

REPORT DOCUMENTATION PAGE

AD-A266 607



Form Approved

OMB No. 0704 0188

of information is estimated to average 1 hour per response, including the time for reviewing instructions, searching existing data sources, gathering the necessary data, reviewing the collection of information, Send comments regarding this burden estimate or any other aspect of this collection of information, including suggestions for reducing this burden, to Washington Headquarters Services, Directorate for Information Operations and Reports, 1215 Jefferson Davis Highway, Suite 1204, Arlington, VA 22202-4302, and to the Office of Management and Budget, Paperwork Reduction Project (0704-0188), Washington, DC 20503.

(blank)

2. REPORT DATE

20 April 1993

3. REPORT TYPE AND DATES COVERED

Final 15 May 66 - 14 Sep 89

5. FUNDING NUMBERS

DAA03-86-G-0084

4. TITLE AND SUBTITLE

Nonparametric Statistics: Weighted Rankings Analysis

6. AUTHOR(S)

Ibrahim A. Salama

7. PERFORMING ORGANIZATION NAME(S) AND ADDRESS(ES)

North Carolina Central University
Durham, NC 277078. PERFORMING ORGANIZATION
REPORT NUMBER

9. SPONSORING/MONITORING AGENCY NAME(S) AND ADDRESS(ES)

U. S. Army Research Office
P. O. Box 12221
RTP, NC 2770910. SPONSORING/MONITORING
AGENCY REPORT NUMBER

ARO 23369.6-MA

11. SUPPLEMENTARY NOTES

DTIC
ELECTE
S JUL 09 1993 B D

12a. DISTRIBUTION/AVAILABILITY STATEMENT

DISTRIBUTION STATEMENT A

Approved for public release
Distribution Unlimited

12b. DISTRIBUTION CODE

13. ABSTRACT (Maximum 200 words)

We examine the credibility hypothesis associated with the method of Weighted Rankings Analysis. We consider a multivariate test for homogeneity based on all nearest neighbors. We extend exact tables for Spearman's footrule by using a Markov chain structure over the permutation group of $1, \dots, n$. We prove the asymptotic permutational normality of certain weighted measures of correlation.

98 7 08 08 4

93-15502



14. SUBJECT TERMS

Weighted Rankings Analysis, Weighted Rank Correlation, Testing Equality of Two Distributions based on nearest neighbors

15. NUMBER OF PAGES

16

16. PRICE CODE

17. SECURITY CLASSIFICATION
OF REPORT

Unclassified

18. SECURITY CLASSIFICATION
OF THIS PAGE

Unclassified

19. SECURITY CLASSIFICATION
OF ABSTRACT

Unclassified

20. LIMITATION OF ABSTRACT

SAR

Research Grant on Nonparametric Statistics

WEIGHTED RANKINGS ANALYSIS

FINAL REPORT

by

Ibrahim A. Salama

April 20, 1993

US ARMY RESEARCH OFFICE

Grant Number: DAALO-86-G-0084

NORTH CAROLINA CENTRAL UNIVERSITY

This report contains a summary of the main results obtained by research supported by this grant.

1. Justification of the basis for weighted rankings analysis

The method of weighted rankings analysis was introduced in Quade (1972, 1979). A brief description is as follows: Let X_{ij} be the yield of the j -th treatment in the i -th block, for $i = 1, \dots, n$ and $j = 1, \dots, m$ and let R_{ij} be the within-block rank of X_{ij} . Let D_i be a measure of apparent variability for the i -th block, Q_i be the corresponding rank. Let $t_1, \dots, t_m$ be any constants, and let $s_1, \dots, s_n$ be constants such that $0 \leq s_1 \leq \dots \leq s_n$. Then a statistic of the form

$$C_n = \frac{\sum_{j=1}^m \left\{ \sum_{i=1}^n s_i Q_i t_{R_{ij}} \right\}^2}{\sum_{i=1}^n s_i^2 \sum_{j=1}^m t_j^2}$$

may be used for testing the hypothesis of no treatment effects.

The method of weighted rankings is based on the idea that if some blocks appear more variable than others, then they are perhaps better referred to as more discriminating. Hence it seems intuitively reasonable that these blocks receive greater weight in the analysis.

In this research we examine this credibility hypothesis and provide objective evidence showing how the more variable blocks do indeed better reflect any true treatment effects. The n -block two treatment exponential case was considered, and the foundation is established to extend the result to the more general cases (see Appendix I).

2. A nonparametric multivariate test for homogeneity based on all nearest neighbors

Let $\{X_{ij}; i = 1, 2; j = 1, \dots, n_i\}$ be two independent random samples of observations in the Euclidean space R^d , where X_{ij} has distribution function F_i . The

For	
1	☒
d	☐
on	☐
per	
form 50	
on/	
ity Codes	
and/or	
total	

problem under consideration is to test the hypothesis $H_0 : F_1 \equiv F_2$, against the completely general alternative $H_a : F_1 \neq F_2$.

Let $R(i,j;i',j')$ be the rank of observation $X_{i',j'}$ with respect to nearness to X_{ij} ; we assume that there will be no ties. Then we define $X_{i',j'}$ as "the" k -th nearest neighbor of X_{ij} if $R(i,j;i',j') = k$, and as "a" k -nearest neighbor if $R(i,j;i',j') \leq k$. Interest in statistical procedures based on such nearest neighbors has grown as high-speed computers have made the application of these techniques practicable, since the idea of making inferences about an object based on nearby objects appears to be a fundamental mechanism of human perception. A review of early work using nearest-neighbor approaches to our problem may be found in Schilling (1986).

Schilling's own approach is as follows. Let

$$\begin{aligned} I_{ij}(r) &= I\{r\text{-th nearest neighbor of } X_{ij} \text{ is in sample } i\} \\ &\text{for } i = 1, 2; j = 1, \dots, n_i; \\ &r = 1, \dots, N-1 \ (N = n_1 + n_2), \end{aligned}$$

where $I\{E\}$ is the indicator function of the event E . Count the number of k -nearest neighbors to X_{ij} which are in the i -th (same) sample, viz.

$$C_{ijk} = \sum_{r=1}^k I_{ij}(r).$$

Summing these counts over all observations yields what may be called the "Schilling total", of order k :

$$T_k = \sum_{ij} C_{ijk}.$$

His test statistic for our problem is then

$$S_k = T_k / Nk,$$

which is "the proportion of all k -nearest neighbor comparisons in which a point and its neighbor are members of the same sample" (1986, p. 800), or "Schilling

proportion" of order k . One would expect S_k to have a larger value under H_a than under H_0 because of a lack of complete mixing of the two samples when the parent distributions are not identical; hence large values of S_k may be considered significant. Schilling shows that the asymptotic distribution of S_k under H_0 is normal, for every positive integer k , and that the test which rejects for large values of S_k is consistent against the general alternative H_a .

Schilling's work suggests that the choice of order is not of great importance; nevertheless, it is arbitrary, and he gives no guidance for choosing it. To resolve this issue we propose to take as test statistic the sum of the Schilling totals, which is equivalent to a certain weighted average of the Schilling proportions:

$$W = \sum T_k = N \sum k S_k.$$

This statistic may have intuitive appeal in that it is equivalent to a sum of N Wilcoxon rank sums, and is also a linear combination of two U -statistics, as we have shown. We study some of the exact properties of W . For the asymptotic properties, we note that W is equivalent to a sum of Schilling proportions, each of which is asymptotically normal; but asymptotic normality of their joint distribution has not been shown. Similarly, we have shown that W is a sum of Wilcoxon statistics, each of which is asymptotically normal; but again asymptotic normality of their joint distribution has not been shown. We might also attempt to show this by means of U -statistics theory, but have not yet been able to work out the necessary conditions on the variance. Nevertheless, it seems reasonable to conjecture that W itself is asymptotically normal. Computer simulation was used to estimate the 95th percentile of W , then this value will be used as the critical value for a test. That is, the null hypothesis will be rejected if W is greater than this critical value. The empirical power of the test based on W was also calculated based on this critical value.

3. Exact tables for Spearman's footrule

Spearman (1904) based the measure of rank correlation as his "footrule" on

$$D(\sigma, \pi) = \sum_{i=1}^n |\sigma(i) - \pi(i)|,$$

where σ and π are any two permutations of S_n , the set of all permutations of the first n integers. Ury and Kleinecke (1979) provided the exact commulative distribution of D for $n = 2(1)10$. Franklin (1988) has extended the tabulation to $n = 11(1)18$. The difficulty in extending these tables is due to the exponential growth rate of time needed to generate all permutations of S_n . By relating D to a Markov chain on S_n , we extend the exact tabulation to $n = 19(1)40$. We also investigate the adequacy of approximation to the normal distribution.

4. The asymptotic permutational normality of certain weighted measures of correlation

Let be given observation (X_i, Y_i) , $i = 1, \dots, n$, on the continuous bivariate random variable (X, Y) . Without loss of generality we may relabel the X 's to have ranks $1, \dots, n$: then let the corresponding ranks of the Y 's be $R_1, \dots, R_n$. Define

$$T_j = \sum_{i=1}^j I\{R_i \leq j\}$$

where $I\{E\}$ is the indicator function of the event E , and let

$$T = \sum_{j=1}^n w_j T_j$$

where $w_1, \dots, w_n$ is a sequence of weights. Then T may be regarded as a weighted measure of rank correlation between the X 's and Y 's. This notion was introduced in Salama and Quade (1982), although they discussed only the special case where $w_j = 1/j$.

In this research, a martingale approach is used to establish the asymptotic normality of a class of weighted correlation statistics.

5. Spherical uniformity and the Cauchy distribution

To motivate the characterization problems (to be posed), we consider the following. Let Y be a stochastic (m -) vector for which $EY = X\beta$, where X is a known $m \times n$ (design) matrix and β is an n -vector of unknown regression parameters. Then $\hat{\beta}$, the least squares estimator (LSE) of β , satisfies the equation: $(X'X)\hat{\beta} = X'Y$. When $X'X$ is not of full rank, additional constraints are imposed on β in order that the LSE may be defined uniquely. In the absence of a physically natural form of such a (linear) constraint, often this choice is made rather arbitrarily. In some random effects or mixed models, when $\text{Rank}(X'X) = n-1$, an additional random constraint is taken as $c'\beta = 0$, where $\|c\| = 1$. In robust regression (cf. Huber (1981, p. 170)), when n is large, often c is chosen at random with respect to the invariant measure on the unit sphere $\|c\| = 1$. In the context of asymptotic normality of estimators, Huber (1981) assumed that $\text{Rank}(X'X) = n$, and we shall see that a different picture emerges when $\text{Rank}(X'X)$ is less than n . Suppose that $\text{Rank}(X'X) = n-1$ and let $c'\beta = 0$ be an additional constraint such that $\|c\| = 1$. Let $\hat{\beta}_c$ be the solution satisfying $(X'X)\hat{\beta}_c = X'Y$ and $c'\hat{\beta}_c = 0$. Let v be the eigenvector corresponding to the null eigenvalue of $X'X$. Then $\hat{\beta}_c$ may be written as

$$\hat{\beta}_c = \hat{\beta}_0 + ((c'\hat{\beta}_0)/(c'v))v, \quad (1)$$

where $\hat{\beta}_0$ is any particular solution of $(X'X)\beta = X'Y$ (and is therefore a random variable). Thus, given $\hat{\beta}_0$, we may write

$$(\hat{\beta}_c - \hat{\beta}_0) = v[(c'\hat{\beta}_0)/(c'v)], \quad (2)$$

and its conditional distribution, given $\hat{\beta}_0$, is generated exclusively by the uniform distribution of c on the spherical surface. Note that v is nonstochastic (as $X'X$ is), so

that the key factor on the right-hand side of (2) is the ratio $(c' \hat{\beta}_0)/(c' v)$, where $\hat{\beta}_0$ is held fixed. This leads us to the following problem.

For two arbitrary points b and v in R^n , let $L(t) = b + tv$, $t \in R$ be a given line. Let $S^{n-1} = \{c: c'c = 1\}$ be the unit sphere, and for every $c \in S^{n-1}$, let $P(c)$ be the hyperplane defined by $Pc = 0$, $P \in R^n$. Let $L(t_c) = b + (t_c)v$ be the point at which L and $P(c)$ intersect, for $c \in S^{n-1}$. Define then

$$X_n = X_n(c) = \text{sign}(t_c) \|L(t_c) - L(0)\| = \text{sign}(t_c) |t_c| \|v\|, \quad c \in S^{n-1}. \quad (3)$$

Assuming that c has a uniform distribution on the sphere S^{n-1} , we show that for every $n \geq 2$, X_n has the same distribution as $aX + d$, where a and d are real numbers and X has the standard Cauchy distribution. Some other related characterization results are also considered in the same view.

6. Topological entropy of countable Markov chains

We consider a symbolic dynamical system (X, σ) on a countable state space. We introduce a kind of topological entropy for such systems, denoted h^* , which coincides with usual topological entropy when X is compact. We use a pictorial approach, to classify a graph Γ (or a chain) as transient, null recurrent, or positive recurrent. We show that given $0 \leq \alpha \leq \beta \leq \infty$, there is a chain whose h^* entropy is β and where Gurevic entropy is α . We compute the topological entropies of some classes of chains, including larger chains built up from smaller ones by a new operation which we call the Cartesian sum.

List of publications under the grant

Salama, I.A. (1988). "Topological entropy and recurrence of countable chains."

Pacific Journal of Mathematics, Vol. 134, No. 2, 325–341.

Barakat, A.S. (1989). "A Nonparametric test for homogeneity based on all nearest neighbors." University of North Carolina Institute of Statistics Mimeo Series No. 1866T. (Ph.D. Dissertation)

Salama, I.A. and Quade, D. (1990). "A note on Spearman's footrule."

Communications in Statistics, Simulation and Computation, Vol. 19, No. 2, 591–601.

Salama, I.A. and Quade, D. (1990). "Remarks on justifying the intuitive basis for the method of weighted rankings." ASA Proceedings, Business and Economic Section, 369–373. (Under revision for journal publication.)

Barakat, A.S., Quade, D., and Salama, I. (1990). "A nonparametric test for homogeneity based on all nearest neighbors." As a Proceedings, Social Statistics Section, 221–225. (Under revision for Journal of Nonparametric Statistics.)

Salama, I.A. and Quade, D. (1990). "On the asymptotic permutational normality of certain weighted measures of correlation." *Statistics and Probability Letters*, Vol. 10, 125–133.

Salama, I.A. and Sen, P.K. (1992). "Spherical uniformity and some characterizations of the Cauchy distribution." *Journal of Multivariate Analysis*, Vol. 41, No. 2, 212–219.

Participating Scientific Personnel

- 1. Ibrahim A. Salama**
- 2. Dana Quade**
- 3. P.K. Sen**
- 4. Azza R. Karmous. (Received Ph.D., Department of Biostatistics, University of North Carolina at Chapel Hill, 1986.)**
- 5. Ali S. Barakat. (Received Ph.D., Department of Biostatistics, University of North Carolina at Chapel Hill, 1989.)**

References

- Franklin, L.A. (1988) "Exact tables of Spearman's footrule for $N = 11(1)18$ with estimate of convergence and errors for the normal approximation." *Statistics and Probability Letters* 6: 399—406.
- Huber, P.J. (1981) *Rubust Statistics*. Wiley, New York.
- Quade, D. (1972) Analyzing randomized blocks by weighted rankings. Report SW 18/72, Mathematical Center Amsterdam.
- Quade, D. (1979) Using weighted rankings in the analysis of complete blocks with additive block effects. *Journal of the American Statistical Association*, 74: 680—783.
- Salama, I.A. and Quade, D. (1982) "A nonparametric comparison of two multiple regressions by means of a weighted measure of correlation." *Communications in Statistics — Theory and Methods*, 11: 1185—1195.
- Schilling, M.F. (1986) "Multivariate two—sample tests based on nearest neighbors." *JASA*, 81, 799—806.
- Spearman, C. (1904) "The proof and measurement of association between two things." *American Journal of Psychology*, 15: 72—101.
- Ury, H.K. and Kleinecke, D.C. (1979) "Tables of the distribution of Spearman's footrule." *Applied Statistics* 78: 271—275.

APPENDIX I

REMARKS ON JUSTIFYING THE INTUITIVE BASIS FOR THE METHOD OF WEIGHTED RANKINGS

Ibrahim A. Salama, School of Business, North Carolina Central University, Durham
Dana Quade, Department of Biostatistics, University of North Carolina at Chapel Hill

1. INTRODUCTION

Let X_{ij} be the observation of the j -th of m treatments in the i -th of n complete blocks, and consider the hypothesis of no treatment effects, specifically

H_0 : $X_{i1}, \dots, X_{im}$ are interchangeable for each i .

(By definition random variables are *interchangeable* if their joint distribution function is invariant under permutation; this implies that they have identical marginal distributions and equal — but not necessarily zero — correlations.) We assume throughout:

(I) Independent Blocks

For $i=1, \dots, n$, the random vectors $X_i = (X_{i1}, \dots, X_{im})'$ (the blocks), are mutually independent.

To simplify the exposition it is convenient also to assume:

(II) No Within-Blocks Ties

$P\{X_{ij} = X_{ij'}\} = 0$ for $j \neq j'$

Thus with probability 1 there will be no ties within blocks.

The alternatives under consideration can be fairly general: however, we have particularly in mind that there may be *additive treatment effects*, as follows:

UNORDERED CASE

$H_a[u]$: There exist quantities $\tau_1, \dots, \tau_m$ (treatment effects) not all equal to zero, such that for $i=1, \dots, n$, $X_{i1} - \tau_1, \dots, X_{im} - \tau_m$ are interchangeable.

ORDERED CASE

$H_a[o]$: The quantities $\tau_1, \dots, \tau_m$ (as above) satisfy $\tau_1 \leq \dots \leq \tau_m$, with $\tau_1 \neq \tau_m$.

Standard nonparametric procedures for attacking this problem are based on within-block

rankings: for example, the tests of Friedman (1937) and Brown and Mood (1951) for $H_a[u]$; and Lyerly (1952), Page (1963) and Jonckheere (1954) for $H_a[o]$. Let $C_{ii'}$ be some measure of rank correlation between block i and block i' . Then

$$\bar{C} = \sum_{i < i'} C_{ii'} / \binom{n}{2}$$

is the *average internal rank correlation*, and we reject H_0 in favor of $H_a[u]$ if $\bar{C}$ is too big. (This is equivalent to Friedman's test if Spearman correlation is used.) Similarly let C_i be the rank correlation between block i and the ordering given by the alternative. Then

$$\bar{C} = \sum_i C_i / n$$

is the *average external rank correlation*, and we reject H_0 in favor of $H_a[o]$ if $\bar{C}$ is too big. (This is equivalent to Page's test if Spearman correlation is used.)

The standard procedure is parametric two-way analysis of variance, but this adds two assumptions:

(III) Additive Block Effects

There exist quantities $\beta_1, \dots, \beta_n$ (block effects) such that the random vectors $(X_{i1} - \beta_i, \dots, X_{im} - \beta_i)'$ are identically distributed.

(IV) Normality

The X_{ij} 's are [jointly] normal.

By Assumption III, comparisons of observations are possible between blocks as well as within, so procedures which use only within-block comparisons waste information. A method of *weighted* within-block rankings, which makes use of Assumption (III) without requiring Assumption (IV), has been introduced by Quade (1972, 1979). The idea behind this method is that blocks in which the observations are more distinct are more likely to reflect any underlying

true ordering of the treatment effects. These blocks, which may be referred to as more *credible*, should receive greater weight in the analysis. (In practice, credibility is measured by *apparent* variability; but note that by Assumption III the *true* variability is the same in all blocks.)

To determine the weight for the i -th block, use some location-free statistic $D_i = D(X_{i1}, \dots, X_{im})$ which measures the credibility of the block with respect to treatment ordering, and let Q_i be the rank of D_i among $D_1, \dots, D_n$. Again for simplicity of exposition, make the (unessential) assumption:

(V) No Between-Block Ties

$$P\{D_i = D_{i'}\} = 0 \text{ for } i \neq i',$$

This assures that there will be no ties in the ranking of the blocks. Let $0 \leq b_1 \leq \dots \leq b_n$, with $0 \neq b_n$, be a fixed set of block scores; and weight the i -th block proportionally to b_{Q_i} .

Then in testing against the unordered alternative we use the *weighted average internal rank correlation*

$$\bar{W} = \frac{\sum_{i \neq i'} b_{Q_i} b_{Q_{i'}} C_{ii'}}{((\sum b_i)^2 - \sum b_i^2)}.$$

Against the ordered alternative we use the *weighted average external rank correlation*

$$\bar{W} = \frac{\sum_{i=1}^n b_{Q_i} C_i}{\sum_{i=1}^n b_i},$$

The purpose of this paper is to examine the notion on which these weighted rank correlation coefficients are based, which we call the *credibility hypothesis*.

2. THE CASE OF TWO TREATMENTS

With $m=2$ treatments, suppose the true ranking is $(1, 2)$: i.e., $\tau_1 < \tau_2$. Let R_{ij} be the rank of X_{ij} within the i -th block, and consider

$$P\{(R_{i1}, R_{i2}) = (1, 2) | Q_i = k\} = \phi_k \text{ (say).}$$

The intuitive notion is that

$$\phi_1 \leq \phi_2 \leq \dots \leq \phi_n,$$

where of course

$$\sum \phi_i / n = \phi \text{ (say)}$$

is the unconditional probability $P\{(R_{i1}, R_{i2}) = (1, 2)\}$. With two treatments we may let $D_i = X_{2i} - X_{1i}$. Then

$$(R_{i1}, R_{i2}) = (1, 2) \Leftrightarrow D_i > 0$$

and

$$Q_i = k \Leftrightarrow |D_i| \text{ has rank } k \text{ among } |D_1|, \dots, |D_n|.$$

In the very special case of 2 treatments and 2 blocks, let us define

$$\phi = P\{D_1 > 0\}$$

and

$$\theta = P\{D_1 + D_2 > 0\}.$$

Then

$$\phi_1 = P\{D_1 > 0 | |D_1| < |D_2|\} = \phi - (\theta - \phi) = 2\phi - \theta$$

$$\phi_2 = P\{D_1 > 0 | |D_1| > |D_2|\} = \phi + (\theta - \phi) = \theta$$

and the credibility hypothesis holds if $\theta > \phi$. Suppose $(X_{i1} - \tau_1, X_{i2} - \tau_2)$ are IID normal with variance σ^2 ; then it is easily seen that

$$\phi = \Phi\left(\frac{\tau_2 - \tau_1}{\sigma\sqrt{2}}\right), \theta = \Phi\left(\frac{\tau_2 - \tau_1}{\sigma}\right)$$

where Φ is the standard normal distribution function, and $\theta > \phi$ if $\tau_1 < \tau_2$. On the other hand, suppose $(X_{i1} - \tau_1, X_{i2} - \tau_2)$ are IID Cauchy; then $\theta = \phi$, and the credibility hypothesis does not hold. Note, by the way, that with only two treatments the weighted rankings procedures are equivalent to signed-rank procedures, which are thoroughly discussed in Chapter 3 of Pratt and Gibbons (1981).

If we do not limit ourselves to additive treatment effects, we may consider the interesting special case where X and Y are exponentially distributed, with parameters (say) λ and μ , respectively. Then it is easily shown that $\phi = \lambda/(\lambda + \mu)$, $\theta = \phi^2(3 - 2\phi)$, and $\theta > \phi$ if and only if $\lambda < \mu$.

3. THE GENERAL CASE

We now turn our attention to the general case where we have m treatments and n blocks.

Write

$(R_{i1}, \dots, R_{im}) = R_i$,
and consider

$$\begin{aligned} & P\{R_i = r_i | Q_i = k\} \\ &= nP\{R_i = r_i, Q_i = k\} \text{ since } P\{Q_i = k\} = 1/n \\ &= nP\{R_i = r_i, \frac{k-1}{D_s} < D_i < \frac{n-k}{D_s}\} \\ &= nP\{R_i = r_i, D_{(k-1)} < D_i < D_{(k)}\} \end{aligned}$$

where $D_{(j)}$ is the j -th order statistic from a sample of $(n-1)$ values of D [that is, all values except D_i].

Theorem:

Let $g_{k-1,k}$ be the joint density function of $D_{(k-1)}$ and $D_{(k)}$. Then

$$\begin{aligned} & P\{R_i = r_i | Q_i = k\} \\ &= n \int_0^{\infty} \int_0^b P\{R_i = r_i, a < D_i < b\} g_{k-1,k}(a, b) da db. \end{aligned}$$

Proof:

$$\begin{aligned} & P\{R_i = r_i | Q_i = k\} \\ &= nP\{R_i = r_i, Q_i = k\} \\ &= nP\{R_i = r_i, \\ & \quad D_{i_1} < \dots < D_{i_{k-1}} < D_i < D_{i_k} < \dots < D_{i_{n-1}}\} \end{aligned}$$

where $(i_1, \dots, i_{n-1})$ is a permutation of $(1, \dots, i-1, i+1, \dots, n)$. This in turn is equal to

$$nP\{R_i = r_i, d_{(k-1)} < D_i < d_{(k)}\}$$

where $d_{(j)}$ is the j -th order statistic of a sample $(n-1)$ observations on D , and thence equal to the integral of the Theorem. QED

This form allows the possibility of actual computation. The density g is well-known given the distribution of D . And the probability expression in the integrand depends only on the m -dimensional joint distribution of the observations within a single block. However, it does involve working with distributions of order statistics from heterogeneous distributions.

In the special cases of $k=1$ and $k=n$ we have

$$P\{R_i = r_i | Q_i = 1\}$$

$$= n \int_0^{\infty} P\{R_i = r_i, 0 \leq D_i < t\} f_{(1)}(t) dt$$

and

$$\begin{aligned} & P\{R_i = r_i | Q_i = n\} \\ &= n \int_0^{\infty} P\{R_i = r_i, t \leq D_i < \infty\} f_{(n-1)}(t) dt, \end{aligned}$$

where $f_{(k)}$ is the density function of $D_{(k)}$.

4. AN EXPONENTIAL SPECIAL CASE

In this section we consider an application of the preceding theorem to a simple and analytically tractable case. We consider two treatments realized by the random variables X and Y , which we assume independent with probability density functions $f_x(x) = e^{-x}$, $f_y(y) = \lambda e^{-\lambda y}$, and $0 \leq X, Y < \infty$. For the n blocks we have independent observations $Z_i = (X_i, Y_i)$, $i=1, \dots, n$. Let $D_i = |X_i - Y_i|$, $i=1, \dots, n$. Assuming that $\lambda > 1$, the "correct" ordering of (X, Y) is given by the permutation $\sigma = (2, 1)$. [The "correct" ordering means $P(X > Y) > P(Y > X)$]. (In this notation the range of the observations in the i -th block is less than that of the $(i+1)$ -th block.)

$$\text{Since } \text{Corr}(R_{(i)}, \sigma) = \begin{cases} 1 & \text{if } R_{(i),1} = 2 \\ -1 & \text{if } R_{(i),1} = 1 \end{cases}$$

it follows that

$$\begin{aligned} E(\text{Corr}(R_{(i)}, \sigma)) &= P\{R_{(i),1} = 2\} - P\{R_{(i),1} = 1\} \\ &= P\{R_{(i),1} = 2\} - [1 - P\{R_{(i),1} = 2\}] \\ &= 2P\{R_{(i),1} = 2\} - 1. \end{aligned}$$

Thus, we compute

$$P_k = P\{R(X_i) = 2 | R(D_i) = k\}, k=1, \dots, n+1$$

and show that $\{P_k\}$ is strictly monotone increasing in k . The following lemma will be used in showing the result.

Lemma:

Let $X_{(1)}, \dots, X_{(n)}$ be the order statistics for a sample of size n from $H(x)$, $0 < x < 1$. Let $d_k = E(X_{(k+1)}) - E(X_{(k)})$. If $h(x)$ is monotone increasing (decreasing), then $\{d_k\}$ is monotone decreasing (increasing) in k .

Proof:

$E(X_{(k)})$

$$\begin{aligned} &= \int_0^1 x \frac{n!}{(k-1)!(n-k)!} H^{k-1}(x) h(x) [1-H(x)]^{n-k} dx \\ &= \sum_{i=k}^n \binom{n}{i} H^i(x) [1-H(x)]^{n-i} x \Big|_0^1 \\ &\quad - \int_0^1 \sum_{i=k}^n \binom{n}{i} H^i(x) [1-H(x)]^{n-i} dx \\ &= 1 - \int_0^1 \sum_{i=k}^n \binom{n}{i} H^i(x) [1-H(x)]^{n-i} dx, \end{aligned}$$

So

$$d_k = E(X_{(k+1)}) - E(X_{(k)})$$

$$= \int_0^1 \binom{n}{k} H^k(x) [1-H(x)]^{n-k} dx.$$

Let $y=H(x)$, so $x=H^{-1}(y)$, and set $\mu(y)=dx/dy$. Then

$$\begin{aligned} d_k &= \int_0^1 \mu(y) \binom{n}{k} y^k (1-y)^{n-k} dy \\ &= \int_0^1 \mu(y) L_k(y) dy \quad (\text{say}). \end{aligned}$$

Clearly, $L_k(0)=L_k(1)=0$, L_k is unimodal, and there exists a y^* such that $0 < y^* < 1$ with $L_k(y^*)=L_{k+1}(y^*)$. Now if $h(x)$ is monotone increasing, then $\mu(y)$ is monotone decreasing and

$$\begin{aligned} d_{k+1} &= \int_{y^*}^1 \mu(y) [L_{k+1}(y) - L_k(y)] dy \\ &\quad - \int_0^{y^*} \mu(y) [L_k(y) - L_{k+1}(y)] dy \\ &< \mu(y^*) \int_{y^*}^1 [L_{k+1}(y) - L_k(y)] dy \\ &\quad - \mu(y^*) \int_0^{y^*} [L_k(y) - L_{k+1}(y)] dy \\ &= \mu(y^*) \int_0^1 [L_{k+1}(y) - L_k(y)] dy \\ &= 0. \end{aligned}$$

$$[\text{Note } \int_0^1 L_{k+1}(y) dy = \int_0^1 L_k(y) dy = \frac{1}{n+1}.]$$

Hence $\{d_k\}$ is monotone increasing in k . The case where $h(x)$ is monotone decreasing is

similar.

Theorem:

Let $(x_1, y_1), \dots, (x_n, y_n)$ be the observations corresponding to a design with two treatments and n blocks. Assume that X and Y are independent, with density functions $f_x(x) = e^{-x}$ and $f_y(y) = \lambda e^{-\lambda y}$, $0 < X, Y < \infty$. Let $P_k = P\{R(X_i) = 2 | R(D_i) = k\}$, where $D_i = |X_i - Y_i|$. If $\lambda > 1$, then $\{P_k\}$ is monotone in k ; that is, $P_1 < P_2 < \dots < P_n$.

Proof:

Let $D(X, Y) = |X - Y|$, then

$$G(t) = P(D \leq t) = 1 - \frac{\lambda}{1+\lambda} e^{-t} - \frac{1}{1+\lambda} e^{-\lambda t},$$

$$g(t) = \frac{\lambda}{1+\lambda} (e^{-t} + e^{-\lambda t}), \quad 0 \leq t < \infty.$$

We also have

$$P\{R(X)=2, 0 \leq D < t\} = \frac{\lambda}{1+\lambda} (1 - e^{-t}),$$

$$P\{R(X)=2, t_1 \leq D < t_2\} = \frac{\lambda}{1+\lambda} (e^{-t_1} - e^{-t_2}),$$

and

$$P\{R(X)=2, t \leq D < \infty\} = \frac{\lambda}{1+\lambda} e^{-t}.$$

By the preceding theorem, for $k=2, \dots, n$ we have

$$\begin{aligned} P_k &= P\{R(X)=2 | R(D)=k\} \\ &= n \int_0^\infty \int_0^{t_2} P\{R(X)=2, t_1 \leq D < t_2\} \\ &\quad g_{(k-1,k)}(t_2, t_1) dt_1 dt_2 \\ &= n \int_0^\infty \int_0^{t_2} \left(\frac{\lambda}{1+\lambda} \right) (e^{-t_1} - e^{-t_2}) g_{(k-1,k)}(t_2, t_1) dt_1 dt_2 \\ &= \frac{n}{1+\lambda} \int_0^\infty \int_0^{t_2} (e^{-t_1} - e^{-t_2}) \frac{(n-1)!}{(k-2)!(n-1-k)!} \\ &\quad [G(T_1)]^{k-2} g(t_1) g(t_2) [1-G(t_2)]^{n-1-k} dt_1 dt_2 \\ &= \frac{n\lambda}{1+\lambda} \int_0^\infty e^{-t} g_{(k-1)}(t) dt - \frac{n}{1+\lambda} \int_0^\infty e^{-t} g_{(k)}(t) dt. \end{aligned}$$

Now,

$$\int_0^\infty e^{-t} g_{(k)}(t) dt =$$

$$\begin{aligned}
& \int_0^{\infty} e^{-t} \frac{(n-1)!}{(k-1)!(n-1-k)!} [G(t)]^{k-1} g(t) [1-G(t)]^{n-1-k} dt \\
&= \int_0^{\infty} e^{-t} \frac{(n-1)!}{(k-1)!(n-1-k)!} \left[1 - \frac{\lambda}{1+\lambda} e^{-t} - \frac{\lambda}{1+\lambda} e^{-\lambda t} \right]^{k-1} \\
&\quad * \left(\frac{\lambda}{1+\lambda} \right) (e^{-t} + e^{-\lambda t}) \left[\frac{\lambda}{1+\lambda} e^{-t} + \frac{\lambda}{1+\lambda} e^{-\lambda t} \right]^{n-1-k} dt \\
&= \int_0^{\infty} X \frac{(n-1)!}{(k-1)!(n-1-k)!} \left[1 - \frac{\lambda}{1+\lambda} X - \frac{1}{1+\lambda} X^{\lambda} \right]^{k-1} \\
&\quad * \left(\frac{\lambda}{1+\lambda} \right) (X + X^{\lambda}) \left[\frac{\lambda}{1+\lambda} X + \frac{1}{1+\lambda} X^{\lambda} \right]^{n-1-k} \left(\frac{1}{X} \right) dx.
\end{aligned}$$

(via the transformation $X = e^{-t}$)

$$\begin{aligned}
&= \int_0^1 X \frac{(n-1)!}{(k-1)!(n-1-k)!} \left[1 - \frac{\lambda}{1+\lambda} X - \frac{1}{1+\lambda} X^{\lambda} \right]^{k-1} \\
&\quad \left(\frac{\lambda}{1+\lambda} \right) (1 + X^{\lambda-1}) \left[\frac{\lambda}{1+\lambda} X + \frac{1}{1+\lambda} X^{\lambda} \right] dx \\
&= \int_0^1 X \frac{(n-1)!}{(k-1)!(n-1-k)!} H^{n-1-k}(x) h(x) [1-H(x)]^{k-1} dx,
\end{aligned}$$

where $H(x)$ is the distribution function given by

$$H(x) = \frac{\lambda}{1+\lambda} X + \frac{1}{1+\lambda} X^{\lambda}, \quad 0 \leq X \leq 1.$$

Thus we have

$$\int_0^{\infty} e^{-t} g_{(k)}(t) dt = E(X_{(n-k)}),$$

where $X_{(k)}$ is the k -th order statistic of a sample of size $n-1$ from $H(x)$. Now

$$P_k = \frac{(n+1)\lambda}{1+\lambda} [E(X_{(n-k+1)}) - E(X_{(n-k)})]$$

and for $\lambda > 1$, $h(x)$ is strictly monotone increasing, hence $P_n > P_{n-1} \dots > P_1$.

BIBLIOGRAPHY

- Brown, G.W., and Mood, A.M. (1951). On median tests for linear hypotheses, *Proceedings of the Second Berkeley Symposium on Mathematical Statistics and Probability*, Berkeley: University of California Press, 159-166.
- Friedman, M. (1937). The use of ranks to avoid the assumption of normality implicit in the analysis of variance. *Journal of the American Statistical Association*, 32:675-701.
- Jonckheere, A.R. (1954). A test of significance for the relation between m rankings and k ranked categories. *British Journal of Statistical Psychology*, 7:93-100.
- Lyerly, S.B. (1952). The average Spearman rank correlation coefficient. *Psychometrika*, 17:421-428.
- Page, E. B. (1963). Ordered hypotheses for multiple treatments: a significance test for linear ranks. *Journal of the American Statistical Association*, 58:216-230.
- Pratt, J.W. and Gibbons, J.D. (1981). *Concepts of Nonparametric Theory*, New York: Springer Verlag.
- Quade, D. (1972). Analyzing randomized blocks by weighted rankings. Report SW 18/72, Mathematical Center Amsterdam.
- Quade, D. (1979). Using weighted rankings in the analysis of complete blocks with additive block effects. *Journal of the American Statistical Association*, 74:680-783.
- Salama, I.A. and Quade, D. (1981). Using weighted rankings to test against ordered alternatives in complete blocks. *Communications in Statistics: Theory and Methods*, A10:385-399.
- Silva, C. and Quade, D. (1980). Evaluation of weighted rankings using expected significance level. *Communications in Statistics: Theory and Methods*, A9:1087-1096.