

AD-A131 204

ROBUST ESTIMATES OF ORDERED PARAMETERS(U) MISSOURI
UNIV-ROLLA R MAGEL ET AL. SEP 82 N00014-80-C-0322

1/1

UNCLASSIFIED

F/G 12/1

NL

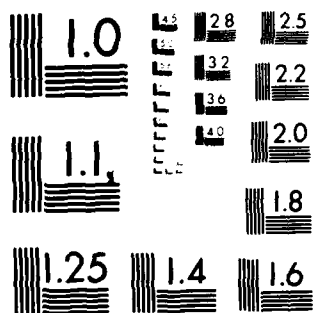
END

DATE

FILED

9 83

DTIC



MICROCOPY RESOLUTION TEST CHART
NATIONAL BUREAU OF STANDARDS 1963-A

12

SECURITY CLASSIFICATION OF THIS PAGE (When Data Entered)

REPORT DOCUMENTATION PAGE		READ INSTRUCTIONS BEFORE COMPLETING FORM
1. REPORT NUMBER	2. GOVT ACCESSION NO.	3. RECIPIENT'S CATALOG NUMBER
	AD-A131	904
4. TITLE (and Subtitle)		5. TYPE OF REPORT & PERIOD COVERED
Robust Estimates of Ordered Parameters		Interim
6. AUTHOR(s)		7. PERFORMING ORG. REPORT NUMBER
Rhonda Magel and F. T. Wright		
8. CONTRACT OR GRANT NUMBER(s)		
N00014-80-C-0322		
9. PERFORMING ORGANIZATION NAME AND ADDRESS		10. PROGRAM ELEMENT, PROJECT, TASK AREA & WORK UNIT NUMBERS
Department of Mathematics and Statistics University of Missouri-Rolla Rolla, MO 65401		
11. CONTROLLING OFFICE NAME AND ADDRESS		12. REPORT DATE
Office of Naval Research Department of the Navy Arlington, VA 22217 -		September, 1982
13. MONITORING AGENCY NAME & ADDRESS (if different from Controlling Office)		14. NUMBER OF PAGES
		18
		15. SECURITY CLASS. (of this report)
		UNCLASSIFIED
		16. DECLASSIFICATION/DOWNGRADING SCHEDULE
17. DISTRIBUTION STATEMENT (of this Report)		
APPROVED FOR PUBLIC RELEASE: DISTRIBUTION UNLIMITED.		
18. DISTRIBUTION STATEMENT (of the abstract entered in Block 20, if different from Report)		
19. SUPPLEMENTARY NOTES		
20. KEY WORDS (Continue on reverse side if necessary and identify by block number)		
Order restricted inference, Robust estimation, Isotonic regression, Cauchy mean value property, Location parameters		
21. ABSTRACT (Continue on reverse side if necessary and identify by block number)		
We consider the estimation of a collection of location parameters when it is believed, <u>a priori</u> , that their ordering is known. The least squares and least absolute deviations estimates subject to this ordering restriction have been studied in the literature. We seek robust estimators which perform well for a broad range of distributions. The results of a Monte Carlo study and a study of computation algorithms are discussed.		

DTIC
ELECTE
AUG 10 1983
S D

AD A 1 3 1 2 0 4

DTIC FILE COPY

DD FORM 1473

EDITION OF 1 NOV 65 IS OBSOLETE

S/N 0102-LF-014-6601

SECURITY CLASSIFICATION OF THIS PAGE (When Data Entered)

88 08 09 00 5

ROBUST ESTIMATES OF ORDERED PARAMETERS⁽¹⁾

Rhonda Magel⁽²⁾ and F. T. Wright

University of Texas-San Antonio and University of Missouri-Rolla

Key Words: Order restricted inference, Robust estimation, Isotonic regression, Cauchy mean value property, Location parameters

ABSTRACT

We consider the estimation of a collection of location parameters when it is believed, a priori, that their ordering is known. The least squares and least absolute deviations estimates subject to this ordering restriction have been studied in the literature. We seek robust estimators which perform well for a broad range of distributions. The results of a Monte Carlo study and a study of computation algorithms are discussed.

(1) This research was sponsored by the Office of Naval Research under ONR contract N00014-80-C-0322.

(2) Parts of this work are taken from this author's doctoral dissertation written at the University of Missouri-Rolla.

Accession For	
NTIS GRA&I	☒
DTIC TAB	☐
Unannounced	☐
Justification	
By	
Distribution/	
Availability, also	
Distribution	
Dist	Special

A



1. INTRODUCTION AND SUMMARY. We consider the estimation of k location parameters when it is believed that they are nondecreasing. Let $\theta_1 \leq \theta_2 \leq \dots \leq \theta_k$ denote the parameters and suppose that independent random samples, X_{ij} , $j=1,2,\dots, n_i$, $i=1,2,\dots,k$, are available. Brunk (1955) obtained the maximum likelihood estimates of nondecreasing normal means, which, of course, minimize

$$\sum_{i=1}^k \sum_{j=1}^{n_i} (X_{ij} - \theta_i)^2 \quad \text{subject to } \theta_1 \leq \theta_2 \leq \dots \leq \theta_k.$$

There are several algorithms for computing these estimates, but we emphasize the pool adjacent violators algorithm (PAVA). (For a detailed discussion of such algorithms, see Section 2.3 of Barlow et al. (1972).) If the sample means, $\bar{X}_i$, are nondecreasing, then they are the restricted least squares estimates. If not there is a violation, that is $\bar{X}_i > \bar{X}_{i+1}$ for some i , then the corresponding samples are pooled and the two sample means are replaced by the mean of the pooled sample. Next the resulting $k-1$ sample means are considered with the understanding that once two samples are pooled they must remain together. This process is continued until a nondecreasing set of means is obtained. If the i th sample has not been pooled with another, then the estimates of θ_i is $\bar{X}_i$, but if the i th sample has been pooled, then the estimate is the mean of the final pooled sample containing the i th sample.

As would be expected, these restricted least squares estimates are unduly affected by extreme observations and so

robust estimators would be desirable in many situations. Robertson and Waltman (1968) derived the least absolute deviations estimates, in particular they obtained the values which minimize

$$\sum_{i=1}^k \sum_{j=1}^{n_i} |x_{ij} - \theta_i| \quad \text{subject to } \theta_1 \leq \theta_2 \leq \dots \leq \theta_k.$$

They also showed that if one uses the sample medians as the initial estimates (in place of $\bar{x}_i$) and computes the median of the pooled samples rather than the mean, then the PAVA can be used to compute these estimates also. It is well-known that the sample median is not very efficient for a normal model, but this does suggest an approach for obtaining robust estimates of the θ_i . Using the notation from Huber (1981), let T be a location statistic, such as the mean, median or an M-estimate, which may be thought of as a functional defined for all empirical distribution functions (EDFs). One could obtain nondecreasing estimates of the θ_i by employing the PAVA and using T to obtain initial estimates as well as the estimates from pooled samples. We denote such estimates by $\bar{\theta}_1 \leq \bar{\theta}_2 \leq \dots \leq \bar{\theta}_k$. (While these estimates depend on T , we do not indicate this in the notation.)

Robertson and Wright (1980) considered computation algorithms for order restricted estimates and found that it is desirable for the estimator, T , to have the Cauchy mean value property (CMVP). T has the CMVP provided its value for a pooled sample is between the T values for the two samples. That is, if F_n and F_m are EDFs based on samples of size n and m respectively, then

$$T(F_n) \wedge T(F_m) \leq T((nF_n + mF_m)/(n+m)) \leq T(F_n) \vee T(F_m).$$

An estimator with this property will be called a Cauchy mean (CM). If T is a CM, then the PAVA yields the same estimate independent of the order in which violators are pooled and these estimates agree with those obtained from the max-min formulas which are discussed there. Robertson and Wright (1974) have shown that if T is a consistent CM, then the estimates obtained from the PAVA also are provided $n_i \rightarrow \infty$ for each i .

In Section 2, several classes of robust estimators are studied to determine if they are CMs. Since so few of those considered in the literature are, the use of non-CMs is discussed further. Section 3 describes the results of a Monte Carlo study of robust estimators for ordered location parameters. If one wants estimators that perform well over the range of normal to Cauchy errors, then Gastwirth's estimator or the trimmed mean which trims 25% on each side are recommended. The trimmed mean is a little more efficient for the normal model and Gastwirth's estimator is more efficient for the Cauchy. The latter is also easier to compute. If very heavy tailed distributions are not a concern, we recommend the Huber with $c = 1.5$ (c is defined in Section 2) which has the advantages of the CMs.

2. ROBUST CMs AND COMPUTATION ALGORITHMS. In this section, we consider the types of location estimators discussed in Andrews et al. (1972) to determine which of those would be

appropriate for computing order restricted estimates.

An M-estimate is a solution, $T_n = T(F_n)$, to an equation of the form

$$\sum_{j=1}^n \psi((x_j - T_n)/s) = 0, \quad (1)$$

where ψ is an odd function and s is estimated independently. (One can also estimate s simultaneously, but, as we shall see, this may yield an estimator which does not have the CMVP.) Hampel and Andrews have proposed some redescending ψ functions (cf. Andrews et al. (1972)), but Hogg (1979) points out that the M-estimates corresponding to these ψ functions, as well as Tukey's biweight, may possess convergence problems when solving iteratively. Since such an iterative procedure must be implemented several times when computing order restricted estimates, these ψ functions were not considered further. Huber (1964) proposed

$$\psi(x) = x \text{ for } |x| \leq c \text{ and } \psi(x) = c \operatorname{sgn}(x) \text{ for } |x| > c, \quad (2)$$

with c a fixed positive constant. A common choice for s is

$$\operatorname{median}(|x_i - \operatorname{median}(x_i)|) / .6754,$$

however, it is not difficult to construct examples to show that if s is computed this way, varying with the sample, then the resulting estimator need not be a CM. (An example is given in Magel (1982).) In many situations, the k populations are assumed to differ only in location and so one could use a fixed value of s , namely,

$$s = \operatorname{median}(|x_{ij} - \operatorname{median}(x_{ij} : j=1, 2, \dots, n_j)| : i=1, 2, \dots, k) / .6754.$$

It is easy to show that the solution set for this ψ must be a nonempty, closed interval (possibly one point) and so we define the estimator to be the midpoint.

Remark. Huber's M-estimator (ψ given by (2)) with a fixed s is a CM.

Proof. Without loss of generality we assume $s = 1$. Let $T_n = T(F_n)$, $T_m = T(F_m)$ and $T_{n+m} = T(F_{n+m})$, where F_n and F_m are the EDFs corresponding to two samples and $F_{n+m} = (nF_n + mF_m)/(n+m)$. Let the solution set for T_n be $[a_1, b_1]$, for T_m be $[a_2, b_2]$ and for T_{n+m} be $[a_3, b_3]$. We also assume that the samples have been labeled so that $a_1 \leq a_2$. Now the left hand side (lhs) of (1) can be written as $n \int \psi(x - T_n) dF_n(x)$ and we see that

$$\begin{aligned} \int \psi(x-t) dF_n(x) \wedge \int \psi(x-t) dF_m(x) &\leq \int \psi(x-t) dF_{n+m}(x) \\ &\leq \int \psi(x-t) dF_n(x) \vee \int \psi(x-t) dF_m(x). \end{aligned} \quad (3)$$

Setting $t = a_3 - \epsilon$, $\epsilon > 0$, the middle term of (3) is positive and so one of the expressions in the rhs of (3) must be positive. Because $a_1 \leq a_2$ this implies that $a_3 - \epsilon \leq a_2$. Letting $\epsilon \rightarrow 0$, yields $a_3 \leq a_2$. Setting $t = a_3$ in the middle term implies that one of the terms in the lhs is nonpositive or $a_3 \geq a_1$. So if T were chosen to be the lh endpoint it would have the CMVP. By symmetry considerations, the same can be shown for the rh endpoints. If $b_1 \leq b_3 \leq b_2$, then $(a_1 + b_1)/2 \leq (a_3 + b_3)/2 \leq (a_2 + b_2)/2$. If $b_2 \leq b_3 \leq b_1$, then $[a_2, b_2] \subset [a_1, b_1]$, in which case $[a_3, b_3] = [a_2, b_2]$, because

$$(n+m) \int \psi(x-t) dF_{n+m}(x) = n \int \psi(x-t) dF_n(x) + m \int \psi(x-t) dF_m(x) = \begin{cases} < 0 & t < a_2 \\ 0 & a_2 \leq t \leq b_2 \\ > 0 & b_2 < t. \end{cases}$$

This remark applies when s is fixed, but in the above definition of Huber's estimate s may vary with the k samples. However, for a fixed set of variables X_{ij} , s has a fixed value and if that value is used to compute the $\bar{\theta}_i$, then the order in which violators are pooled will not affect $\bar{\theta}_i$, the max-min formulae will give the same $\bar{\theta}_i$, and the norm reducing property of Robertson and Wright (1974) will imply that

$$\max_{1 \leq i \leq k} |\bar{\theta}_i - \theta_i| \leq \max_{1 \leq i \leq k} |T_{n_i} - \theta_i|.$$

So if the error distribution, ie. the distribution of $X_{ij} - \theta_i$, is symmetric about zero, then $\bar{\theta}_i \rightarrow \theta_i$ for each i provided $n_i \rightarrow \infty$ for each i .

Another large class of location estimators are the linear combinations of order statistics, or L-estimators. This class includes trimmed means, adaptive trimmed means, Gastwirth's estimator and Tukey's Trimean, as well as the mean and median. Leurgans (1981) has shown that there are only three basic types of L-estimators which possess the CMVP: the mean; weighted midranges, that is a weighted average of the smallest and largest order statistics, $wx_{(1)} + (1-w)x_{(n)}$ where $x_{(1)} \leq x_{(2)} \leq \dots \leq x_{(n)}$ are the order statistics from a sample of size n and $0 \leq w \leq 1$; and weighted percentiles, $x_{([np]+1)}$ if np is not an integer, and $wx_{(np)} + (1-w)x_{(np+1)}$

if np is an integer. Hence, the commonly used robust L-estimators are not CMs. Hogg (1967) proposed an adaptive estimator which uses various L-estimators depending on the value of the sample kurtosis. Two of these, the trimmed mean and the outer mean are not CMs.

Estimators derived from rank tests for a shift are called R-estimators. The Hodges-Lehmann estimator, a popular member of this class, is defined to be the median of pairwise averages, $\text{med}((x_i + x_j)/2)$. Magel (1982) gives an example which shows that this estimator does not possess the CMVP and the same example shows that this is true if one considers all n^2 pairs, just those with $i \leq j$, or just those with $i < j$. The folded medians comprise a closely related class of estimators which depend on the averages of symmetrically placed order statistics, ie. $(x_{(1)} + x_{(n)})/2$, $(x_{(2)} + x_{(n-1)})/2$, etc. The Bickel-Hodges estimator, the median of these numbers, does not have the CMVP (Magel (1982)). It is also shown there that the multiply folded medians (cf. Andrews et al. (1972) for a description) are not CMs.

Magel (1982) has also shown that the skipped means and the one step Hubers do not possess the CMVP. (Andrews et al. (1972) also discuss these estimators.)

Since so few robust estimators of location have the CMVP, we consider algorithms that do not require this property. Leurgans (1982) considered a SLOCOM algorithm with linear functions of order statistics, but noted that the resulting estimators need not be nondecreasing. Robertson and Wright

(1974, Example 5) considered the estimation of ordered modes and observed that the max-min formula applied to consistent estimators which are not CMs may produce estimator which are not consistent. That example also shows that the lower sets algorithm (cf. Barlow et al. (1972)) has this same difficulty. The problem stems from the fact that, with estimators that are not CMs, the initial estimates may be nondecreasing, but yet these algorithms may modify them substantially. This does not happen with the PAVA. In fact, it is clear that if $\theta_1 < \theta_2 < \dots < \theta_k$, if the estimator T is consistent and if $\min_{1 \leq i \leq k} n_i \rightarrow \infty$, then θ_i is consistent for θ_i for $i=1,2,\dots,k$. So we recommend the PAVA when computing order restricted estimates based on initial estimators which are not CMs. However, there is one difficulty with this approach. The estimates may depend on the order in which violators are pooled. In the Monte Carlo study that was conducted, pooling was always from left to right, starting with the first violators, $\bar{X}_i > \bar{X}_{i+1}$, the i th and $i+1$ st samples are pooled, then $(n_i \bar{X}_i + n_{i+1} \bar{X}_{i+1}) / (n_i + n_{i+1})$ is compared with $\bar{X}_{i+2}$, etc.

3. MONTE CARLO RESULTS. Since the small sample properties of order restricted estimators have proved to be quite intractable, even in the case of normal means, a Monte Carlo study was conducted to assess the performance of such estimators for various choices of T . As was mentioned in the last section, even for moderate values of k , several

values of T must be computed when employing the PAVA and so we have not considered some estimators because of computational complexities. We have studied the following T : the mean; median; Tukey's trimean; Gastwirth's estimator (Gast); symmetrically trimmed means, trimming 15% and 25% on each side (Trim 15, Trim 25); and Huber's M-estimator with $c = 1.5$ and 2.0 (H-1.5, H-2.0). We have assumed that the $X_{ij} - \theta_i$ are iid with the distributions studied in Andrews et al. (1972). A key to the distributions used is given in Table I. Many of these distributions are mixtures and $N(0,1)/U(0,1)$ represents the distribution of the ratio of independent variables, one a standard normal and the other a uniform variable on $(0,1)$. For $k = 2, 3, 5$ and various choices of $\theta_1 \leq \theta_2 \leq \dots \leq \theta_k$; $n_1 = n_2 = \dots = n_k = n = 10$ or 20 ; each distribution in Table I; and each T given above, the mean square error (MSE) associated with estimating θ_i , ie. $E(\bar{\theta}_i - \theta_i)^2$, was estimated based on 5000 iterations. The total MSE, $\sum_{i=1}^k E(\bar{\theta}_i - \theta_i)^2$, was estimated by summing these k values and with $k, \theta_1 \leq \theta_2 \leq \dots \leq \theta_k$, and an error distribution fixed, an estimated relative efficiency for a particular T was computed by taking the reciprocal of the ratio of its total MSE to the smallest total MSE for all T in the study. These estimated relative efficiencies are given in Tables II through V.

In Tables II and IV, with k and n fixed, the dispersions in the θ vector, $\theta_k - \theta_1$, are varied. However it is interesting

to note that there are no unusual changes in the efficiencies. This would lead us to believe that the conclusions that follow are valid for a reasonable range of dispersions in θ . (Since the estimators, T , considered are location statistics, the $\bar{\theta}_i - \theta_i$ have distributions that are invariant under shifts by a constant vector.)

First, we consider the estimators in groups. The median, trimean and Gast are easily computed L-statistics. Gast dominates the median for all but the very heavy tailed distributions (11 and in some cases 7). However, the loss in efficiency is at most 10 percentage points in all the cases considered and the gain is as much as 24 percentage points in some light tailed cases. The trimean is on the other extreme, performing even better for light tailed distributions (a gain of 5 percentage points or so in efficiency), but considerably worse for heavy tailed distributions, with a loss of approximately 20 percentage points. In this group we recommend Gastwirth's estimator unless the user is quite certain that very heavy tailed errors are not present, then the trimean could be used.

Andrews et al. (1972) suggest that Hubers, with c in the range we have considered, might be used in practical situations. (They consider the Cauchy distribution to be unreasonable.) In the cases considered here, the efficiency of H-2.0 for a normal model was about 3% more than H-1.5, but that is reversed in even the lightest contamination considered and there is as much as 10 percentage points

difference for moderate tailed errors. So we recommend $H-1.5$ in this family. If very heavy tailed distributions (7,11) can be ruled out, an experimenter may prefer $H-1.5$ over the other T in this study. It holds its efficiency in the normal model, performs well for all but the very heavy tailed distributions and has the advantages of a CM.

Among the trimmed means we recommend a trimming proportion around 25% on each side. The Trim 15 has 7 percentage points more efficiency for the normal model, but the losses for moderate and heavy tailed distributions (4-6 and 8-10) can be far greater.

If one wants an estimator that performs well over the range of situations considered here, then the Trim 25 and the Gast should be considered. The efficiency of the Gast for a normal model is about 4 percentage points smaller, but is up to 7 percentage points larger for the Cauchy distribution (cf. $k=5$). Gastwirth's estimator is easier to compute. If one feels it is not necessary to guard against very heavy tailed distributions (7,11), then $H-1.5$ is recommended. It is about 95% efficient in the normal case and has the advantage of a CM.

While it is not possible to consider enough k and θ to draw definitive conclusions, the fact that the recommendations given above are supported by each case considered strengthens their credibility. Also for each k and θ considered, the Monte Carlo experiment was carried out for $n = 10$ and 20 , and for $k = 5$ additional θ , such

as $(0,0,0,0,1)$ and $(2,4,8,16,32)$, were considered. Each case substantiated the recommendations above.

REFERENCES

- Andrews, D. F., Bickel, P. J., Hampel, F. R., Huber, P. J., Rogers, W. H. and Tukey, J. W. (1972). Robust Estimates of Location. Princeton University Press, Princeton, N.J.
- Barlow, R. E., Bartholomew, D. J., Bremner, J. M. and Brunk, H.D. (1972). Statistical Inferences Under Order Restrictions. Wiley, New York.
- Brunk, H. D. (1955). Estimating monotone parameters. Ann. Math. Statist. 26, 607-616.
- Hogg, R. V. (1967). Some observations on robust estimation. J. Am. Statist. Assoc. 62, 1179-1186.
- Hogg, R. V. (1979). Statistical robustness: One view of its use in applications today. Am. Statistician 33, 108-115.
- Huber, P. J. (1964). Robust estimation of a location parameter. Ann. Math. Statist. 35, 73-101.
- Huber, P. J. (1981). Robust Statistics. Wiley, New York.
- Leurgans, Sue E. (1981). The Cauchy mean value property and linear functions of order statistics. Ann. Statist. 9, 905-908.
- Leurgans, Sue E. (1982). Asymptotic distributions of slope-of-greatest-convex-minorant-estimators. Ann. Statist. 10, 287-296.
- Magel, Rhonda (1982). Topics in Isotonic Regression. Ph.D. Dissertation, University of Missouri-Rolla.
- Robertson, Tim and Waltman, P. (1968). On estimating monotone parameters. Ann. Math. Statist. 39, 1030-1039.
- Robertson, Tim and Wright, F. T. (1974). A norm reducing property for isotonized Cauchy mean value functions. Ann. Statist. 2, 1302-1307.
- Robertson, Tim and Wright, F. T. (1980). Algorithms in order restricted statistical inference and the Cauchy mean value property. Ann. Statist. 8, 645-651.

TABLE I. KEY TO THE ERROR DISTRIBUTION USED

1. $N(0,1)$
2. 90% $N(0,1)$ and 10% $N(0,9)$
3. Double Exponential
4. 50% $N(0,1)$ and 50% $N(0,9)$
5. 25% $N(0,1)$ and 75% $N(0,9)$
6. 90% $N(0,1)$ and 10% $N(0,100)$
7. 75% $N(0,1)$ and 25% $N(0,100)$
8. 90% $N(0,1)$ and 10% $N(0,1)/U(0,1)$
9. 75% $N(0,1)$ and 25% $N(0,1)/U(0,1)$
10. 90% $N(0,1)$ and 10% $N(0,1)/U(0,\frac{1}{3})$
11. Cauchy

TABLE II: RELATIVE EFFICIENCIES: $k=2$, $n=20$ $\theta = (0,0)$

Distribution	Mean	Median	Trimean	Gast	H-1.5	H-2.0	Trim 15	Trim 25
1	1.000	.675	.865	.819	.959	.987	.916	.845
2	.730	.747	.944	.899	1.000	.969	.989	.951
3	.503	.639	.954	.883	.953	.825	1.000	.925
4	.642	.914	.928	.986	.856	.773	.901	1.000
5	.871	.937	.955	.987	.929	.904	.948	1.000
6	.136	.827	.996	.979	.984	.882	.984	1.000
7	.100	.968	.724	1.000	.763	.584	.566	.984
8	<.1	.770	.950	.906	1.000	.973	.988	.958
9	<.1	.828	.987	.965	.982	.900	1.000	.991
10	<.1	.808	.987	.955	.991	.904	1.000	.984
11	<.1	1.000	.696	.904	.647	.525	.591	.892

 $\theta = (-.3, .3)$

1	1.000	.703	.882	.838	.960	.988	.926	.859
2	.751	.769	.953	.914	1.000	.971	.991	.959
3	.528	.683	.964	.927	.967	.860	1.000	.955
4	.671	.931	.936	.987	.871	.795	.909	1.000
5	.884	.948	.963	.989	.936	.914	.954	1.000
6	.158	.845	.988	.982	.977	.882	.977	1.000
7	.116	.982	.744	1.000	.785	.614	.589	.982
8	<.1	.782	.951	.914	1.000	.974	.988	.959
9	<.1	.845	.991	.970	.987	.911	1.000	.990
10	<.1	.839	.993	.971	.994	.918	1.000	.991
11	<.1	1.000	.715	.901	.673	.556	.615	.889

TABLE III: RELATIVE EFFICIENCIES: $k=3$, $\theta=(0,0,1)$

n = 10								
Distribution	Mean	Median	Trimean	Gast	H-1.5	H-2.0	Trim 15	Trim 25
1	1.000	.729	.892	.833	.954	.981	.939	.870
2	.762	.810	.955	.921	1.000	.978	.985	.899
3	.539	.800	.974	.962	1.000	.890	.941	.970
4	.723	.954	.945	1.000	.917	.851	.920	.992
5	.893	.942	.958	.985	.949	.928	.975	1.000
6	.173	.893	.931	.984	1.000	.978	.809	.981
7	.128	1.000	.576	.915	.767	.610	.436	.764
8	<.1	.811	.956	.912	1.000	.986	.957	.938
9	<.1	.883	.968	.976	1.000	.935	.842	.997
10	<.1	.871	.954	.968	1.000	.936	.844	.971
11	<.1	1.000	.638	.904	.713	.597	.423	.825
n = 20								
1	1.000	.682	.873	.820	.961	.991	.932	.860
2	.736	.761	.949	.910	1.000	.971	.991	.938
3	.511	.667	.978	.914	.997	.858	1.000	.934
4	.667	.914	.923	.982	.863	.787	.904	1.000
5	.837	.919	.923	.949	.894	.868	.950	1.000
6	.163	.841	.992	.978	1.000	.904	.984	.997
7	.121	.961	.772	1.000	.806	.633	.585	.963
8	<.1	.760	.944	.895	1.000	.981	.962	.913
9	<.1	.783	.929	.906	.939	.870	1.000	.988
10	<.1	.808	.974	.960	1.000	.927	.996	.915
11	<.1	1.000	.724	.919	.687	.572	.617	.886

TABLE IV: RELATIVE EFFICIENCIES, $k=3$, $n=20$

$$\theta = (-2, 0, 2)$$

Distribution	Mean	Median	Trimean	Gast	H-1.5	H-2.0	Trim 15	Trim 25
1	1.000	.681	.873	.824	.962	.990	.928	.858
2	.734	.762	.951	.913	1.000	.970	.991	.959
3	.513	.681	.979	.932	1.000	.865	.987	.940
4	.647	.905	.909	.965	.850	.772	.899	1.000
5	.855	.934	.950	.971	.915	.888	.952	1.000
6	.147	.845	.984	.981	.993	.895	.965	1.000
7	.116	.989	.706	1.000	.807	.622	.529	.957
8	<.1	.759	.943	.899	1.000	.982	.972	.922
9	<.1	.804	.954	.929	.963	.891	1.000	.991
10	<.1	.818	.980	.959	1.000	.924	.986	.978
11	<.1	1.000	.694	.903	.680	.560	.590	.888

$$\theta = (-.33, -.09, .33)$$

1	1.000	.714	.885	.840	.963	.988	.937	.872
2	.762	.785	.955	.919	1.000	.974	.991	.944
3	.555	.699	.980	.932	1.000	.875	1.000	.944
4	.677	.926	.932	.989	.871	.796	.908	1.000
5	.825	.911	.913	.943	.883	.856	.948	1.000
6	.171	.858	.995	.984	1.000	.913	.981	.987
7	.118	.963	.795	1.000	.815	.646	.602	.970
8	<.1	.785	.950	.908	1.000	.927	.955	.912
9	<.1	.808	.944	.922	.953	.889	1.000	.989
10	<.1	.834	.983	.967	1.000	.933	1.000	.990
11	<.1	1.000	.728	.923	.688	.574	.615	.876

TABLE V: RELATIVE EFFICIENCIES, $k=5$, $n=10$

$$\theta = (-.8, -.3, 0, .9, 1.5)$$

Distribution	Mean	Median	Trimean	Gast	H-1.5	H-2.0	Trim 15	Trim 25
1	1.000	.757	.903	.849	.962	.985	.949	.889
2	.782	.828	.958	.924	1.000	.982	.987	.954
3	.512	.668	.979	.916	1.000	.861	.954	.990
4	.725	.962	.936	.991	.911	.846	.925	1.000
5	.884	.945	.951	.973	.938	.916	.971	1.000
6	.197	.901	.951	.978	1.000	.926	.831	.975
7	.145	1.000	.645	.945	.823	.672	.509	.854
8	<.1	.826	.957	.921	1.000	.988	.961	.940
9	<.1	.888	.967	.971	.987	.929	.920	1.000
10	<.1	.878	.948	.961	1.000	.945	.886	.987
11	<.1	1.000	.724	.919	.687	.572	.430	.856

$$\theta = (-2.2, -.7, 0, 0, .7)$$

1	1.000	.742	.896	.839	.958	.985	.938	.875
2	.776	.826	.955	.925	1.000	.983	.982	.953
3	.543	.796	.954	.938	1.000	.891	.910	.942
4	.723	.959	.935	.990	.912	.846	.921	1.000
5	.889	.947	.956	.980	.942	.921	.975	1.000
6	.192	.888	.942	.973	1.000	.924	.801	.958
7	.147	1.000	.630	.936	.823	.670	.500	.846
8	<.1	.820	.955	.917	1.000	.989	.957	.943
9	<.1	.878	.957	.960	.978	.920	.912	1.000
10	<.1	.871	.943	.958	1.000	.941	.857	.969
11	<.1	1.000	.642	.912	.743	.632	.431	.834

DATE
ILME