisson, but generally did so more conservatively and consequently displayed reduced power. With underdispersion, the poisrank method partially corrected the inflated Type I error of untransformed Poisson regression, but improvement was modest. In short, the poisrank procedure improved the performance compared with the untransformed Poisson when the distribution shows extra-Poisson variation, but it does not perform as well as the quasi-Poisson approach. The simplicity of the quasi-Poisson approach is also an advantage as it just requires multiplying the standard errors by the dispersion coefficient.

In summary, while the normrank procedure has been shown to work well in a variety of settings [3,7], it is important to explore whether similar approaches are worth recommending in other settings. The poisrank approach had advantages compared with the Poisson regression on under and overdispersed data, but it did not perform better than the quasi-Poisson regression. As such, we do not recommend using the rank Poisson approach.

#### 4.1. Limitations

The results highlight important limitations of rank-based transformations for count outcomes. Because count variables are discrete, ties are common, and ties remain a central challenge for any ranking procedure. In the present setting, different original counts can map onto the same transformed value, so the transformation is not invertible in a practically meaningful sense. More broadly, ranking reduces the role of metric spacing and therefore changes the interpretation of regression coefficients. After a ranking transformation, coefficients are no longer naturally interpreted on the original count scale, which limits the method's usefulness when substantive interpretation of effect size is important. These limitations do not make the procedure invalid, but they do constrain the contexts in which it is likely to be attractive.

Furthermore, the present findings suggest that making a variable look more Poisson-like does not necessarily improve statistical inference compared with other alternatives. The poisrank transformation often reduced the discrepancy between the observed marginal distribution and a Poisson distribution, and produced better performance than using a Poisson regression on the untransformed data, but it did not improve upon the quasi-Poisson regression. A likely reason is that the inferential problem created by non-Poisson count data is not solely a matter of marginal shape. Overdispersion, heaping, and other irregularities affect the variance structure of the data, and model-based approaches such as quasi-Poisson directly target that variance misspecification. By contrast, a rank-based transformation changes the outcome scale and redistributes values without explicitly modeling the mechanism that produced the non-Poisson behavior. As a result, the transformed outcome may appear more compatible with the Poisson distribution while still failing to preserve the information needed for efficient testing.

#### 4.2. Future Directions

For applied researchers, the practical implications are fairly straightforward. When under and overdispersion are present, quasi-Poisson and negative binomial models remain the approaches that should be used because they address the inferential consequences of count-data misspecification more directly. When excess zeros, censoring, or other structural features are theoretically meaningful, specialized models tailored to those mechanisms

are preferable to a generic transformation. The present findings therefore do not support routine use of `poisrank` as a replacement for established count-data methods. At most, the transformation may be useful as a secondary robustness or sensitivity analysis in settings where the source of non-Poisson behavior is unclear and the interpretability of the original count scale is not the main priority.

**Supplementary Materials:** The aforementioned supporting information can be downloaded at <https://github.com/dbrookswr/poisrank> (accessed on 13 May 2026).

**Author Contributions:** Conceptualization, D.B.W.; Methodology, D.B.W. and S.N.S.; Software, D.B.W.; Formal analysis, D.B.W.; Writing—original draft, D.B.W.; Writing—review and editing, D.B.W. and S.N.S. All authors have read and agreed to the published version of the manuscript.

**Funding:** This research received no external funding.

**Data Availability Statement:** The data and code that support the findings of this study are available at <https://github.com/dbrookswr/poisrank> (accessed on 13 May 2026).

**Conflicts of Interest:** The authors declare no conflicts of interest.

## References

1. Cameron, A.C.; Trivedi, P.K. *Regression Analysis of Count Data*, 2nd ed.; Cambridge University Press: Cambridge, UK, 2013.
2. Hilbe, J.M. *Negative Binomial Regression*, 2nd ed.; Cambridge University Press: Cambridge, UK, 2011.
3. van der Waerden, B.L. Order tests for the two-sample problem and their power. *Indag. Math. (Proc.)* **1952**, *55*, 453–458. [\[CrossRef\]](#)
4. van der Waerden, B.L. Order tests for the two-sample problem (second communication). *Indag. Math. (Proc.)* **1953**, *56*, 303–310. [\[CrossRef\]](#)
5. van der Waerden, B.L. Order tests for the two-sample problem (third communication). *Indag. Math. (Proc.)* **1953**, *56*, 311–316. [\[CrossRef\]](#)
6. Beasley, T.M.; Erickson, S.; Allison, D.B. Rank-based inverse normal transformations are increasingly used, but are they merited? *Behav. Genet.* **2009**, *39*, 580–595. [\[CrossRef\]](#) [\[PubMed\]](#)
7. Wright, D.B. Normrank correlations for testing associations and for use in latent variable models. *Open Educ. Stud.* **2024**, *6*, 20240003. [\[CrossRef\]](#)
8. Wilcox, R.R. *Introduction to Robust Estimation and Hypothesis Testing*, 4th ed.; Academic Press: San Diego, CA, USA, 2017.
9. Cox, S.; West, S.G.; Aiken, L.S. The analysis of count data: A gentle introduction to poisson regression and its alternatives. *J. Personal. Assess.* **2009**, *91*, 121–136. [\[CrossRef\]](#) [\[PubMed\]](#)
10. Branch, T.A.; Côté, I.M.; David, S.R.; Drew, J.A.; LaRue, M.; Márquez, M.C.; Parsons, E.C.M.; Rabaiotti, D.; Shiffman, D.; Steen, D.A.; et al. Controlled experiment finds no detectable citation bump from Twitter promotion. *PLoS ONE* **2024**, *19*, e0292201. [\[CrossRef\]](#) [\[PubMed\]](#)
11. Wright, D.B. *Statistics for Education and Social Science Using R*; Springer-Nature: Berlin/Heidelberg, Germany, 2027, *in press*.
12. Hastie, T.J.; Pregibon, D. Generalized linear models. In *Statistical Models in S*; Chambers, J.M., Hastie, T.J., Eds.; Wandsworth & Brooks/Cole: Pacific Grove, CA, USA, 1992; pp. 195–247.
13. Venables, W.N.; Ripley, B.D. *Modern Applied Statistics with S*, 4th ed.; Springer: New York, NY, USA, 2002.
14. Wright, D.B.; Strubler, K.A.; Vallano, J.P. Statistical techniques for juror and jury research. *Leg. Criminol. Psychol.* **2011**, *16*, 90–125. [\[CrossRef\]](#)

**Disclaimer/Publisher’s Note:** The statements, opinions and data contained in all publications are solely those of the individual author(s) and contributor(s) and not of MDPI and/or the editor(s). MDPI and/or the editor(s) disclaim responsibility for any injury to people or property resulting from any ideas, methods, instructions or products referred to in the content.