

AD-A101 733

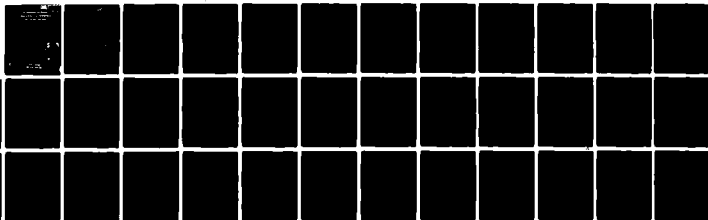
PENNSYLVANIA STATE UNIV UNIVERSITY PARK DEPT OF STAT--ETC F/G 12/1
A GEOMETRIC INTERPRETATION OF INFERENCES BASED ON RANKS IN THE --ETC(U)
JUL 81 T P HETTMANSPERGER, J W MCKEAN N00014-80-C-0741

UNCLASSIFIED

TR-36

NL

1 OF 1
AD-A101 733



END

DATE

FORMED

8-8-81

DTIC

LEVEL II

42

The Pennsylvania State University
Department of Statistics
University Park, Pennsylvania

AD A101733

DTIC FILE COPY

DTIC
ELECTE
JUL 22 1981

E



DISTRIBUTION STATEMENT A
Approved for public release
Distribution Unlimited

81 7 21 0 66

LEVEL II

(12)

DEPARTMENT OF STATISTICS

The Pennsylvania State University
University Park, PA 16802, U.S.A.

9/15/81 kept

TECHNICAL REPORTS AND PREPRINTS

Number 36: July 1981

⑥ A GEOMETRIC INTERPRETATION OF INFERENCES
BASED ON RANKS IN THE LINEAR MODEL.

⑩ Thomas P. Hettmansperger*
The Pennsylvania State University

Joseph W. McKean**

Western Michigan University

⑪
(12, 42)

⑭ 71F-36

DTIC
ELECTE
S JUL 22 1981
E

*Research partially supported by ONR contract N00014-80-C-0741

**Research partially supported by ARO contract DRXRO-MA-16064-M

Accession For	
NTIS GRA&I	☒
DTIC TAB	☐
Unannounced	☐
Justification	
By _____	
Distribution/	
Availability Codes	
Dist	Avail and/or Special
A	

1. Introduction and Summary

There are at least three separate approaches to testing hypotheses based on ranks in the linear model. Three of these methodologies, scattered through the literature, are described by McKean and Hettmansperger (1976), Sen and Puri (1977) and Adichie (1978). In addition, we will introduce a fourth approach in this paper based on a suggestion of Bickel (1976) in the context of M-estimation. All of these tests have the same approximating distribution under the null hypothesis and the same asymptotic efficiency. Other than to note the asymptotic similarities there has been no previous attempt to study the similarities and differences among these tests for small to moderate sample sizes. In particular, there have been no suggestions for users on which of these methods is the more practical. Recommendations on the implementation of these methods can be found at the end of Section 5.

In Sections 2 through 4 we provide a unified discussion of all four tests in the context of the geometry of the linear model. By considering the geometry of the statistics we can quite easily describe differences and similarities in the tests. In addition to providing a comparison of the rank tests, the geometry suggests a comparison with the classical F-test. There are three algebraically equivalent forms of the F statistic and the rank tests can be identified with these different forms. The rank tests, however, are not algebraically equivalent and this is one source of their small sample differences. For an excellent account of the geometry of the linear model, see Arnold (1981).

In Section 5 we investigate in a Monte Carlo study the small sample levels and powers of these tests along with the F-test. The study includes several designs and error distributions. On the basis of this study we conclude that some of the tests seem to be unusually sensitive to the design and to the error structure.

The new approach to testing described here is based on Bickel's (1976) idea of pseudo-observations. The pseudo-observations are constructed from rank estimates of the parameters in the linear model in such a way that the F-test calculated from these pseudo-observations is a normalized quadratic form in the rank estimates which can be used to carry out hypothesis tests. A plot of the pseudo-observations versus the data illustrates the effect which robust methods have on the observations. This plot is included with the example in Section 7.

Both rank and signed rank tests can be constructed and this may be a source of some confusion. Although they differ numerically, their asymptotic theory is the same. To avoid more notation we present only the signed rank versions. If, in the formulas for each test described in Section 3, we replace the signed rank score of the absolute value of the residual (defined under 2.8) by the centered rank score of the regular residual and replace the design matrix by the mean centered design matrix we have the corresponding rank score version of the test. In the case of rank scores the intercept parameter is handled separately. One estimate of the intercept is a location estimate computed from the residuals. McKean and Hettmansperger (1976) and Sen and Puri (1977) discuss only the rank score tests while Adichie (1978) discusses both

rank and signed rank score tests. The Monte Carlo study includes both types in order that their small sample behavior can be compared.

A final note needed for Sections 2 and 3 concerns estimability. In this paper, a function, $\lambda'\beta$, of regression parameters is estimable provided λ lies in the row space of the design matrix; see McKean and Schrader (1980) for a discussion of estimability in terms of robust estimation.

2. Estimation in the Linear Model

2.1 Notation and Assumptions

Let Y denote an $n \times 1$ vector of observations. Assume it follows the linear model

$$Y = X\beta + e \quad (2.1)$$

where X is an $n \times r$ matrix of known constants, β is an unknown $r \times 1$ vector of parameters, and e is an $n \times 1$ vector of iid random errors, symmetrically distributed about 0 with density function $f(x)$. Let Ω denote the subspace of R^n which is spanned by the columns of X . Assume the dimension of Ω is $p \leq r$. Then alternately we can write the model (2.1) as

$$Y = \theta + e, \theta \in \Omega. \quad (2.2)$$

When expectations exist $EY = \theta$; in any case, Y is symmetrically distributed about θ .

For vectors $y, z \in R^n$ let $\langle y, z \rangle$ denote the usual inner product; so $\langle y, z \rangle = \sum y_i z_i$. Two vectors are orthogonal when their inner product is 0. Let $\Omega^\perp$ denote the orthogonal complement of Ω , the collection of vectors in R^n which are orthogonal to all vectors in Ω .

Denote the usual Euclidean norm by $\|y\|_{LS} = \langle y, y \rangle^{1/2}$. Least squares procedures are then based on this norm. Procedures based on ranks use another norm which involves a set of scores $a(1), \dots, a(n)$. These are often generated by a non-negative, non-decreasing, square integrable function $\phi(u)$, $0 < u < 1$, by setting $a(i) = \phi(i/(n+1))$. Without loss of generality, $\int \phi^2 = 1$. Scores that are frequently used are the sign scores, $\phi(u) = 1$, and the Wilcoxon scores, $\phi(u) = 3^{1/2}u$. Let $R|y_1|$ denote the rank of $|y_1|$ among $|y_1|, \dots, |y_n|$; then for a given set of scores

the function on R^n defined by

$$\begin{aligned} ||y||_R &= \langle a(R|y|), |y| \rangle \\ &= \sum a(R|y_i|) |y_i| \end{aligned} \quad (2.3)$$

is a norm on R^n ; see McKean and Schrader (1980). Note that for sign scores, $||y||_R$ is the L_1 -norm. In general we refer to (2.3) as the weighted least absolute deviation norm (WLAD-norm), and the weights depend on the size of the absolute residuals.

2.2 Prediction and Estimation

Let $||\cdot||$ represent any norm on R^n . A prediction of Y or estimate of θ is defined as a point $\hat{\theta}$ in Ω closest to y , that is, $\hat{\theta}$ satisfies

$$||y - \hat{\theta}|| = \min ||y - \theta||, \theta \in \Omega. \quad (2.4)$$

Such a point always exists. For the Euclidean norm $||\cdot||_{LS}$, $\hat{\theta}$ is the least squares projection of y onto Ω which we will denote by $\hat{\theta}_{LS}$. For the norm $||\cdot||_R$, we will call $\hat{\theta}$ a best rank or R -prediction of Y and denote it by $\hat{\theta}_R$. Computation of predictions is discussed in Section 4.

When the gradient $\nabla ||y - X\beta||$ exists, $\hat{\theta} = X\hat{\beta}$ is determined by the equation

$$\nabla ||y - X\hat{\beta}|| = 0. \quad (2.5)$$

In the case of least squares (2.5) represents the linear normal equations

$$-2X'(y - X\hat{\beta}) = 0$$

which, in the full rank case, results in the estimate $\hat{\beta}_{LS} = (X'X)^{-1}X'y$.

Note that the least squares residual vector is orthogonal to Ω , i.e.

$$(y - \hat{\theta}_{LS}) \in \Omega^\perp. \quad (2.7)$$

In the case of the weighted least absolute deviations norm the gradient exists at all but a finite number of points. The corresponding non-linear equation is

$$\nabla \|y - X\beta\|_R = -X' a^+(R|y - X\beta|) = 0 \quad (2.8)$$

where $a^+(R|y_i - \theta_i|) = a(R|y_i - \theta_i|) \operatorname{sgn}(y_i - \theta_i)$ for $i = 1, \dots, n$.

From (2.8) it follows that a best R-prediction is determined so that the vector of signed rank residuals $a^+(R|y - \hat{\theta}_R|)$ is orthogonal to Ω , i.e.

$$a^+(R|y - \hat{\theta}_R|) \in \Omega^\perp. \quad (2.9)$$

For the R-estimates, the minimum distance of y to Ω , $\|y - \hat{\theta}\|_R$ is unique. Although the WLAD-estimate is not unique, under regularity conditions the diameter of the solution set tends to 0 in probability quite rapidly, see Jaeckel (1972).

The gradient (2.8) consists of signed rank statistics appropriate for testing the various components of β . The gradient equation yields estimates which can be considered as extensions of the rank estimates of location proposed by Hodges and Lehmann (1963). If $\phi(u) = 3^{1/2}u$ and X is the vector of n ones then (2.8) becomes the Wilcoxon signed rank statistic and $\hat{\beta}$ is the median of the $n(n+1)/2$ pairwise averages of the observations. A motivation for replacing the LS norm with the WLAD norm is that the influence of an aberrant y point is bounded in the case of $\hat{\beta}_R$ and unbounded in the case of $\hat{\beta}_{LS}$. This type of robustness is discussed by Hampel (1974) and does not extend to protection against aberrant design points. Another motivation is the increased estimation efficiency discussed next.

2.3 Asymptotic Theory for Estimation

For the full rank case, under suitable regularity conditions (see Huber (1973), Jaeckel (1972), Jureckova (1971), Kraft and Van Eeden (1972)) the estimates $\hat{\beta}$ derived from the LS or WLAD norms are approximately normally distributed as

$$\hat{\beta} \sim \text{MVN}(\beta, K^2(X'X)^{-1}). \quad (2.10)$$

For least squares $K^2 = \sigma^2$, the variance of the error distribution. For the WLAD estimate $K^2 = \tau^2$ where

$$\tau^{-1} = -\int_0^1 \phi(u) f'[F^{-1}(\{2^{-1}(u+1)\})]/f[F^{-1}(\{2^{-1}(u+1)\})]du. \quad (2.11)$$

Following Bickel (1964) the efficiency of $\hat{\beta}_R$ relative to $\hat{\beta}_{LS}$ is σ^2/τ^2 . In the case of Wilcoxon scores $\sigma^2/\tau^2 = 12 \sigma^2 (\int f^2(x) dx)^2$ which is bounded below by .864 and may be arbitrarily large, depending on F . When F is the normal cdf the efficiency for Wilcoxon scores is .955. See Lehmann (1975 Sections 2.4 and 4.3) for a further discussion of the efficiency. Hence we find that the WLAD norm produces rank estimates which are generally more efficient than the least squares estimate, at least for error distributions with tails heavier than those of the normal distribution.

3. Testing in the Linear Model

3.1 General Linear Hypotheses

For the model (2.1) we are interested in general linear hypotheses of the form

$$H_0: H\beta = 0 \text{ versus } H_A: H\beta \neq 0 \quad (3.1)$$

where $H\beta$ is a collection of q linearly independent estimable functions.

Let Ω_0 be the $(p-q)$ dimensional subspace of Ω constrained by the hypothesis $H\beta = 0$. In terms of the model (2.2) we are testing

$$H_0: \theta \in \Omega_0 \text{ versus } H_A: \theta \in \Omega - \Omega_0. \quad (3.2)$$

We will call the model (2.1), $\theta \in \Omega$, the full model and the model with the constraint $\theta \in \Omega_0$ the reduced model.

3.2 Tests based on minimum distances

Given a norm $||\cdot||$ on R^n , tests of (3.2) can be naturally constructed by comparing the distances between y and the two subspaces Ω and Ω_0 . If $\hat{\theta}_0, \hat{\theta}$ are the best reduced and full model predictions of y , for the norm, then the minimum distances are $||y - \hat{\theta}_0||$ and $||y - \hat{\theta}||$.

If the norm is $||\cdot||_{LS}$ then the comparison is between the square of the minimum distances, $||y - \hat{\theta}_0||_{LS}^2 - ||y - \hat{\theta}||_{LS}^2$. Asymptotic distribution theory suggests standardizing this difference by an estimate of σ^2 , the variance of the error distribution. The usual F-test is

$$F_{LS} = \frac{||Y - \hat{\theta}_0||_{LS}^2 - ||Y - \hat{\theta}||_{LS}^2}{q\hat{\sigma}^2} \quad (3.3)$$

where $\hat{\sigma}^2 = ||y - \hat{\theta}||_{LS}^2 / (n - p)$. Under H_0 and regularity conditions, see Huber (1973), qF_{LS} has an asymptotic $\chi^2(q)$ distribution. The usual small sample test compares F_{LS} with $F(q, n - p)$ -critical values, rejecting H_0 at approximate level α if $F \geq F(\alpha, q, n - p)$. Note also that although $\hat{\theta}_0$ and $\hat{\theta}$ are derived from $||\cdot||_{LS}$ the notation has been suppressed.

If the norm $||\cdot||_R$ is used then a natural comparison is between the minimum distances, $||y - \hat{\theta}_0||_R - ||y - \hat{\theta}||_R$. The asymptotic distribution theory developed by McKean and Hettmansperger (1976) leads to standardizing this difference by an estimate of τ , (2.11), resulting in the test statistic,

$$D = \frac{||Y - \hat{\theta}_0||_R - ||Y - \hat{\theta}||_R}{q(\hat{\tau}/2)} \quad (3.4)$$

where $\hat{\tau}$ is an estimate of τ , for instance (4.2). Under H_0 and regularity conditions (see McKean and Hettmansperger (1976)) qD also has an asymptotic $\chi^2(q)$ distribution. For small samples, the level of the test seems to be more stable if the $F(\alpha, q, n - p)$ critical point is used; see Hettmansperger and McKean (1977) and Section 5 of the present paper. Hence, the test rejects H_0 at approximate level α if $D \geq F(\alpha, q, n - p)$.

Because of the close relationship between the inner product $\langle \cdot, \cdot \rangle$ and the norm $||\cdot||_{LS}$ the F statistic (3.2) can be written in two other algebraically equivalent forms. Other rank tests for the linear model can be identified with these alternative forms of the F; however, since they are based on the norm $||\cdot||_R$, they are not algebraically equivalent. We next discuss these other forms of the F-test and the corresponding rank tests. For convenience we will assume X has full column rank.

3.3 Tests based on the full model estimates

This form of the F statistic can be derived from (3.3) using the Pythagorean theorem. Let $\Omega|\Omega_0$ be the orthogonal complement of Ω_0 in Ω ; then $\|y - \hat{\theta}_0\|_{LS}^2 - \|y - \hat{\theta}\|_{LS}^2 = \|P_{\Omega_0^\perp} y\|_{LS}^2 - \|P_{\Omega^\perp} y\|_{LS}^2 = \|P_{\Omega|\Omega_0} y\|_{LS}^2$,

where $P_{\Omega|\Omega_0} y$ denotes the projection of y onto $\Omega|\Omega_0$. Hence

$$F = \frac{\|P_{\Omega|\Omega_0} y\|_{LS}^2}{q\hat{\sigma}^2} \quad (3.5)$$

From (3.1), Ω_0 is determined by $H\beta = 0$. Assume X is a basis matrix for Ω then $\beta = (X'X)^{-1}X'\theta$ and $H(X'X)^{-1}X'\theta = 0$ for every $\theta \in \Omega_0$. It then follows that $Z = X(X'X)^{-1}H'$ is a basis matrix for $\Omega|\Omega_0$ and $P_{\Omega|\Omega_0} y = Z(Z'Z)^{-1}Z'y$.

When the definition of Z is combined with $P_{\Omega|\Omega_0} y$, (3.5) becomes

$$F = \frac{(H\hat{\beta})'[H(X'X)^{-1}H']^{-1}(H\hat{\beta})}{q\hat{\sigma}^2} \quad (3.6)$$

This is the coordinatized version of (3.3) and expresses the numerator in terms of the full model least squares estimates.

For the corresponding rank test, replace the least squares estimate in (3.6) with the rank estimate $\hat{\beta}_R$ and replace $\hat{\sigma}^2$ by $\hat{\tau}^2$. Standard asymptotic theory based on (2.10) shows that q times this test statistic is approximately chi squared with q degrees of freedom. Preliminary simulations indicate that the probability of a type I error is better controlled if $\hat{\tau}^2$ is estimated by $\hat{\tau}^2$, (4.2), which includes the bias correction similar to that used in least squares. The test is carried out by referring

$$B = \frac{(H\hat{\beta}_R)'[H(X'X)^{-1}H']^{-1}(H\hat{\beta}_R)}{q\hat{\tau}^2} \quad (3.7)$$

to an $F(\alpha, q, n - p)$ critical point. In the numerator $\hat{\beta}_R$ can be replaced by $\hat{\beta}^{(k)}$, (4.1), the k -step rank estimate.

Bickel (1976) defined a pseudo-observation based on an M -estimate and described how these pseudo-observations could be used in a least squares program to yield a robust test of hypothesis; see Schrader and Hettmansperger (1980). We now adapt this idea to the WLAD or rank estimate approach.

Given the full model estimate $\hat{\theta}_R = X\hat{\beta}_R$, define the pseudo-observation $\tilde{y}$ by:

$$\begin{aligned}\tilde{y} &= X\hat{\beta}_R + \lambda a^+(R|y - X\hat{\beta}_R|) \\ &= \hat{\theta}_R + \lambda a^+(R|y - \hat{\theta}_R|)\end{aligned}\tag{3.8}$$

where λ is a constant to be determined. Note that since $a^+(R|y - \hat{\theta}_R|) \in \Omega^1$ from (2.9), the least squares estimate of β based on $\tilde{y}$ is $\hat{\beta}_R$. The least squares variance estimate based on $\tilde{y} - \hat{\theta}_R$ is

$$\begin{aligned}\hat{\sigma}^2 &= \frac{1}{n-p} ||\tilde{y} - \hat{\theta}_R||^2 \\ &= \frac{\lambda^2}{n-p} \sum_{i=1}^n a_i^2\end{aligned}\tag{3.9}$$

where $\sum_{i=1}^n a_i^2$ is a known constant. Hence for $\hat{\tau}$ as in (4.2), if we take

$$\lambda = \hat{\tau} / (n^{-1} \sum_{i=1}^n a_i^2)^{1/2}\tag{3.10}$$

then $\hat{\sigma}^2 = \hat{\tau}^2$, the proper denominator in (3.7). In the case of Wilcoxon scores, $\lambda = \hat{\tau} [(n+1)/(n-1)]^{1/2}$.

Hence we can compute B in (3.6) quite easily, as follows: 1. find $\hat{\beta}_R$ (or $\hat{\theta}_R$) and $\hat{\tau}$, 2. construct $\tilde{y}$ and 3. use $\tilde{y}$ in a least squares AOV program.

The numerator of B is the appropriate line of the AOV table and the denominator is the error line. The pseudo observations also have diagnostic value for data analysis. A plot of $\hat{y}$ against y shows the effect of "robustification" on the data. This is illustrated in the example of Section 7.

3.4 Aligned Tests

The aligned tests are most conveniently described for the case $H = (0, I)$ where I is the $q \times q$ identity. The model (2.4) is partitioned as

$$Y = X_1 \beta_1 + X_2 \beta_2 + e \quad (3.11)$$

where X_1 and X_2 are of order $n \times (p - q)$ and $n \times q$ while β_1 and β_2 are $(p - q) \times 1$ and $q \times 1$ vectors, respectively. The null hypothesis (3.1) becomes $H_0: \beta_2 = 0$ where β_1 is treated as a nuisance vector.

The least squares F test can be derived by first removing the effects of the nuisance vector β_1 by projecting both y and the columns of X_2 onto Ω_0 . Hence consider $P_{\Omega_0} y = y - P_{\Omega_0} y$ and $P_{\Omega_0} X_2 = X_2 - P_{\Omega_0} X_2$ where $P_{\Omega_0} = X_1 (X_1' X_1)^{-1} X_1'$. Now using $P_{\Omega_0} X_2$ as the design, project the vector $P_{\Omega_0} y$ to construct the F test for $H_0: \beta_2 = 0$. The numerator in this case is the squared length of this projection and the resulting form of the F statistic is

$$F = \frac{(P_{\Omega_0} y)' X_2 (X_2' P_{\Omega_0} X_2)^{-1} X_2' (P_{\Omega_0} y)}{q \hat{\sigma}^2} \quad (3.12)$$

This can be derived algebraically from (3.3) and expresses the numerator as a quadratic form in the reduced model residuals. Draper and Smith

(1966, Section 4.1) discuss this method for fitting a regression equation with two independent variables as a sequence of simple regression fits.

Now let $\hat{\theta}_0 = X_1 \hat{\beta}_1$ denote the reduced model rank estimate. Hence (similar to (2.9))

$$a^+(R|y - \hat{\theta}_0|) \in \Omega_0^1. \quad (3.13)$$

Using the numerator in (3.12) we replace the reduced model residuals by the reduced model signed ranks of the residuals and define

$$Q = [a^+(R|y - \hat{\theta}_0|)]' X_2 (X_2' P_{\Omega_0^1} X_2)^{-1} X_2' a^+(R|y - \hat{\theta}_0|). \quad (3.14)$$

The vector

$$S_2 = X_2' a^+(R|y - \hat{\theta}_0|) \quad (3.15)$$

is called the vector of aligned rank statistics and

$$Q = S_2' (X_2' P_{\Omega_0^1} X_2)^{-1} S_2 \quad (3.16)$$

is called the aligned rank test statistic. See Hodges and Lehmann (1962).

The results of Sen and Puri (1977) when specialized to the univariate linear model show that Q has an approximate chi squared distribution with q degrees of freedom. Hence the asymptotic test of $H_0: \beta_2 = 0$ rejects when $Q > \chi^2(\alpha, q)$

We next derive an equivalent rank test. From (3.13) we can replace $a^+(R|y - \hat{\theta}_0|)$ by $P_{\Omega_0^1} a^+(R|y - \hat{\theta}_0|)$ in (3.14). Using the matrix identity

$$P_{\Omega_0^1} X_2 (X_2' P_{\Omega_0^1} X_2)^{-1} X_2' P_{\Omega_0^1} = P_{\Omega_0^1} - P_{\Omega_0^1} \quad (3.17)$$

we can write Q in (3.16) as

$$Q = [a^+(R|y - \hat{\theta}_0|)]' [P_{\Omega_0^\perp} - P_{\Omega_1}] a^+(R|y - \hat{\theta}_0|). \quad (3.18)$$

This version of Q is introduced and discussed in detail by Adichie (1978). In his discussion the estimate $\hat{\theta}_0$ need not be a rank estimate to obtain the asymptotic distribution. To our knowledge no study has been made of the sensitivity of Q to the type of estimate used for $\hat{\theta}_0$. Since Q is a rank test it would seem most natural to use a rank estimate and under this condition the aligned rank test (3.14) is equivalent to Adichie's test (3.18). For the remainder of this paper the term aligned rank tests will refer to either construction.

All four tests: (3.4), (3.7), (3.14) and (3.18) are very similar in their asymptotic properties. They all have an asymptotic chi squared distribution under the null hypothesis and they have the same asymptotic efficiency. We have further shown in this section that though they are not algebraically identical, each is closely associated with one of the various equivalent forms of the F test. In Section 5 we consider some of the small sample power properties via Monte Carlo simulations.

4. Computation

The R-estimates, predicted values, and hence minimum distances can be obtained by the k-step procedure discussed by McKean and Hettmansperger (1978a). It is an algorithm based on approximating the dispersion function $\|y - X\beta\|_R$ by a quadratic function. For a general design matrix X the kth step estimate of θ is

$$\hat{\theta}^{(k)} = \hat{\theta}^{(k-1)} + P_{\Omega} \hat{\tau}^{(k-1)} a(R|y - \hat{\theta}^{(k-1)}|);$$

while for the full rank case the estimate of β is

$$\begin{aligned} \hat{\beta}^{(k)} &= \hat{\beta}^{(k-1)} + \hat{\tau}^{(k-1)} (X'X)^{-1} V \|y - X\hat{\beta}^{(k-1)}\| \\ &= \hat{\beta}^{(k-1)} + \hat{\tau}^{(k-1)} (X'X)^{-1} X' a^+(R|y - X\hat{\beta}^{(k-1)}|). \end{aligned} \quad (4.1)$$

In general this algorithm will not converge and, hence, to obtain fully iterated estimates algorithms such as steepest descent must be used. The asymptotic distribution theory for inference based on these k-step estimates, however, is the same as that of the fully iterated estimates. Furthermore in a simulation study by the authors (1978a), small sample results were practically the same for estimates of 2 or 3 steps as that of the fully iterated ones. For starting values resistant to outliers, we recommend ℓ_1 -estimates. Recent algorithms such as Barrodale and Roberts (1973) and Bartels and Conn (1980) make ℓ_1 -starts computationally feasible.

Note from (4.1), that the k-step residuals, $r^{(k)}$, can be written as

$$r^{(k)} = r^{(k-1)} - P_{\Omega} (\hat{\tau}^{(k-1)} a^+(R|r^{(k-1)}|))$$

where P_{Ω} is the projection transformation onto Ω . Hence a convenient, and numerically stable algorithm, is based on first obtaining a QR-decomposition of the design matrix X, using, for instance, the collection of algorithms LINPACK by Dongara et al. (1979). A QR-decomposition results

in an orthonormal basis for Ω from which projections can readily be formed. Another advantage is that the design matrix X need not be of full column rank. For fully iterated estimates, steepest descent in terms of an orthonormal basis is simply a search along the direction specified by $\hat{\theta}^{(k)} - \hat{\theta}^{(k-1)}$. Finally, reduced model design matrices for natural hypotheses of the form $H\beta = 0$ can readily be obtained using QR-decompositions; see McKean and Schrader (1981).

A consistent estimate of τ proposed by the authors (1976), (1978a), is

$$\hat{\tau} = (n^{1/2}(U - L)/2Z_{\alpha/2})[n/(n - p)]^{1/2}, \quad (4.2)$$

where $Z_{\alpha/2}$ is the $(1 - \alpha/2)$ percentile point of the standard normal distribution and U and L are solutions C of the equations $n^{-1/2} \sum a^+(R|r_i - C) \doteq Z$ for $Z = -Z_{\alpha/2}$ and $Z_{\alpha/2}$, respectively. In the case of Wilcoxon scores U and L are ordered Walsh averages of the residuals. Even in this case, though, the equations are solved by an iterative algorithm similar to the one discussed by McKean and Ryan (1977).

As the authors (1977) discussed, small sample corrections are necessary for this estimation of τ . One such correction is given by the term in the brackets of (4.2) which corresponds to the usual least squares degree of freedom correction for the estimate of variance. Another correction which has proven useful in the Wilcoxon case is to modify the Z in the above equations so as to eliminate the p smallest, in absolute value, ordered Walsh averages. A similar idea was used by Hill and Holland (1977) for a scale estimate based on L_1 -residuals. The estimate of τ used in Sections 5 and 7 employed both of these small sample corrections with $Z_{\alpha/2} = 1.645$.

5. Monte Carlo Study

As discussed in Section 3 the different rank tests have the same asymptotic null distribution and relative efficiency. Monte Carlo simulations are needed to investigate their small sample behavior. Simulation studies by the authors (1977), (1978a), (1978b) have indicated that for the statistic D , (3.4), the probability of a type I error is quite stable over a variety of underlying distributions. This stability is confirmed in the study discussed below. Small sample properties of the rank test B , (3.7), and the aligned rank tests, (3.14) and (3.18), have not previously been investigated.

Before turning to the comparative simulation, we will briefly discuss two recent simulations of related aligned rank tests proposed by Sen (1969) and Adichie (1974) for testing parallelism of several regression lines. Sen ranks within the individual data sets and avoids estimating the intercept; Adichie ranks the combined data set but assumes the intercepts are all equal and avoids estimation of the common intercept.

Lo, Simkin, and Worthley (1978) did a simulation study of these tests of parallelism for the case of equal intercepts. The aligned rank tests performed uniformly worse than the least squares F-test on the distributions considered. The power of the aligned tests was always low and in some cases almost nonexistent.

Smit (1979) performed a similar simulation study. His results also indicated that least squares generally performed better than the aligned tests. There is, however, one glaring difference with the first study.

In several cases the F test and Sen's test had very similar power but Adichie's test had little or no power. A careful comparison of the two studies yields contradictory conclusions. In the first study the aligned rank tests behaved similarly and were inferior to the F test. In the second, Sen's test and the F test often behaved similarly and were superior to Adichie's test.

The aligned tests described in Section 3 were developed for the general linear hypotheses, quite analogous to least squares, and differ from these earlier tests. With the above studies in mind, we decided to investigate the tests discussed in Section 3 on essentially the same model as considered by Lo, Simkin, and Worthley (1978): three regressions with unconstrained intercepts, equal sample sizes, and uniformly spaced x 's. We considered sample sizes of 5 and 10, hence, total sample sizes of 15 and 30. The dimension of Ω for this design is 6; the ratios of observations to parameters are 2.5 and 5. We label this design A.

A second parallel regression problem, design B, consists of two regressions with common intercepts and x 's placed at 1, 2, 3, 4, 5, 10 for the first sample and at 7, 8, 9, 10, 11, 12 for the second. This design contains a point of moderate leverage corresponding to the point 10 in the first sample. This is a valuable point of the design and should not be confused with points of leverage combined with discrepant observations, see Hoaglin and Welsch (1978). Least squares and the rank tests D (3.4) and B (3.7) are not affected by this design but, as shown below and discussed in the next section, the aligned tests are adversely affected. The observation-parameter ratio for this design is 12 to 3. As in the first design, we are testing for parallelism.

In order to study the performance of these tests on both moderate and heavy tailed error structure, we selected normal and Cauchy errors. With Lo, Simkin and Worthley's study in mind, we also included the double exponential distribution for design A. The normals were obtained using the transformation proposed by Marsaglia and Bray (1964) on a pair of uniform variates. The uniforms were generated by the algorithm UNI developed by Gross (1976). The double exponential and the Cauchy observations were generated in the form normal over an independent variable as noted in Simon (1976). The tests simulated were F, (3.3); D, (3.4); B, (3.7); and Q, (3.14). Both signed rank and rank scores were used. The results are all based on 500 simulations.

Empirical 5% and 10% levels for the test statistics on all the designs and distributions are displayed in Table 5.1. Least squares (3.3) and the rank tests based on drop in dispersion (3.4) are fairly stable over almost all the situations. The tests based on R-pseudo-observations, (3.7), appear to be conservative for the small sample sizes in the normal model and tend to be liberal for Cauchy errors on design A with $n_1 = 10$ and design B. This behavior for the statistic B on Cauchy errors confirms similar findings on tests derived from M-pseudo-observations (Schrader and Hettmansperger (1980)). It is also predictable in the light of robust regression estimators which seem to have larger small sample variances than their asymptotic counterparts for heavy tailed error structures; see Huber (1973) and McKean and Hettmansperger (1978a). Small sample corrections for the tests based on pseudo-observations seem to need some measure of tail weight of the underlying error structure.

- Table 5.1 about here -

The null levels for the aligned tests are more erratic. They are quite liberal for the small sample size on design A. This improves somewhat for the larger sample size. For design B they seem to be conservative, especially the aligned test based on rank scores. As shown in Table 5.4, their power was adversely affected by this design; for instance, the least squares test is as powerful on Cauchy errors as the aligned tests based on rank scores. The aligned rank tests exhibited even worse behavior for other designs which included points of moderate leverage. A partial explanation of this behavior is found in the next section.

Simple small sample corrections, such as using F-critical values, that would improve the behavior of the aligned tests on design A, would be quite detrimental to their behavior on design B. Due to their sensitivity to design, small sample corrections for the aligned rank tests appear to be more complicated than those for the tests based on the drop in dispersion or pseudo-observations. Corrections for the aligned tests need more information involving the design matrix.

The results of the power study for the tests at level .05 appear in Tables 5.2 - 5.4. The alternatives were selected separately for each distribution in order to achieve a reasonable range of powers. For valid comparisons, the empirical levels should be close to .05. Since this is true for the least squares test and the rank test based on D, (3.4), these tests can be compared (other than LS at Cauchy errors on design B). For all designs, least squares dominates on normal errors, while D dominates on Cauchy errors. On double exponential errors note that least squares is slightly more powerful for samples of size 5, whereas D is more powerful for the samples of size 10.

The test B based on R-pseudo-observations is slightly dominated by the test D.

- Tables 5.2 - 5.4 about here.

When comparable (design A, $n_i = 10$, double exponential or Cauchy errors) the power results of the aligned tests are similar to D. Certainly the results on aligned tests are much improved over the earlier tests of Sen (1969) and Adichie (1974) considered by the two studies mentioned above. The adverse behavior of aligned tests on design B was noted above. The results for the rank tests other than the aligned tests seem to be similar for rank or signed-rank scores. Neither type of score dominated the other.

Our general recommendation is to use the WLAD or rank estimate $\hat{\beta}$ along with D. This approach combines the estimate which minimizes D, a data fitting criterion, with the test that considers the reduction in D due to fitting the various models under consideration. The asymptotic theory and corresponding small sample adjustments combine to provide an effectively distribution free, robust test, and a robust estimate. The level and power of D appear quite stable over a variety of underlying error distributions.

6. The Effect of Leverage in the Design

In order to understand the behavior of the test statistics on the design containing a point of moderate leverage, consider a simple model containing two predictor variables with observations taken at the points (x_1, x_2) as shown in Table 6.1. The point $(0, x)$ is an extreme point

- Table 6.1 about here -

in the x_2 -direction. Let y denote the observation corresponding to $(0, x)$. Consider the full model $Y = \beta_1 X_1 + \beta_2 X_2 + e$, and the hypothesis $H_0: \beta_2 = 0$ versus $H_1: \beta_2 \neq 0$.

If y is on the surface then it is a valuable point in determining the fit. The reduced model residual at y will be large and its value incorporated into F and D ; Q however down weights the residuals and consequently loses power.

In order to see this consider a perfect model: $Y = X_1 + X_2 + 0$; that is all the points lie on the surface. Since the denominators of F and D are then zero, we will only consider their numerators. A straight forward calculation shows

$$F = \frac{3(3xy + 20)^2}{10(3x^2 + 20)} \quad (6.1)$$

$$Q = \frac{3(3.78x + 18.89)^2}{10(3x^2 + 20)} \quad (6.2)$$

In computing Q we supposed that $y > 1$ so that the reduced model residual has rank 10. When y is on the surface we have $y = x$. As x increases, the leverage increases, F also increases but Q decreases! In fact Q may

decrease below the critical value and fail to detect $H_1: \beta_2 > 0$. The formulas (6.1) and (6.2) continue to hold when y is not on the surface. In (4.2) we see that Q is only dependent on y through the rank. Again when $y = x$ since the full model fit is perfect, the full model dispersion is 0 and D only depends on the reduced model dispersion. Now for the example, when $y = x$ and $x > 0$, the numerator of D is (for rank scores),

$$D = 2(5.67 + 1.42y) \quad (6.3)$$

Here x does not enter the formula since it does not appear in the reduced model. Note that D increases with y as the leverage increases.

This simple example provides some explanation for the adverse effect of moderate leverage on aligned tests noted in the simulations. It further suggests that the other tests will not be so effected.

7. An Example

We consider a 3×4 factorial experiment discussed by Box and Cox (1964). An experiment was carried out on 48 animals to study the relative effectiveness of 3 poisons and 4 treatments. There were 4 animals at each poison x treatment combination and the data consists of the 48 survival times. The data is reproduced in Table 7.1

- Table 7.1 about here -

The least squares AOV is given in Table 7.2.

- Table 7.2 about here -

Note that the F test for interaction fails to achieve significance at the 5% level. A glance at a plot of the cell means shows that there are crossing means in Poisons 1 and 2 while Poison 3 is almost indifferent to which treatment is applied. These sorts of patterns are highly suggestive of interaction. See Brown (1975). If we plot the standardized full model residuals against the full model predicted values, a fan shaped plot appears. See Figure 7.1.

- Figure 7.1 about here -

There is clear heterogeneity of cell variances and the four circled observations are noted for their large standardized residual. Note the residuals were standardized by dividing by the pooled estimate of their standard deviations.

We now consider how a parallel, robust analysis based on the ranks of the residuals can enhance the data analysis. We will use Wilcoxon scores.

A robust AOV table based on (3.4) is shown in Table 7.3.

- Table 7.3 about here -

Note the test for interaction is now significant at the 5% level.

The pseudo-observations (3.8) provide some insight into how the rank tests are operating on the data. In Figure 7.2 we plot the pseudo-observations against the actual observations. The 4 observations with large standardized residuals are marked. Notice how they are "brought back in."

- Figure 7.2 about here -

In Figure 7.3 we plot the standardized residuals of the pseudo-observations against the robust predicted values. There are no longer any extreme standardized residuals.

- Figure 7.3 about here -

Further, it can be seen how the ranking process is working to equalize the cell standard deviations. Compare Figures 7.1 and 7.3.

Table 7.3 was based on the drop in dispersion (3.4). If the pseudo-observations were used with a least squares program then a similar table based on (3.7) could be constructed. The results are quite similar for the two approaches.

- Adichie, J. N. (1974). "Rank Score Comparison of Several Regression Lines". Ann. of Statist., 2, 396-402.
- Adichie, J. N. (1978). "Rank Tests of Subhypotheses in the General Linear Regression." Annals of Statist., 6, 1021-1026.
- Arnold, S. (1981). The Theory of Linear Models and Multivariate Analysis. New York: John Wiley and Sons.
- Barrodale, I. and Roberts, F. D. K. (1974). Algorithm 478: Solution of an overdetermined system of equations in the L_1 norm. Comm. Assoc. Comput. Mach., 17, 319-20.
- Bartels, R. H. and Conn, A. R. Linearly constrained discrete l_1 problems. Trans. Math. Soft., 6, 594-608.
- Bickel, P. J. (1964). "On Some Alternative Estimates for Shift in the p-Variate One Sample Problem," Ann. Math. Statist., 35, 1079-1090.
- Bickel, P. J. (1976). Another look at robustness: A review of reviews and some new developments (reply to discussant). Scand. J. Statist., 3, 167.
- Box, G. E. P. and Cox, D. R. (1964). An analysis of transformations (with discussion). J. R. Statist. Soc., B2, 211-52.
- Brown, M. B. (1975). Exploring interaction effects in the ANOVA. Appl. Statist. 24, 288-98.
- Dongara, J. J., Bunch, J. R., Moler, C. B. and Stewart, G. W. (1979). LINPACK Users Guide. Philadelphia: Society for Industrial and Applied Mathematics.
- Draper, N. R. and Smith, H. (1966). Applied Regression Analysis. New York: John Wiley and Sons.
- Gross, A. M. (1976). Portable random number generation. Unpublished memorandum, Bell Laboratories, Murray Hill, NJ.
- Hampel, F. R. (1974). The influence curve and its role in robust estimation. J. Amer. Statist. Assoc., 69, 383-93.
- Hettmansperger, T. P. and McKean, J. W. (1977). A Robust Alternative Based on Ranks to Least Squares in Analyzing Linear Models. Technometrics, 19, 275-284.
- Hill, R. W. and Holland, P. W. (1977). Two robust alternatives to least-squares regression. J. Am. Statist. Assoc., 72, 828-33.

- Hoaglin, D. C. and Welsch, R. E. (1978). The Hat Matrix in Regression and ANOVA. The American Statistician, 32, 17-22.
- Hodges, J. L., Jr. and Lehmann, E. L. (1962). Rank methods for combination of independent experiments in analysis of variance. Ann. Math. Statist., 33, 482-497.
- Hodges, J. L., Jr. and Lehmann, E. L. (1963). Estimates of location based on rank tests. Ann. Math. Statist., 34, 598-611.
- Huber, P. J. (1973). "Robust Regression: Asymptotics, Conjectures and Monte Carlo," Ann. Statist., 1 799-821.
- Jaekel, L. A. (1972). Estimating regression coefficients by minimizing the dispersion of the residuals. Ann. Math. Statist., 43, 1449-58.
- Jureckova, J. (1971). Nonparametric estimate of regression coefficients. Ann. Math. Statist., 42, 1328-38.
- Kraft, C. H. and Van Eeden, C. (1972). Linearized rank estimates and signed-rank estimates for the general linear hypothesis. Ann. Math. Statist., 43, 42-57.
- Lehmann, E. L. (1975). Nonparametrics: Statistical Methods Based on Ranks. San Francisco: Holden-Day, Inc.
- Lo, L. C., Simkin, M. E. and Worthley, R. E. (1978). A small sample comparison of rank score tests for parallelism of several regression lines. Journal of Am. Statist. Assoc., 73, 666-669.
- Marsaglia, G. and Bray, T. A. (1964). A convenient method for generating normal variables. SIAM Rev., 6, 260-4.
- McKean, J. W. and Hettmansperger, T. P. (1976). "Tests of Hypothesis in the General Linear Model Based on Ranks," Comm. in Statist., A5(8), 693-709.
- McKean, J. W. and Hettmansperger, T. P. (1978a). Statistical Inference Based on Ranks. Psychometrika, 43, 69-79.
- McKean, J. W. and Hettmansperger, T. P. (1978b). A Robust Analysis of the General Linear Model Based on One Step R-estimates. Biometrika, 65, 571-579.
- McKean, J. W. and Ryan, T. A., Jr. (1977). An algorithm for obtaining confidence intervals and point estimates based on ranks in the two sample location problem. Trans. Math. Software, 3(2), 183-185.
- McKean, J. W. and Schrader, R. M. (1980). The Geometry of Robust Procedures in Linear Models. JRSS(B), 42, 366-371.
- McKean, J. W. and Schrader, R. M. (1981). The Use and Interpretation of Robust Analysis of Variance. Proceedings of ARO Conference on Modern Data Analysis, New York: Academic Press.

- Schrader, R. M. and Hettmansperger, T. P. (1980). Robust Analysis of Variance Based Upon a Likelihood Ratio Criterion. Biometrika, 67, 93-101.
- Sen, P. K. (1969). On a Class of Rank Order Tests for the Parallelism of Several Regression Lines, Ann. Math. Statist., 40, 1668-1683.
- Sen, P. K. and Puri, M. L. (1977). Asymptotically Distribution-Free Aligned Rank Order Tests for Composite Hypotheses for General Multivariate Linear Models, Z. Wahrscheinlichkeitstheorie und Verw. Gebiete, 39, 175-186.
- Simon, G. (1976). Computer simulation swindles, with applications to estimates of location and dispersion. JRSS(C), 25, 266-74.
- Smith, C. F. (1979). An Empirical Comparison of Several Tests for Parallelism of Regression Lines, Comm. in Statist. (B), 8, 61-74.

TABLE 5.1
EMPIRICAL LEVELS
PARALLEL DESIGN A($n_1 = 5$)

DIST	NOMINAL LEVEL	LS	D		B		Q	
			SR	R	SR	R	SR	R
NORMAL	05	050	036	036	032	030 ⁻	076 ⁺	056
	10	098	080	070	056 ⁻	048 ⁻	182 ⁺	142 ⁺
DEXP	05	060	048	044	040	046	094 ⁺	076 ⁺
	10	100	096	082	074 ⁻	078	178 ⁺	148 ⁺
CAUCHY	05	042	044	036	060	050	098 ⁺	068
	10	114	108	082	096	086	200 ⁺	160 ⁺

PARALLEL DESIGN A($n_1 = 10$)

NORMAL	05	048	058	056	046	042	064	058
	10	090	098	094	094	082	126 ⁺	120
DEXP	05	038	050	048	058	050	050	034
	10	074 ⁻	098	094	090	088	120	100
CAUCHY	05	062	076 ⁺	062	094 ⁺	096 ⁺	068	064
	10	116	138 ⁺	128 ⁺	154 ⁺	138 ⁺	142 ⁺	132 ⁺

PARALLEL DESIGN B($n_1 = 5$)

NORMAL	05	048	034	032	030 ⁻	028 ⁻	034	016 ⁻
	10	088	072	068 ⁻	060 ⁻	056 ⁻	100	090
CAUCHY	05	092 ⁺	062	046	104 ⁺	102 ⁺	030 ⁻	024 ⁻
	10	110	122	110	128 ⁺	130 ⁺	102	078

Note: A -(+) is attached if the empirical level is below(above) (.031, .070) or (.077, .123), the 95% intervals around .05 and .10 respectively.

TABLE 5.2
EMPIRICAL POWER FOR .05 TESTS
PARALLEL DESIGN A($n_1 = 5$)

		LS	D		B		Q	
			SR	R	SR	R	SR ⁺	R
NORMAL	NULL	050	036	036	032	030 ⁻	076 ⁺	056
	1	126	098	084	076	068	226	160
	2	372	298	276	224	214	462	380
ALTS	3	646	592	546	472	452	724	626
	4	976	942	930	870	860	986	968
	5	1.000	998	1.00	998	1.00	1.00	1.00
DEXP	NULL	060	048	044	040	046	094 ⁺	076 ⁺
	1	094	072	070	058	054	168	114
	2	220	192	174	156	142	324	262
ALTS	3	406	376	340	298	284	556	472
	4	782	758	730	678	668	860	802
	5	974	972	958	930	930	980	962
CAUCHY	NULL	042	044	036	060	050	098 ⁺	068
	1	062	088	076	080	080	202	136
	2	218	300	278	244	230	476	394
ALTS	3	376	598	582	524	504	692	610
	4	536	762	740	740	716	788	748
	5	598	814	800	794	786	832	786

TABLE 5.3
EMPIRICAL POWER FOR .05 TESTS
PARALLEL DESIGN A($n_1 = 10$)

		LS	D		B		Q	
NORMAL			SR	R	SR	R	SR	R
	NULL	048	058	056	046	042	064	058
	1	078	082	076	066	064	096	086
ALTS	2	450	432	398	364	374	476	452
	3	888	884	858	836	832	904	874
	4	994	990	990	990	984	990	990
	5	1.000	998	998	998	998	998	998
DEXP	NULL	042	050	048	058	050	050	034
	1	134	164	150	148	142	176	152
	2	424	530	504	486	476	536	506
	3	772	832	814	784	800	838	794
	4	986	984	984	990	990	988	986
	5	998	998	996	1.000	1.000	994	992
CAUCHY	NULL	062	076 ⁺	062	094 ⁺	096 ⁺	068	064
	1	070	096	086	126	112	124	090
	2	146	400	386	376	380	410	390
	3	248	738	734	744	734	710	678
	4	380	894	892	904	902	844	824
	5	480	958	960	964	966	912	908

TABLE 5.4
EMPIRICAL POWER FOR .05 TESTS
PARALLEL DESIGN B

		LS	D		B		Q	
			SR	R	SR ⁻	R ⁻	SR	R ⁻
NORMAL	NULL	048	034	032	030 ⁻	028 ⁻	034	016 ⁻
	1	252	196	178	154	150	204	124
	2	624	486	424	422	428	448	328
ALTS	3	894	790	760	760	764	684	546
	4	990	974	954	936	930	842	718
	5	996	988	986	982	974	884	768
CAUCHY	NULL	092 ⁺	062	046	104 ⁺	102 ⁺	030 ⁻	024 ⁻
	1	350	460	428	410	406	430	300
	2	634	812	780	774	772	720	572
ALTS	3	696	858	852	846	844	772	660
	4	750	902	884	886	888	818	722
	5	798	926	922	924	926	868	778

TABLE 6.1

DESIGN WITH HIGH LEVERAGE POINT

x_1	-1	-1	-1	1	1	1	0	0	0	0
x_2	-1	0	1	-1	0	1	-1	0	1	x

TABLE 7.1

Survival times (unit, 10 hr) of animals in a 3 x 4 factorial experiment

Poison	Treatment			
	A	B	C	D
I	031	082	043	045
	045	110	045	071
	046	088	063	066
	043	072	076	062
II	036	092	044	056
	029	061	035	102
	040	049	031	071
	023	124	040	038
III	022	030	023	030
	021	037	025	036
	018	038	021	031
	023	029	022	033

TABLE 7.2

AOV					
<u>SOURCE</u>	<u>df</u>	<u>SS</u>	<u>MS</u>	<u>F</u>	<u>5% Pt.</u>
TREAT.	3	.92	.31	15.5	2.88
POISON	2	1.03	.52	26.0	3.28
T x P	6	.25	.04	2.0	2.38
ERROR	36	.80	.02		
TOTAL	47	3.00			

TABLE 7.3

<u>SOURCE</u>	<u>df</u>	<u>D</u>	<u>"AOV"</u>		<u>5% Pt.</u>
			<u>D/df</u>	<u>F_R</u>	
TREAT.	3	2.9	.97	20.9	2.88
POISON	2	3.6	1.8	38.8	3.28
T x P	6	.85	.14	3.1	2.38
ERROR	36		.046		

ERROR = $\hat{\tau}/2$ IS COMPUTED FROM THE FULL MODEL

FIGURE 7.1
STANDARDIZED RESIDUALS VS. PREDICTED VALUES

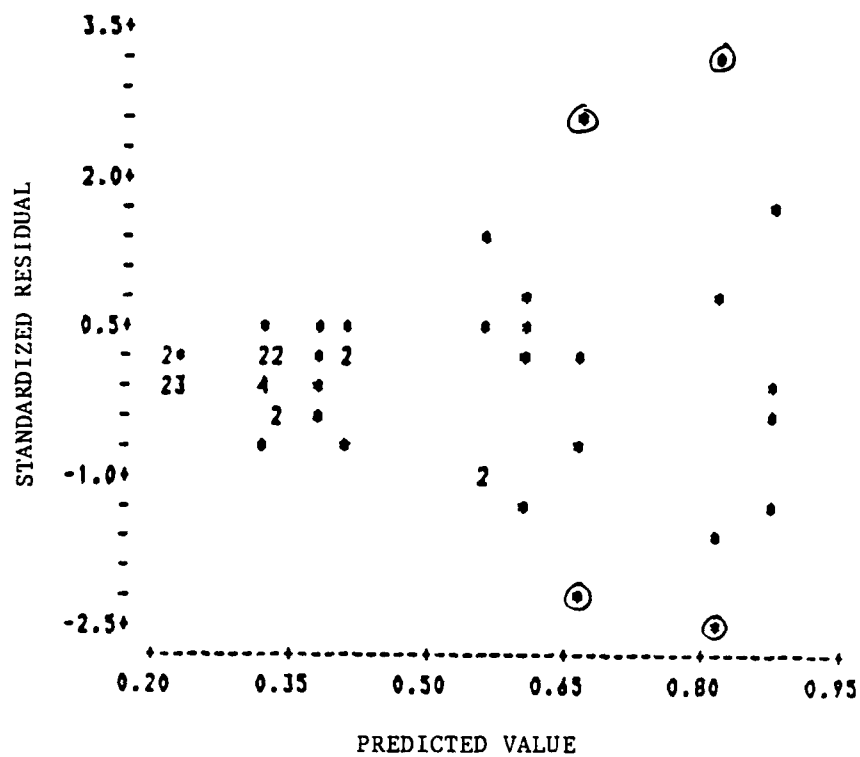


FIGURE 7.2
PSEUDO-OBSERVATIONS VS. OBSERVATIONS

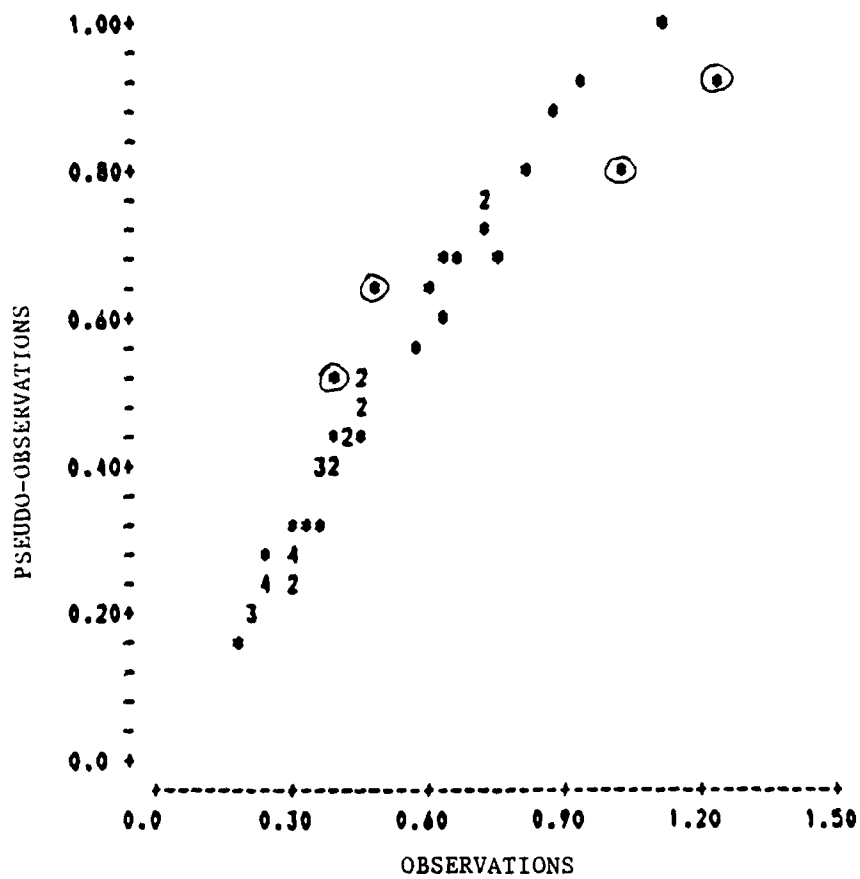
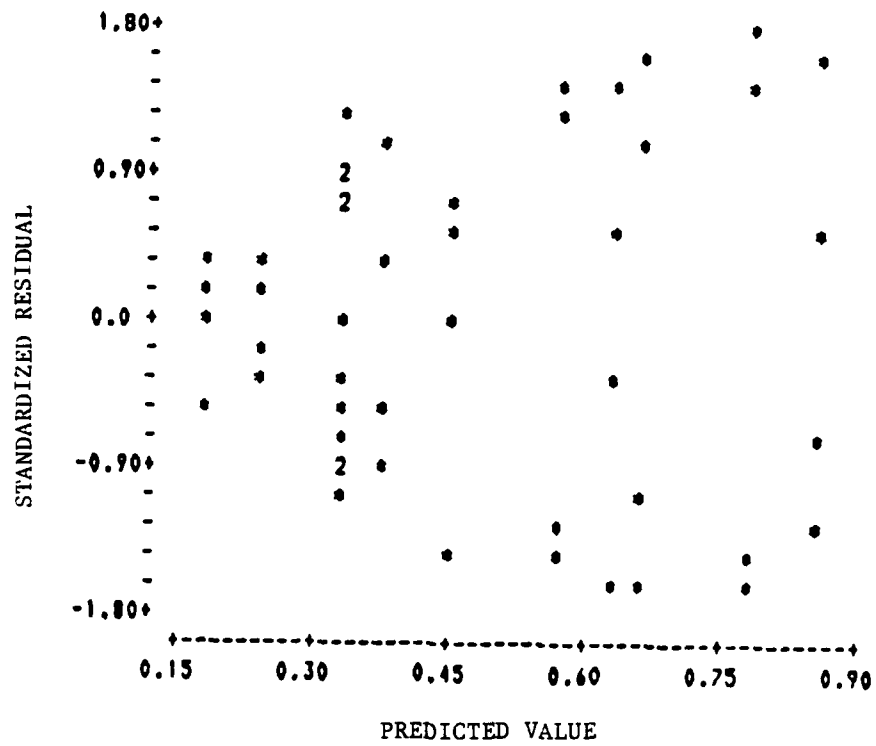


FIGURE 7.3
STANDARDIZED RESIDUALS VS. PREDICTED VALUES
BASED ON PSEUDO-OBSERVATIONS



Unclassified

SECURITY CLASSIFICATION OF THIS PAGE (When Data Entered)

REPORT DOCUMENTATION PAGE		READ INSTRUCTIONS BEFORE COMPLETING FORM
1. REPORT NUMBER Penn State Tech. Rpt. #36	2. GOVT ACCESSION NO. AD-A101 733	3. RECIPIENT'S CATALOG NUMBER
4. TITLE (and Subtitle) A Geometric Interpretation of Inferences Based on Ranks in the Linear Model		5. TYPE OF REPORT & PERIOD COVERED
		6. PERFORMING ORG. REPORT NUMBER
7. AUTHOR(s) Thomas P. Hettmansperger, Penn State University Joseph W. McKean, Western Michigan University		8. CONTRACT OR GRANT NUMBER(s) N00014-80-C-0741
9. PERFORMING ORGANIZATION NAME AND ADDRESS Department of Statistics Penn State University University Park, PA 16802		10. PROGRAM ELEMENT, PROJECT, TASK AREA & WORK UNIT NUMBERS NR 042-446
11. CONTROLLING OFFICE NAME AND ADDRESS Office of Naval Research Statistics and Probability Program Code 436 Arlington, VA 22217		12. REPORT DATE July 1981
		13. NUMBER OF PAGES 39
14. MONITORING AGENCY NAME & ADDRESS (if different from Controlling Office)		15. SECURITY CLASS. (of this report) Unclassified
		15a. DECLASSIFICATION/DOWNGRADING SCHEDULE
16. DISTRIBUTION STATEMENT (of this Report) APPROVED FOR PUBLIC RELEASE: DISTRIBUTION UNLIMITED.		
17. DISTRIBUTION STATEMENT (of the abstract entered in Block 20, if different from Report)		
18. SUPPLEMENTARY NOTES		
19. KEY WORDS (Continue on reverse side if necessary and identify by block number) R-estimates, rank tests, robust inference, nonparametric tests		
20. ABSTRACT (Continue on reverse side if necessary and identify by block number) Four different approaches to testing and estimation based on ranks are unified through the geometry of the linear model. The various tests are identified with various algebraically equivalent forms of the classified F-statistic. Small sample differences are investigated via a Monte Carlo study using both rank and signed rank tests.		

DD FORM 1 JAN 73 1473 EDITION OF 1 NOV 65 IS OBSOLETE
S. N 0102- LF-014-6601

Unclassified
SECURITY CLASSIFICATION OF THIS PAGE (When Data Entered)

