

22948865

MICHIGAN STATE UNIVERSITY LIBRARIES



3 1293 00600 4661



This is to certify that the

dissertation entitled

An Empirical Study of the Type I Error Rate
and Power for Some Selected Normal-Theory
and Nonparametric Tests of the Independence
of Two Sets of Variables

presented by

Abdul Razak Habib

has been accepted towards fulfillment
of the requirements for

Ph. D. degree in Education

Betty Jane Becker
Major professor

Date August 10th, 1988

PLACE IN RETURN BOX to remove this checkout from your record.
TO AVOID FINES return on or before date due.

DATE DUE	DATE DUE	DATE DUE
JUL 28 1992		
Jan 8, 1992		
Feb 6, 1992		
May 4 1992 (111)		
JUN 9 1992		

MSU Is An Affirmative Action/Equal Opportunity Institution

AN EMPIRICAL STUDY OF THE TYPE I ERROR RATE AND POWER
FOR SOME SELECTED NORMAL-THEORY AND NONPARAMETRIC TESTS
OF THE INDEPENDENCE OF TWO SETS OF VARIABLES

By

Abdul Razak Habib

A DISSERTATION

Submitted to
Michigan State University
in partial fulfillment of the requirements
for the degree of

DOCTOR OF PHILOSOPHY

Department of Counseling, Educational
Psychology and Special Education

1988

5368819

ABSTRACT

AN EMPIRICAL STUDY OF THE TYPE I ERROR RATE AND POWER FOR SOME SELECTED NORMAL-THEORY AND NONPARAMETRIC TESTS OF THE INDEPENDENCE OF TWO SETS OF VARIABLES

By

Abdul Razak Habib

The present study empirically examined the effect of non-normality, sample size, number of variables, and degree of dependency on the Type I error and power properties of five normal-theory and nonparametric tests of the independence of two sets of variables. Simulated data representing light-, moderately heavy-, and heavy-tailed distributions, three sample sizes, three sets of correlations-among-variables, and three sets of numbers-of-variables were included.

This study yielded the following results. The Type I error rates of the normal-theory Bartlett and Rao F tests increase substantially for the moderately heavy- and heavy-tailed distributions, whereas the Type I error rates of the nonparametric rank-transform Rao F and the pure- and mixed-rank tests are not affected by the form of a parent distribution for moderately-small and moderately-large samples. The Type I error rates of the Bartlett and Rao F tests increase with increases in the correlation among predictor and/or dependent variables, and with increases in the number-of-variables for heavy-tailed distributions. The Type I error rate of the rank- transform Rao F test is not affected by the within-set-correlation and the number-of-variables factors, while those of the pure- and mixed-rank tests are not affected by the within-set-correlation factor but decrease as the number of variables increases for all distributions.

The power values of the normal-theory Bartlett and Rao F tests increase substantially only for extremely heavy-tailed distributions. The power values of all three nonparametric tests increase with increases in the kurtosis values. The power values of all five tests increase with increases in the sample size and the correlation among the predictor variables, and decrease with increases in the correlation among the dependent variables for all distributions. The increments due to the sample size are higher for the three nonparametric tests. The power values of the Bartlett, Rao F , and the rank-transform Rao F tests are not affected by the number of variables, while those of the pure- and mixed-rank tests decrease as the number of variables increases for all distributions. However, the reduction in the power values tends to be compensated for by increases in the sample size.

To Faisal and Hilmy

ACKNOWLEDGEMENTS

I wish to express my appreciation to all the individuals who made the completion of this program possible. To the members of my doctoral committee: Dr. Betsy J. Becker (Chairperson), Dr. Dennis C. Gilliland, Dr. Michael R. Harwell, Dr. Perry E. Lanier, and Dr. Stephen W. Raudenbush for their constant encouragement and valuable suggestions and comments.

I wish to extend my appreciation to the Vice Chancellor of the Malaysian National University (MNU) for granting the study leave; to members of the Faculty of Education at MNU for carrying extra work during my absence; to the Chairpersons (past and present) and faculty members of the Department of Counseling, Educational Psychology, and Special Education at Michigan State University (MSU) for hiring me as a teaching and research assistant, research consultant, and instructor; to the President of the Sage Foundation and the MSU Coordinator of the Sage Foundation Fund for awarding me a dissertation grant; and to all friends who supported me financially and emotionally during my struggle to complete this program.

Finally, I wish to express my deep gratitude to members of my family: Pn. Rafeah, En. Habib, Pn. Embun, Pn. Fatimah, Pn. Maimun, Ad. Mahadzir, Ad. Yahya, Ad. Rohani, Ad. Bashah, Ad. Basuri, Ad. Nazri, Ad. Salmah, Ad. Rokiah, Ad. Habibah, Ad. Mohammad Nasir, An. Faisal, and An. Hilmy for their continued love.

May God blesses everyone.

TABLE OF CONTENT

	Page
LIST OF TABLES	ix
LIST OF FIGURES	xi
 CHAPTER	
I. STATEMENT OF THE PROBLEM	1
Role of Multivariate Analysis in Educational Research	1
Normal-Theory and Nonparametric-Multivariate Tests in Educational Research	4
Purpose of the Study	6
Factors Which Influence the Choice of a Normal-Theory or Nonparametric-Multivariate Test	7
Research Questions and Hypotheses	10
Role of Simulation in Distributional Studies	13
Significance of the Study	13
Limitations of the Study	15
II. REVIEW OF THE LITERATURE	16
Comparing the Normal Distribution and Some Non-normal Distributions	16
Methods of Dealing With Non-Normal Data	20
Nonparametric Methods	20
Rank-Transform Methods	22
Robustness and Power Results for Some Normal-Theory and Nonparametric-Multivariate Tests	23
Analytic Results - Univariate Case	24
Empirical Results - Univariate Case	25

Empirical Results - Normal-Theory-Multivariate Case . . .	28
Analytic Results - Nonparametric-Multivariate Case . . .	30
Empirical Results - Nonparametric-Multivariate Case . . .	31
Dependency Among Variables	35
Summary	35
III. METHODOLOGY	37
Multivariate Statistical Models	37
Test Statistics and Their Assumptions	41
Bartlett Statistic	42
Rao F Statistic	43
Rank-Transform Rao F Statistic	44
Pure-Rank Statistic	44
Mixed-Rank Statistic	47
Data Generation	48
Simulation Conditions	52
Presentation of Simulation Results	54
IV. RESULTS	57
Characteristics of the Generated Data	57
Type I Error and Power Conditions	62
Overall Type I Error and Power Results	63
Main Effects of Simulation Conditions	66
Distribution	66
Sample Size	69
Within-Set Correlation	72
Number-of-Variables	74

Interaction Effects of Simulation Conditions	76
Distribution and Sample Size	76
Distribution and Within-Set Correlation	81
Distribution and Number-of-Variables	85
Sample Size, Within-Set Correlation, and Number-of- Variables	89
V. SUMMARY	92
Research Questions	92
Methodology	92
Findings	94
Conclusions	96
Implications for Data Analysis in Educational Research . .	98
Recommendations for Further Study	100
APPENDICES	103
A. Definition of Terms	103
B. Procedures and Algorithms	105
C. Pure-Rank Statistic	113
D. Computer Programming	117
E. Tables	127
BIBLIOGRAPHY	150

LIST OF TABLES

Table	Page
1. Skewness and Kurtosis Values of Some Univariate Distributions	18
2. Multivariate Skewness and Kurtosis Values of Some Bivariate Distributions	19
3. Asymptotic Relative Efficiencies of Some Nonparametric and Normal-Thery Tests of Location	24
4. Simulation Factors	52
5. Average Mean, Variance, Skewness, Kurtosis, and Within-Set Correlations of the Simulated Data	59
6. Sample Measures of Multivariate Skewness and Kurtosis For the Normal Distribution	61
7. Regression Coefficients Used for Power Simulations and the Resulting Between-Set Correlations	62
8. Overall Empirical Type I Error and Power Values	64
9. Overall Number of Conservative and Liberal Type I Errors	66
10. Average Type I Error and Power Values and Number of Conservative and Liberal Error Rates by Distribution	67
11. Average Type I Error and Power Values and Number of Conservative and Liberal Error Rates by Sample Size	70
12. Average Type I Error and Power Values and Number of Conservative and Liberal Error Rates by Within-Set Correlation	72
13. Average Type I Error and Power Values and Number of Conservative and Liberal Error Rates by Number-of-Variables.	74
14. Average Type I Error and Power Values and Number of Conservative and Liberal Error Rates by Distribution and Sample Size	78
15. Average Type I Error and Power Values and Number of Conservative and Liberal Error Rates by Distribution and Within-Set Correlation	82

16. Average Type I Error and Power Values and Number of Conservative and Liberal Error Rates by Distribution and Number-of-Variables	86
17. Average Type I Error and Power Values and Number of Conservative and Liberal Error Rates by Sample Size and and Within-Set Variables For Rao F Test	89
18. Average Type I Error and Power Values and Number of Conservative and Liberal Error Rates by Sample Size and Number-of-Variables For Rao F and Pure- and Mixed-Rank Tests	90

LIST OF FIGURES

Figure	Page
1. Frequency-Distribution Curves of Simulated Data.	58
2. Type I Error and Power Curves by Alpha Level	64
3. Type I Error and Power Curves by Distribution	68
4. Type I Error and Power Curves by Sample Size	71
5. Type I Error and Power Curves by Within-Set Correlation . .	73
6. Type I Error and Power Curves by Number-of-Variables	75
7. Type I Error Curves by Distribution and Sample Size	79
8. Power Curves by Distribution and Sample Size	80
9. Type I Error Curves by Distribution and Within-Set Correlation	83
10. Power Curves by Distribution and Within-Set Correlation . .	84
11. Type I Error Curves by Distribution and Number-of-Variables	87
12. Power Curves by Distribution and Number-of-Variables. . . .	88

CHAPTER I

STATEMENT OF THE PROBLEM

The present study used computer-simulated data to assess the distributional behavior (i.e., Type I error rate and power) of selected normal-theory and nonparametric tests of the independence of two sets of variables. The focus of the investigation was the behavior of the tests in the presence of non-normal skewness and kurtosis values. This chapter discusses the (a) role of multivariate analysis in educational research, (b) normal-theory and nonparametric-multivariate tests in educational research, (c) purpose of the study, (d) factors which influence the choice of a normal-theory or nonparametric-multivariate test, (e) research questions and hypotheses, (f) role of simulation in distributional studies, (g) significance of the study, and (h) limitations of the study. The definitions of some statistical terms are given in Appendix A.

Role of Multivariate Analysis in Educational Research

Multivariate analysis refers to a collection of descriptive and inferential methods that have been developed for situations where one or more sets of correlated variables are treated as outcome measures, predictors, or both (Harris, 1975, p. 5). More specifically, multivariate methods allow researchers to simultaneously analyze the interrelationships among many variables. In contrast, univariate analyses are carried out separately for each outcome variable. One potential shortcoming of univariate methods is that they may lead to an

incomplete description of the data since they ignore interrelationships among predictor and outcome variables.

Multivariate methods have found widespread use in educational research. A primary reason for their popularity is the interest educational researchers show in testing theories that are multivariate in character, which implies the use of multiple variables. Because these variables are chosen to be consistent with the theory under test, they form a multidimensional system and are expected to be correlated (Takeuchi, Yanai, & Mukherjee, 1982, p. 54). Testing theories by collecting data on several variables leads quite naturally to multivariate data-analytic methods.

Among the inferential multivariate methods used in educational research are multivariate analysis of variance (MANOVA), factor analysis, discriminant analysis, canonical-correlation analysis, and multivariate-multiple-regression. These methods have served as important explanatory tools for researchers attempting to summarize the information in a data set containing multiple (correlated) outcome variables.

As noted above, many studies in education involve the analysis of relationships between multiple outcome and predictor variables. Canonical-correlation analysis and multivariate-multiple-regression represent general data-analytic methods that may be used to study such relationships. The fundamental difference between the two approaches lies in the nature of the measured relationship. Canonical-correlation analysis assesses the degree of relationship among two sets of random variables (Takeuchi et al., 1982, p. 225). Although researchers often

refer to one set of variables as predictors and the other as outcomes, the mathematical model underlying canonical correlation makes no such distinction (Gittins, 1985, p. 19).

The multivariate-multiple-regression model, on the other hand, simultaneously assesses the degree of relationship between each of the random outcome variables and the set of fixed and known predictor variable values (Takeuchi et al., 1982, p. 116). However, predictors are rarely fixed and known in practice and regression analysis is routinely performed for predictors that, in essence, are random variables. An important consequence of this practice is that data-analytic inferences are limited to predictor values appearing in the sample (Rogosa, 1980).

As an example of the differing applications of canonical-correlation and multivariate-multiple-regression in educational research, consider a study of the relationship between school organizational climate and teacher job satisfaction. Two well known instruments in this area are the Teacher Job Satisfaction Questionnaire (Lester, 1983), which measures nine identified factors of teacher job satisfaction, and the Organizational Climate Description Questionnaire (Kottkamp, Mohlern, & Hoy, 1985), which measures five dimensions of organizational climate. If the research question focuses on the interdependence between these two sets of (random) variables, the relationship between the job satisfaction factors and the organizational climate dimensions is most properly examined using canonical-correlation analysis. If one set of variables is conceptualized as outcomes and the other as predictors, then

multivariate-multiple-regression would be appropriate. It is important to emphasize that these models are conceptually different, yet are identical with respect to making statistical inferences about the relationship between two sets of variables. Both of these procedures are important explanatory tools for educational researchers interested in testing theories that are multivariate in character.

Normal-Theory and Nonparametric-Multivariate Tests in Educational Research

Historically, researchers opting for multivariate methods have been confronted with the problem of fitting the observed data into the framework of multivariate-normal-theory procedures. Such methods have collectively been labelled parametric, and are identified by their reliance on the assumption that the population distribution of observations follows a multivariate-normal density function (Puri & Sen, 1971, p. 1). Yet in many data-analytic situations there is little doubt that the observations can be characterized as moderately or even distinctly non-normal (Puri & Sen, 1971, p. 1). Under the assumption of random sampling from a specified population, this casts doubt on the normality of the population distribution. For example, educationally-oriented variables such as the number of days absent from school are likely to produce (non-normal) data that are badly skewed or slightly or heavily kurtic.

One approach for dealing with non-normality is to transform the original data to a form more closely resembling a normal distribution (see Box & Cox, 1964) and then employ normal-theory methods. This requires that the underlying distribution of the original variable(s)

be known before deciding which transformation is most appropriate. In many cases, the distribution of the original variable is not known, and transformations of this type may be problematic (Kendall & Stuart, 1969, V. 2, p. 487). Another issue is that the transformed variable may not be interpretable.

A second approach is to transform the original data to ranks (or some other monotonic transformation) and then employ nonparametric methods. These methods do not require that the form of the underlying distribution be known, and are characterized by their relaxation of the normality assumption. However, the underlying distribution must be continuous (Kendall & Stuart, 1969, V. 2, p. 487).

In surveying the literature, a number of nonparametric alternatives to normal-theory, univariate methods are available (e.g., Conover, 1980; Gibbons, 1971; Marascuilo & McSweeney, 1977). This is not true for the multivariate case, where nonparametric alternatives to normal-theory multivariate methods exist only in certain areas of statistical inference.

A primary source of the development of nonparametric-multivariate methods is the work of Puri and Sen (1969, 1971, 1985). Of special importance are the tests these authors generated for hypotheses subsumed under the multivariate general linear model. One is the pure-rank procedure, in which values of all variables are ranked prior to any analysis, and the other is the mixed-rank procedure, in which some but not all variables are ranked. Another nonparametric approach that is closely related to classical nonparametric methods is the rank-transform procedure due to Iman and Conover (Conover & Iman, 1981;

Iman, 1974b; Iman & Conover, 1979). This procedure, which likewise does not require a normality assumption, involves transforming the original data to their corresponding ranks and then applying normal-theory procedures. The pure- and mixed-rank procedures as illustrated by Puri and Sen, and the rank-transform procedure of Conover and Iman, represent the primary nonparametric alternatives to normal-theory multivariate analysis.

Purpose of the Study

The purpose of the present study was to compare the Type I error and power properties of two normal-theory (Bartlett, Rao F) and three nonparametric (rank-transform Rao F , pure-and mixed-rank) tests of the independence of two sets of variables (i.e., tests of no canonical correlation or no regression). The distributional properties of these tests hold exactly only for the asymptotic case (i.e., for very large samples and/or a parent normal distribution), and hence the focus was on the behavior of these tests for small samples under a variety of non-normal skewness and kurtosis conditions.

Other factors examined included the sample size, the correlation within the set of predictors and within the set of dependent variables, the correlation among the sets of predictor and dependent variables, and the number of variables. Since the effects of such factors are difficult to evaluate analytically (Ito & Schull, 1964; Zwick, 1984, p. 2), a simulation study was performed to investigate the behavior of the tests. It is anticipated that the results of the present study will provide educational researchers with guidelines for choosing

between normal-theory tests of the hypothesis of no relationship among two sets of variables and their nonparametric counterparts under a variety of non-normal data and sample-size conditions.

Factors Which Influence the Choice of a Normal-Theory
or Nonparametric-Multivariate Test

As a result of theoretical and computational advances a variety of normal-theory and nonparametric multivariate tests are available to educational researchers. The question arises of how best to choose among the two kinds of tests. The application of a normal-theory procedure to test a statistical hypothesis requires some statistical assumptions on the observations and the population distribution. For example, the omnibus test that all population regression coefficients equal zero in multivariate-multiple-regression assumes that (a) the population of outcomes, conditional on the predictors, is normally distributed, (b) the outcomes, conditional on the predictors, have a common covariance matrix, and (c) the residuals for a given outcome variable are independent.

Violations of one or more of these assumptions have been shown to have adverse effects on the Type I error probability and power of normal-theory multivariate tests under a variety of data conditions (Ito & Schull, 1964; Mardia, 1971; Olson, 1976). Thus, the use of a normal-theory test when assumptions are violated may lead to an incorrect conclusion. For example, a true statistical hypothesis may be rejected, not because the statistical hypothesis is false but because one or more of the underlying statistical assumptions are violated (Conover, 1980, p. 84). Thus, the effect of violating underlying

statistical assumptions is an important factor in choosing a normal-theory or nonparametric-multivariate test.

It is important to emphasize that some violations of the assumptions underlying a statistical test will always occur. This points to a need for criteria that define the "best" test with respect to distributional properties when underlying assumptions are violated. Gibbons (1971, p. 16) defines the "best" test as the test which is most successful in correctly distinguishing between the conditions as stated in the null and alternative hypotheses. An equivalent and more technical definition of the "best" test is given by Ito (1980), who argues that it is the one which is robust (i.e., insensitive to the violation of test assumptions) with respect to the Type I error probability and also most powerful among its competitors.

Unfortunately, the search for the "best" test is complicated by the variety of assumption violations that can affect the distributional properties of a test. Fundamental to the comparison of normal-theory and nonparametric tests is the assumption of normality. As noted earlier, the application of a nonparametric procedure in testing statistical hypotheses does not require a normality assumption, and hence the choice of a normal-theory or nonparametric test in general depends on the tenability of the normality assumption. It is important to emphasize that this is the fundamental difference between the normal-theory and nonparametric tests considered in the present study. While nonparametric methods may be less sensitive to other assumption violations than their normal-theory counterparts (e.g., heteroscedasticity of variance), it is primarily the lack of a normality requirement

that distinguishes the two methodologies, and serves as an important factor influencing the choice of one kind of test rather than another.

Factors other than the tenability of the normality assumption also influence the Type I error rate and power of tests, and hence should be considered in the choice of a normal-theory or nonparametric test. These include the number of variables, their degree of dependency, and the sample size.

As noted earlier it would be desirable to analytically examine the effects of all of the above factors on the distributional properties of these tests. Such analyses are extremely difficult if not impossible in the multivariate case because the analytic methods rely on specific statistical assumptions about the underlying distributions and on the asymptotic distribution of sample statistics. Hence investigations of the effects of these factors on normal-theory and nonparametric-multivariate tests have primarily been empirical.

The results of a number of studies suggest that the form of the underlying distribution plays a key role on the Type I error rate and power performance of multivariate tests (Arnold, 1964; Chase & Bulgren, 1971; Davis, 1982b; Harwell & Serlin, 1985; Mardia, 1970). The number of dependent variables has also been found to affect the distributional behavior of multivariate tests (Ito, 1980; Olson, 1974), in that the tests tend to become less robust as the number of outcome variables increases. This result may be explained by the fact that the degree of non-normality in the joint distribution of the dependent variables is likely to increase as their number increases (Puri & Sen, 1971, p. 2).

As implied by Puri and Sen (1971, p. 176), the degree and pattern of dependency among variables would also be expected to affect the power of nonparametric-multivariate tests. For example, high correlations among variables would (other factors being equal) tend to produce less powerful tests than low correlations among variables. Sample size is also an important factor influencing the choice of a normal-theory or nonparametric test. This occurs because most of the normal-theory multivariate methods are based on sampling distributions of test statistics that are derived from large samples. However, the same is true for nonparametric tests, and hence both normal-theory and nonparametric tests are expected to be less robust for small samples.

In addition to the form of the underlying distribution, the number of variables, their degree of dependency, and the sample size all have been cited as influencing the choice of a normal-theory or nonparametric test. Consequently, any investigation of normal-theory and nonparametric-multivariate tests should consider these factors.

Research Questions and Hypotheses

In comparing the Type I error and power properties of normal-theory and nonparametric tests of the independence of two sets of variables, special attention was given to the influence of skewness and kurtosis. Such attention is justified by previous research in both the univariate and multivariate cases, and has indicated the importance of these two characteristics in the performance of a test (Chase & Bulgren, 1971; Harwell & Serlin, 1985; Mardia, 1970; Olson, 1976). The questions of particular interest and their related hypotheses are the

following:

1. Do the skewness and kurtosis values affect the Type I error and power values of normal-theory and nonparametric tests? Previous empirical results suggest that increasing skewness results in normal-theory tests that become liberal (i.e., rejecting a true null hypothesis more often than expected), while increasing kurtosis results in tests that become conservative (i.e., rejecting a true null hypothesis less often than expected) (Chase & Bulgren, 1971; Mardia, 1970; Olson, 1974). Increasing skewness or kurtosis may also reduce the power of normal-theory tests (Harwell & Serlin, 1985; Olson, 1974). Such effects would not be expected for nonparametric tests, since these procedures in general do not depend on the form of the underlying distribution.
2. Does sample size influence the effects of skewness and kurtosis on the Type I error and power values of normal-theory and nonparametric tests? The sampling distributions of most normal-theory and nonparametric test statistics are derived for large samples. Thus, for small samples neither normal-theory or nonparametric tests would be expected to be robust with respect to Type I error rate and power. In particular, for small and moderate samples departures from normality would be expected to noticeably affect the Type I error rate and power of the normal-theory tests (Olson, 1974; Zwick, 1984, p. 2). Similarly, the nonparametric tests would not be expected to perform well for small samples, but would be expected to do well for moderate samples. The power values

of all tests would be expected to increase with increases in the sample size.

3. Does the degree of dependency among variables influence the effects of skewness and kurtosis on the Type I error and power values of normal-theory and nonparametric tests? As implied by Puri and Sen (1971, p. 176), the degree of dependency among variables would be expected to affect the power of some multivariate tests. However, the results of the study by Harwell and Serlin (1985) suggested that in general the degree of dependency among variables would not affect the power of normal-theory tests, although for extremely skewed data a high degree of dependency among variables tended to slightly reduce the power of nonparametric tests. Thus, a high degree of dependency among variables might be expected to slightly reduce the power of the nonparametric tests for extremely skewed data.
4. Does the number of variables influence the effects of skewness and kurtosis on the Type I error and power values of normal-theory and nonparametric tests? Previous empirical results suggest that normal-theory procedures become less robust with respect to Type I errors and power as the number of outcome variables increases (Ito, 1980; Olson, 1974). Puri and Sen (1971, p. 2) pointed out that as the number of variables increases, the degree of non-normality of their joint distribution might be expected to increase simply because of the increase in the dimensionality of the distribution.

Role of Simulation in Distributional Studies

It has been pointed out that analytic studies of the Type I error probability and power efficiency of multivariate tests for small samples are very difficult if not impossible to carry out. An alternative to analytic methods is the application of computer simulation in assessing the performance of various statistical tests. Hartley (1976) argues that computer simulation has become an important technique for verifying analytic results. For example, such techniques can be used to find the exact sampling distribution of most statistics using data that have been drawn from any parent distribution (Tracy & Conley, 1982, p. 262). Fawcett and Salter (1987) stressed that a distribution study should not be regarded as complete without the inclusion of computer simulation for finding the exact distribution of the statistic used. Although the present study will present analytic expressions when they exist, it will rely on simulated data to answer the research questions by examining the distributional behavior of the selected test statistics under various data conditions.

Significance of the Study

The significance of the present study is related to the multivariate character of educational research questions and the empirical tests of these questions. As noted earlier, studies in education often generate research questions that are multivariate in character and lead to the use of multiple (correlated) variables. The multivariate methods used in the analysis of this data typically assume normality of the underlying distribution. In many practical

situations, however, there is evidence that the underlying distributions are non-normal (Puri & Sen, 1971, p. 1). Because such non-normality can have a deleterious effect on the distributional properties of normal-theory tests, particularly for studies that employ small to moderate samples, the use of normal-theory methods may not be appropriate. Under these circumstances, educational researchers should consider a multivariate-nonparametric alternative.

In studying normal-theory and nonparametric-multivariate tests, canonical-correlation analysis/multivariate-multiple-regression seem a natural starting point. Their importance as a general system of statistical inference has been demonstrated by several authors. For example, Knapp (1978) showed that many of the commonly used normal-theory tests can be treated as special cases of the canonical-correlation model. The same is true in the nonparametric case.

It should be emphasized that the Bartlett, Rao F , and rank-transform Rao F tests were developed for an omnibus test involving canonical correlations, while the pure- and mixed-rank tests to be examined were developed for an omnibus test involving multivariate-multiple-regression. However, as shown in the next chapter the models underlying canonical-correlation and multivariate-multiple-regression are identical, and hence these tests are equivalent with respect to concluding whether two sets of variables are independent (Gittins, 1985, p. 19).

The choice of the "best" test depends on a number of factors, including the form of the underlying distribution, the number of variables, their degree of dependency, and the sample size. These

factors were examined in the present simulation study. It is expected that the results will provide educational researchers with guidelines for performing either normal-theory or nonparametric canonical-correlation/multivariate-multiple-regression analysis for a variety of data conditions.

Limitations of the Study

The findings of the present study are valid only if the between-set correlation matrix containing zeros is a sufficient indication of the independence of two sets of non-normal variates. The generalizability of the results is also limited by the range of simulation conditions investigated.

CHAPTER II

REVIEW OF THE LITERATURE

A review of the literature pertinent to the present study is presented in this chapter. The review includes the following areas (a) defining the normal distribution and some non-normal distributions on the basis of skewness and kurtosis, (b) methods of dealing with non-normal data, including transforming the original data to ranks, and (c) robustness and power results for some normal-theory and nonparametric-multivariate tests.

Comparing the Normal Distribution and Some Non-Normal Distributions

Because of the critical role of skewness and kurtosis in the present study, these characteristics of a distribution are defined and illustrated for the normal and some non-normal distributions. In theory, the shape of a distribution is defined by its probability distribution function, which is an algebraic expression indicating the distribution of a variable across all of its possible values. In practice, an approximate distribution of a variable can be characterized by its first four central moments (i.e., mean, variance, skewness, and kurtosis) (Fleishman, 1978). In this scheme, the center and dispersion of a distribution are determined by the mean and variance, and its symmetry and tailedness by the skewness and kurtosis values. Assuming the observations are standardized to have a known mean and variance, this permits distributions to be classified according to their skewness and kurtosis values.

Skewness and kurtosis are defined through central moments. For a continuous random variable X , the r (th) central moment (μ_r) is defined as (Kendall & Stuart, 1969, V. 1, p. 55):

$$\mu_r = \int_{-\infty}^{\infty} (X - \mu_1)^r f(X) dX, \quad (1)$$

where μ_1 is the population mean and $f(X)$ is the probability distribution function of X . The population skewness (γ_1) and kurtosis (γ_2) can be defined in terms of the second (μ_2), third (μ_3), and fourth (μ_4) central moments (Kendall & Stuart, 1969, V. 1, p. 85):

$$\gamma_1 = \mu_3 / \mu_2^{3/2}, \quad (2)$$

$$\gamma_2 = \mu_4 / \mu_2^2 - 3. \quad (3)$$

According to expressions (2) and (3), the normal distribution has skewness and kurtosis values equal to zero. Non-normal distributions can then be defined as those having skewness and/or kurtosis values other than zero. Expressions (2) and (3) were used to define the normal distribution and a variety of non-normal distributions in the present simulation study.

In general, a distribution is characterized as mesokurtic if its kurtosis is zero, platykurtic if its kurtosis is negative, and leptokurtic if its kurtosis is positive (Kendall & Stuart, 1969, V. 1, p. 86). Platykurtic distributions are flatter and have lighter tails (i.e., less extreme values) than the normal distribution. Leptokurtic distributions on the other hand, are sharply peaked and have heavier tails (i.e., more extreme values) than the normal distribution. A distribution is said to be symmetric if its skewness is zero, or

asymmetric (skewed) if it has a negative or positive skewness value.

To illustrate the relationship between various skewness and kurtosis combinations and the shape of the distribution they reflect, the skewness and kurtosis values of some common (standardized) univariate distributions are shown in Table 1. Most of these distributions are symmetric or leptokurtic (i.e., normal, uniform, logistic, double exponential, t). In general, skewness and kurtosis

Table 1
Skewness and Kurtosis Values of Some
Univariate Distributions^a

Distribution	γ_1	γ_2	Shape of Distribution
normal	.0	.00	symmetric, mesokurtic
uniform	.0	-1.12	symmetric, platykurtic
logistic	.0	1.20	symmetric, leptokurtic
double-exp.	.0	3.00	symmetric, leptokurtic
t	.0	$6/(\nu-4)$ ^b	symmetric, leptokurtic
exponential	2.0	6.00	asymmetric, leptokurtic

^a Johnson and Kotz (1970, Vols. 1 and 2)

^b ν = degrees of freedom ($\nu > 4$).

values that deviate substantially from zero indicate a greater degree of non-normality, although it should be emphasized that the effects of both skewness and kurtosis must be considered.

Another distribution that is important in simulation studies is the symmetric, extremely heavy-tailed Cauchy, since it reflects an extreme in non-normality that can occur in practice. Theoretically, the Cauchy distribution has an infinite variance, and hence does not possess a finite kurtosis value. However, in empirical studies a pseudo-Cauchy distribution can be generated using a zero skewness

value and a large positive kurtosis value (Harwell & Serlin, 1985).

Although univariate measures of skewness and kurtosis are relatively unambiguous in their interpretation, such is not the case for their multivariate counterparts. Even so, these measures have been useful in identifying a particular member of a family of distributions, in developing a test of normality, and in investigating the robustness of normal-theory procedures (Mardia, 1970).

Mardia (1970, 1974) developed measures of skewness and kurtosis for multivariate distributions, details of which appear in Chapter III. For illustrative purposes, multivariate measures of skewness and kurtosis for some bivariate distributions in which both variables are standardized are given in Table 2 (Mardia, 1970). Notice that the skewness values are zero for all symmetric distributions. In particular, Mardia (1974) showed that the multivariate-normal distribution has skewness and kurtosis values of zero. Just as in the univariate case, any multivariate distribution is considered to be non-normal for non-zero skewness and/or kurtosis values.

Table 2

Multivariate Skewness and Kurtosis Values of Some
Bivariate Distributions (Mardia, 1970)

Distribution	Skewness	Kurtosis
normal	0.00	0.00
uniform	0.00	-2.24
double-exponential	0.00	6.00
exponential	8.00	12.00

Methods of Dealing With Non-Normal Data

As noted earlier, the normality assumption underlying normal-theory tests is of prime importance in the correct use and interpretation of these procedures. At the same time, there is a general recognition that data obtained in a variety of settings, including education, are frequently at least moderately non-normal and hence the use of normal-theory tests is problematic. This has led to the emergence of two approaches for fitting such data into the framework of statistical theory (a) transforming the data to an approximate normal form and then applying a normal-theory procedure (see Box & Cox, 1964; Kaskey, Kolman, Krishnaiah, & Steinberg, 1980), or (b) transforming the data to their corresponding ranks, which removes the normality requirement on the form of the underlying distribution, and employing nonparametric methods. The focus here is on (b), and in the following sections the applications of rank methods in hypothesis testing are discussed.

Nonparametric Methods

Nonparametric methods have a long history in both theoretical and applied statistics (Noether, 1984). These methods require a transformation of the original scores such that the resulting transformed values have known distributional properties. These tests are often classified as distribution-free, since the methods are based on (sample) statistics whose sampling distributions do not depend on the form of the parent distribution from which the sample was drawn (Gibbons, 1971, p. 3). The valid use of these tests requires that the distributions underlying the data are continuous and that all

observations are independently and identically distributed (Puri & Sen, 1985, p. 307).

The principle underlying nonparametric procedures is that under a postulated statistical hypothesis the joint distribution of the random variables is invariant under appropriate groups of transformations (Puri & Sen, 1985, p. 7). It is this invariance that produces what are called genuinely distribution-free tests. However, some distribution-free tests are not genuinely distribution-free, and are usually classified as asymptotically or permutationally distribution-free. In general, an asymptotically distribution-free test can be defined as one which is distribution-free given that the sample size is infinite (i.e., by virtue of the central limit theorem) (Conover & Iman, 1981; Hollander & Wolfe, 1973, p. 437), whereas a permutationally distribution-free test depends only on the set of permutations of the observations associated with testing a hypothesis, and not on the underlying distribution function (Puri & Sen, 1985, p. 149).

Applications of nonparametric methods in testing univariate hypotheses are described in detail by Gibbons (1971), Conover (1980), and Marascuilo and McSweeney (1977). A primary source of nonparametric methods in testing multivariate hypotheses is the work of Puri and Sen (1969, 1971, 1985). Multivariate tests presented by these authors include the single-sample location problem (e.g., sign, signed-rank, extended signed-rank), the multi-sample location problem (e.g., median, rank sum), and a number of tests for hypotheses subsumed by the general linear model.

Rank-Transform Methods

A related set of nonparametric procedures that are used to test univariate and multivariate hypotheses are rank-transform methods. Transforming the original data to their corresponding ranks and applying the usual normal-theory procedures is an idea championed by Conover and Iman (1981). This approach generates a class of rank-transform methods that in many ways are comparable to well known univariate nonparametric procedures such as the Wilcoxon-Mann-Whitney, Kruskal-Wallis, Wilcoxon-signed ranks, and Friedman tests (Conover & Iman, 1976, 1980a, 1982, 1980b; Iman, 1974a, 1974b, 1976; Iman & Conover 1976, 1978, 1979, 1980a, 1980b). Other applications of the rank-transform approach include correlation and regression analysis (Boyer, Palachek, & Schucany; 1983; Hogg & Randles, 1975; Iman & Conover, 1979). An extension of this approach to tests based on the multivariate general linear model is possible by ranking each quantitative variable separately, and then applying the usual normal-theory methods. One advantage of this approach is that the resulting tests can be performed using existing statistical computer packages.

Although the rank-transform approach seems promising, one important limitation should be noted. The rank-transform methods rely on the distributions of the normal-theory test statistics as approximations to the actual distributions of the rank transformation statistics. To date, the theoretical distributions of such statistics have not been established. Consequently their distributional properties can only be determined empirically through computer simulation studies and are always restricted by the conditions of a particular simulation study.

In sum, the two primary nonparametric methods for handling non-normal data are the pure- and mixed-rank procedures illustrated in Puri and Sen (1985, pp. 307-328), and the rank-transform method of Conover and Iman (1981). Both remove the underlying normality requirement. The major difference between them is that the pure- and mixed-rank procedures have a known theoretical substructure, which permits analytic statements about the distributional properties of a test, while the rank-transform procedure permits no such statements.

Robustness and Power Results for Some Normal-Theory and Nonparametric-Multivariate Tests

In comparison to the univariate case, surprisingly little is known about the robustness and power of multivariate tests when underlying assumptions are violated (e.g., non-normality). A natural starting point is the robustness and power properties of univariate procedures, which have often been found to parallel those of their multivariate counterparts. A large number of studies comparing the distributional properties of univariate procedures are available. Some of these studies have been analytic, involving asymptotic approximations, but most have been empirical. The following review will summarize the robustness and power of some nonparametric-univariate and -multivariate tests as compared to their normal-theory counterparts. In reporting these results, the focus will be on the violation of the normality assumption. Analytic results, where available, will be presented first, followed by empirical results. The sparseness of information on multivariate tests clearly indicates the need for further work in this area.

Analytic Results - Univariate Case

The power of a particular test relative to a competitor is usually reflected by its asymptotic relative efficiency (A.R.E.), which is the limiting ratio of the sample size required by one test relative to that required by a second test such that they have equal power for the same alternative hypothesis. A test is said to be more powerful and efficient than another test if the A.R.E. is greater than 1. Details of the computation of the A.R.E. appear in Appendix B. As an example, some A.R.E. results for nonparametric tests of location relative to the normal-theory t and F tests are shown in Table 3 (Marascuilo & McSweeney, 1977, p. 87). The nonparametric "normal scores" tests in Table 3 refer to a rank test applied to normal scores, which are obtained by transforming the ranks to their standard normal scale counterparts.

Table 3
Asymptotic Relative Efficiencies of Some Nonparametric
and Normal-Theory Tests of Location.

Test/Distribution	normal	uniform	logistic	double- exponential
One-sample				
Sign	0.637	0.333	0.750	2.000
Wilcoxon	0.955	1.000	1.047	1.500
Two-sample				
Median	0.637	0.333	0.750	2.000
Mann-Whitney	0.955	1.000	1.047	1.500
Normal-scores	1.000			1.273
K-sample				
Median	0.637	0.333	0.750	2.000
Kruskal-Wallis	0.955	1.000	1.047	1.500
Normal-scores	1.000			1.273

As an example, consider the A.R.E. of the nonparametric two-sample Mann-Whitney test relative to the t test for a parent normal distribution. The A.R.E. of .955 implies that the Mann-Whitney test requires a sample size of 100 in order to have a power equal to that of the t test using a sample size of approximately 95. In this case, the t test is said to be more efficient than the rank test since it requires a smaller sample size in order to have the same power.

It should be noted that although all four distributions in Table 3 are symmetric, their kurtosis varies from moderate-negative to large-positive. The results indicate that for the normal, uniform, and logistic distributions most of the nonparametric tests are as efficient as their normal-theory counterparts, while for the double exponential distribution the nonparametric tests are more efficient than the corresponding normal-theory tests. These findings suggest that these nonparametric alternatives are, in general, almost as powerful as the corresponding normal-theory procedures for normal and near-normal parent distributions, and certainly more powerful for leptokurtic distributions.

Empirical Results - Univariate Case

A.R.E. is a large-sample property of a test which may not be valid for small to moderate samples (Gibbons, 1971, p. 19). As an alternative, simulation studies may be used to assess the performance of two or more tests by comparing their empirical Type I error and power values for various underlying distributions, alternative hypotheses, and sample sizes. This section summarizes the results of a number of empirical studies carried out to investigate the effects

of non-normality on the robustness and power of univariate normal-theory and nonparametric tests. The studies cover a variety of tests and were chosen because of their representativeness in comparing normal-theory and nonparametric tests for parent non-normal distributions.

In general, empirical studies have suggested that univariate normal-theory tests are robust to moderate non-normality for large samples, especially when the underlying distribution is symmetric (Gaito, 1970; Glass, Peckham, & Sanders, 1972, Kendall & Stuart, V. 2, p. 484, Scheffe', 1959, p. 347). However, a number of studies have suggested that distinct departures from normality in combination with small samples affect univariate normal-theory tests. For example, Feir-Walsh and Toothaker (1974) compared the performance of the normal-theory F test and the nonparametric Kruskal-Wallis test when samples were drawn from an exponential (positively skewed) population. While the Type I error rates were somewhat conservative for both tests, the Kruskal-Wallis procedure was found to be more powerful than the F test. Srisukho (1974), in a similar study, found the power of the Kruskal-Wallis test to be greater than that of the F test when all samples were drawn from a double exponential (symmetric, leptokurtic) population, and less than the F test when all samples were drawn from a uniform (symmetric, platykurtic) population. The Type I error rates for the Kruskal-Wallis test tended to be closer to the nominal value than those of the F test for both the double-exponential and uniform populations.

Studies that examined the distributional behavior of normal-theory and rank-transform statistics have shown favorable results for the latter tests. Boyer, Palachek, and Schucany (1983) studied the distributional behavior of the Williams' (1959) test of the equality of dependent correlations (i.e., $H_0: \rho_{yx_1} = \rho_{yx_2}$ given that X_1 and X_2 are correlated), and its rank-transform alternative. The results indicated that although the power values of both procedures were similar, the Type I error rates of the rank-transform test were closer to the nominal alpha level than those of Williams' test for data that were drawn from a parent lognormal distribution. Williams' test, however, produced higher power values for data that were drawn from a parent normal distribution. Based on these results, the authors recommended the use of Williams' test when normality can be assumed, and the rank-transform version of Williams' test when the normality assumption is not tenable.

Iman (1974b) examined the Type I error and power properties of the normal-theory F and rank-transform F tests for a two-way ANOVA problem. The results indicated that the Type I error rate of the rank-transform F test was similar to that of the F test, and that the rank-transform F was more powerful when the underlying distribution was non-normal.

In sum, the smattering of empirical results for univariate tests presented above suggests that the Type I error rate of normal-theory tests is generally not affected by moderate departures from normality for large samples, particularly when the underlying distribution is symmetric and light-tailed (e.g., uniform). However, for distinctly

non-normal populations (e.g., exponential, double-exponential), nonparametric tests appear to be robust with respect to Type I error rate, and produce higher power values than their normal-theory counterparts.

Empirical Results - Normal-Theory-Multivariate Case

Normal-theory-multivariate tests depend on the assumption that the observations are governed by the multivariate-normal density function. Since departures from this assumption are very difficult to investigate analytically (Ito & Schull, 1964), empirical methods are used. The following review summarizes some empirical studies that have examined the effects of non-normality on the Type I error and power properties of some normal-theory- multivariate tests.

A number of studies for the one- and two-independent groups case have been carried out examining the effects of non-normal skewness and kurtosis and sample size on Hotelling's T^2 statistic. The one-sample T^2 , like the univariate one-sample t statistic, is (a) not affected by small departures from normality, (b) more sensitive to non-normal skewness than to non-normal kurtosis, (c) produces liberal Type I error rates for a large skewness, and (d) produces conservative Type I error rates for large kurtosis (Chase & Bulgren, 1971; Davis, 1982a; Mardia, 1970). Similar results for the two-sample location problem were obtained by Davis (1980, 1982b) for Wilks's likelihood ratio and Roy's largest root tests. Other results have suggested that non-normal kurtosis has no substantial effect on the two-sample T^2 statistic for large samples (Hopkins & Clay, 1963; Ito, 1980).

Olson (1974) empirically studied the effects of non-normality and heterogeneity of covariance matrices on six multi-sample, normal-theory MANOVA tests (Roy, Hotelling-Lawley, Wilks, Pillai-Bartlett, Gnanadesikan, and Gnanadesikan-alternative) using small-to-large sample sizes (5, 10, 50). The empirical Type I error results indicated that moderate departures from a kurtosis of zero had mild effects on three tests (Hotelling-Lawley, Wilks, and Pillai-Bartlett), and severe effects on the remaining tests. The direction of the effect of positive kurtosis was generally toward conservatism. The results of this study also suggested that the Gnanadesikan and Gnanadesikan-alternative tests tended to produce liberal Type I error rates with increases in the number of outcome variables for non-zero kurtosis values, especially for small samples crossed with a large number of groups. The power results indicated that all six tests suffer under moderate departures from a kurtosis of zero, and that increases in the number of outcome variables tended to decrease the power of all of the tests.

In general, the results of the studies that used large samples suggest that normal-theory tests are robust to non-normality (Zwick, 1984, p. 2). This result is expected for many tests because of the role of the multivariate analog of the central limit theorem (Morrison, 1976, p. 85). This theorem states that any statistic which can be represented as a linear combination of the observations has a sampling distribution that can be approximated by the normal distribution for large samples. Since many test statistics are derived from some linear function of the observations, the sampling

distribution of the test statistic can be approximated by the normal distribution as sample size increases. Consequently the test would be robust for large samples; this is not necessarily the case for small or moderate samples.

In sum, these studies suggest that (a) small-to-moderate departures from normality have only minor effects on normal-theory-multivariate tests, (b) such effects are more pronounced for small samples than for large samples, (c) increasing skewness tends to result in liberal Type I error rates, (d) increasing kurtosis tends to results in conservative Type I error rates, and (e) the distributional behavior of normal-theory-multivariate tests is affected more by non-normal skewness than by non-normal kurtosis.

Analytic Results - Nonparametric-Multivariate Case

Analytic results for the distributional properties of nonparametric-multivariate tests are available for a few special cases. Available analytic studies on the asymptotic efficiency of nonparametric-multivariate tests relative to their normal-theory counterparts have shown results similar to those of the univariate case (Puri & Sen, 1971, p. 177). For example, the A.R.E. of the multivariate-nonparametric one-sample test of location with normal-scores relative to Hotelling's T^2 is equal to 1.00 for a multivariate normal distribution, and sometimes greater than 1.00 for other multivariate distributions. (Puri & Sen, 1971, p. 177). Recall that similar results were reported by Marascuilo and McSweeney (1977, p. 87) for the one-sample, normal-scores and t tests in the univariate case (Table 3).

Additional results for multivariate tests of location are available for special cases: (a) the A.R.E. of the rank procedure to the normal-theory, two-sample test of location is always less than 1.00 in the bivariate normal case, (b) the normal-scores test is more efficient than the normal-theory test for any mixture of multivariate normal distributions (i.e., a combination of two normal deviates) and heavy-tailed multivariate distributions, and (c) the normal-scores test is more efficient for a multivariate distribution with marginal densities that have light tails, a result that parallels the univariate case (see Zwick, 1984, p. 9). In general, the nonparametric rank tests for hypotheses subsumed under the multivariate-general-linear model are asymptotically power-equivalent to the normal-theory likelihood-ratio test for a parent normal distribution (Puri & Sen, 1985, p. 184). The nonparametric rank tests for location are asymptotically as efficient as their normal-theory counterparts for a parent normal distribution and more efficient for a parent non-normal distribution (Zwick, 1984, p. 27).

The analytic results indicated that multivariate-nonparametric rank and normal scores tests are more efficient than their normal-theory counterparts for non-normal distributions, and are almost as efficient for the normal distribution.

Empirical Results - Nonparametric-Multivariate Case

Surprisingly few simulation studies have been done comparing the Type I error rate and power of multivariate nonparametric tests against their normal-theory counterparts. The available results are presented in some detail since they have important implications for

the conduct of the present study.

Tiku and Singh (1982) studied the Type I error rate and power of the two-sample Hotelling's T^2 and rank tests using samples of size 20 drawn from six bivariate distributions [normal, t , two chi-square ($\nu = 2, 4$), and two mixed-normal]. In all cases the two outcome variables had a correlation of .5. Their results indicated that the rank test was robust with respect to the Type I error rate for three distributions [normal, chi-square ($\nu = 2$), and one mixed-normal], and was conservative for the remaining distributions. The T^2 test was robust with respect to Type I error rate for the normal and t distributions, and conservative for the chi-square and mixed-normal distributions. The rank test proved more powerful than the T^2 test for all distributions except the normal. The results of Tiku and Singh suggest that the two-sample rank test should be the procedure of choice for testing the equality of mean vectors for even moderately non-normal distributions.

Zwick (1984) studied the empirical Type I error rate and power of the multivariate-nonparametric two-sample rank and normal-scores alternatives to the T^2 test under mild non-normality and heterogeneity of variance-covariance conditions. Just as in the univariate case, the Type I error rate was affected mainly by variance-sample size combinations and not by the parent distribution. Power was affected by both the variance-sample size combination and parent distribution, with all tests producing approximately the same power values. In summarizing the results, Zwick recommended that (a) under the conditions of normality and homogeneity of variance the normal-theory test was the

best procedure with respect to both Type I error and power, (b) under normality and heterogeneity of variance with equal sample sizes or when the larger group had the larger variance the rank test was the best choice, and (c) for negatively-skewed distributions the normal-scores test appeared to be the best overall choice except when the smaller group had the larger variance.

Harwell and Serlin (1985) examined the Type I error rate and power of the Rao F (1951), the nonparametric rank-transform Rao F (Conover & Iman, 1981), and the pure- and mixed-rank tests illustrated by Puri and Sen (1985, p. 312). The simulation conditions included in this study were form of distribution (normal, uniform, double-exponential, exponential, Cauchy), sample size (20, 40, 100), correlation within each set of variables (.3, .7), and correlation among the two sets of variables (Type I error, power). The Cauchy was represented by a symmetric distribution with a kurtosis value of twenty.

The Type I error results suggested that the Rao F test was robust with respect to the Type I error for the normal and uniform distributions, became liberal for the Cauchy distribution, and produced mixed results for the double exponential and exponential distributions. There was no clear pattern for the liberal Type I error rates with respect to sample size and within-set correlation. As a measure of the Type I error behavior of these tests, the Rao F overall produced 38% liberal Type I error rates, taking into account sampling error, while the rank-transform Rao F produced only 7% liberal Type I error rates. The latter test performed most satisfactorily for

extremely non-normal distributions and poorly for the smallest sample size and within-set correlation = .3 conditions. In contrast, the pure- and mixed-rank tests did not produce a single liberal Type I error rate across all simulation conditions.

With respect to power, the results indicated that under a normal or uniform distribution the Rao F test was most powerful across all within-set correlation and sample size conditions. In general, the rank-transform Rao F produced the largest power values among the three nonparametric tests, especially for small samples. However, all four tests produced similar power values for the sample size of 100. The power values of all four tests for the double-exponential were comparable for a sample size of 100, and slightly less for a sample size of 40. In general the mixed-rank power values were slightly higher than those of the pure-rank procedure.

The power results for the Cauchy and exponential distributions showed that the pure- and mixed-rank tests performed poorly for the sample size of 20 and the .01 level of significance. Once again the power values for all three nonparametric tests were quite similar for larger sample sizes. The three nonparametric tests overall produced power values substantially larger than the reported values of the Rao F test for the sample size of 100. In general, the mixed-rank test produced slightly lower power values than the pure-rank test.

Based on the Type I error and power results, the authors recommended that the Rao F test be used for symmetric, light-tailed distributions, the rank-transform Rao F for small samples for any of the non-normal distribution investigated, and the pure- and mixed-rank

statistics for larger samples and moderate to distinctly non-normal distributions. The results also confirmed earlier findings that a normal-theory test is affected by moderate-to-large skewness and by a large kurtosis.

Dependency Among Variables

Studies that examined the effects of the degree of dependency among variables on the distributional behavior of multivariate tests suggest that such dependency affects the power of nonparametric tests. Bhattacharyaa, Johnson, and Neave (1971) examined the power of the two-sample Hotelling's T^2 and nonparametric rank-sum tests. The A.R.E.'s of the rank test relative to T^2 test for correlation values of .0, .3, .6, and .9 are .955, .947, .924, and .884, respectively. These results suggest that the A.R.E. of the rank test to T^2 decreases as the degree of dependency among variables increases. Similar results were found by Puri and Sen (1971, p. 176) for negative correlation values. The results of Harwell and Serlin (1985) showed that the nonparametric pure- and mixed-rank tests produced slightly lower power values for extremely non-normal distributions (e.g., exponential, Cauchy) as the correlation among outcome and predictor variables increased from .3 to .7. These studies suggest that a high degree of dependency among outcome and predictor variables slightly reduces the power of some nonparametric tests.

Summary

The present review of the literature leads to the following conclusions with respect to the effects of skewness and kurtosis on Type I error probability and power of normal-theory and nonparametric-

multivariate statistical tests: (a) the effects of non-normality on normal-theory-multivariate tests appear to parallel those in the univariate case, (b) for large samples slight departures from normality have negligible effects on the Type I error rate and power of most normal-theory tests, (c) for small samples moderate to large departures from normality affect the Type I errors and power of these tests, (d) increasing skewness results in normal-theory tests that become liberal while increasing kurtosis results in tests that become conservative (except for the Harwell & Serlin 1985 study in which increasing kurtosis results in normal-theory tests that become liberal), (e) normal-theory tests are more sensitive to non-normal skewness than to non-normal kurtosis, (f) nonparametric tests are superior at controlling Type I errors within nominal levels and are asymptotically more efficient compared to their normal-theory competitors when the underlying distributions are at least moderately non-normal, and (g) a high degree of dependency among variables slightly decreases the power of some nonparametric tests.

The review of the literature indicates that most of the studies on distributional properties of multivariate tests were confined to the MANOVA procedure. The present study of the tests for canonical-correlation/multivariate-multiple-regression complements previous work. The present study also extends the results of Harwell and Serlin (1985) by including the number-of-variables factor and some additional parent distributions and within-set correlations. The focus was the interaction of the form of parent distribution and the sample size, the within-set correlation, and the number-of-variables.

CHAPTER III

METHODOLOGY

This chapter presents the methodology employed in the present study. The following topics are discussed (a) multivariate statistical models, (b) test statistics and their assumptions, (c) data generation method, (d) simulation conditions, and (e) presentation of simulation results.

Multivariate Statistical Models

This section describes the multivariate-multiple-regression and canonical-correlation models and their relationship to the multivariate general linear model. The term general linear model refers to a family of algebraic models characterized by the linearity of the parameters of the equations specifying the models (Gittins, 1985, p. 19). The multivariate-multiple-regression model is a member of one such family. Let $\mathbf{Y}$ be an $N \times p$ ($i=1,2,\dots,N; j=1,2,\dots,p$) data matrix of N observations on p outcome variables, $\mathbf{X}$ an $N \times q$ ($k=1,2,\dots,q$) matrix of regression constants, $\boldsymbol{\beta}$ a $p \times q$ matrix of unknown parameters of the model, and $\mathbf{E}$ an $N \times p$ matrix of unobserved random errors. The multivariate-multiple-regression model can be written:

$$\begin{matrix} \mathbf{Y} & - & \mathbf{X} & \boldsymbol{\beta} & + & \mathbf{E} & . \\ N \times p & & N \times q & q \times p & & N \times p & \end{matrix} \quad (4)$$

A canonical-correlation model may be conceived of as a special case of the multivariate-general-linear model (Gittins, 1985, pp. 19-20). Recall that the focus of the present study is testing whether there is a linear relationship among two sets of variables, and that testing the hypothesis of independence among two sets of variables is equivalent to testing the hypothesis of no regression. This linear relationship may be represented and studied using a canonical-correlation model (Gittins, 1985, p. 19). The following paragraph introduces canonical correlation and canonical variables.

Let $Y_1, Y_2, \dots, Y_p$ and $X_1, X_2, \dots, X_q$ be two sets of random variables. Define

$$\hat{Y} = h_1 Y_1 + h_2 Y_2 + \dots + h_p Y_p \quad (5)$$

as a weighted linear combination of the Y_j variables, and

$$\hat{X} = m_1 X_1 + m_2 X_2 + \dots + m_q X_q \quad (6)$$

as a weighted linear combination of the X_k variables. Define also $\underline{h}' = (h_1, h_2, \dots, h_p)$ and $\underline{m}' = (m_1, m_2, \dots, m_q)$ as the vectors of constants that maximize the correlation between the Y_j and X_k variables. The correlation between the canonical variates $\hat{Y}$ and $\hat{X}$ is the canonical correlation and $\underline{h}$ and $\underline{m}$ are the canonical weights. Given two sets of variables a total of $s = \text{minimum}(p, q)$ pairs of linear combinations can be constructed, and hence s canonical correlations can be obtained. The canonical correlations are found as solutions of a determinantal equation and the canonical weights as solutions of an eigen equation. The process of obtaining these

quantities is outlined below (Morrison, 1976, pp. 254-257).

Let $\hat{\rho}$ represents a sample canonical correlation, S_{xx} the sample covariance matrix of the X_j variables, S_{yy} the sample covariance matrix of the Y_j variables, and S_{yx} the sample covariance matrix of the Y_j and X_k variables. By definition ρ^2 is given by

$$\rho^2 = \frac{(h' S_{yx} m)^2}{(h' S_{yy} h)(m' S_{xx} m)} \quad (7)$$

Since we wish to maximize the correlation between the linear combinations $\hat{Y}$ and $\hat{X}$, the problem can be solved by obtaining the values of h and m that maximizes expression (7). To simplify the "maximization" process while assuring the uniqueness of h and m the variance of both linear combinations is set equal to 1:

$$h' S_{yy} h - m' S_{xx} m = 1 \quad (8)$$

Hence we need only to maximize $(h' S_{yx} m)^2$ subject to the constraint in (8). This problem can be solved by introducing Lagrangian multipliers λ and θ as follows:

$$(h' S_{yx} m)^2 - \lambda(h' S_{yy} h - 1) - \theta(m' S_{xx} m - 1) \quad (9)$$

The first partial derivatives of expression (9) with respect to h and m are then taken. Setting these equations equal to zero produces a homogenous system of two simultaneous matrix equations, namely

$$\begin{aligned} -\lambda S_{yy} h + (h' S_{yx} m) S_{yx} m &= 0, \\ (h' S_{yx} m) S_{yx} h - \theta S_{xx} m &= 0. \end{aligned} \quad (10)$$

Premultiplication of the first equation by $\underline{h}'$ and the second equation by $\underline{m}'$ produces

$$\lambda - \theta - (\underline{h}' \underline{S}_{yx} \underline{m})^2. \quad (11)$$

Hence each of the Lagrangian multipliers is equal to the squared maximum correlation between $\hat{Y}$ and $\hat{X}$. In order for the equations in (10) to have a nontrivial solution their determinant must vanish. This leads to the determinantal equation

$$|\underline{S}_{yy}^{-1} \underline{S}_{yx} \underline{S}_{xx}^{-1} \underline{S}'_{yx} - \lambda \underline{I}| = 0, \quad (12)$$

where $\underline{I}$ is an identity matrix. Morrison (1976, p. 257) shows that the eigenvalues of equation (12) may also be obtained by replacing the covariance matrices in equation (12) with their corresponding sum-of-cross-product (SCP) matrices:

$$|\underline{A}_{yy}^{-1} \underline{A}_{yx} \underline{A}_{xx}^{-1} \underline{A}'_{yx} - \lambda \underline{I}| = 0, \quad (13)$$

where $\underline{A}_{xx}$ is the SCP matrix of the centered X_k variables, $\underline{A}_{yy}$ is the SCP matrix of the centered Y_j variables, and $\underline{A}_{yx}$ is the SCP matrix of the centered Y_j and X_k variables. The maximum $\hat{\rho}^2$ is the largest eigenvalue (λ) of equation (12) or (13) (Morrison, 1976, p. 256). The canonical correlations are ordered $1 > \hat{\rho}_1 > \hat{\rho}_2 > \dots > \hat{\rho}_s > 0$. The canonical weights are the solutions to the associated eigen equations (Morrison, 1976, p. 263). Note that each pair of canonical variates is orthogonal to all other pairs and that the total number of pairs equals s .

Test Statistics and Their Assumptions

Five test statistics will be used to test the hypothesis of no linear relationship among the two sets of variables (Y_j, X_k) . These tests can be conceived of as a test of the matrix of the regression parameters β against zero, or, synonymously, as an omnibus test that all squared population canonical correlations $(\rho_1^2, \rho_2^2, \dots, \rho_s^2)$ are simultaneously equal to zero (Gittins, 1985, p. 57). Hence the null hypothesis for the canonical problem can be written as

$$H_0: \rho_1^2 = \rho_2^2 = \dots = \rho_s^2 = 0. \quad (14)$$

Retaining H_0 is equivalent to concluding that there is no Y_j, X_k relationship, while rejecting H_0 implies the existence of such a relationship. Each test is performed using the eigen values (squared canonical correlations) obtained from expression (13).

Two normal-theory and three nonparametric omnibus tests of the hypothesis of expression (14) will be employed. The normal-theory tests are the Bartlett (1938) and Rao F (1951) procedures. The nonparametric tests are the pure- and mixed-rank procedures discussed in Puri and Sen (1985, pp. 307-328) and the rank-transform Rao F (Conover & Iman, 1981). Although all five tests provide tests of an omnibus statistical hypothesis it is important to emphasize the differences in the nature of the variables upon which the canonical correlations are computed. For the normal-theory tests the canonical correlations are obtained using the original values of the outcome and predictor variables; for the pure-rank and rank-transform Rao F the canonical correlations are obtained using the ranks of the original

values of the outcome and predictor variables; and for the mixed-rank test the canonical correlations are obtained using the original values of the predictors and the ranks of the outcome variables. Note, however, that the hypothesis being tested involves the raw score canonical correlations or regression coefficients. The test statistics, their computations, and assumptions are described below.

Bartlett Statistic

Given two sets of random variables, there exist several measures that summarize the strength of the relationship among them. The best known is Wilk's lambda (Λ), which is defined as (Anderson, 1958, p. 233)

$$\Lambda = \prod_{r=1}^s (1 - \lambda_r), \quad (r=1, 2, \dots, s), \quad (15)$$

where the λ_r 's are the solutions (eigenvalues) of expression (13). The values of Λ have a range between 0 and 1, with smaller values indicating a strong relationship between the X_k and Y_j variables and larger values indicating a weak relationship (Marascuilo & Levin, 1983, p. 185).

Given that two sets of variables are independent, Wilk's Λ has been shown to follow Wilk's Λ distribution (Anderson, 1958, p. 242). Unfortunately, tables of the exact Λ distribution are needed to perform the test. To obviate the need for these tables, Bartlett (1938) introduced a large-sample chi-square approximation to the exact Wilks's Λ distribution. Under the truth of the hypothesis of expression (14), Bartlett (1938) showed that Bartlett statistic (BAR) is asymptotically distributed as a central chi-square variable with pq

degrees of freedom. The BAR statistic can be computed using the following formula (Marascuilo & Levin, 1983, p. 185):

$$\text{BAR} = -[(N-1) - (p + q + 1)/2] \log_e \hat{\Lambda} \sim \chi_{pq}^2 \quad (16)$$

If BAR exceeds the $100(1 - \alpha)$ percentile of the chi-square distribution with pq degrees of freedom, the hypothesis of expression (14) is rejected (Bartlett, 1938).

The Bartlett procedure assumes that the observations are independently and identically distributed random variables (i.i.d.r.v.'s) with a common (multivariate-normal) distribution function (Gittins, 1985, p. 242).

Rao F Statistic

A more precise approximation to the exact Wilk's Λ was developed by Rao (1951) (Marascuilo & Levin, 1983, p. 185). In testing the independence of two sets of variables, Rao's procedure yields an exact test when the smaller of the two sets contains two or less variables (Marascuilo & Levin, 1983, p. 187). In contrast, no exact test is possible with the Bartlett test. The Rao F statistic (RAO) can be computed using the following formula (Marascuilo & Levin, 1983, p. 186):

$$\text{RAO} = \frac{(1 - \hat{\Lambda}^{1/b})/\nu_1}{\hat{\Lambda}^{1/b}/\nu_2} \sim F_{\nu_1, \nu_2} \quad (17)$$

where $\nu_1 = pq$, $\nu_2 = 1 + ab - pq/2$, $a = (N - 1) - (p + q + 1)/2$, and $b = [(p^2 q^2 - 4)/(p^2 + q^2 - 5)]^{1/2}$. If RAO exceeds the $100(1 - \alpha)$ percentile of the F_{ν_1, ν_2} distribution the hypothesis of expression

(14) is rejected (Rao, 1951). The assumptions of the Rao F procedure are the same as those of the Bartlett.

Rank-Transform Rao F

The rank-transform approach (Conover & Iman, 1981) involves transforming the original values of the outcome and predictor variables into their corresponding ranks and then applying the normal-theory Rao F procedure. It is important to emphasize that the theoretical F distribution is used as an approximation to the unknown distribution of the rank-transform Rao F statistic (RTF). The decision rule for the rank-transform Rao F is the same as that of the normal-theory Rao F test. However, unlike the previous two tests the ρ_r^2 of expression (13) are based on the ranks of the X_k and the Y_j variables. The rank-transform procedure assumes that the observations are i.i.d.r.v.'s with a common distribution function (Conover & Iman, 1981).

Pure-Rank Statistic

The nonparametric pure- and mixed-rank statistics illustrated in Puri and Sen (1985, pp. 307-328), and discussed by Harwell and Serlin (1985), are used when both predictor and outcome variables are random. It should be noted that the model originally presented by Puri and Sen (1969) requires the predictor variables to be known regression constants. However, since this condition rarely obtains in practice interest centers on models in which the predictors are assumed to be random. The pure-rank test is particularly useful when only the ranks of the predictor or outcome variables are available.

In the pure-rank model all p outcomes and q predictors are assumed to be i.i.d.r.v.'s. To represent the pure-rank test in a canonical correlation context, let $G(Y|X)$ be the conditional distribution function of the i (th) subject's vector of outcomes Y_i , given a vector of predictor values X_i . Recalling the multivariate-multiple-regression model illustrated in expression (4), the conditional distribution function of the Y_i given the X_i can be written as:

$$G(Y_i|X_i) = G_0(Y_i - X_i \beta), \quad (18)$$

where G_0 is some continuous distribution function. Expression (18) implies that the conditional distribution function for each subject is identical, and that this function depends on the observed predictor values. As noted earlier, this implies that inferences are limited to subpopulations having the same configuration of predictor values as those in the sample.

Puri and Sen (1985, pp. 307-328) used the form given in expression (18) to write the hypothesis of no relationship among the two sets of variables as

$$H_0: G(Y_i|X_i) = G_0(Y_i). \quad (19)$$

Retention of the above hypothesis implies the two sets of variables are independent, while rejection implies that they are related. In order to compute the pure-rank statistic (PUR), the N observations for each of the p outcomes and q predictors must be separately ranked. Let R_{ji} and R_{ki} represent the rank of the i (th) subject on the j (th)

outcome and k (th) predictor variables, respectively. Since $G(Y_i|X_i)$ is assumed to be continuous the theoretical probability of tied R_{ji} or R_{ki} values is zero. In practice, as long as the proportion of ties is small, assigning midranks to tied values will have a negligible effect on the test statistic (Lehmann, 1975, p. 18).

Puri and Sen (1985, pp. 307-312) presented a large-sample form of the pure-rank statistic based on the SCP matrix $\underline{S}$ of the centered R_j and R_k values, with elements

$$s_{jk} = \frac{1}{N} \sum_{i=1}^N (R_{ji} - \bar{R}_j)(R_{ki} - \bar{R}_k), \quad (20)$$

where $\bar{R}_j$ and $\bar{R}_k$ are the rank means for the j (th) dependent and k (th) predictor variables, respectively. In the construction of the pure-rank test, Puri and Sen show that the $E(\underline{S}) = \underline{0}$ and that the elements of $\underline{S}$ are asymptotically multivariate-normal given that the sets of the Y_j and the X_k variables are independent. Details of the construction of this statistic appear in Appendix C.

The form of the pure-rank statistic as presented by Puri and Sen does not easily permit the use of existing computer software packages. Harwell and Serlin (1985) provide a form of the pure-rank test that allows existing computing packages to be used. Details of the derivation of the alternative form also appear in Appendix C. Assuming the X_k and Y_j have been separately ranked, the ranks are submitted to a standard computing package (e.g., SAS, 1982), and the canonical correlations obtained from the output. The form of the PUR statistic presented by Harwell and Serlin (1985) is

$$PUR = (N-1) \sum_{i=1}^N \theta_r, \quad (21)$$

where θ_r represents the (squared) canonical correlation (eigenvalue) between the ranks of the X_k and the Y_j . These are obtained by replacing the SCP matrices based on the original values of the X_k and Y_j with the SCP matrices based on their ranks in expression (13). Puri and Sen (1985, p. 312) showed that under the truth of the hypothesis of expression (14), the pure-rank statistic is asymptotically distributed as a central chi-square variable with pq degrees of freedom. The decision rule for the pure-rank test is to reject the hypothesis of expression (14) if PUR exceeds the $100(1 - \alpha)$ percentile of the chi-square distribution with pq degrees of freedom (Puri & Sen, 1985, p. 312). Rejection of the hypothesis of expression (14) implies that the population regression coefficients are not all simultaneously equal to zero, or, synonymously, that the population canonical correlations are not all equal to zero.

The pure-rank test assumes that the X_k and Y_j observations are i.i.d.r.v.'s whose common distribution function is $G(Y_i|X_i)$. The difference in assumptions between the normal-theory and pure-rank procedures is that normality of $G(Y_i|X_i)$ is not required for the pure-rank test.

Mixed-Rank Statistic

As presented by Puri and Sen (1985, pp. 307-328), and discussed by Harwell and Serlin (1985), the mixed-rank statistic is computed using the original X_k values and the ranks of the Y_j . This test assumes that all outcomes and predictors are i.i.d.r.v.'s and provides a test

of the hypothesis that all regression coefficients equal to zero, or, synonymously, all squared canonical correlations simultaneously equal to zero. A procedure similar to that of the pure-rank test was employed by Puri and Sen to obtain a large-sample form of the mixed-rank statistic. In the mixed-rank case, however, the original X_k values are used instead of their ranks.

The mixed-rank (MIX) test statistic has exactly the same form as that of the pure-rank test given in expression (21). The decision rule for the mixed-rank test is the same as that of the pure-rank test, and rejection of the hypothesis of expression (14) implies that the two sets of variables are related. The assumptions of the mixed-rank test are the same as those of the pure-rank test.

Data Generation Method

This section outlines the method that was used to generate the multivariate data for the present study. The data generation and analysis were performed on an IBM 3090-180 computer at Michigan State University. The program was coded in FORTRAN V and incorporated a number of subroutines from the International Mathematical and Statistical Libraries (IMSL) (1983). A summary of the IMSL subroutines used in the present study is given in Appendix D. In all cases the data were in standard form (i.e., $\mu = 0$, $\sigma^2 = 1$). The Basic Uniform Number Generator (GGUBS) subroutine of the International Mathematical and Statistical Libraries (IMSL, 1983) was used to generate random uniform deviates in the range of (0, 1). GGUBS has been extensively tested and has been found to produce deviates with good statistical

properties (Learmonth & Lewis, 1973).

The resulting uniform deviates were transformed into normal deviates using the Box-Muller (1958) approach. This procedure transforms a pair of uniform deviates (u_1, u_2) into a pair of standard normal deviates (z_1, z_2) using the following transformations:

$$\begin{aligned} z_1 &= (-2 \log_e u_1)^{1/2} \cos(2 \pi u_2), \\ z_2 &= (-2 \log_e u_2)^{1/2} \sin(2 \pi u_2). \end{aligned} \quad (22)$$

The resulting variable has (approximately) a mean of 0 and a variance of 1.

In generating multivariate data the following structure was assumed to underlie the $\underline{Y}$ values:

$$\underline{Y}_{p \times N} = \underline{\beta}_{p \times q} \underline{X}_{q \times N} + \underline{E}_{p \times N}, \quad (23)$$

where the $\underline{X}$ (predictor) and $\underline{E}$ (residual) matrices contain random deviates from a specified distribution. These deviates were generated by specifying population correlations among the X_k variables and among the residuals and using the method described below. The Y_j values were then obtained using specified values of the $\underline{\beta}$ in expression (23).

In generating the multivariate data, a $(p+q) \times N$ matrix of standard normal deviates was initially generated using the transformation given in (22). The first p rows of this matrix represented the uncorrelated residuals ($\underline{E}$), and the remaining q rows the predictor values for a sample of size N . In all cases the X_k and Y_j variables had the same distribution. The two matrices of uncorrelated normal deviates were then separately transformed such that the resulting deviates were

correlated with a specified distribution. To generate these correlated deviates the procedure due to Vale and Maurelli (1983) was used. Details of this procedure, which combines the approaches of Kaiser and Dickman (1962) and Fleishman (1978), are presented in Appendix B. The same procedure was followed separately for the predictor and residual correlation matrices.

The Vale and Maurelli procedure begins by using the Kaiser and Dickman (1962) method to generate a sample of multivariate-normal deviates using a matrix decomposition of the desired population correlation matrix, say $\underline{P}$. A matrix $\underline{Z}$ of multivariate-normal deviates can be obtained using the following transformation:

$$\underline{Z}_{(p+q) \times N} = \underline{F}_{(p+q) \times (p+q)} \underline{z}_{(p+q) \times N}, \quad (24)$$

where $\underline{F}$ is a matrix of principal components (or some other decomposition) of the population correlation matrix $\underline{P}$ and $\underline{z}$ is a matrix of uncorrelated standard normal deviates. Multiplication of $\underline{F}$ and $\underline{z}$ produces variables with a mean and variance approximately equal to 0 and 1, respectively, and intervariable correlations approximately equal to those in $\underline{P}$.

To generate non-normal deviates the Vale and Maurelli procedure combines the Fleishman (1978) and Kaiser-Dickman methods. Fleishman (1978) developed a technique for generating a (univariate) non-normal variable, say w_i , by finding the first four central moments (mean, variance, skewness, and kurtosis) of the distribution of the variable. The technique uses a polynomial involving the first three powers of a standard normal deviate z_i :

$$w_i = a + bz_i + cz_i^2 + dz_i^3, \quad i = 1, 2, \dots, N, \quad (25)$$

where a , b , c , and d are the so-called Fleishman power function constants. These constants were computed using the nonlinear equation-solving routine NEQNF (IMSL, 1983). Fleishman's (1978) procedure has been shown to produce non-normal deviates with the desired distributional properties (i.e., mean, variance, skewness, kurtosis) (Fleishman, 1978).

In generating multivariate non-normal random deviates the processes of decomposition of the population correlation matrix and the Fleishman transformation interact, which leads to non-normal deviates with correlations different from those of the desired population. Vale and Maurelli (1983) developed a method to counteract this effect such that the resulting non-normal deviates would possess (approximately) the desired correlations. Essentially, this method involves the creation of an intermediate correlation matrix, P^* , from the desired correlation matrix P . The P^* matrix is then factored to obtain the F matrix of expression (24), and the matrix of multivariate-normal deviates is transformed to non-normal multivariate deviates using expression (25).

It should be noted that the data-generation process is not based on the probability density function of any theoretical multivariate distribution (e.g., multivariate-exponential), and hence the method does not actually produce data from such a distribution. Rather, the method produces data that have the same marginal skewness and kurtosis values as those of a theoretical multivariate distribution. However,

the Vale and Maurelli (1983) procedure has been shown to produce multivariate data with (asymptotically) the expected marginal (univariate) mean, variance, skewness, kurtosis, and correlations (Vale & Maurelli, 1983). The present study used five univariate summary measures (i.e., mean, variance, skewness, kurtosis, correlation) to determine if the simulated data actually possessed the desired distributional properties. In addition, Mardia's (1974) measures of multivariate skewness and kurtosis were computed.

Simulation Conditions

As noted earlier there are several factors which are expected to influence the distributional behavior of the tests. They include differences in the parent distributions as specified by γ_1 and γ_2 , numbers-of-variables, between-set correlations, within-set correlations, and sample sizes. The simulation factors and their levels are shown in Table 4 and are discussed in detail below.

Table 4

Simulation Factors

Factor	Level
skewness and kurtosis [γ_1 , γ_2]	[0, 0], [0, -1.12], [0, 3], [0, 20], [.5, 0], [1, .5], [1, 3], [2, 6]
number-of-variables (p, q)	(2, 2), (3, 3), (4, 4)
between-set correlation	.0 (Type I error), or >.0 (power)
within-set correlation (ρ_y , ρ_x)	(.3, .3), (.3, .7), (.7, .7)
sample size (N)	25, 50, 100

Data were generated to represent observations from eight selected distributions representing a range of skewness and kurtosis values. The (univariate) skewness and kurtosis values of four known distributions (normal [0,0], uniform [0, -1.12], double-exponential [0, 3], and exponential [2, 6]), and four additional distributions were included. Pairings of skewness and kurtosis values allowed an examination of the effects of this factor over a broad range of non-normal conditions. Specifically, the {[0, -1.12], [.5, 0], [1, .5]} pairings represent three mildly non-normal distributions, {[0, 3], [1, 3]} two moderately non-normal distributions, and {[2, 6], [0, 20]} two extremely non-normal distributions. The three combinations of the number-of-variables used [(2,2), (3,3), (4,4)] were chosen to examine the effects of increased dimensionality. Recall that the Rao F test is exactly distributed as an F variate under the null hypothesis when the smaller variable set contains two or fewer variables. The (3, 3) and (4, 4) combinations also reflected the increasing non-normality that tends to be associated with increasing numbers of variables.

The three sets of values of the correlations within the set of outcome and predictor variables [(.3, .3), (.3, .7), (.7, .7)] represented a range of correlations encountered in practice. These combinations of within-set correlations permitted an examination of the effects of equal and unequal within-set correlation on the distributional behavior of the tests. For example, the power value of nonparametric rank tests appears to decrease slightly as the (absolute) within-set correlation increases for extremely non-normal data. The within-set correlation values allowed the behavior of the

tests under these conditions to be examined.

Three different sample sizes (25, 50, 100) were also included in the present study. This range permitted an examination of the effect of varying sample size on the tests. Marascuilo and Levin (1983, p. 204) recommended that samples of larger than $10(p+q)$ should be used in multivariate studies. According to this recommendation, the chosen sample sizes represent small, small-moderate, and moderate-large samples, depending on the numbers of variables used.

Finally, a range of between-set correlations ($.0 \leq \rho_{xy} \leq .366$) were included to examine the Type I error and power properties of the tests. A zero between-set correlation corresponds to the Type I error case while a non-zero between-set correlation corresponds to the power case. The between-set correlation for the power case was obtained analytically using a procedure due to Muller and Peterson (1984) and the tabled power values of the F test due to Pearson and Hartley (1951). This procedure uses an approximation involving the non-centrality parameter of the F distribution. The non-centrality value was found such that a power of .8 would be achieved at an alpha level of .05 for a sample size of 100 and a multivariate-normal distribution.

Presentation of Simulation Results

The $8 \times 3 \times 2 \times 3 \times 3$ fully-crossed design employed in the present study generated a total of 432 simulation conditions. Three thousand replications were carried out for each condition, and the five test statistics (Bartlett, Rao F , rank-transform Rao F , pure- and mixed-rank) were calculated for each replication. The resulting Type I error

and power values were tabulated at three levels of significance (.01, .05, .10).

The Type I error probability was estimated by the proportion of the number of rejections of the null hypothesis of expression (14) when the null condition was true (i.e. data were sampled from a multivariate population with a between-set correlation equal to zero). The robustness of the Type I error probability of the five tests was determined using a 95% confidence interval of the Type I error probability (i.e., $\alpha \pm 1.96/[\alpha(1 - \alpha)/3000]$, where α is the nominal Type I error probability). The 95% confidence interval of the average Type I error probability was obtained using the standard error of the average empirical Type I error rate (i.e., $\alpha \pm 1.96/[\alpha(1 - \alpha)/3000n]$, where n is the number of the Type I errors involved in computing the average). A test was considered robust with respect to the Type I error probability if its empirical Type I error rate fell inside the confidence interval. Otherwise, the test was considered either conservative or liberal. The empirical power value was estimated by the proportion of the number of rejections of the null hypothesis of expression (14) when the null condition was false (i.e., data were sampled from a multivariate population with a nonzero between-set correlation).

Evidence that the data generation method was actually producing data with the desired distributional characteristics was obtained by computing five (marginal) summary measures: average mean, variance, skewness, kurtosis, and correlations. In addition to the marginal measures, multivariate measures of skewness and kurtosis proposed by

Mardia (1974) were computed and used to examine the skewness and kurtosis of the data with a parent multivariate-normal distribution. A brief description of these measures is given below with a detailed presentation of the statistics provided in Appendix C.

Let $\underline{U}_1, \underline{U}_2, \dots, \underline{U}_N$ be N vectors of random observations on $t = (p + q)$ variables, $\underline{U}$ the vector of sample means, and $\underline{V}$ the matrix of sample covariances:

$$\underline{U}_i = \begin{pmatrix} u_{1i} \\ u_{2i} \\ \vdots \\ u_{ti} \end{pmatrix} \quad \bar{\underline{U}} = \begin{pmatrix} \bar{u}_1 \\ \bar{u}_2 \\ \vdots \\ \bar{u}_t \end{pmatrix} \quad \underline{V} = \begin{pmatrix} v_{11} & v_{12} & \cdots & v_{1t} \\ v_{21} & v_{22} & \cdots & v_{2t} \\ \vdots & \vdots & \ddots & \vdots \\ v_{t1} & v_{t2} & \cdots & v_{tt} \end{pmatrix} \quad (26)$$

Mardia (1974) proposed the following sample measures of multivariate skewness ($\gamma_{1,t}$) and kurtosis ($\gamma_{2,t}$) for a multivariate distribution with t dimensions:

$$\gamma_{1,t} = N^{-2} \sum_{i=1}^N \sum_{j=1}^N [(\underline{U}_i - \bar{\underline{U}})' \underline{V}^{-1} (\underline{U}_j - \bar{\underline{U}})]^3 \quad (27)$$

$$\gamma_{2,t} = N^{-1} \sum_{i=1}^N [(\underline{U}_i - \bar{\underline{U}})' \underline{V}^{-1} (\underline{U}_i - \bar{\underline{U}})]^2 - t(t + 2). \quad (28)$$

Mardia (1974) showed that these multivariate skewness and kurtosis values are zero for any multivariate-normal distribution and nonzero for any other multivariate distribution. However, nonzero values of expressions (26) and (27) do not permit identification of a particular non-normal distribution.

CHAPTER IV

RESULTS

The purpose of this study was to empirically evaluate the distributional behavior of some selected normal-theory and nonparametric tests of the hypothesis of no relationship among two sets of variables when the normality assumption is violated. The study also examined the effects of sample size, within-set correlation, and number-of-variables on the Type I error and power values of two normal-theory and three nonparametric tests. The results of the simulation study are reported in this chapter.

Specifically, this chapter discusses the (a) characteristics of the simulated data, (b) Type I error and power conditions, (c) main effects of the parent distribution, sample size, within-set correlations, and number-of-variables factors, and (d) interaction effects of these factors. These results are summarized in tables and figures in this chapter. Detailed results of the simulation study are presented in Tables E6 through E13 of Appendix E.

Characteristics of the Simulated Data

Simulated data representing observations from eight multivariate distributions were used. Frequency distributions for these distributions were generated using 10,000 deviates, and are presented in Table E14 of Appendix E and displayed in Figure 1.

Five sample (univariate) marginal statistics (mean, variance, skewness, kurtosis, and variable intercorrelations) were computed to

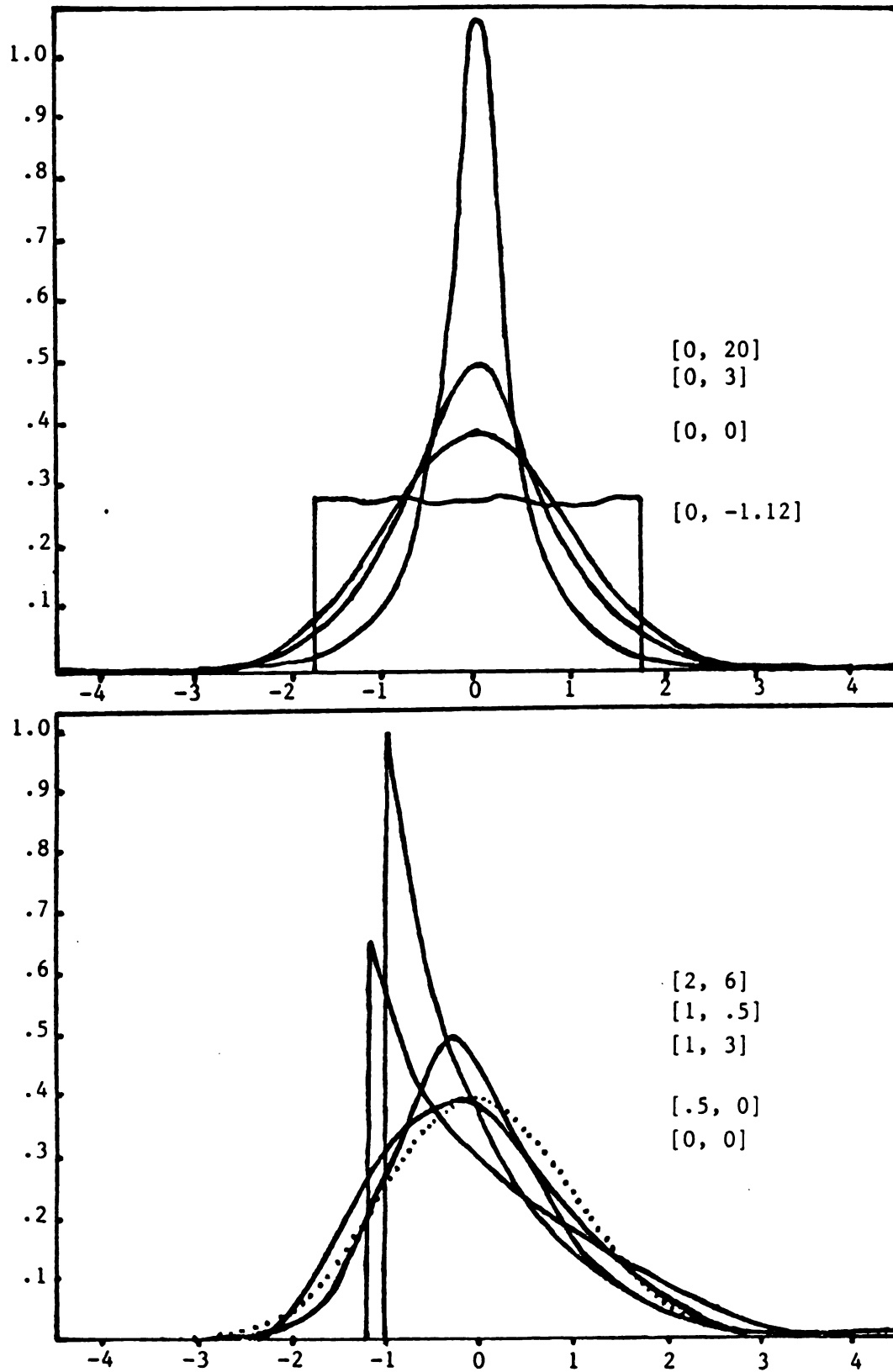


Figure 1. Frequency Distributions of Simulated Data

Table 5

Average Mean, Variance, Skewness, Kurtosis, and
Within-Set Correlation of the Simulated Data^a

$[\gamma_1, \gamma_2]$	μ	σ^2	γ_1	γ_2	$\rho = .3$	$\rho = .7$
[0, 0]						
Ave.	.0017	.9851	.0025	.0772	.3008	.7005
Std.	.0009	.0071	.0050	.0432	.0009	.0006
Min.	.0006	.9773	-.0055	.0137	.2992	.6996
Max.	.0030	.9959	.0079	.1233	.3020	.7014
[0, -1.12]						
Ave.	.0003	.9951	-.0011	-1.1156	.2979	.6936
Std.	.0008	.0087	.0024	.0367	.0010	.0006
Min.	-.0007	.9856	-.0040	-1.1625	.2966	.6927
Max.	.0015	1.0065	.0033	-1.0630	.2998	.6944
[.5, 0]						
Ave.	.0010	.9845	.5025	.0751	.2994	.6996
Std.	.0010	.0080	.0034	.0497	.0011	.0007
Min.	-.0009	.9756	.4976	-.0076	.2975	.6986
Max.	.0021	.9948	.5084	.1235	.3013	.7007
[1, .5]						
Ave.	.0009	.9863	1.0094	.5864	.3010	.7002
Std.	.0007	.0082	.0054	.0548	.0019	.0008
Min.	-.0005	.9771	1.0012	.5120	.2988	.6990
Max.	.0019	.9978	1.0155	.6426	.3038	.7011
[0, 3]						
Ave.	.0006	.9861	.0066	3.1495	.3001	.7000
Std.	.0014	.0089	.0129	.0861	.0019	.0013
Min.	-.0015	.9772	-.0156	2.9914	.2980	.6981
Max.	.0029	1.0012	.0203	3.2639	.3038	.7023
[1, 3]						
Ave.	.0010	.9867	1.0110	3.1497	.3002	.7000
Std.	.0013	.0078	.0149	.1292	.0012	.0008
Min.	-.0001	.9780	.9804	2.9300	.2984	.6988
Max.	.0044	1.0009	1.0268	3.2697	.3023	.7014
[2, 6]						
Ave.	-.0004	.9832	2.0288	6.2224	.2989	.6994
Std.	.0009	.0086	.0205	.1109	.0009	.0006
Min.	-.0016	.9720	1.9951	6.0160	.2978	.6986
Max.	.0013	.9971	2.0550	6.3820	.3000	.7004
[0, 20]						
Ave.	.0007	.9872	-.0116	20.2620	.2997	.7000
Std.	.0015	.0107	.0335	.5450	.0014	.0012
Min.	-.0019	.9704	-.0555	19.6140	.2969	.6976
Max.	.0030	1.0053	.0453	21.2950	.3017	.7016

^a Ave = average, Std = standard deviation, Min = minimum, Max = maximum value of the sample mean, variance, skewness, kurtosis, and within-set correlations.

assess how well the generated data actually represented the eight parent distributions. These values were examined for their similarity to the known population values. The average, standard deviation, minimum, and maximum values of the five sample statistics are given in Table 5. The average values for the sample mean, variance, skewness, and kurtosis were obtained using sample values across all sample-size, within-set-correlation, and number-of-variables conditions. The average values for the sample within-set correlations were obtained using sample values for all within-set-correlation conditions among the Y_j and the X_k variables across all sample-size and number-of-variables conditions. The values of the five sample statistics for various sample-size and number-of-variables conditions are presented in Table E2 of Appendix E.

The statistics in Tables 5 and E2 show excellent agreement between the mean, variance, and intercorrelations of the resulting variables and their population values. Similarly, both tables show good agreement between the skewness values and their population counterparts across all distributions, sample sizes, and numbers-of-variables. However, as seen in Table E2 the agreement between the kurtosis values and their known theoretical counterparts for non-normal distributions was somewhat poor, even for moderately-large samples.

As a further check on the adequacy of the data generation, multivariate measures of skewness and kurtosis (Mardia, 1974) were computed for a sample of $N = 100$. The sample values proved to be quite different from their known population counterparts. As an additional check a number of runs were made using $N = 300$ and the multivariate-

normal distribution. The results appear in Table 6. The multivariate measures of skewness and kurtosis of Mardia (1974) are equal to zero for the multivariate-normal distribution, and hence deviations from zero of the sample skewness and kurtosis values suggest non-normality.

Even for a sample size of 300 the results indicate less than perfect agreement between the sample measures of the multivariate skewness and kurtosis values and their theoretical values of zero, especially as the number of variables increases. However tests of

Table 6
Sample Measures of Multivariate Skewness and
Kurtosis For the Normal Distribution^a

v ^b	Measure	N ^c	Within-set correlation		
			(.3 .3)	(.3 .7)	(.7 .7)
4	Skewness	300	0.3919	0.3919	0.3919
	Kurtosis	300	-0.3094	-0.3094	-0.3093
6	Skewness	300	1.0867	1.0867	1.0867
	Kurtosis	300	-0.6064	-0.6063	-0.6062
8	Skewness	300	2.3257	2.3257	2.3257
	Kurtosis	300	-1.0318	-1.0317	-1.0316

^a The tabled values represent the multivariate skewness and kurtosis statistics of the simulated data based on a sample of size 300 and 3000 replications, ^b V = number of variables, ^c N = sample size.

multivariate normality (Mardia, 1970) on each of the data sets of Table 6 were not significant (at $\alpha = .05$). These results suggest that the simulated data may, for a small number-of-variables, be assumed to approximate that obtained from a multivariate-normal distribution. For larger numbers-of-variables the results of Table 6 suggest that the simulated data tended to be positively-skewed and more kurtic than that

of a normal distribution, and hence the assumption that these data represent (approximately) a multivariate-normal distribution is less tenable.

Type I Error and Power Conditions

Two between-set correlation values were used to generate Type I error and power conditions. A between-set correlation of zero (i.e., $\beta = 0$) was used to establish the Type I error case, meaning each observation was sampled from a population in which there was no correlation between the predictors and the dependent variables. Hence each rejection counted toward the empirical Type I error rate. For the power case, a non-zero β matrix was obtained analytically using a procedure due to Muller and Peterson (1984) and the tabled power values of the F test due to Pearson and Hartley (1951). The regression coefficients of expression (23) were computed such that a power of .8 would be achieved at an alpha level of .05 using a sample of size 100 and a parent normal distribution. The values of the regression coefficients and the resulting non-zero, between-set correlations are presented in Table 7 (see Appendix C for computational details).

Table 7

Regression Coefficients Used for Power Simulations and the Resulting Between-Set Correlations

v^a	β_{jk}	Within-set correlation		
		(.3 .3)	(.3 .7)	(.7 .7)
4	.180	.234	.306	.306
6	.138	.221	.331	.331
8	.118	.224	.366	.366

^a v = number of variables.

In the present study a total of eight distributions x three sample sizes x three within-set correlations x three numbers-of-variables = 216 conditions were separately generated for the Type I error and power cases. For each of these conditions, observations were selected from populations possessing the requisite properties and the empirical Type I error rates and power values tabulated for the various tests across 3000 replications. The robustness of the Type I error probabilities was determined using a 95% confidence interval. The intervals for the alpha levels of .01, .05, and .10 were (.0064, .0136), (.0422, .0578), and (.0893, .1107), respectively. The 95% confidence intervals for the average Type I error rates (at .05 alpha level) are presented and displayed (as broken lines) in the following tables and figures. Type I error rates exceeding the upper limit of these intervals were considered to be liberal, and values below the lower limit were considered to be conservative.

Overall Type I Error and Power Results

A summary of the empirical Type I error and power values across the distributions, sample sizes, within-set correlations, and numbers-of-variables for the five tests [Bartlett (BAR), Rao F (RAO), rank-transform Rao F (RTF), pure-rank (PUR), mixed-rank (MIX)] are presented in Table 8. Because the Type I error and power values for the Bartlett and Rao F tests were virtually identical, only the Rao F test results are reported. The overall Type I error and power curves of the various tests are displayed in Figure 2.

Table 8
Overall Empirical Type I Error and Power Values^a

Test	Type I Error			Power		
	.01	.05	.10	.01	.05	.10
BAR	.0154	.0587	.1082	.3485	.5265	.6250
RAO	.0151	.0579	.1077	.3472	.5250	.6240
RTF	.0104	.0505	.1010	.3674	.5460	.6435
PUR	.0052	.0376	.0864	.3023	.4948	.6081
MIX	.0046	.0362	.0839	.2986	.4868	.5995

^a The tabled values represent the average Type I error and power values across all distributions, sample sizes, within-set correlations, and numbers-of-variables ($n=216$, $[\.0495, .0505]$).

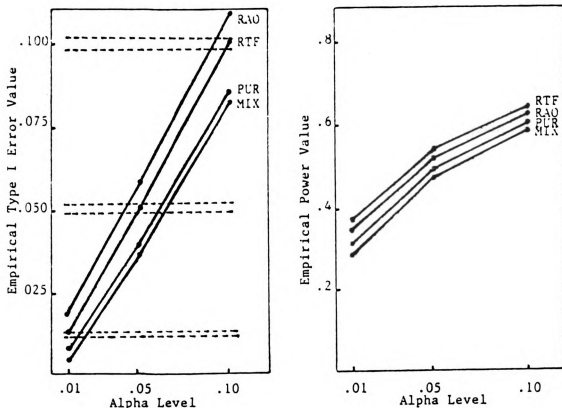


Figure 2. Type I Error and Power Curves by Alpha Level

The Type I error results of Table 8 and Figure 2 indicate that the overall Type I error rate was highest for the Rao F test, and lowest for the mixed-rank test. The rank-transform Rao F produced empirical Type I error rates closest to the nominal alpha values. Under the same simulation conditions, the average power values were highest for the rank-transform Rao F , followed by the Rao F , pure-rank, and mixed-rank tests. This pattern persisted for all three alpha levels.

The Type I error and power results of Table 9 and Figure 2 indicate that the (a) average Type I error rates of the four tests increased at approximately the same rate as the alpha level increased from .01 to .05, and from .05 to .10, (b) average power values of all four tests increased at approximately the same rate as the alpha level increased, and (c) rate of change of the power value was greater for .01 to .05 than for .05 to .10. These results indicate that there was no interaction of alpha level and test statistic on the Type I error and power values. On the basis of these results, subsequent Type I error and power curves will be displayed only for the .05 level.

The total number of conservative and liberal Type I errors for the four tests across all simulation conditions are presented in Table 9. These results indicate that the Rao F test had a higher number of liberal Type I errors as compared to the rank-transform Rao F , the pure-rank, and the mixed-rank tests across all alpha levels. The percentage of liberal Type I errors for the Rao F , rank-transform Rao F , pure-rank, and mixed-rank tests were 28.2%, 4.5%, .3%, and 0%, respectively. The corresponding percentages of the number of conservative Type I errors for these tests were .8%, 2%, 60%, and

66.5%, respectively.

In general, the (a) Rao F proved to be more liberal than it's nonparametric competitors, (b) pure- and mixed-rank tests were more conservative than the rank-transform Rao F and Rao F tests, and (c) number of liberal Type I errors for the Rao F test and the number of conservative Type I errors for the pure- and mixed-rank tests decreased as the alpha level increased.

Table 9
Overall Number of Conservative and Liberal
Type I Errors^a

Test	Alpha level							
	.01		.05		.10		Total	
	C	L	C	L	C	L	C	L
RAO	1	72	2	59	2	52	5	183
RTF	3	9	4	10	6	10	13	29
PUR	140	0	138	0	111	2	389	2
MIX	159	0	143	0	129	0	431	0

^a The tabled values represent overall frequencies of the conservative (C) and liberal (L) Type I errors across all distributions, sample sizes, within-set correlations, and numbers-of-variables (216 cases).

Main Effects of Simulation Conditions

The next four sections present the Type I error and power results categorized by distribution, sample size, within-set correlation, and number-of-variables.

Distribution

A total of eight parent distributions were included in the present study to examine the effects of varying skewness and kurtosis

Table 10

Average Type I Error and Power Values and Number of
Conservative and Liberal Errors by Distribution^a

$[\gamma_1, \gamma_2]$	Type I Error			(C L)	Power		
	.01	.05	.10		.01	.05	.10
[0, 0]							
RAO	.0099	.0508	.1000	(0 0)	.3286	.5124	.6159
RTF	.0100	.0499	.1010	(5 2)	.3088	.4932	.5980
PUR	.0045	.0369	.0869	(52 0)	.2512	.4442	.5634
MIX	.0047	.0375	.0848	(53 0)	.2543	.4469	.5639
[0, -1.12]							
RAO	.0100	.0509	.1020	(1 6)	.3183	.5015	.6084
RTF	.0103	.0509	.1016	(0 2)	.2880	.4704	.5779
PUR	.0051	.0379	.0868	(45 0)	.2330	.4231	.5434
MIX	.0045	.0382	.0873	(50 0)	.2345	.4263	.5467
[.5, .0]							
RAO	.0112	.0507	.1008	(0 3)	.3268	.5077	.6099
RTF	.0108	.0500	.1012	(1 2)	.3101	.4926	.5950
PUR	.0056	.0377	.0860	(46 0)	.2512	.4446	.5606
MIX	.0051	.0367	.0847	(45 0)	.2496	.4423	.5578
[1, .5]							
RAO	.0104	.0500	.1000	(3 3)	.3350	.5149	.6143
RTF	.0095	.0500	.1008	(0 2)	.3737	.5539	.6511
PUR	.0044	.0377	.0863	(52 0)	.3070	.5029	.6152
MIX	.0043	.0365	.0836	(54 0)	.2920	.4788	.5888
[0, 3]							
RAO	.0121	.0524	.1015	(1 9)	.3441	.5258	.6268
RTF	.0108	.0501	.0985	(2 2)	.3615	.5401	.6397
PUR	.0059	.0378	.0848	(45 0)	.2973	.4898	.6033
MIX	.0050	.0359	.0816	(50 0)	.3065	.4967	.6115
[1, 3]							
RAO	.0119	.0542	.1045	(0 14)	.3455	.5249	.6237
RTF	.0098	.0507	.1007	(1 4)	.3573	.5392	.6381
PUR	.0049	.0370	.0859	(57 0)	.2951	.4875	.6034
MIX	.0043	.0359	.0847	(62 0)	.2943	.4832	.5952
[2, 6]							
RAO	.0198	.0651	.1135	(0 67)	.3682	.5345	.6255
RTF	.0110	.0521	.1025	(0 9)	.4415	.6153	.7031
PUR	.0058	.0388	.0878	(41 0)	.3670	.5603	.6669
MIX	.0046	.0350	.0838	(55 0)	.3318	.5070	.6122
[0, 20]							
RAO	.0358	.0895	.1393	(0 81)	.4112	.5782	.6673
RTF	.0110	.0505	.1015	(4 6)	.4984	.6628	.7452
PUR	.0056	.0374	.0864	(51 2)	.4165	.6061	.7085
MIX	.0041	.0335	.0804	(62 0)	.4257	.6133	.7198

^a The tabled values represent the average Type I error and power values across all sample sizes, within-set correlations, and numbers-of-variables ($n=27$, $[\gamma_1, \gamma_2] = [.0485, .0515]$). The number of conservative (C) and liberal (L) Type I errors are the total across three alpha levels.

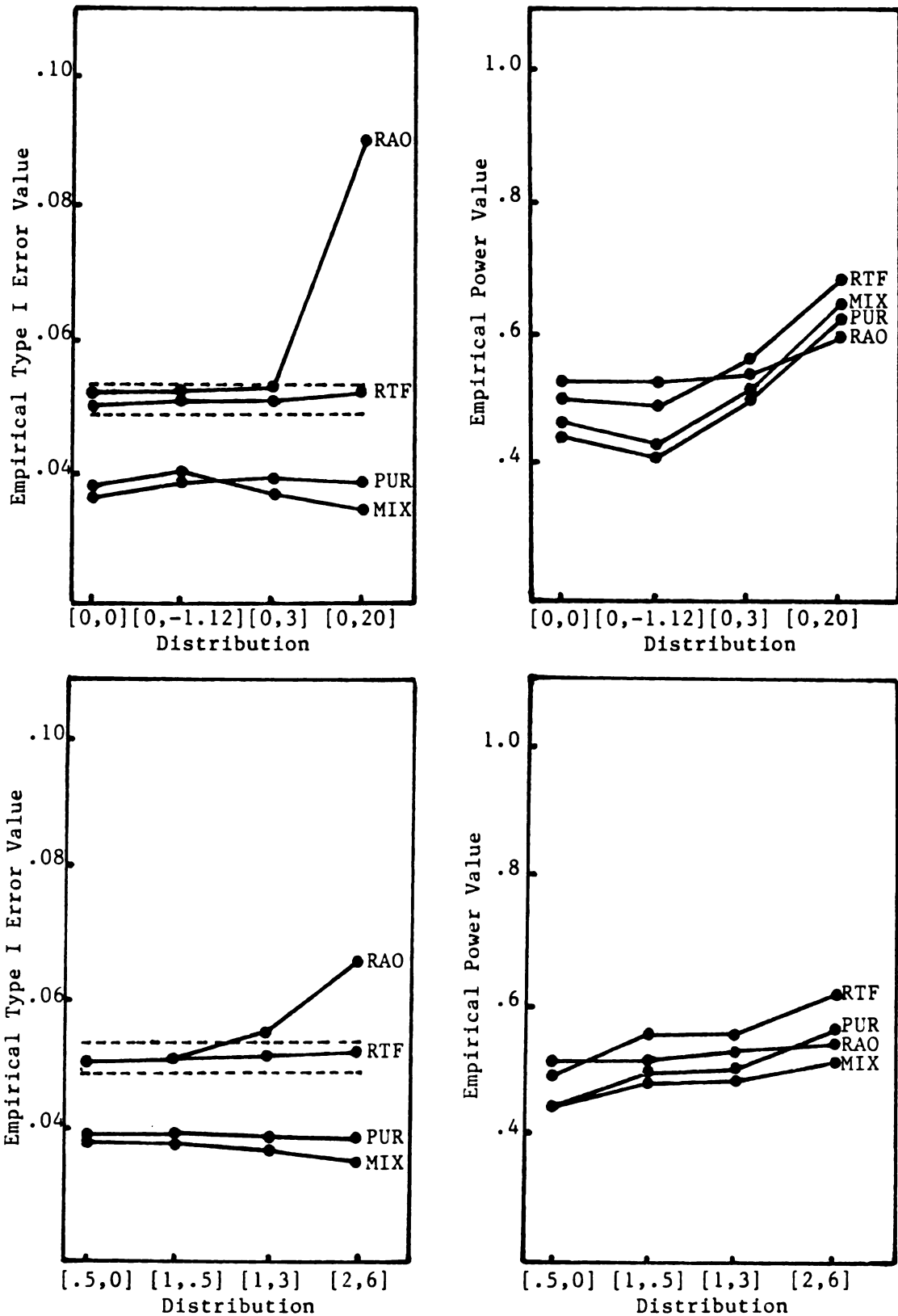


Figure 3. Type I Error and Power Curves by Distribution ($\alpha = .05$)

values on the Type I error and power values of the four tests. The average Type I error and power values and the total number of conservative and liberal Type I errors are presented in Table 10 and displayed in Figure 3. The Type I error results of Table 10 and Figure 3 indicate that the (a) Type I error rate of the Rao F test increased substantially only for the [1, 3], exponential, and Cauchy distributions, and (b) Type I error rates of the rank-transform Rao F and the pure- and mixed-rank tests were not affected by distribution.

The conservative and liberal Type I error results of Table 11 indicate that the Rao F test tended to produce more liberal Type I errors for the [1, 3], exponential, and Cauchy distributions. The percentage of liberal Type I errors increased from 0% for the normal distribution to 18% for the [1, 3] distribution, 83% for the exponential distribution, and 100% for the Cauchy distribution. In contrast, the number of liberal and conservative Type I errors for the rank-transform Rao F and the pure-rank tests did not seem to be affected by the degree of non-normality.

The power results of Table 10 and Figure 3 indicate that the (a) power of the Rao F test tended to increase only for the Cauchy distribution, and (b) power of the rank-transform Rao F , the pure-rank, and the mixed-rank tests tended to increase with increases in the kurtosis value of the parent distributions.

Sample Size

Three sample sizes were used to examine the effects of varying sample sizes on the Type I error and power values of the four tests. The average Type I error and power values and the total number of

conservative and liberal Type I errors are presented in Table 11 and displayed in Figure 4.

The Type I error results of Table 11 and Figure 4 indicate that the (a) Type I error rate of the Rao F tended to shrink toward the nominal alpha level as the sample size increased, (b) Type I error rate of the rank-transform Rao F test did not seem to be affected by increases in the sample size, and (c) Type I error rates of the pure- and mixed-rank tests tended to increase substantially toward the nominal alpha level with increases in the sample size.

Table 11

Average Type I Error and Power Values and Number of
Conservative and Liberal Errors by Sample Size^a

N	Type I Error			(C L)	Power		
	.01	.05	.10		.01	.05	.10
25							
RAO	.0159	.0593	.1088	(1 75)	.0967	.2395	.3500
RTF	.0108	.0515	.1012	(0 11)	.0938	.2437	.3581
PUR	.0026	.0287	.0747	(199 0)	.0299	.1611	.2922
MIX	.0020	.0264	.0702	(211 0)	.0184	.1302	.2566
50							
RAO	.0150	.0581	.1073	(2 50)	.2755	.4943	.6182
RTF	.0101	.0498	.0986	(6 1)	.2957	.5253	.6493
PUR	.0054	.0390	.0873	(144 0)	.2098	.4685	.6159
MIX	.0045	.0378	.0862	(159 0)	.1881	.4591	.6144
100							
RAO	.0145	.0564	.1070	(2 58)	.6694	.8411	.9037
RTF	.0103	.0503	.1031	(7 17)	.7128	.8689	.9231
PUR	.0077	.0452	.0970	(46 2)	.6671	.8549	.9162
MIX	.0073	.0442	.0952	(61 0)	.6892	.8711	.9276

^a The tabled values represent the average Type I error and power values across all distributions, within-set correlations, and numbers-of-variables (n=72, [.0491, .0509]). The number of conservative (C) and liberal (L) Type I errors are the total across three alpha levels (216 cases).

The conservative and liberal Type I error results of Table 12 indicate that the (a) number of liberal Type I errors of the Rao F decreased as the sample size increased, (b) number of liberal Type I errors of the rank-transform Rao F was largest for the moderate-large sample size, and (c) number of liberal Type I errors of the pure- and mixed-rank tests decreased substantially as the sample size increased.

The power results of Table 12 and Figure 4 indicate that (a) the power of all four tests increased substantially with increases in the sample size, and (b) the increment in the power values was higher for the pure- and mixed-rank tests as the sample size increased from small-moderate to moderate-large.

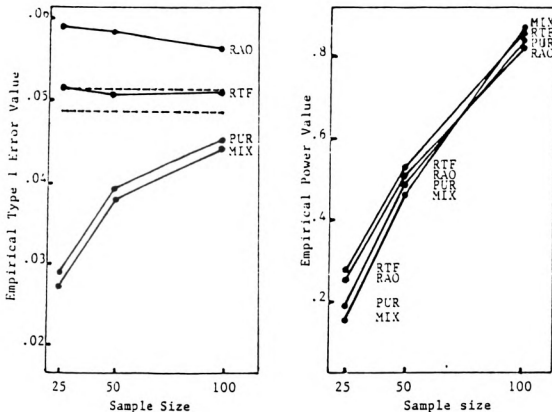


Figure 4. Type I Error and Power Curves by Sample Size ($\alpha = .05$)

Within-Set Correlation

Three combinations of the within-set correlation among the Y_j variables and among the X_k variables were included to examine the effect of varying correlations on the Type I error and power values of the four tests. The average Type I error and power values and the total number of conservative and liberal Type I errors are presented in Table 12 and displayed in Figure 5.

The Type I error results of Table 12 and Figure 5 indicate that the (a) Type I error rate of the Rao F tended to increase slightly as the within-set correlation among the Y_j variables and among the X_k variables increased, and (b) Type I error rates of the rank-transform

Table 12
Average Type I Error and Power Values and Number of Conservative
and Liberal Errors by Within-Set Correlations^a

(ρ_y, ρ_x)	Type I Error			(C L)	Power		
	.01	.05	.10		.01	.05	.10
(.3 .3)							
RAO	.0144	.0568	.1063	(3 56)	.3007	.4823	.5862
RTF	.0106	.0507	.1019	(3 14)	.3269	.5079	.6092
PUR	.0054	.0378	.0869	(127 1)	.2672	.4584	.5738
MIX	.0046	.0361	.0838	(139 0)	.2737	.4615	.5759
(.3 .7)							
RAO	.0150	.0579	.1080	(1 61)	.4354	.6072	.6970
RTF	.0103	.0506	.1007	(7 8)	.4537	.6251	.7137
PUR	.0053	.0376	.0863	(132 1)	.3781	.5698	.6771
MIX	.0045	.0359	.0838	(150 0)	.3716	.5594	.6666
(.7 .7)							
RAO	.0160	.0591	.1088	(1 66)	.3056	.4854	.5887
RTF	.0102	.0504	.1004	(3 7)	.3217	.5049	.6077
PUR	.0050	.0375	.0858	(130 0)	.2616	.4563	.5733
MIX	.0047	.0365	.0841	(142 0)	.2505	.4396	.5560

^a The tabled values represent the average Type I error and power values across all distributions, sample sizes, and numbers-of-variables ($n=72$, [.0491, .0509]). The number of conservative (C) and liberal (L) Type I errors is the total across three alpha levels (216 cases).

Rao F and the pure-and mixed-rank tests were not affected by the increment in the within-set correlations.

The conservative and liberal Type I error results of Table 12 indicate that the (a) number of liberal Type I errors of the Rao F was slightly higher for the largest within-set correlation, and (b) number of liberal Type I errors of the rank-transform Rao F and the pure- and mixed-rank tests varied slightly with the values of the within-set correlation.

The power results of Table 12 and Figure 5 indicate that the (a) increment in the correlation among the X_k variables tended to increase the power of all tests, (b) the increment in the correlation among the

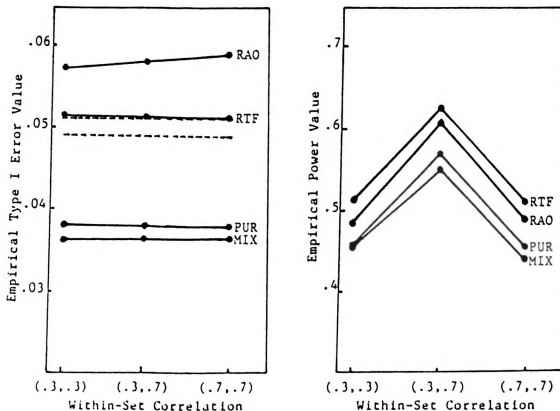


Figure 5. Type I Error and Power Curves by Within-Set Correlation ($\alpha = .05$)

Y_j variables tended to decrease the power of all tests, and (c) change in the power values for all tests was approximately the same.

Number of Variables

Three combinations of the number of Y_j and X_k variables were included to examine the effects of varying numbers-of-variables on the Type I error and power values of the four tests. The average Type I error and power values and the total number of conservative and liberal Type I errors are presented in Table 13 and displayed in Figure 6.

Table 13

Average Type I Error and Power Values and Number of Conservative Liberal Errors by Number-of-Variables^a

v^b	Type I Error				Power		
	.01	.05	.10	(C L)	.01	.05	.10
4							
RAO	.0137	.0556	.1031	(2 54)	.3452	.5302	.6291
RTF	.0107	.0506	.1002	(6 15)	.3650	.5493	.6476
PUR	.0065	.0425	.0934	(90 2)	.3216	.5222	.6342
MIX	.0057	.0414	.0913	(100 0)	.3342	.5376	.6491
6							
RAO	.0155	.0584	.1088	(2 65)	.3423	.5203	.6204
RTF	.0102	.0515	.1025	(3 7)	.3627	.5431	.6408
PUR	.0051	.0381	.0874	(135 0)	.2974	.4911	.6057
MIX	.0045	.0355	.0842	(153 0)	.2938	.4814	.5965
8							
RAO	.0162	.0599	.1112	(1 64)	.3542	.5245	.6224
RTF	.0102	.0496	.1002	(4 7)	.3746	.5455	.6422
PUR	.0041	.0323	.0783	(164 0)	.2878	.4711	.5844
MIX	.0036	.0316	.0761	(178 0)	.2678	.4415	.5529

^a The tabled values represent the average Type I error and power values across all distributions, sample sizes, and within-set correlations ($n=72$, $[\cdot0491, \cdot0509]$). The number of conservative (C) and liberal (L) Type I errors is the total across three alpha levels.

^b V = number of variables.

The Type I error results of Table 13 and Figure 6 indicate that the (a) Type I error rate of the Rao $\underline{F}$ tended to increase with increases in the number of variables, (b) Type I error rate of the rank-transform Rao $\underline{F}$ did not seem to be affected by the number of variables, and (c) Type I error rates of the pure- and mixed-rank tests tended to decrease as the number of variables increased.

The conservative and liberal Type I error results of Table 13 indicate that (a) the number of liberal Type I errors of the Rao $\underline{F}$ increased as the number of variables increased, (b) no particular pattern was found for the number of liberal Type I errors of the rank-transform Rao $\underline{F}$, (c) the number of conservative Type I errors of the

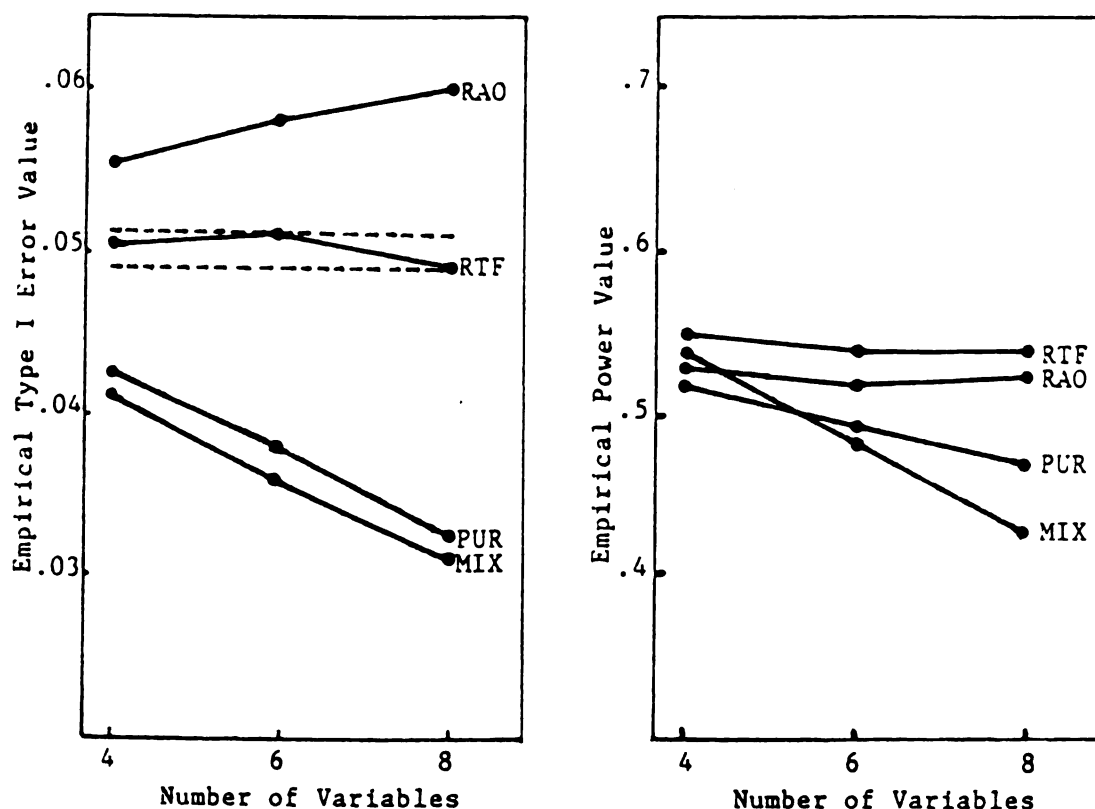


Figure 6. Type I Error and Power Curves by Number-of-Variables
($\alpha = .05$)

pure-and mixed-rank tests increased substantially as the number of variables increased.

The power results of Table 13 and Figure 6 indicate that the power of the Rao F and rank-transform Rao F tests did not vary with increases in the number of variables, and that the power of the pure-and mixed-rank tests tended to decrease as the number of variables increased.

Interaction Effects of Simulation Conditions

The next three sections present the Type I error and power results categorized by distribution and sample size, distribution and within-set correlation, and distribution and number-of-variables. Complete results appear in Tables E3, E4, and E5 of Appendix E.

Distribution and Sample Size

The average Type I error and power values and the total number of conservative and liberal Type I errors categorized by distribution and sample size for the .05 alpha level are presented in Table 14. The average Type I error and power values for the normal (thin-tailed), double-exponential (moderately non-normal/moderate-tailed), and exponential and Cauchy distributions (extremely non-normal/heavy-tailed) are displayed in Figures 7 and 8, respectively.

The Type I error results of Table 14 indicate that the (a) Type I error rate of the Rao F varied only slightly for the normal and mildly and moderately non-normal distributions and decreased substantially for the extremely non-normal distributions as the sample size increased, (b) Type I error rate of the rank-transform Rao F did

not vary much with sample size for all distributions, and (c) Type I error rates of the pure-and mixed-rank tests increased substantially toward the nominal alpha level as the sample size increased for all distributions.

The power results of Table 14 indicate that as the sample size increased the (a) power values of all four tests increased substantially for all distributions, (b) power values of the rank-transform Rao F and the pure- and mixed-rank tests increased at higher rates for all non-normal distributions other than the uniform and the $[-.5, 0]$ distributions, (c) power of the Rao F was largest for the normal, uniform, and the $[-.5, 0]$ distributions across all sample sizes, (d) power value of the rank-transform Rao F was largest for the $[1, .5]$, double-exponential, $[1, 3]$, exponential, and Cauchy distributions and small and small-moderate sample sizes, and (e) the mixed-rank test produced the highest power value for the double-exponential, $[1, 3]$, exponential, and Cauchy distributions and moderately-large sample size.

Table 14.
Average Type I Error and Power Values and Number of Conservative
and Liberal Errors by Distribution and Sample Size ($\alpha = .05$)^a

$[\gamma_1, \gamma_2]$	Type I Error			Power		
	N=25(C L)	N=50(C L)	N=100(C L)	N=25	N=50	N=100
[0, 0]						
RAO	.0474(0 0)	.0539(0 0)	.0511(0 0)	.2090	.4773	.8509
RTF	.0499(0 1)	.0500(0 0)	.0500(2 0)	.2047	.4548	.8202
PUR	.0274(9 0)	.0394(6 0)	.0440(3 0)	.1325	.3961	.8039
MIX	.0256(9 0)	.0423(6 0)	.0446(2 0)	.1253	.3995	.8159
[0, -.12]						
RAO	.0504(0 1)	.0532(0 0)	.0490(0 0)	.1951	.4674	.8420
RTF	.0528(0 1)	.0515(0 0)	.0483(0 0)	.1879	.4291	.7924
PUR	.0302(7 0)	.0404(5 0)	.0431(3 0)	.1219	.3740	.7734
MIX	.0299(7 0)	.0409(4 0)	.0439(3 0)	.1237	.3787	.7764
[.5, .0]						
RAO	.0521(0 0)	.0487(0 0)	.0512(0 0)	.2041	.4754	.8435
RTF	.0517(0 0)	.0484(0 0)	.0499(0 0)	.2023	.4563	.8191
PUR	.0287(9 0)	.0385(5 0)	.0458(1 0)	.1299	.4026	.8012
MIX	.0277(9 0)	.0366(6 0)	.0459(1 0)	.1182	.3980	.8106
[1, .5]						
RAO	.0514(1 1)	.0503(1 0)	.0483(0 0)	.2245	.4802	.8401
RTF	.0507(0 0)	.0501(0 0)	.0493(0 0)	.2417	.5315	.8885
PUR	.0282(9 0)	.0399(7 0)	.0450(3 0)	.1598	.4751	.8739
MIX	.0272(9 0)	.0385(6 0)	.0437(2 0)	.1164	.4352	.8849
[0, 3]						
RAO	.0535(0 0)	.0516(0 0)	.0519(0 0)	.2279	.5037	.8459
RTF	.0494(0 0)	.0521(0 0)	.0488(0 0)	.2319	.5184	.8700
PUR	.0279(9 0)	.0414(3 0)	.0439(4 0)	.1515	.4606	.8573
MIX	.0259(9 0)	.0376(6 0)	.0440(1 0)	.1339	.4735	.8826
[1, 3]						
RAO	.0566(0 3)	.0503(0 0)	.0557(0 4)	.2340	.4993	.8415
RTF	.0529(0 1)	.0472(1 0)	.0520(0 2)	.2372	.5139	.8665
PUR	.0291(9 0)	.0357(9 0)	.0463(3 0)	.1551	.4564	.8511
MIX	.0264(9 0)	.0355(9 0)	.0458(3 0)	.1253	.4474	.8768
[2, 6]						
RAO	.0665(0 9)	.0687(0 9)	.0603(0 5)	.2799	.5005	.8230
RTF	.0515(0 1)	.0516(0 0)	.0533(0 1)	.2951	.6144	.9366
PUR	.0287(8 0)	.0399(8 0)	.0479(0 0)	.1980	.5553	.9277
MIX	.0250(9 0)	.0361(9 0)	.0440(2 0)	.1074	.4707	.9427
[0, 20]						
RAO	.0961(0 9)	.0884(0 9)	.0840(0 9)	.3417	.5509	.8420
RTF	.0535(0 1)	.0471(1 0)	.0510(0 2)	.3469	.6839	.9577
PUR	.0296(7 0)	.0366(9 0)	.0459(2 0)	.2397	.6277	.9510
MIX	.0235(9 0)	.0351(9 0)	.0419(4 0)	.1918	.6698	.9784

^a The tabled values represent the average and the number of conservative (C) and liberal (L) Type I error rates and the average power values across all within-set correlations and numbers-of-variables ($n=9$, [.0474, .0526]).

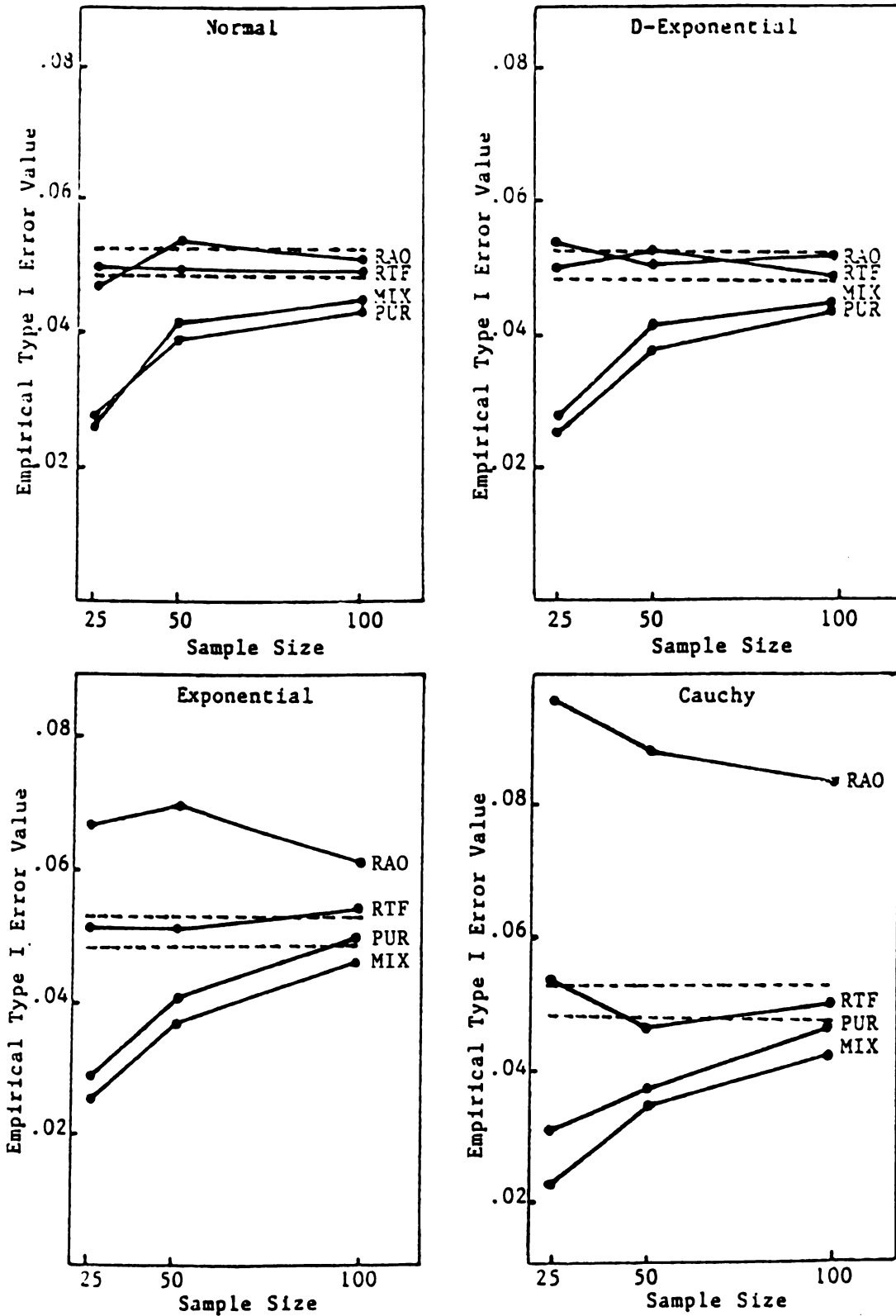


Figure 7. Type I Error Curves by Distribution and Sample Size ($\alpha = .05$)

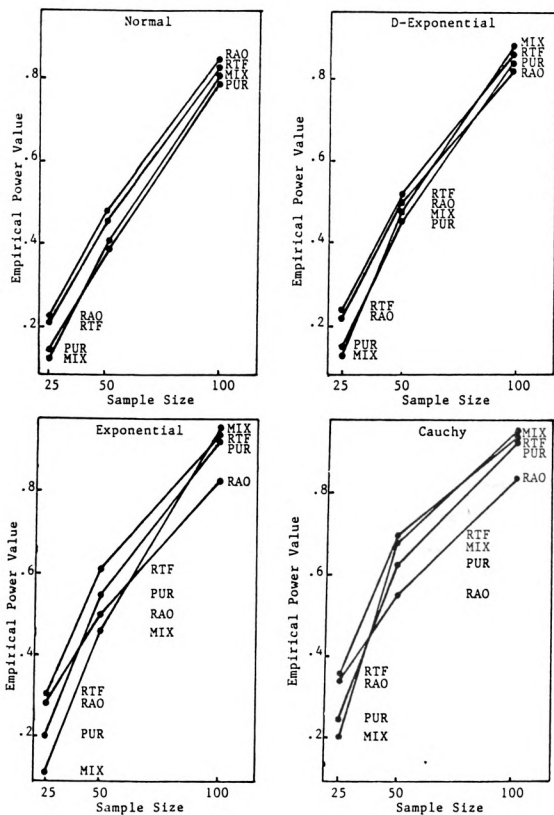


Figure 8. Power Curves by Distribution and Sample Size ($\alpha = .05$)

Distribution and Within-Set Correlation

The average Type I error and power values and the total number of conservative and liberal Type I errors categorized by distribution and within-set correlation for the .05 alpha level are presented in Table 15. The average Type I error and power values for the normal, double-exponential, exponential, and Cauchy distributions are displayed in Figures 9 and 10, respectively.

The Type I error results of Table 16 indicate that (a) the Type I error rates of the rank-transform Rao F and the pure- and mixed-rank tests were not affected by the within-set-correlation values for all distributions, and (b) except for the Cauchy in which the Type I error rate increased with increases in the within-set correlation, the Type I error rate of the Rao F was minimally affected by the within-set-correlation values.

The power results of Table 15 indicate that the power of all four tests increased as the within-set correlation among the predictor variables increased from .3 to .7, and decreased as the within-set correlation among the outcome variables increased from .3 to .7 for all distributions.

Table 15
Average Type I Error and Power Values and Number of Liberal
Errors by Distribution and Within-Set Correlation ($\alpha = .05$)^a

$[\gamma_1, \gamma_2]$	Type I Error			Power		
	(.3 .3)(C L)	(.3 .7)(C L)	(.7 .7)(C L)	(.3 .3)	(.3 .7)	(.7 .7)
[0, 0]						
RAO	.0508(0 0)	.0508(0 0)	.0508(0 0)	.4687	.5999	.4686
RTF	.0499(1 0)	.0504(0 1)	.0494(1 0)	.4472	.5814	.4511
PUR	.0369(6 0)	.0369(6 0)	.0369(6 0)	.4003	.5383	.4039
MIX	.0370(6 0)	.0370(6 0)	.0384(5 0)	.4070	.5288	.4050
[0, -1.12]						
RAO	.0514(0 0)	.0502(0 0)	.0510(0 1)	.4583	.5897	.4565
RTF	.0513(0 0)	.0507(0 0)	.0506(0 1)	.4257	.5560	.4296
PUR	.0385(5 0)	.0373(6 0)	.0379(4 0)	.3817	.5038	.3838
MIX	.0387(4 0)	.0377(6 0)	.0382(4 0)	.3837	.5084	.3866
[.5, .0]						
RAO	.0505(0 0)	.0505(0 0)	.0510(0 0)	.4653	.5919	.4658
RTF	.0496(0 0)	.0511(0 0)	.0493(0 0)	.4461	.5789	.4526
PUR	.0378(6 0)	.0382(4 0)	.0370(5 0)	.3997	.5251	.4090
MIX	.0369(6 0)	.0367(5 0)	.0366(5 0)	.4027	.5236	.4006
[1, .5]						
RAO	.0488(2 0)	.0495(0 0)	.0516(0 1)	.4710	.6013	.4725
RTF	.0502(0 0)	.0494(0 0)	.0505(0 0)	.5118	.6389	.5111
PUR	.0383(5 0)	.0374(7 0)	.0374(7 0)	.4636	.5822	.4631
MIX	.0355(5 0)	.0370(6 0)	.0369(6 0)	.4504	.5573	.4288
[0, 3]						
RAO	.0520(0 0)	.0517(0 0)	.0533(0 0)	.4831	.6087	.4857
RTF	.0506(0 0)	.0496(0 0)	.0501(0 0)	.5004	.6208	.4992
PUR	.0374(5 0)	.0381(6 0)	.0377(5 0)	.4521	.5652	.4520
MIX	.0359(5 0)	.0359(6 0)	.0357(5 0)	.4688	.5697	.4515
[1, 3]						
RAO	.0543(0 2)	.0542(0 2)	.0541(0 3)	.4821	.6084	.4842
RTF	.0511(0 2)	.0507(0 0)	.0503(1 1)	.4976	.6203	.4997
PUR	.0373(7 0)	.0367(8 0)	.0370(6 0)	.4482	.5646	.4498
MIX	.0357(7 0)	.0355(7 0)	.0365(7 0)	.4559	.5571	.4365
[2, 6]						
RAO	.0623(0 7)	.0666(0 8)	.0666(0 8)	.4922	.6123	.4989
RTF	.0514(0 0)	.0527(0 0)	.0523(0 2)	.5868	.6853	.5740
PUR	.0385(6 0)	.0389(5 0)	.0391(5 0)	.5318	.6272	.5220
MIX	.0354(7 0)	.0350(7 0)	.0348(6 0)	.5019	.5708	.4482
[0, 20]						
RAO	.0841(0 9)	.0898(0 9)	.0946(0 9)	.5381	.6453	.5513
RTF	.0512(0 2)	.0502(1 1)	.0503(0 0)	.6479	.7190	.6216
PUR	.0378(7 0)	.0371(5 0)	.0372(6 0)	.5894	.6617	.5672
MIX	.0334(8 0)	.0321(8 0)	.0351(6 0)	.6214	.6591	.5594

^a The tabled values represent the average and the number of conservative (C) and liberal (L) Type I error rates and the average power values across all sample sizes and numbers-of-variables ($n=9$, [.0474, .0526]).

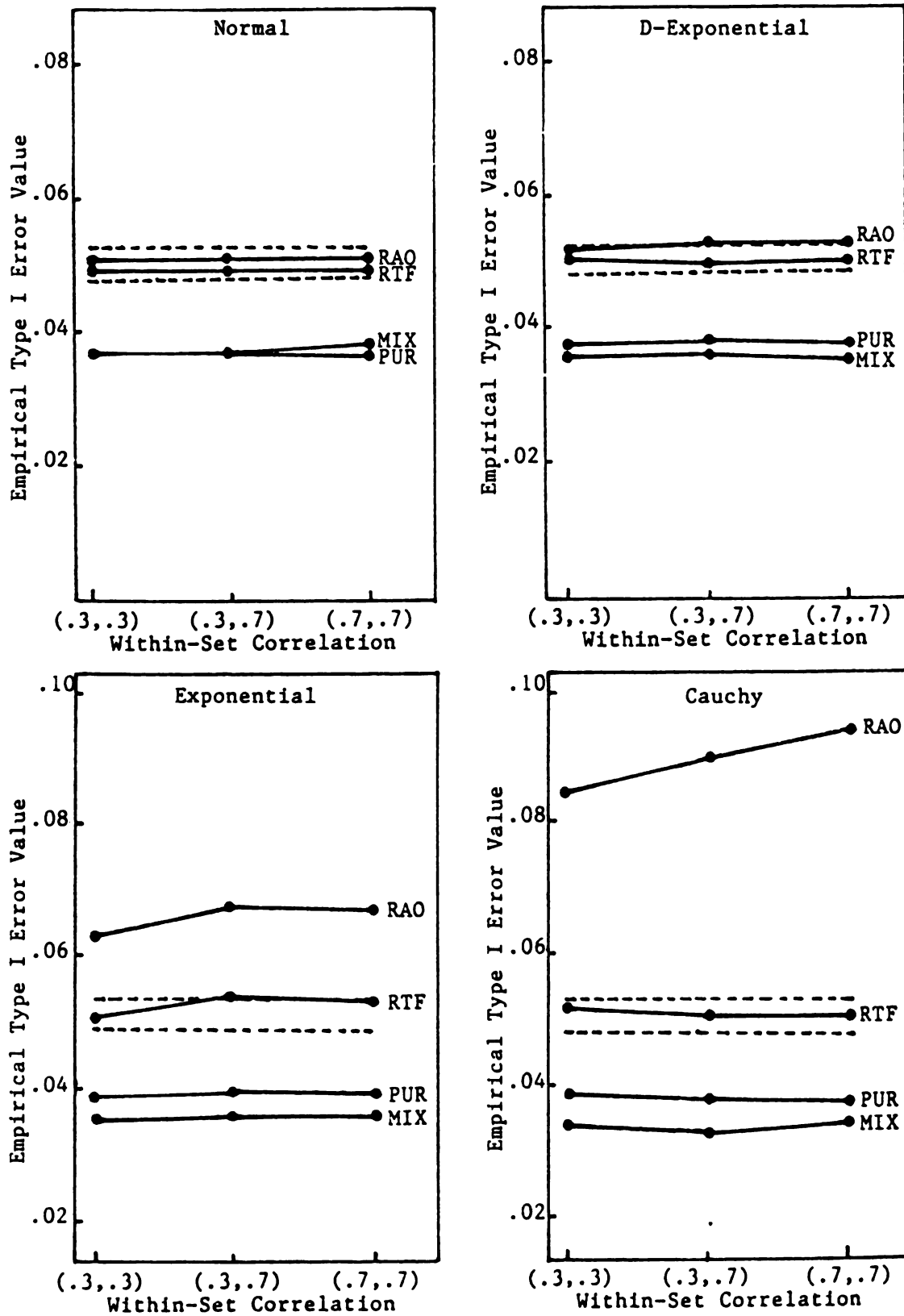


Figure 9. Type I Error Curves by Distribution and Within-Set Correlation ($\alpha = .05$)

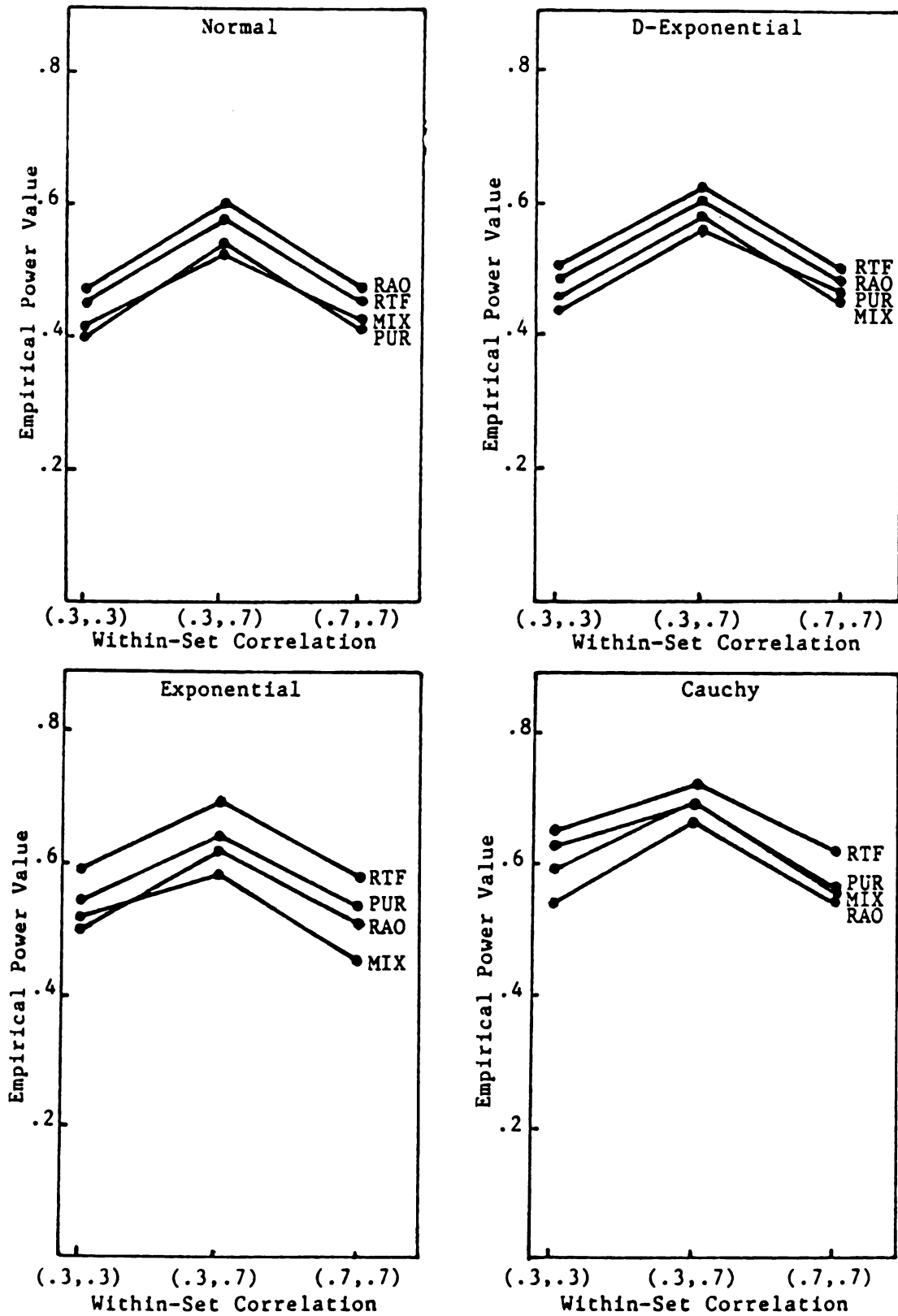


Figure 10. Power Curves by Distribution and Within-Set Correlation ($\alpha = .05$)

Distribution and Number-of-Variables

The average Type I error and power values and the total number of conservative and liberal Type I errors categorized by distribution and number-of-variables for the .05 alpha level are presented in Table 16. The average Type I error and power values for the normal, double-exponential, exponential, and Cauchy distributions are displayed in Figures 11 and 12, respectively.

The Type I error results of Table 16 indicate that (a) only for the exponential and Cauchy distributions did the Type I error rate of the Rao F increased as the number of variables increased, (b) the Type I error rate of the rank-transform Rao F was not affected by the number-of-variables factor for all distributions, and (c) the Type I error rates of the pure-and mixed-rank tests decreased as the number of variables increased for all distributions.

The power results of Table 16 indicate that the (a) power values of the Rao F and rank-transform Rao F tests were not affected by the number-of-variables factor for all distributions, and (b) power values of the pure- and mixed-rank tests decreased with increases in the number of variables for all distributions.

Table 16
Average Type I Error and Power Values and Number of Conservative and Liberal Errors by Distribution and Number-of-Variables ($\alpha = .05$)^a

$[\gamma_1, \gamma_2]$	Type I Error			Power		
	V=4(C L)	V=6(C L)	V=8(C L)	V=4	V=6	V=8
[0, 0]						
RAO	.0494(0 0)	.0528(0 0)	.0502(0 0)	.5273	.5031	.5067
RTF	.0468(2 0)	.0526(0 1)	.0504(0 0)	.4984	.4850	.4963
PUR	.0386(6 0)	.0389(6 0)	.0332(6 0)	.4719	.4340	.4267
MIX	.0409(5 0)	.0386(6 0)	.0329(6 0)	.4829	.4354	.4224
[0, -.12]						
RAO	.0531(0 1)	.0512(0 0)	.0483(0 0)	.5149	.4960	.4937
RTF	.0518(0 1)	.0519(0 0)	.0489(0 0)	.4784	.4646	.4683
PUR	.0439(2 0)	.0390(4 0)	.0308(9 0)	.4526	.4170	.3998
MIX	.0444(1 0)	.0392(4 0)	.0311(9 0)	.4578	.4186	.4023
[.5, .0]						
RAO	.0515(0 0)	.0504(0 0)	.0501(0 0)	.5141	.5045	.5044
RTF	.0500(0 0)	.0526(0 0)	.0474(0 0)	.4906	.4908	.4963
PUR	.0426(3 0)	.0393(5 0)	.0310(7 0)	.4640	.4402	.4295
MIX	.0422(3 0)	.0363(6 0)	.0317(7 0)	.4746	.4397	.4126
[1, .5]						
RAO	.0503(1 0)	.0485(1 0)	.0512(0 1)	.5223	.5067	.5159
RTF	.0494(0 0)	.0520(0 0)	.0487(0 0)	.5633	.5487	.5497
PUR	.0422(5 0)	.0387(6 0)	.0323(8 0)	.5363	.4970	.4755
MIX	.0426(3 0)	.0357(6 0)	.0311(8 0)	.5386	.4717	.4262
[0, 3]						
RAO	.0480(0 0)	.0553(0 0)	.0539(0 0)	.5335	.5233	.5207
RTF	.0487(0 0)	.0540(0 0)	.0476(0 0)	.5393	.5394	.5417
PUR	.0410(5 0)	.0402(3 0)	.0320(8 0)	.5130	.4879	.4685
MIX	.0400(3 0)	.0360(7 0)	.0316(6 0)	.5373	.4929	.4599
[1, 3]						
RAO	.0541(0 3)	.0531(0 3)	.0553(0 1)	.5260	.5237	.5250
RTF	.0510(1 2)	.0504(0 0)	.0507(0 1)	.5355	.5365	.5456
PUR	.0426(6 0)	.0362(8 0)	.0323(7 0)	.5083	.4859	.4685
MIX	.0419(6 0)	.0323(9 0)	.0334(6 0)	.5244	.4820	.4430
[2, 6]						
RAO	.0590(0 6)	.0662(0 9)	.0702(0 8)	.5320	.5330	.5385
RTF	.0526(0 1)	.0509(0 0)	.0529(0 1)	.6247	.6169	.6044
PUR	.0441(4 0)	.0378(6 0)	.0346(6 0)	.5954	.5604	.5251
MIX	.0394(7 0)	.0344(6 0)	.0314(7 0)	.5970	.5025	.4213
[0, 20]						
RAO	.0793(0 9)	.0894(0 9)	.0999(0 9)	.5714	.5723	.5909
RTF	.0543(0 3)	.0474(1 0)	.0500(0 0)	.6639	.6629	.6617
PUR	.0453(4 0)	.0344(7 0)	.0323(7 0)	.6630	.6067	.5756
MIX	.0397(6 0)	.0314(8 0)	.0294(8 0)	.6881	.6081	.5439

^a The tabled values represent the average and the number of conservative (C) and liberal (L) Type I error rates and the average power values across all sample sizes and within-set correlations ($n=9$, $[\gamma_1, \gamma_2]$).

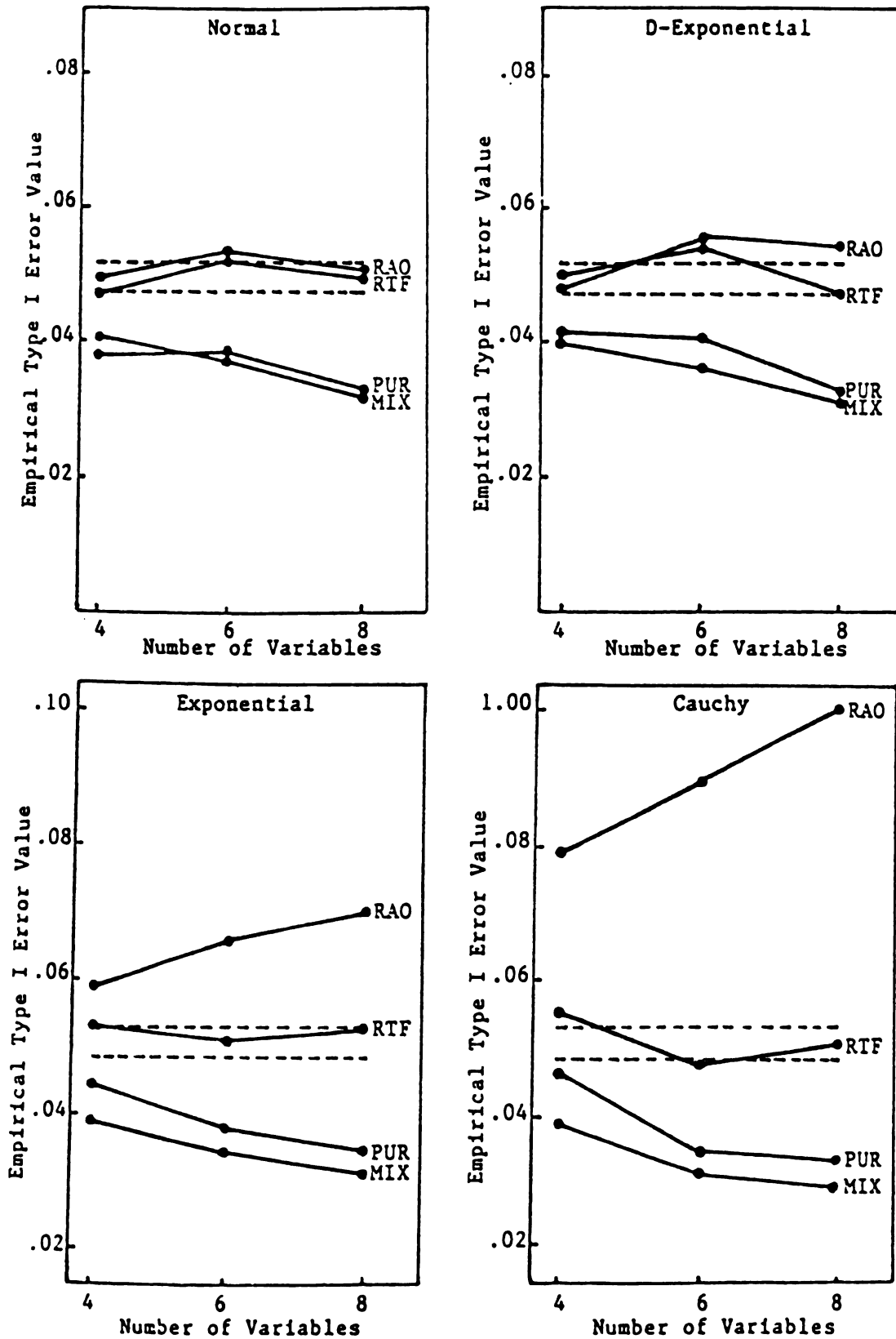


Figure 11. Type I Error Curves by Distribution and Number-of-Variables ($\alpha = .05$)

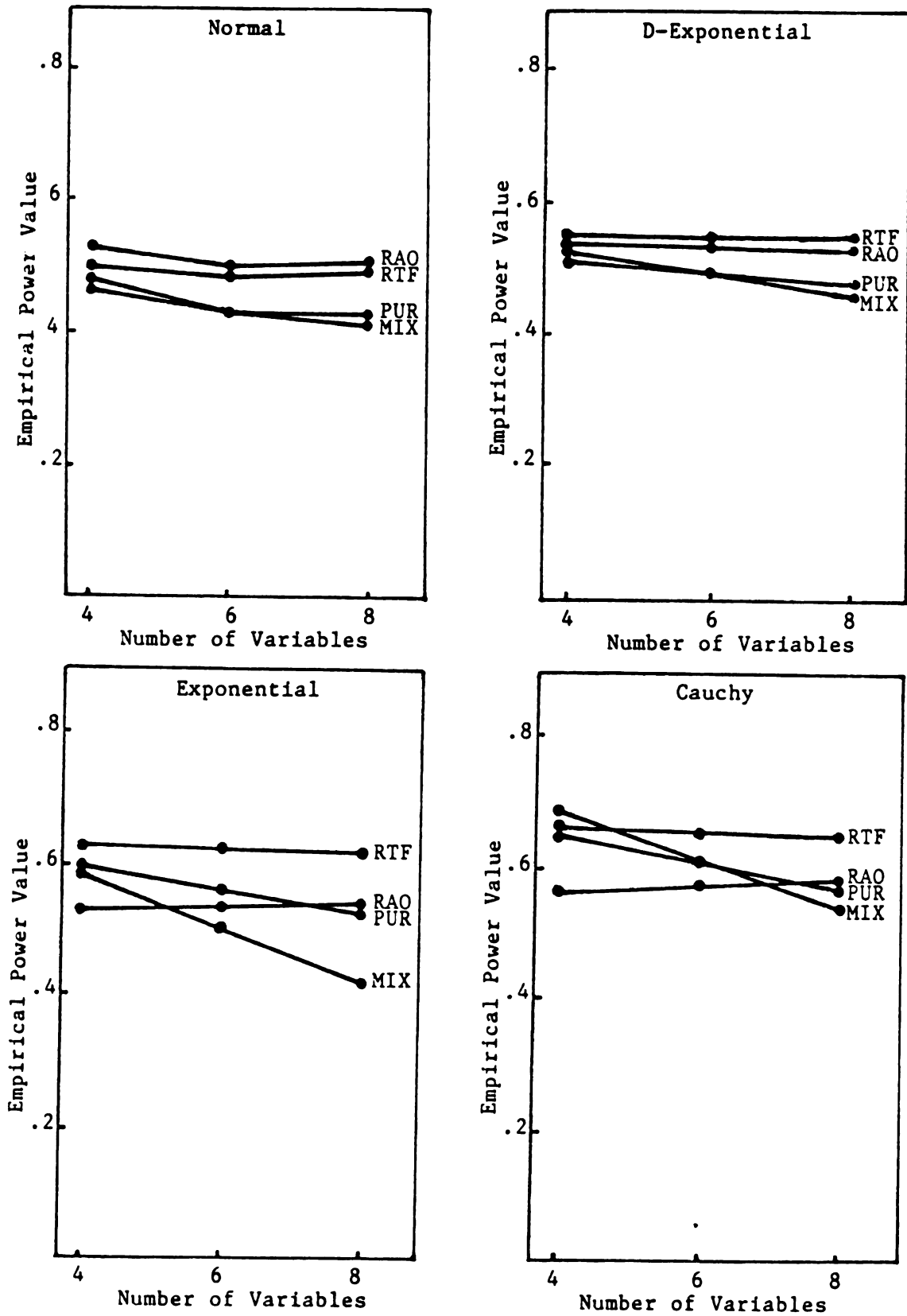


Figure 12. Power Curves by Distribution and Number-of-Variables ($\alpha = .05$)

Sample Size, Within-Set Correlation and Number-of-Variables

The results from previous sections indicate that the (a) Type I error rate of the Rao F increased with increases in the correlation among the Y_j and the X_k variables and the number-of-variables, particularly for the extremely non-normal distributions, and (b) Type I error and power values of the pure- and mixed-rank tests decreased with increases in the number-of-variables for all distributions. A further analysis was carried out to examine the interactions of sample size, within-set correlation, and number-of-variables.

The average and the total number of conservative and liberal Type I error rates for the Rao F test categorized by sample size and within-set-correlation for the .05 alpha level are presented in Table 17. The average and the total number of conservative and liberal Type I error rates and the average power values for the Rao F and the pure- and mixed-rank tests categorized by sample size and number-of-variables for the .05 alpha level are presented in Table 18.

Table 17

Average Type I Error and Power Values and Number of Conservative and Liberal Errors by Sample Size and Within-Set Correlation For Rao F Test ($\alpha = .05$)^a

N	(.3,.3)(C L)	(.3,.7)(C L)	(.7,.7)(C L)
25	.0579 (1 7)	.0589 (0 7)	.0609 (0 9)
50	.0570 (1 6)	.0584 (0 6)	.0591 (0 6)
100	.0554 (0 5)	.0565 (0 6)	.0574 (0 7)

^a The tabled values represent the average and the number of conservative (C) and liberal (L) Type I error rates across all distributions and numbers-of-variables ($n=24$, [.0484, .0516]).

The results of of Tables 17 and 18 indicate that the Type I error rate of the Rao F was only slightly reduced as the sample size increased for all combinations of the within-set-correlations and the number-of-variables. In both cases, the number of liberal Type I errors was only slightly reduced, even for the largest sample size.

Table 18

Average Type I Error and Power Values and Number of Conservative and Liberal Errors by Sample Size and Number-of-Variables
For Rao F and Pure- and Mixed-Rank Tests ($\alpha = .05$)^a

N	Type I Error			Power		
	V=4 (C L)	V=6 (C L)	V=8 (C L)	V=4	V=6	V=8
25						
RAO	.0570(1 7)	.0592(0 9)	.0615(0 7)	.2517	.2381	.2288
PUR	.0378(19 0)	.0281(24 0)	.0202(24 0)	.2086	.1572	.1174
MIX	.0349(22 0)	.0253(24 0)	.0190(24 0)	.1955	.1200	.0753
50						
RAO	.0565(0 6)	.0575(1 6)	.0605(0 6)	.5060	.4884	.4886
PUR	.0429(11 0)	.0400(17 0)	.0339(24 0)	.5061	.4640	.4354
MIX	.0430(9 0)	.0369(22 0)	.0336(24 0)	.5380	.4540	.3853
100						
RAO	.0533(0 6)	.0583(0 6)	.0576(0 6)	.8328	.8345	.8559
PUR	.0468(5 0)	.0461(4 0)	.0427(10 0)	.8519	.8522	.8607
MIX	.0462(3 0)	.0443(6 0)	.0422(9 0)	.8792	.8701	.8638

^a The tabled values represent the average and the number of conservative (C) and liberal (L) Type I error rates and the average power values across all distributions and within-set correlations ($n=24$, [.0484, .0516]).

The results of of Table 18 indicate that the (a) Type I error rates of the pure- and mixed-rank tests substantially increased as the sample size increased, particularly for the largest number of variables, whereas the number of conservative Type I errors decreased substantially for the largest sample size, and (b) power values of the

pure- and mixed-rank tests increased substantially with increases in the sample size, particularly for the largest number of variables.

These results suggest that increments in the sample size result in the pure- and mixed-rank tests being less conservative and more powerful as the number of variables increased, but it did not reduce the number of liberal Type I errors of the Rao F test for increasing within-set correlations and numbers-of-variables.

CHAPTER V

SUMMARY

Educational researchers opting for multivariate methods have historically employed multivariate-normal-theory procedures. The valid use of these procedures depends on the tenability of the underlying statistical assumptions (e.g., normality). The importance of satisfying these underlying assumptions can not be overemphasized, since violations have been shown to affect the distributional properties of normal-theory tests. The present study focused on the effect of non-normality of the observations on the Type I error and power properties of some selected normal-theory and nonparametric-multivariate tests. This chapter summarizes the (a) research questions, (b) methodology, (c) findings, and (d) implications for data analysis in educational research. Recommendations for further research in this area are also presented.

Research Questions

The following research questions were formulated (a) do varying skewness and kurtosis values affect the Type I error rate and power of normal-theory and nonparametric-multivariate tests, and (b) do sample size, within-set correlation, and number-of-variables influence the effects of skewness and kurtosis on the Type I error rate and power of normal-theory and nonparametric-multivariate tests?

Methodology

The present study used a computer simulation to empirically examine the Type I error and power values of the Bartlett, Rao F,

rank-transform Rao F , and the pure- and mixed-rank tests. Computer simulated data representing skewness and kurtosis values of eight multivariate distributions were used. Three combinations of correlations within the sets of predictors and dependent variables, three numbers-of-variables, and three sample sizes were used. All combinations of simulation factors were investigated.

The sample statistics of the simulated data showed excellent agreement between the average marginal mean, variance, intercorrelations, and skewness values of the resulting variables and their population counterparts. The average kurtosis values, however, indicated that the simulated data came from slightly more kurtic parent distributions than expected. In addition, Mardia's (1974) multivariate measure of skewness indicated that for a parent normal distribution increasing numbers-of-variables produced increasingly skewed data. Thus caution must be exercised in categorizing the data as multivariate-normal for a large number of variables. In this case, approximately-multivariate-normal may be a more appropriate description of the data.

The robustness of Type I error probabilities was determined using a 95% confidence interval about a particular (nominal) alpha level. Empirical Type I error rates exceeding the upper limit of these intervals were considered to be liberal and the values below the lower limit were considered to be conservative. Empirical power values were reported for all simulation conditions, even those with a liberal Type I error rate.

Findings

A total of five tests were included in the present study. However, the Type I error and power results of the Bartlett and Rao F tests were virtually identical, and hence only the results of the Rao F test were reported. The overall Type I error results showed that the Rao F test produced the highest proportion of liberal Type I errors, which were mostly confined to the (1, 3), exponential, and Cauchy distributions. The mixed-rank test produced the highest proportion of conservative Type I errors. These were mostly confined to the .01 alpha level for the small-sample-size and the largest number-of-variables conditions.

The overall power results showed that the average power value was highest for the rank-transform Rao F test, followed by the Rao F , pure-rank, and mixed-rank tests. However, the difference in the average power values among the tests was less than 7%. This pattern of Type I error and power results persisted across all three alpha levels.

The results of the present study showed that the Type I error rate of the normal-theory Rao F test increased substantially for the [1, 3], exponential, and Cauchy distributions. As expected, the Type I error rates of the three nonparametric tests were not affected by the degree of non-normality. The power of the normal-theory Rao F test increased substantially only for the Cauchy distribution. The power of the rank-transform Rao F and the pure- and mixed-rank tests increased with increases in the kurtosis value.

The results of previous studies have suggested that the Type I error rates of normal-theory tests increase with increases in the skewness value, but decrease with increases in the kurtosis value. The present study found that the Type I error rate of the normal-theory Rao F test increased substantially for the distributions with large skewness and kurtosis values. Similar results were found by Harwell and Serlin (1985). As for power, increasing skewness and/or kurtosis values have been shown to reduce the power of normal-theory tests.

Sample size has also been shown to affect the distributional properties of both normal-theory and nonparametric tests. The results of the present study showed that the Type I error rates of all normal-theory and nonparametric tests moved toward the nominal alpha level as the sample size increased. The power values of all five tests increased substantially with the sample size. All three nonparametric tests had a higher increment rate of the power values for distributions other than the normal, uniform, and the .5 skewness and 0 kurtosis combination (i.e., [.5, 0] distribution).

With respect to the within-set correlations, the results of the present study showed that the Type I error rate of the Rao F test increased as the correlation among the predictors and/or among the dependent variables increased only for the Cauchy distribution. The Type I error rates of the rank-transform Rao F and the pure- and mixed-rank tests were not affected by the increment in the within-set correlations for all distributions. The results also showed that the power values of all normal-theory and nonparametric tests increased with increases in the correlation among the predictors, and decreased

with increases in the correlation among the dependent variables for all distributions. Previous studies suggested that high correlation among dependent variables tended to reduce the power of only nonparametric- multivariate tests.

The results of the present study showed that Type I error rate of the normal-theory Rao F test increased with increases in the number of variables for the [1, 3], exponential, and Cauchy distributions. However, its power value was not affected by the number of variables for all distributions. The Type I error and power values of the rank-transform Rao F test were not affected by the number of variables for all distributions. The Type I error and power values of the pure-and mixed-rank tests decreased as the number of variables increased for all distributions.

Conclusions

The purpose of this study was to empirically evaluate the Type I error and power values of five selected normal-theory and nonparametric tests of the hypothesis of no relationship among two sets of variables under a variety of distribution, sample size, correlation among the predictor and dependent variables, and number-of-variables conditions. Recalling the research questions given earlier, the results of the present study lead to a number of conclusions. The generalization of these conclusions is limited to the simulation conditions examined.

(a) The Type I error rates of the normal-theory Bartlett and Rao F tests increase substantially for moderately-heavy and heavy-tailed distributions. Although the Type I error rates decrease with increases in the sample size, they remain liberal for heavy-tailed distributions.

The Type I error rates of the Bartlett and Rao F tests also increase with increases in the correlation among predictor and/or dependent variables, and with increases in the number-of-variables for heavy-tailed distributions.

(b) The Type I error rates of the nonparametric rank-transform Rao F and the pure- and mixed-rank tests are not affected by the form of the parent distribution for moderately-small and moderately-large samples. The Type I error rate of the rank-transform Rao F is slightly liberal for extremely non-normal distributions, while those of the pure- and mixed-rank tests are conservative for a small sample across all distributions. The Type I error rates of the pure- and mixed-rank tests move toward nominal alpha levels as the sample size increases.

(c) The Type I error rate of the rank-transform Rao F test is not affected by the within-set-correlation and the number-of-variables factors. The Type I error rates of the pure- and mixed-rank tests are not affected by the within-set-correlation factor, but decrease as the number of variables increases for all distributions.

(d) The power values of the normal-theory Bartlett and Rao F tests increase substantially only for extremely heavy-tailed distributions. The power values of the three nonparametric tests increase with increases in the kurtosis value. The power values of all five tests also increase with increases in the sample size and correlation among predictor variables, and decrease with increases in the correlation among dependent variables for all distributions. The increments due to the sample size are higher for the three nonparametric tests.

(e) The power values of the Bartlett and Rao F tests are largest for the light-tailed distributions across all sample sizes. The power value of the rank-transform Rao F test is largest for moderately-heavy and heavy-tailed distributions and small to moderately-small samples. The power value of the mixed-rank test is largest for moderately-heavy and heavy-tailed distributions for moderately-large samples.

(f) The power values of the Bartlett, Rao F , and the rank-transform Rao F tests are not affected by the number of variables for all distributions. The power values of the pure- and mixed-rank tests decrease as the number of variables increases for all distributions. However, the reduction in the power values tends to be compensated for by increases in the sample size.

Implications for Data Analysis in Educational Research

Based on the Type I error and power results, the following recommendations are suggested as guidelines for educational researchers in choosing the "best" test among the normal-theory Bartlett and Rao F tests and the nonparametric rank-transform Rao F and pure- and mixed-rank tests in testing the hypothesis of no relationship among two sets of variables. The criteria used for recommending a test as "best" are that the Type I error rate of the test is not liberal and its power is higher than its competitors. In the spirit of the Neyman-Pearson lemma the "best" test is a test that minimizes both Type I and Type II errors.

Note that the within-set correlation and the number-of-variables are not included as factors in determining for the "best" test because they have similar effects on Type I error and power

values of the five tests, or their effects depend on the type of distribution or the sample size. The recommendations are as follows:

- (a) The Bartlett and Rao F tests are recommended for all light-tailed distributions and any sample size.
- (b) The rank-transform Rao F test is recommended for moderately-heavy and heavy-tailed distributions for small and moderately-small samples.
- (c) The pure- or mixed-rank test is recommended for moderately and extremely heavy-tailed distributions and moderately-large (or larger) samples.

Using the percentage of the number of extreme observations (3 standard deviations away from the mean) as a measure of "tailedness" of a distribution, the simulated data indicated the following order of "tailedness" ("light" to "heavy"): Uniform (0%), normal (.28%), [.5, 0] (.48%), [1, .5] (.92%), double-exponential (1.41%), [1, 3] (1.45%), exponential (2.07%), and Cauchy (2.34%).

The results of the present study suggest that the Type I error rates of the Bartlett and Rao F tests were robust for the light-tailed distributions (i.e., uniform, normal, [.5, 0], [1, .5]). For the "heavy-tailed" distributions, the Bartlett and Rao F tests tended to produce liberal Type I error rates and smaller power values than that of the rank-transform Rao F test. For a moderate-large sample, the power values of all three nonparametric tests were larger than those of the two normal-theory tests.

As noted earlier, educationally-oriented variables are likely to produce data that are skewed or kurtic. For small to moderate-large samples (i.e. 25 to 100) the use of the nonparametric rank-transform Rao F test is recommended, while for larger samples (i.e. 100 or more) the use of the pure- or mixed-rank test is recommended when the observed data have more than 1% of the extreme values.

Recommendations for Further Study

The results of a simulation study are limited in their generalization to the conditions examined in the study. The present study confirmed many results from previous simulation studies and generated a more comprehensive set of guidelines for the appropriate use of the five tests. However, there is a need to investigate the effects of data conditions which were not examined in the present study. Based on the results and the limitations of the methodology used in the present study, the following recommendations for further research are suggested:

- (a) The present study was limited to the violation of the normality assumption. The assumptions of homogeneity of the elements of the covariance matrix and independence of observations have not been examined. Previous studies have suggested that the normal-theory and nonparametric procedures are not robust to violations of these assumptions. Knowledge of the behavior of the normal-theory and nonparametric tests (of the independence between two sets of variables) under violations of these assumptions would further help to determine the utility of these tests.

- (b) The distributional properties of the normal-theory and nonparametric tests should be examined for non-normal distributions with large skewness values and a zero kurtosis value. The data representing these distributions could not be generated for the present study.
- (c) The present study examined the behavior of test statistics using variables with the same (univariate) marginal distribution. A further study should be conducted to examine the behavior of these tests using variables with different marginal distributions.
- (d) The distributional properties of the Rao F and the pure- and mixed-rank tests should be examined using a larger variety of sample sizes. The results of the present study indicate that the Type I error rates of these tests move toward the nominal alpha levels as the sample size increases. However, it would be useful to examine the behavior of the tests for a wider range of sample sizes.
- (e) The distributional properties of the normal-theory and nonparametric tests should be examined for normal and non-normal distributions using various combinations of unequal within-set correlations among the predictor and dependent variables. The present study used only a single unequal, within-set correlation.

In conclusion, normal-theory tests of the hypothesis of no relationship among two sets of variables require the assumptions of independence, homogeneity of covariance, and normality. When the assumption of normality is not tenable the use of these tests may result in falsely accepting or rejecting a null hypothesis. The

development of nonparametric hypothesis-testing frameworks due to Conover and Iman and Puri and Sen provides alternatives to normal-theory procedures when the data do not meet the normality assumption. The results of the present study indicate that the Type I error rates of these nonparametric tests are not liberal and, for moderate samples, their power values are almost equal to those of their normal-theory alternatives for light-tailed distributions and superior for moderately and extremely heavy-tailed distributions. Hence, the three nonparametric tests examined should be routinely used by educational researchers.

APPENDICES

APPENDIX A

DEFINITION OF TERMS

Asymptotic relative efficiency. The asymptotic relative efficiency of test T_1 with respect to test T_2 is the limiting ratio of sample sizes n_2/n_1 , where n_2 and n_1 are the sample sizes required by test T_2 and test T_1 respectively such that both tests achieve equal power against equal alternatives that are "close to" the null hypothesis (Hollander and Wolfe, 1973, p. 439).

Canonical-Correlation analysis. A canonical-correlation analysis refers to a procedure employed to find correlation values (largest to smallest) among a set of one or more linear functions of two sets of random variables (Timm, 1975, p. 348). It is also used as a data-reduction method.

Conservative test. A test is conservative if the actual level of significance is smaller than the stated level of significance (Conover, 1980, p. 90).

Consistent test. A sequence of tests is consistent against all alternatives in the class H_1 if the power of the test approaches 1.0 as the sample sizes approaches infinity, for each fixed alternative possible under H_1 (Conover, 1980, p. 86).

Equivalent tests. Two statistical tests of the hypothesis H_0 are equivalent if, for each possible set of observations, the decision reached by one test agrees with the decision reached by the other test (Hollander and Wolfe, 1973, p. 447).

Hypothesis test. A hypothesis (significance) test is a decision rule which, on the basis of sample observations, either accepts or rejects the null hypothesis (Hollander & Wolfe, 1973, p. 450).

Liberal test. A test is liberal if the actual level of significance is greater than the stated level of significance.

Multivariate analysis of variance (MANOVA). A MANOVA refers to a procedure which is used to simultaneously compare a set of means of several outcome variables between several treatment populations (Timm, 1975, p. 369).

Multivariate-multiple-regression analysis. A multivariate-multiple-regression analysis refers to a procedure which is used to simultaneously explain the relationships among a set of outcome variables and a set of predictor variables, and predict a set of outcome values for a given set of predictor values (Timm, 1975, p. 307).

Most powerful test. A test is said to be most powerful for a specified alternative if no other test, at the same level of significance, has greater power against the same alternative (Gibbons, 1971, p. 15).

Power of a test. The power of a test against a specified alternative is the probability of (correctly) rejecting the null hypothesis when in fact the alternative is true (Hollander and Wolfe, 1973, p. 456).

Robust. A statistical procedure is said to be robust, with respect to a particular postulated assumption, if the procedure is relatively insensitive to (slight) departures from the assumption (Hollander & Wolfe, 1973, p. 460).

Test statistic. A test statistic is a statistic that determines the critical region of a hypothesis test (Hollander & Wolfe, 1973, p. 463).

Type I error. A Type I error is a false acceptance of the alternative hypothesis, that is, a rejection of the null hypothesis when in fact it is true (Hollander & Wolfe, 1973, p. 463).

Unbiased test. An unbiased test is a test in which the probability of rejecting H_0 when H_0 is false is always greater than or equal to the probability of rejecting H_0 when H_0 is true (Conover, 1980, p. 86).

Uniformly most powerful test. A test is uniformly most powerful against a class of alternatives if it is most powerful with respect to each specific alternative within the class of alternatives (Gibbons, 1971, p. 16).

APPENDIX B

PROCEDURES AND ALGORITHMS

Fleishman Procedure

Fleishman (1978) developed a technique for generating a non-normal deviate using a function involving the first three powers of a standard normal deviate. The procedure for obtaining the power-function constants is outlined here.

Let w_1 be a non-normal deviate and a , b , c , and d the power-function constants [see expression (25) of Chapter III]. Fleishman showed that for any distribution the expected value and variance of w_1 are given by:

$$E(w) = a + c , \quad (29)$$

$$\text{Var } (w) = b^2 + 6bd + 2c^2 + 15d^2 . \quad (30)$$

Assuming that the distribution is standardized expressions (29) and (30) become:

$$a + c = 0 , \quad (31)$$

$$b^2 + 6bd + 2c^2 + 15d^2 = 1 . \quad (32)$$

After considerable algebraic manipulation the skewness (γ_1) and kurtosis (γ_2) values for a desired distribution can be expressed in terms of b , c , and d (Fleishman, 1978):

$$\gamma_1 = 2c(b^2 + 24bd + 105d^2 + 2) , \quad (33)$$

$$\gamma_2 = 24[bd + c^2(1 + b^2 + 28bd) + d^2(12 + 48bd + 141c^2 + 225d^2)] . \quad (34)$$

The values for b, c, and d can be obtained by simultaneously solving equations (32), (33), and (34). The value for a can be obtained using expression (31).

Vale and Maurelli Procedure

Vale and Maurelli (1983) developed a procedure for computing the intermediate correlations among non-normal variables. The procedure is outlined below.

Let two standardized variables, W_1 and W_2 , be distributed as a bivariate non-normal distribution with a specified population correlation ($\rho_{W_1 W_2}$). Assume also that W_1 and W_2 have a common density function and thus possess common skewness and kurtosis values. Hence, the Fleishman power-function constants for W_1 and W_2 are identical. Let z_1 and z_2 be two standard normal deviates and Z_1 and Z_2 the vectors containing these deviates to the powers of zero through three:

$$Z_1 = [1 \quad z_1 \quad z_1^2 \quad z_1^3], \quad (35)$$

$$Z_2 = [1 \quad z_2 \quad z_2^2 \quad z_2^3]. \quad (36)$$

Let D be the vector containing the Fleishman power-function constants

$$D' = [-c \quad b \quad c \quad d]. \quad (37)$$

Using Fleishman's (1978) power function of expression (25), the non-normal deviates w_1 and w_2 can be defined as

$$w_1 = \underset{1 \times 4}{D'} * \underset{4 \times 1}{Z_1}, \quad w_2 = \underset{1 \times 4}{D'} * \underset{4 \times 1}{Z_2}. \quad (38)$$

Since W_1 and W_2 are standardized their correlation is equal to their

expected cross product

$$\begin{aligned}\rho_{W_1 W_2} &= E(W_1 W_2) \\ &= E(D' Z_1 Z_2' D) \\ &= D' P D\end{aligned}\quad (39)$$

where P is the expected matrix product of Z_1 and Z_2 :

$$P = E(Z_1 Z_2') = \begin{bmatrix} 1 & 0 & 1 & 0 \\ 0 & \rho_{z_1 z_2} & 0 & 3\rho_{z_1 z_2} \\ 1 & 0 & 2\rho_{z_1 z_2}^2 + 1 & 0 \\ 0 & 3\rho_{z_1 z_2} & 0 & 6\rho_{z_1 z_2}^3 + 9\rho_{z_1 z_2} \end{bmatrix} \quad (40)$$

Returning to the vector and matrix product of expression (40) the correlation $\rho_{W_1 W_2}$ is given by the scalar expression

$$\rho_{W_1 W_2} = \rho_{z_1 z_2} (b^2 + 6bd + 9d^2 + 2c^2 \rho_{z_1 z_2} + 6d^2 \rho_{z_1 z_2}^2). \quad (41)$$

The values of $\rho_{z_1 z_2}$ obtained by solving the polynomial of expression (41) provides the intermediate correlation matrix P^* .

Multivariate Skewness and Kurtosis

The sample measures of Mardia (1974) multivariate skewness and kurtosis values are given in expressions (27) and (28) of Chapter III.

The following algorithms were used to compute these values:

$$\gamma_{1,t} = N^{-2} \sum_{i=1}^N \sum_{j=1}^N g_{ij}^3, \quad (42)$$

$$\gamma_{2,t} = N^{-1} \sum_{i=1}^N g_{ii}^2 - t(t+2), \quad (43)$$

where g_{ij} is the (ij) th element of the matrix

$$\left(\begin{matrix} [U_1 - \bar{U}] & [U_2 - \bar{U}] & \dots & [U_N - \bar{U}] \end{matrix} \right) \underset{\text{Nxt txt}}{V^{-1}} \left(\begin{matrix} [U_1 - \bar{U}] & [U_2 - \bar{U}] & \dots & [U_N - \bar{U}] \end{matrix} \right), \quad (44)$$

txN

where $U_1, U_2, \dots, U_N$ and $\bar{U}$ are defined in expression (26) of Chapter III.

Tests of Multivariate Normality

Mardia (1970) introduced two tests of multivariate normality based on the sampling distributions of the multivariate skewness and kurtosis statistics. The null hypotheses for the tests are that both the population skewness and kurtosis values are equal to zero. The test statistics are:

$$A = (N\gamma_{1,t})/6 \sim \chi^2_{t(t+1)(t+2)/6} \quad (45)$$

$$B = \gamma_{2,t}/\sqrt{[8t(t+2)/N]} \sim N(0, 1), \quad (46)$$

where t is the number of variables.

The sample measures of the multivariate skewness and kurtosis and the corresponding test-statistic values are given in Table B1.

Table B1
Tests For Multivariate Normality (Mardia, 1974)

t^a	Measure	N^b	Value	Statistic	df^c	cv^d	dc^e
4	skewness	300	.3919	19.595	20	31.41	NS
	kurtosis	300	-.3094	-.387		-1.96	NS
6	skewness	300	1.0867	54.335	56	74.45	NS
	kurtosis	300	-.6064	-.536		-1.96	NS
8	skewness	300	2.3257	116.285	120	124.34	NS
	kurtosis	300	-1.0318	-.706		-1.96	NS

^a t = number of variables, ^b N = sample size, ^c df = degrees of freedom, ^d cv = critical value, ^e dc = decision, NS = not significant.

Computation of Power Value

The power value of the Rao F test was computed analytically using a method due to Muller and Peterson (1983) and the tabled power values of the F test due to Pearson and Hartley (1951). Given the matrices of regression coefficients (β) and the within-set correlations among the dependent variables (R_{11}) and among the predictor variables (R_{22}), the power value of the Rao F test can be determined using the following procedures:

- (1) Let R be the matrix of intercorrelations among the dependent and predictor variables

$$R = \begin{bmatrix} R_{11} & R_{12} \\ R_{21} & R_{22} \end{bmatrix} \quad (47)$$

R_{12} , the matrix of between-set correlations of the dependent and predictor variables, is given by (Timm, 1975, p. 309):

$$R_{12} = \beta R_{22}. \quad (48)$$

- (2) Wilks' lambda (Λ) can be obtained using the ratio of the determinants of R , R_{11} , and R_{22} (Anderson, 1958, p. 233)

$$\Lambda = |R| / (|R_{11}| |R_{22}|). \quad (49)$$

- (3) The non-centrality parameter for the Rao F can be obtained using the formulas given by Muller and Peterson (1983) and Pearson and Hartley (1951):

$$\lambda_{\Lambda} = (1 - \Lambda^{1/b}) / (\Lambda^{1/b} / \nu_2), \quad (50)$$

$$\phi = (\lambda_{\Lambda} / \nu_1 + 1)^{1/2} \quad (51)$$

where b , ν_1 , and ν_2 are given by expressions (17) of Chapter III.

- (4) The power value is obtained using the power charts of the F test (Pearson & Hartley, 1951) and the non-centrality parameter of expression (51). The power values of the Rao F test are given in Table B2.

Table B2
Theoretical and Empirical Power Values of the Rao F Test^a

Measure	Within-set correlation		
	(.3, .3)	(.3, .7)	(.7, .7)
β	.180	.180	.180
R_{12}	.234	.306	.306
$\hat{\lambda}$	.870	.831	.870
λ	13.798	18.680	13.798
$\phi^{\wedge}$	1.660	1.933	1.660
Theoretical power	.840	.940	.840
Empirical power	.815	.910	.815

^a Tabled power values were based on $N=100$ and $t = 4$. The empirical power values were obtained using 3,000 replications for a normal distribution.

Asymptotic Relative Efficiency

The efficiencies of two test statistics can be compared using their power values for various alternative hypotheses and sample sizes. However, a single measure of their relative efficiency can be obtained using the so-called asymptotic relative efficiency (A.R.E.). The computation of the A.R.E. of the pure-rank test to the normal-theory likelihood-ratio test for regression (Puri & Sen, 1985, pp. 316-317) is outlined here.

Let $(\underline{Z}' = (\underline{Y}', \underline{X}') = (Y_1, Y_2, \dots, Y_p, X_1, X_2, \dots, X_q))$ be a set of random variables with a normal density function (d.f.) $F(\underline{z})$, $F_1(\underline{x})$

be the marginal d.f. for the $\mathbf{X}$ variables, and $G_0(\mathbf{y} - \boldsymbol{\beta}_0 - \boldsymbol{\beta}'\mathbf{x})$ be the conditional d.f. of $\mathbf{Y}$ given $\mathbf{X} = \mathbf{x}$. Let $f(\mathbf{z})$, $f_1(\mathbf{x})$, and $g_0(\mathbf{y})$ be the probability distribution functions (p.d.f.'s) corresponding to the F , F_1 , and G_0 , respectively. Then $f(\mathbf{z})$ can be written in terms of the conditional and marginal p.d.f.'s as follows:

$$f(\mathbf{z}) = g_0(\mathbf{y} - \boldsymbol{\beta}_0 - \boldsymbol{\beta}'\mathbf{x})f_1(\mathbf{x}) \quad (52)$$

Under $H_0: \boldsymbol{\beta} = \mathbf{0}$, expression (52) becomes

$$f(\mathbf{z}) = g_0(\mathbf{y} - \boldsymbol{\beta}_0)f_1(\mathbf{x}) \quad (53)$$

The likelihood function can then be written as

$$L(\mathbf{Z}_1, \mathbf{Z}_2, \dots, \mathbf{Z}_n) = \prod_{i=1}^n f(\mathbf{Z}_i) = \prod_{i=1}^n g_0(\mathbf{Y}_i - \boldsymbol{\beta}_0 - \boldsymbol{\beta}'\mathbf{X}_i)f_1(\mathbf{X}_i). \quad (54)$$

Denoting the maximum likelihood estimator of $\boldsymbol{\beta}$ by $\hat{\boldsymbol{\beta}}$, the likelihood-ratio statistic is given by

$$\lambda_n = \prod_{i=1}^n \frac{g_0(\mathbf{Y}_i)}{g_0(\mathbf{Y}_i - \hat{\boldsymbol{\beta}}\mathbf{X}_i)} \quad (55)$$

Under H_0 (i.e., $\boldsymbol{\beta} = \mathbf{0}$) $-2 \log(\lambda_n)$ is asymptotically distributed as the central chi-square distribution with pq degrees of freedom. Under H_1 , the statistic has a noncentral chi-square distribution with pq degrees of freedom and a noncentrality parameter, say Δ_λ . The A.R.E. of the pure-rank test (L) with respect to λ_n is

$$e(L, \lambda) = \frac{\Delta_L}{\Delta_\lambda} \quad (56)$$

where ΔL is the noncentrality parameter for the chi-square distribution for the pure-rank test under H_1 . It has been shown that for a parent multivariate normal population the pure-rank test and the normal-theory likelihood-ratio test are asymptotically power equivalent (Puri & Sen, 1969). For other forms of continuous distributions the A.R.E. of the pure-rank test to the likelihood ratio-test is bounded below by 0.864.

APPENDIX C

The Pure-Rank Statistic: Puri and Sen Form

The construction of the pure-rank statistic (Puri & Sen, 1985, pp. 307-312) is outlined in this appendix. Let $[Y_i, X_i] = [Y_{1i}, Y_{2i}, \dots, Y_{pi}, X_{1i}, X_{2i}, \dots, X_{qi}]$, $i = 1, 2, \dots, N$, be a vector of random observations for the i (th) subject on $Y_1, Y_2, \dots, Y_p$ dependent and $X_1, X_2, \dots, X_q$ predictor variables having an identical $(p + q)$ -variate continuous distribution function. Let R_{ji} and R_{ki} represent the rank of the i (th) subject on the j (th) dependent and k (th) predictor variables, respectively. The original and the rank values of the Y_j and X_k can be represented as matrices H and R , respectively:

$$H = \begin{bmatrix} Y_{11} & Y_{12} & \dots & Y_{1N} \\ Y_{21} & Y_{22} & \dots & Y_{2N} \\ \vdots & \vdots & \ddots & \vdots \\ Y_{p1} & Y_{p2} & \dots & Y_{pN} \\ X_{11} & X_{12} & \dots & X_{1N} \\ X_{21} & X_{22} & \dots & X_{2N} \\ \vdots & \vdots & \ddots & \vdots \\ X_{q1} & X_{q2} & \dots & X_{qN} \end{bmatrix} \quad R = \begin{bmatrix} R_{11} & R_{12} & \dots & R_{1N} \\ R_{21} & R_{22} & \dots & R_{2N} \\ \vdots & \vdots & \ddots & \vdots \\ R_{p1} & R_{p2} & \dots & R_{pN} \\ R_{p+1,1} & R_{p+1,2} & \dots & R_{p+1,N} \\ R_{p+2,1} & R_{p+2,2} & \dots & R_{p+2,N} \\ \vdots & \vdots & \ddots & \vdots \\ R_{p+q,1} & R_{p+q,2} & \dots & R_{p+q,N} \end{bmatrix} \quad (57)$$

Since the N vectors of observations are independent each row of R represents a permutation of integers $1, 2, \dots, N$ (assuming no ties by virtue of continuity of the distribution function of the Y_j and X_k), with a total of $N!$ permutations. Since R contains $(p + q)$ rows, under the truth of the hypothesis of independence of the Y_j and X_k , the total number of possible realizations of R is $(N!)^{(p+q)}$.

Following Chatterjee and Sen (1964), two rank matrices are said to be permutationally equivalent if it is possible to obtain the second matrix by permutations of the columns of the first matrix. Suppose the columns of R are rearranged in such a way that the first row has the elements in the natural order $1, 2, \dots, N$, and denote the corresponding matrix by R^* . R is said to be permutationally equivalent to R^* if it is possible to obtain R^* by permutations of the columns of R . Since there are N columns in R , for a given R^* there will be a total of $N!$ possible realizations of R which are permutationally equivalent to R^* .

In general, the probability distribution of R over $(N!)^{(p+q)}$ possible realizations of R depends on the joint distribution of the Y_j and X_k . However, given a particular R^* the conditional distribution of R over the $N!$ permutations of the columns of R^* is uniform under the truth of the hypothesis of independence of the Y_j and X_k . An exact test of independence of Y_j and X_k may be computed using the distribution of R . However, the arithmetic is excessive and a large-sample approximation is of principal interest.

Puri and Sen (1985, pp. 307-312) presented a large-sample test based on the sum-of-cross-products vector $\underline{S}$ of the centered R_j and R_k and the covariance matrices $\underline{M}$ of the R_j and $\underline{C}$ of the R_k with elements

$$s_{jk} = \sum_{i=1}^N (R_{ji} - \bar{R}_j)(R_{ki} - \bar{R}_k), \quad (58)$$

$$m_{jj'} = N^{-1} \sum_{i=1}^N (R_{ji} - \bar{R}_j)(R_{j'i} - \bar{R}_{j'}) \quad j, j' = 1, 2, \dots, p, \quad (59)$$

$$c_{kk'} = N^{-1} \sum_{i=1}^N (R_{ki} - \bar{R}_k)(R_{k'i} - \bar{R}_{k'}) \quad k, k' = 1, 2, \dots, q, \quad (60)$$

where $\bar{R}_j$, $\bar{R}_{j'}$, $\bar{R}_k$, and $\bar{R}_{k'}$, are the rank means for the j (th) and j' (th) dependent and k (th) and k' (th) predictor variables, respectively. Puri and Sen showed that for a large N the expected values of $\underline{S}$ and $\underline{S}'\underline{S}$ are:

$$E(\underline{S}) = \underline{0}, \quad (61)$$

$$E(\underline{S}'\underline{S}) = N(\underline{M} \otimes \underline{C}), \quad (62)$$

where $(\underline{M} \otimes \underline{C})$ represents the Kronecker product of $\underline{M}$ and $\underline{C}$ (see Anderson, 1958, p. 347). The large-sample pure-rank statistic (L) is given by

$$L = N^{-1}[\underline{S} (\underline{M} \otimes \underline{C})^{-1} \underline{S}']. \quad (63)$$

Puri and Sen showed that the L statistic is distributed as a chi-square variable with pq degrees of freedom when the Y_j and X_k variables are independent. The L statistic is genuinely distribution-free for $p = q = 1$, but only (conditionally) permutationally distribution-free for $p > 1$ or $q > 1$.

The Pure-Rank Test: Harwell and Serlin Form

Harwell and Serlin (1985) derived a simpler form of the Puri and Sen L statistic using canonical correlations among the R_j and R_k values. Defining $\hat{\underline{\beta}} = \underline{S} \underline{C}^{-1}$, where $\hat{\underline{\beta}}$ is a matrix of sample regression coefficients based on ranks, Harwell and Serlin showed that the L statistic has the definitional form:

$$L = \sum_{j=1}^N \sum_{j'=1}^N \sum_{k=1}^N \sum_{k'=1}^N d^{jk,j'k'} s_{jk} s_{j'k'}, \quad (64)$$

$$L = \sum_{j=1}^N \sum_{j'=1}^N \sum_{k=1}^N \sum_{k'=1}^N m^{jj'} c^{kk'} s_{jk} s_{j'k'},$$

where $d^{jk,j'k'}$ represent the elements of the $(\underline{M}^{-1} \otimes \underline{C}^{-1})$, and $m^{jj'}$ and $c^{kk'}$ represent the elements of $\underline{M}^{-1}$ and $\underline{C}^{-1}$, respectively.

In this definitional form the elements of $\underline{S}$ appear as part of a quadruple sum across products of the elements of $\underline{M}$ and $\underline{C}$, thus making the computation of expression (64) extremely difficult. In the derivation of the simpler form of the L statistic, Harwell and Serlin showed that the summation of expression (64) may be written as the matrix product:

$$\begin{aligned} L &= (N-1) \text{Tr} (\hat{\underline{S}} \underline{M}^{-1} \hat{\underline{S}} \underline{C}) \\ &= (N-1) \text{Tr} (\underline{S} \underline{C}^{-1} \underline{C}^{-1} \underline{S}' \underline{C} \underline{M}^{-1}) \\ &= (N-1) \text{Tr} (\underline{S} \underline{C}^{-1} \underline{S}' \underline{M}^{-1}) \\ &= (N-1) (\text{sum of eigenvalues of resulting matrix}). \end{aligned} \tag{65}$$

The eigenvalues (squared canonical correlations) of the resulting matrix are the eigenvalues among the sets of the R_j and R_k values, θ_r ($r = 1, 2, \dots, s$). The L statistic can then be written in the form

$$L = (N-1) \sum_{r=1}^s \theta_r, \quad r = 1, 2, \dots, s. \tag{66}$$

The L statistic, as shown by Puri and Sen, is asymptotically distributed as a central chi-square variable with pq degrees of freedom when the Y_j and X_k variables are independent.

APPENDIX D

COMPUTER PROGRAMMING

This appendix describes the subroutines of the Statistical Package for the Social Sciences (SPSSX 2.2) and the International Mathematical and Statistical Libraries (1983) used in the present study. The subroutine FACTOR of the SPSSX was used to obtain the principal component weights of the intermediate correlation matrices. The following subroutines of the IMSL were used for data generation and computation of test statistics:

NEQNF	To solve a set of simultaneous nonlinear equations for Fleishman power function constants
GGUBS	To generate uniform random deviates
VMULFF	To multiply matrices
VMULFM	To multiply the transpose of a matrix A by a matrix B
VMULFP	To multiply a matrix A by the transpose of a matrix B
LINV1F	To compute the inverse of a matrix
NMRANK	To rank the generated deviates
EIGRF	To compute eigenvalues of a matrix
LINV3F	To compute determinant of a matrix

The complete listing of the computer program, which was coded in FORTRAN V, is given below.

```

PROGRAM NONPAR
CC-----
CC TYPE I ERROR RUN (NV=4,NP=2,NQ=2,CY=0.3,CX=0.3,NS=25,NL=3000).
CC DISTRIBUTION: NORMAL AND UNIFORM.
CC-----
      INTEGER NL,NS,TS,NP,NQ,NV,AL,NT,WA,BC
      PARAMETER (NL=3000,NV=4,NS=25,TS=100,NP=2,NQ=2,AL=3,NT=5,WA=8,
&              BC=1)
      INTEGER DGT,EJOB,DJOB,DIS,TST,REJB(AL),REJR(AL),REJT(AL),
&              REJP(AL),REJM(AL),LOOP,IER,TN,IR(NS)
      REAL     RB(NQ,NP),PE(NP,NP),PE1(NP,NP),PE2(NP,NP),PX(NQ,NQ),
&              PX1(NQ,NQ),PX2(NQ,NQ),PH,EPS,D1,AA,BB,CC,DD,SUMD(NV),
&              SUMD1(NV),SUMD2(NV),SUMD3(NV),SUMD4(NV),AVD(NV),
&              AVD1(NV),AVD2(NV),AVD3(NV),AVD4(NV),SUMMS,SUMMK,SUMS,
&              SUMK,UD(TS),ND(TS),UNE(NP,NS),UNX(NQ,NS),MNE(NP,NS),
&              MNX(NQ,NS),MNNE(NP,NS),MNNX(NQ,NS),MBX(NP,NS),CY,CX,
&              MNNY(NP,NS),DATD(NV,NS),VAR(NV),DEVD(NV,NS),AAA(NV,NV),
&              COVN(NV,NV),COVI(NV,NV),WK(WA),DCV(NS,NV),DCD(NS,NS),
&              MULS,MULK,TMEAN,TVARN,TSKEW,TKURT,ACOVN(NV,NV),VARN(NV),
&              SKEW(NV),KURT(NV),OMEAN,OVARN,OUSKW,OUKUR,OMSKW,OMKUR,
&              CORR(NV,NV),CORL(NV,NV),BCR(BC),D2,DCOR,RCP(NV,NV),
&              SUMCC(NV,NV)
      REAL     AYY(NP,NP),AYX(NP,NQ),AXX(NQ,NQ),DTA(NS),R(NS),DTR(NS),
&              S,T,DATR(NV,NS),AVR,DEVR(NV,NS),RRR(NV,NV),RYY(NP,NP),
&              RYX(NP,NQ),RXX(NQ,NQ),DVYR(NP,NS),DVXD(NQ,NS),
&              MYX(NP,NQ),AYYI(NP,NP),AXXI(NQ,NQ),RYYI(NP,NP),
&              RXXI(NQ,NQ),MATD1(NP,NQ),MATD2(NQ,NP),MATD(NP,NP),A1,
&              MATR1(NP,NQ),MATR2(NQ,NP),MATR(NP,NP),MATM1(NP,NQ),B1,
&              MATM2(NQ,NP),MATM(NP,NP),PED,PER,SER,SEM,B2,V1,V2,V3,
&              BAR,RAO,RTF,PRN,MRN,CS1,CS2,CS3,CF1,CF2,CF3
      COMPLEX ED(NP),ER(NP),Z(NP,NP)
      DOUBLE   PRECISION DSEED
      DSEED    =66901.D0
CC-----
CC (1) TO SPECIFY SIMULATION CONDITIONS.
CC-----
CC (1A) TO SPECIFY REGRESSION COEFFICIENT MATRIX (RB), PRINCIPAL
CC COMPONENT WEIGHTS FOR ERRORS (PE) AND PREDICTORS (PX), WITHIN-
CC SET CORRELATION (CY,CX), PHI(PH), TIE VALUE (EPS) FOR RANKING,
CC PARAMETER VALUES FOR THE INVERSE (DGT), EIGENVALUE (IZ,IJOB),
CC AND DETERMINAT (D1,DJOB) PROCEDURES.
      OPEN(20,FILE='TESTS')
      DATA RB / .0, .0, .0, .0 /
      DATA PE1/ .80623, .80623, -.59161, .59161/
      DATA PX1/ .80623, .80623, -.59161, .59161/
      DATA PE2/ .81475, .81475, -.57981, .57981/
      DATA PX2/ .81475, .81475, -.57981, .57981/
      CY=0.3
      CX=0.3
      TN=NS*NL
      PH=3.1428571
      EPS=0.000001
      DGT=0
      EJOB=0
      DJOB=4
      D1=1.0
CC (1B) TO OBTAIN A1, B1, B2, V1, V2 FOR COMPUTING TEST STATISTICS
CC AND TO SPECIFY CRITICAL VALUES: F .01, .05, .10 (CF1,CF2,CF3)
CC AND CHI-SQUARE .01, .05, .10 (CS1,CS2,CS3).
      A1=(NS-1.)-(NP+NQ+1.)/2.

```

```

B1=SQRT((NP*NP*NQ*NQ-4.)/(NP*NP+NQ*NQ-5.))
B2=1./B1
V1=NP*NQ
V2=A1*B1-V1/2.+1.
V3=V2/V1
IF (V1 .EQ. 4) THEN
  CS1=13.277
  CS2= 9.488
  CS3= 7.779
  IF (NS .EQ. 25) THEN
    CF1= 3.800
    CF2= 2.590
    CF3= 2.080
  ENDIF
  IF (NS .EQ. 50) THEN
    CF1= 3.526
    CF2= 2.466
    CF3= 2.008
  ENDIF
  IF (NS .EQ. 100) THEN
    CF1= 3.416
    CF2= 2.422
    CF3= 1.972
  ENDIF
ENDIF
IF (V1 .EQ. 9) THEN
  CS1=21.666
  CS2=16.919
  CS3=14.684
  IF (NS .EQ. 25) THEN
    CF1= 2.820
    CF2= 2.090
    CF3= 1.770
  ENDIF
  IF (NS .EQ. 50) THEN
    CF1= 2.579
    CF2= 1.976
    CF3= 1.689
  ENDIF
  IF (NS .EQ. 100) THEN
    CF1= 2.491
    CF2= 1.924
    CF3= 1.657
  ENDIF
ENDIF
IF (V1 .EQ. 16) THEN
  CS1=32.000
  CS2=26.296
  CS3=23.542
  IF (NS .EQ. 25) THEN
    CF1= 2.2359
    CF2= 1.841
    CF3= 1.604
  ENDIF
  IF (NS .EQ. 50) THEN
    CF1= 2.145
    CF2= 1.727
    CF3= 1.529
  ENDIF
  IF (NS .EQ. 100) THEN

```

```

        CF1= 2.065
        CF2= 1.682
        CF3= 1.494
        ENDIF
    ENDIF
CC      (1B) TO SECIFY FLEISHMAN'S CONSTANTS FOR NORMAL AND UNIFORM DIS.
        DIS=1
        WRITE(20,1)
1        FORMAT(/,'----- 1. NORMAL (0, 0) -----')
        AA=0.0
        BB=1.0
        CC=0.0
        DD=0.0
        DO 2 I=1,NP
        DO 3 J=1,NP
        PE(I,J)=PE1(I,J)
3        CONTINUE
2        CONTINUE
        DO 4 I=1,NQ
        DO 5 J=1,NQ
        PX(I,J)=PX1(I,J)
5        CONTINUE
4        CONTINUE
        GOTO 55
6        WRITE(20,7)
7        FORMAT('----- 2. UNIFORM (0, -1.12) -----')
        AA=0.0
        BB=1.34891701
        CC=0.0
        DD=-.13265955
        DO 8 I=1,NP
        DO 9 J=1,NP
        PE(I,J)=PE2(I,J)
9        CONTINUE
8        CONTINUE
        DO 10 I=1,NQ
        DO 11 J=1,NQ
        PX(I,J)=PX2(I,J)
11       CONTINUE
10       CONTINUE
CC      (1C) TO SET THE NUMBER OF REJECTIONS TO ZERO: BARTLETT (REJB),
CC      RAO F (REJR), RANK-TRANSFORM (REJT), PURE-RANK (REJP), AND
CC      MIXED-RANK (REJM).
55       DO 56 I=1,AL
        REJB(I)=0
        REJR(I)=0
        REJT(I)=0
        REJP(I)=0
        REJM(I)=0
56       CONTINUE
CC      (1C) TO SET SUMS TO ZERO FOR MULTIVARIATE SKEWNESS (SUMMS) AND
CC      KURTOSIS (SUMMK), RAW-SCORE CROSS PRODUCTS (SUMCC), RAW SCORES
CC      (SUMD1), RAW-SCORE SQUARES (SUMD2), RAW-SCORE CUBES (SUMD3), AND
CC      RAW-SCORE DUADS (SUMD4).
        SUMMS=0.0
        SUMMK=0.0
        DO 57 I=1,NV
        DO 58 J=1,NV
        SUMCC(I,J)=0.0
58       CONTINUE

```

```

SUMD1(I)=0.0
SUMD2(I)=0.0
SUMD3(I)=0.0
SUMD4(I)=0.0
57    CONTINUE
CC-----
CC (2) TO GENERATE DATA
CC-----
CC (2A) TO SET SUM OF LOOP RAW SCORES TO ZERO (SUMD), GENERATE
CC UNIVARIATE RANDOM UNIFORM DEVIATES (UD), TRANSFORM UD INTO RANDOM
CC NORMAL DEVIATES (ND), AND FORM A MATRIX OF UNIVARIATE RANDOM
CC NORMAL ERRORS (UNE) AND PREDICTORS (UNX).
      LOOP=0
62    LOOP=LOOP+1
      IF (LOOP .GT. NL) GOTO 250
      DO 64 I=1,NV
      SUMD(I)=0.0
64    CONTINUE
      SUMS=0.0
      SUMK=0.0
      CALL GGUBS(DSEED,TS,UD)
      DO 65 I=1,TS,2
      ND(I)  =SQRT(-2.*LOG(UD(I)))*COS(2.*PH*UD(I+1))
      ND(I+1)=SQRT(-2.*LOG(UD(I)))*SIN(2.*PH*UD(I+1))
65    CONTINUE
      DO 69 I=1,NP
      DO 70 J=1,NS
      K=(I-1)*NS+J
      UNE(I,J)=ND(K)
70    CONTINUE
69    CONTINUE
      DO 71 I=1,NQ
      DO 72 J=1,NS
      K=(NP+I-1)*NS+J
      UNX(I,J)=ND(K)
72    CONTINUE
71    CONTINUE
CC (2B) TO OBTAIN MULTIVARIATE NORMAL ERRORS (MNE) AND PREDICTORS
CC (MNX) BY MULTIPLYING PRINCIPAL COMPONENT WEIGHTS (PE,PX) WITH
CC NORMAL ERRORS AND PREDICTORS (UNE,UNX).
      CALL VMULFF (PE,UNE,NP,NP,NS,NP,NP,MNE,NP,IER)
      CALL VMULFF (PX,UNX,NQ,NQ,NS,NQ,NQ,MNX,NQ,IER)
CC (2C) TO OBTAIN MULTIVARIATE NON-NORMAL DATA (MNNE, MNNX) BY
CC MULTIPLYING EACH SCORE WITH FLEISHMAN CONSTANTS.
      DO 81 I=1,NP
      DO 82 J=1,NS
      MNNE(I,J)=AA+BB*MNE(I,J)+CC*(MNE(I,J)**2)+DD*(MNE(I,J)**3)
82    CONTINUE
81    CONTINUE
      DO 83 I=1,NQ
      DO 84 J=1,NS
      MNNX(I,J)=AA+BB*MNX(I,J)+CC*(MNX(I,J)**2)+DD*(MNX(I,J)**3)
84    CONTINUE
83    CONTINUE
CC (2D) TO OBTAIN DEPENDENT DEVIATES (MNNY):  $Y = B'X + E$ .
      CALL VMULFM (RB,MNNX,NQ,NP,NS,NQ,NQ,MBX,NP,IER)
      DO 85 I=1,NP
      DO 86 J=1,NS
      MNNY(I,J) = MBX(I,J) + MNNE(I,J)
86    CONTINUE

```

```

85      CONTINUE
CC      (2E) TO FORM A COMBINED DEPENDENT-PREDICTOR DATA MATRIX (DATD).
      DO 87 I=1,NP
      DO 88 J=1,NS
      DATD(I,J) = MNXY(I,J)
88      CONTINUE
87      CONTINUE
      DO 89 I=1,NQ
      DO 90 J=1,NS
      K=NP+I
      DATD(K,J)=MNNX(I,J)
90      CONTINUE
89      CONTINUE
CC-----
CC (3) TO OBTAIN DESCRIPTIVE STATISTICS OF THE DATA
CC-----
CC      (3A) TO OBTAIN SUM OF RAW SCORES (SUMD,SUMD1), SUM OF RAW-SCORE
CC      SQUARES (SUMD2), SUM OF RAW-SCORE CUBES (SUMD3), SUM OF RAW-SCORE
CC      QUADS (SUMD4), MEAN (AVD), AND DEVIATION FROM THE MEAN (DEVD).
      DO 95 I=1,NV
      DO 96 J=1,NS
      SUMD(I) =SUMD(I) +DATD(I,J)
      SUMD1(I)=SUMD1(I)+DATD(I,J)
      SUMD2(I)=SUMD2(I)+DATD(I,J)**2
      SUMD3(I)=SUMD3(I)+DATD(I,J)**3
      SUMD4(I)=SUMD4(I)+DATD(I,J)**4
96      CONTINUE
      AVD(I) =SUMD(I)/NS
95      CONTINUE
      DO 98 I=1,NV
      DO 99 J=1,NS
      DEVD(I,J)=DATD(I,J)-AVD(I)
99      CONTINUE
98      CONTINUE
CC      (3B) TO FIND SUM OF RAW-SCORE CROSS PRODUCT MTRIX (RCP), SUM OF
CC      CROSS PRODUCT MATRX (AAA), SUM OF RAW-SCORE CROSS PRODUCTS (SUMCC;
CC      AND COVARIANCE (COVN).
      CALL VMULFP (DATD,DATD,NV,NS,NV,NV,NV,RCP,NV,IER)
      CALL VMULFP (DEVD,DEVD,NV,NS,NV,NV,NV,AAA,NV,IER)
      DO 108 I=1,NV
      DO 109 J=1,NV
      SUMCC(I,J)=SUMCC(I,J)+RCP(I,J)
      COVN(I,J)=AAA(I,J)/(NS-1.)
109      CONTINUE
108      CONTINUE
CC      (3C) TO OBTAIN MULTIVARIATE SKEWNESS (MULS) AND KURTOSIS (MULK)
CC      INVERSE OF COVARIANCE MATRIX (COVI) AND PRODUCT OF THE MATRICES
CC      DEVD'*COVI*DEVD (DCD).
      CALL LINV1F (COVN,NV,NV,COVI,DGT,WK,IER)
      CALL VMULFM (DEVD,COVI,NV,NS,NV,NV,NV,DCV,NS,IER)
      CALL VMULFF (DCV,DEVD,NS,NV,NS,NS,NV,DCD,NS,IER)
      DO 120 I=1,NS
      DO 121 J=1,NS
      SUMS=SUMS+DCD(I,J)**3
121      CONTINUE
      SUMK=SUMK+DCD(I,I)**2
120      CONTINUE
      MULS=SUMS/(NS**2)
      MULK=SUMK/NS-NV*(NV+2.)
      SUMMS=SUMMS+MULS

```

```

SUMMK=SUMMK+MULK
CC-----
CC (4) TO OBTAIN SUM OF CROSS PRODUCTS MATRIX AND SUBMATRICES
CC-----
CC (4A) ORIGINAL DATA: TO OBTAIN SCP SUBMATRIX FOR Y AND X VARIABLES
CC (AYY,AXX) FROM AAA.
      DO 127 I=1,NP
      DO 128 J=1,NP
      AYY(I,J)=AAA(I,J)
128   CONTINUE
127   CONTINUE
      DO 132 I=1,NP
      DO 133 J=1,NQ
      AXX(I,J)=AAA(I,NP+J)
133   CONTINUE
132   CONTINUE
      DO 136 I=1,NQ
      DO 137 J=1,NQ
      AXX(I,J)=AAA(NP+I,NP+J)
137   CONTINUE
136   CONTINUE
CC (4B) RANKED DATA: TO RANK ORIGINAL DATA AND OBTAIN MEAN RANK
CC (AVR), SCP MATRIX (RRR) AND SCP SUBMATRICES FOR Y,YX,X VARIABLES
CC (RYY,RYX,RXX).
      DO 140 I=1,NV
      DO 141 J=1,NS
      DTA(J)=DATD(I,J)
141   CONTINUE
      CALL NMRANK(DTA,NS,EPS,IR,R,DTR,S,T)
      DO 142 J=1,NS
      DATR(I,J)=DTR(J)
142   CONTINUE
140   CONTINUE
      AVR=(NS+1.)/2.
      DO 145 I=1,NV
      DO 146 J=1,NS
      DEVR(I,J)=DATR(I,J) - AVR
146   CONTINUE
145   CONTINUE
      CALL VMULFP(DEVR,DEVR,NV,NS,NV,NV,NV,RRR,NV,IER)
      DO 155 I=1,NP
      DO 156 J=1,NP
      RYY(I,J)=RRR(I,J)
156   CONTINUE
155   CONTINUE
      DO 160 I=1,NP
      DO 161 J=1,NQ
      RYX(I,J)=RRR(I,NP+J)
161   CONTINUE
160   CONTINUE
      DO 165 I=1,NQ
      DO 166 J=1,NQ
      RXX(I,J)=RRR(NP+I,NP+J)
166   CONTINUE
165   CONTINUE
CC (4C) MIXED DATA: TO OBTAIN SCP FOR RANKED Y AND ORIGINAL X (MYX).
      DO 170 I=1,NP
      DO 171 J=1,NS
      DVYR(I,J)=DEVR(I,J)
171   CONTINUE

```

```

170     CONTINUE
        DO 175 I=1,NQ
        DO 176 J=1,NS
            K=NP+I
            DVXD(I,J)=DEV(D,K,J)
176     CONTINUE
175     CONTINUE
        CALL VMULFP(DVYR,DVXD,NP,NS,NQ,NP,NQ,MYX,NP,IER)
CC-----
CC (5) TO OBTAIN WILKS' LAMBDA AND SUM OF EIGENVALUES
CC-----
CC (5A) TO OBTAIN THE INVERSES OF AYY,AXX,RYX,RXX.
        CALL LINVIF(AYY,NP,NP,AYYI,DGT,WK,IER)
        CALL LINVIF(AXX,NQ,NQ,AXXI,DGT,WK,IER)
        CALL LINVIF(RYY,NP,NP,RYYI,DGT,WK,IER)
        CALL LINVIF(RXX,NQ,NQ,RXXI,DGT,WK,IER)
CC (5B) TO OBTAIN PRODUCT OF MATRICES USING ORIGINAL DATA(MATD),
CC RANKED DATA(MATR), AND MIXED DATA(MATM).
        CALL VMULFF(AYYI,AYX,NP,NP,NQ,NP,NP,MATD1,NP,IER)
        CALL VMULFP(AXXI,AYX,NQ,NQ,NP,NQ,NP,MATD2,NQ,IER)
        CALL VMULFF(MATD1,MATD2,NP,NQ,NP,NP,NQ,MATD,NP,IER)
        CALL VMULFF(RYYI,RXX,NP,NP,NQ,NP,NP,MATR1,NP,IER)
        CALL VMULFP(RXXI,RXX,NQ,NQ,NP,NQ,NP,MATR2,NQ,IER)
        CALL VMULFF(MATR1,MATR2,NP,NQ,NP,NP,NQ,MATR,NP,IER)
        CALL VMULFF(RYYI,MYX,NP,NP,NQ,NP,NP,MATM1,NP,IER)
        CALL VMULFP(AXXI,MYX,NQ,NQ,NP,NQ,NP,MATM2,NQ,IER)
        CALL VMULFF(MATM1,MATM2,NP,NQ,NP,NP,NQ,MATM,NP,IER)
CC (5C) TO OBTAIN EIGENVALUES OF MATD (ED) AND MATR (ER), THE PRODUC
CC OF (1-EIGENVALUE) (PED,PER) AND SUM OF THE EIGENVALUES (SER, SEM)
        CALL EIGRF(MATD,NP,NP,EJOB,ED,Z,NP,WK,IER)
        CALL EIGRF(MATR,NP,NP,EJOB,ER,Z,NP,WK,IER)
        PED=1.0
        PER=1.0
        SER=0.0
        SEM=0.0
        DO 230 I=1,NP
            PED=PED*(1.0-ED(I))
            PER=PER*(1.0-ER(I))
            SER=SER+MATR(I,I)
            SEM=SEM+MATM(I,I)
230     CONTINUE
CC-----
CC (6) TO COMPUTE TEST STATISTICS AND TO COUNT REJECTIONS.
CC-----
CC (6A) TO COMPUTE BARTLETT (BAR), RAO F (RAO), RANK-
CC TRANSFORM (RTF), PURE-RANK (PRN), MIXED-RANK (MRN).
        BAR = -A1*LOG(PED)
        RAO = ((1.-PED**B2)/(PED**B2))*V3
        RTF = ((1.-PER**B2)/(PER**B2))*V3
        PRN = (NS-1.)*SER
        MRN = (NS-1.)*SEM
CC (6C) NUMBER OF REJECTIONS FOR ALPHA = .01
        IF(BAR .GE. CS1) REJB(1)=REJB(1)+1
        IF(RAO .GE. CF1) REJR(1)=REJR(1)+1
        IF(RTF .GE. CF1) REJT(1)=REJT(1)+1
        IF(PRN .GE. CS1) REJP(1)=REJP(1)+1
        IF(MRN .GE. CS1) REJM(1)=REJM(1)+1
CC (6D) NUMBER OF REJECTIONS FOR ALPHA = .05
        IF(BAR .GE. CS2) REJB(2)=REJB(2)+1
        IF(RAO .GE. CF2) REJR(2)=REJR(2)+1

```

```

      IF(RTF .GE. CF2) REJT(2)=REJT(2)+1
      IF(PRN .GE. CS2) REJP(2)=REJP(2)+1
      IF(MRN .GE. CS2) REJM(2)=REJM(2)+1
CC   (6E) NUMBER OF REJECTIONS FOR ALPHA = .10
      IF(BAR .GE. CS3) REJB(3)=REJB(3)+1
      IF(RAO .GE. CF3) REJR(3)=REJR(3)+1
      IF(RTF .GE. CF3) REJT(3)=REJT(3)+1
      IF(PRN .GE. CS3) REJP(3)=REJP(3)+1
      IF(MRN .GE. CS3) REJM(3)=REJM(3)+1
      GOTO 62
CC-----
CC (7) TO OBTAIN AVERAGE AND OVERALL DESCRIPTIVE STATISTICS
CC-----
CC (7A) TO OBTAIN AVERAGE AND OVERALL MEAN (AVD1, OMEAN), VARIANCE
CC (AVARN, OVARN), UNIVARIATE SKEWNESS (SKEW, OSKEW) AND KURTOSIS
CC (KURT,OSKEW), MULTIVARIATE SKEWNESS (OMSKW) AND KURTOSIS (OMKUR),
CC CORRELATION MATRIX (CORR) AND DETERMINANT OF CORR (DCOR).
250   TMEAN=0.0
      TVARN=0.0
      TSKEW=0.0
      TKURT=0.0
      DO 251 I=1,NV
      AVD1(I)=SUMD1(I)/TN
      AVD2(I)=SUMD2(I)/TN
      AVD3(I)=SUMD3(I)/TN
      AVD4(I)=SUMD4(I)/TN
      VARN(I)=AVD2(I)-AVD1(I)**2
      SKEW(I)=(AVD3(I)-3.*AVD1(I)*AVD2(I)+2.*AVD1(I)**3)/VARN(I)**1.5
      KURT(I)=((AVD4(I)-4.*AVD1(I)*AVD3(I)+6.*(AVD1(I)**2)*AVD2(I)-
&        3.*AVD1(I)**4)/(VARN(I)**2)) - 3.0
      TMEAN=TMEAN+AVD1(I)
      TVARN=TVARN+VARN(I)
      TSKEW=TSKEW+SKEW(I)
      TKURT=TKURT+KURT(I)
251   CONTINUE
      OMEAN=TMEAN/NV
      OVARN=TVARN/NV
      OUSKW=TSKEW/NV
      OUKUR=TKURT/NV
      OMSKW=SUMMS/NL
      OMKUR=SUMMK/NL
CC   (7B) TO OBTAIN CORRELATION MATRIX (CORR) AND DETERMINANT
CC OF CORRELATION MATRIX (DCOR).
      DO 260 I=1,NV
      DO 261 J=1,NV
      ACOVN(I,J)=(SUMCC(I,J)/TN)-(AVD1(I)*AVD1(J))
261   CONTINUE
260   CONTINUE
      DO 262 I=1,NV
      DO 263 J=1,NV
      CORR(I,J)=ACOVN(I,J)/SQRT(ACOVN(I,I)*ACOVN(J,J))
      CORL(I,J)=CORR(I,J)
263   CONTINUE
262   CONTINUE
      CALL LINV3F(CORL,BCR,DJOB,NV,NV,D1,D2,WK,IER)
      DCOR=D1*2.**D2
      WRITE(20,270) NV,NS,NL,CY,CX
270   FORMAT(/,' NV =' ,I1,' NS =' ,I3,' NL =' ,I4,' CY =' ,F3.1,
&        ' CX =' ,F3.1)
      WRITE(20,272) AVD1(1),AVD1(2),AVD1(3),AVD1(4),OMEAN

```

```

272  FORMAT(/, 'MEAN', 3X, 5F11.6)
      WRITE(20, 274) VARN(1), VARN(2), VARN(3), VARN(4), OVARN
274  FORMAT(/, 'VARN', 3X, 5F11.6)
      WRITE(20, 276) SKEW(1), SKEW(2), SKEW(3), SKEW(4), OUSKW
276  FORMAT(/, 'USKEW', 2X, 5F11.6)
      WRITE(20, 278) KURT(1), KURT(2), KURT(3), KURT(4), OUKUR
278  FORMAT(/, 'UKURT', 2X, 5F11.6)
      DO 280 I=1, NV
      WRITE(20, 281) CORR(I, 1), CORR(I, 2), CORR(I, 3), CORR(I, 4)
281  FORMAT(/, 'CORR', 3X, 4F11.6)
280  CONTINUE
      WRITE(20, 282) OMSKW, OMKUR, DCOR
282  FORMAT(/, 'MSKEW =', F11.6, ' MKURT =', F11.6, ' DETCOR = ', F8.6)
      WRITE(20, 284)
284  FORMAT(/, 'ALFA(REJ)  BART  RAOF  RAOR  PURR  MIXR')
      WRITE(20, 286) REJB(1), REJR(1), REJT(1), REJP(1), REJM(1)
286  FORMAT(/, '0.01( 30)', 5(3X, I4))
      WRITE(20, 288) REJB(2), REJR(2), REJT(2), REJP(2), REJM(2)
288  FORMAT(/, '0.05(150)', 5(3X, I4))
      WRITE(20, 290) REJB(3), REJR(3), REJT(3), REJP(3), REJM(3)
290  FORMAT(/, '0.10(300)', 5(3X, I4), /)
      DIS=DIS+1
      IF (DIS .EQ. 2) GOTO 6
      STOP
      END

```

APPENDIX E

TABLES

Table E1. Fleishman Constants Used for Data Generation^e

γ_1	γ_2	a	b	c	d
0.00	0.00	0.00	1.00	0.00	0.00
0.00	-1.12	0.00	1.348917	0.00	-0.132660
0.50	0.00	-0.092624	1.039946	0.092624	-0.016461
1.00	0.50	-0.258525	1.114655	0.258525	-0.066013
0.00	3.00	0.00	0.782356	0.00	0.067905
1.00	3.00	-0.128397	0.832216	0.128397	0.048032
2.00	6.00	-0.313749	0.826324	0.313749	0.022707
0.00	20.00	0.00	0.338712	0.00	0.184461

^e γ_1 = skewness; γ_2 = kurtosis; a, b, c, d = Fleishman constants.

Table E2. Average Mean, Variance, Skewness, Kurtosis, and Within-Set Correlations of the Generated Data^a

N	V	γ_1	γ_2	μ	σ^2	γ_1	γ_2	$\rho(.3)$	$\rho(.7)$
25	4	0.00	0.00	.002	.993	.006	.031	.301	.701
		0.00	-1.12	-.001	1.006	.000	-1.149	.299	.694
		0.50	0.00	.001	.995	.498	-.008	.299	.699
		1.00	0.50	.001	.998	1.006	.522	.303	.701
		0.00	3.00	.003	1.001	.020	2.991	.304	.702
		1.00	3.00	.001	.992	.980	2.930	.298	.699
		2.00	6.00	.000	.997	2.015	6.196	.300	.700
		0.00	20.00	.003	1.000	.045	20.551	.297	.698
	6	0.00	0.00	.003	.992	-.006	.026	.299	.700
		0.00	-1.12	.001	1.005	-.003	-1.163	.297	.694
		0.50	0.00	.002	.995	.502	.025	.300	.700
		1.00	0.50	.002	.996	1.001	.512	.299	.701
		0.00	3.00	.001	.999	.020	3.264	.301	.701
		1.00	3.00	.004	1.001	1.027	3.156	.299	.700
		2.00	6.00	-.001	.991	2.005	6.159	.299	.700
		0.00	20.00	.002	1.005	-.040	20.031	.301	.700
	8	0.00	0.00	.002	.996	.005	.014	.302	.701
		0.00	-1.12	.000	1.001	.000	-1.162	.299	.694
		0.50	0.00	.002	.994	.500	.013	.299	.699
		1.00	0.50	.001	.996	1.002	.518	.303	.699
		0.00	3.00	.001	.992	-.016	3.061	.298	.699
		1.00	3.00	.001	.995	.994	2.932	.300	.700
		2.00	6.00	-.001	.991	1.995	6.016	.298	.699
		0.00	20.00	-.002	.991	-.027	19.985	.300	.701
50	4	0.00	0.00	.003	.983	.007	.084	.302	.701
		0.00	-1.12	.000	.995	.001	-1.124	.300	.694
		0.50	0.00	.002	.981	.506	.089	.298	.699
		1.00	0.50	.002	.986	1.008	.593	.304	.700
		0.00	3.00	.001	.984	.001	3.228	.302	.701
		1.00	3.00	.000	.984	1.012	3.190	.301	.700
		2.00	6.00	.001	.986	2.034	6.219	.298	.699
		0.00	20.00	.000	.985	.000	20.922	.302	.702
	6	0.00	0.00	.001	.984	-.003	.090	.301	.701
		0.00	-1.12	-.001	.993	.003	-1.113	.297	.693
		0.50	0.00	.000	.984	.508	.080	.299	.699
		1.00	0.50	.001	.984	1.012	.606	.300	.701
		0.00	3.00	.002	.983	.011	3.207	.300	.700
		1.00	3.00	.000	.986	1.012	3.270	.302	.701
		2.00	6.00	-.001	.982	2.045	6.382	.300	.700
		0.00	20.00	-.001	.986	-.056	20.037	.300	.700

Table E2 (continued)

N	V	γ_1	γ_2	μ	σ^2	γ_1	γ_2	$\rho(.3)$	$\rho(.7)$
50	8	0.00	0.00	.003	.984	.008	.085	.300	.700
		0.00	-1.12	.002	.992	-.004	-1.103	.298	.694
		0.50	0.00	.002	.982	.505	.076	.300	.700
		1.00	0.50	.001	.984	1.001	.604	.301	.699
		0.00	3.00	-.002	.983	-.009	3.176	.298	.698
		1.00	3.00	.001	.985	1.010	3.212	.300	.700
		2.00	6.00	.000	.980	2.024	6.185	.298	.699
		0.00	20.00	.003	.979	.019	19.879	.299	.699
100	4	0.00	0.00	.002	.978	.006	.021	.301	.700
		0.00	-1.12	.001	.986	-.003	-1.077	.298	.694
		0.50	0.00	.000	.976	.502	.124	.299	.699
		1.00	0.50	.001	.979	1.015	.640	.300	.700
		0.00	3.00	.000	.979	.009	3.104	.298	.699
		1.00	3.00	.001	.980	1.018	3.250	.301	.700
		2.00	6.00	-.002	.972	2.040	6.192	.298	.699
		0.00	20.00	.001	.970	.011	19.614	.299	.699
	6	0.00	0.00	.001	.978	-.003	.120	.301	.700
		0.00	-1.12	.001	.987	-.004	-1.063	.298	.694
		0.50	0.00	-.001	.977	.502	.120	.301	.701
		1.00	0.50	.001	.978	1.014	.640	.300	.701
		0.00	3.00	.000	.977	.020	3.181	.300	.700
		1.00	3.00	.001	.978	1.021	3.234	.300	.699
		2.00	6.00	.001	.977	2.047	6.299	.300	.700
		0.00	20.00	.001	.981	-.047	20.647	.300	.701
	8	0.00	0.00	.001	.977	.003	.123	.300	.700
		0.00	-1.12	-.001	.986	-.001	-1.087	.297	.693
		0.50	0.00	.002	.977	.500	.121	.300	.700
		1.00	0.50	-.001	.977	1.016	.643	.300	.700
		0.00	3.00	-.001	.978	.003	3.136	.300	.700
		1.00	3.00	.001	.980	1.013	3.176	.301	.701
		2.00	6.00	-.001	.974	2.055	6.354	.299	.699
		0.00	20.00	.000	.987	-.010	21.295	.300	.700

^a The tabled values represent the average mean, variance, skewness, kurtosis, and within-set correlation values based on 9,000 replications.

Table E3. Average of the Type I Error and Power Values by Distribution and Sample Size^a

N	Type I Error						Power					
	$\alpha = .01$			$\alpha = .10$			$\alpha = .01$			$\alpha = .10$		
	25	50	100	25	50	100	25	50	100	25	50	100
[0, 0]												
RAO	010	011	009	097	101	102	065	247	675	323	613	913
RTF	010	010	009	099	098	106	069	227	630	313	587	894
PUR	002	005	007	073	088	100	021	153	580	254	551	885
MIX	002	005	008	068	087	099	016	153	594	243	556	893
[0, -1.12]												
RAO	011	011	008	101	103	102	063	232	660	311	603	912
RTF	011	011	009	101	102	102	063	210	592	297	563	874
PUR	002	006	007	074	091	095	019	141	539	239	528	864
MIX	002	005	006	074	092	096	020	142	541	243	532	865
[.5, 0]												
RAO	012	011	011	101	099	102	071	249	661	318	605	907
RTF	011	010	011	104	095	105	070	234	626	312	583	890
PUR	003	005	009	075	084	099	021	159	574	251	550	881
MIX	002	005	008	073	085	096	017	149	583	233	551	889
[1, .5]												
RAO	012	010	010	100	099	100	086	259	661	334	605	905
RTF	010	010	009	101	099	103	090	296	736	361	657	936
PUR	002	005	006	074	088	097	028	209	684	292	624	930
MIX	002	004	007	071	085	095	015	169	692	234	593	940
[0, 3]												
RAO	014	011	012	102	101	101	083	273	677	344	631	905
RTF	011	011	012	099	098	099	087	287	711	351	645	924
PUR	003	006	009	074	087	093	028	202	663	283	610	917
MIX	002	005	008	070	084	091	019	195	706	270	629	935
[1, 3]												
RAO	013	010	013	105	100	108	089	279	669	346	623	903
RTF	011	009	010	102	098	103	088	282	702	348	643	923
PUR	003	004	008	076	086	096	030	198	657	283	610	917
MIX	002	004	007	072	087	096	019	179	686	249	604	933
[2, 6]												
RAO	021	020	018	116	118	107	138	309	658	379	610	887
RTF	010	011	011	101	101	105	124	374	826	414	729	966
PUR	003	007	008	074	090	099	040	272	789	342	697	962
MIX	002	004	008	070	085	097	013	182	801	221	640	975
[0, 20]												
RAO	036	037	035	149	136	133	181	358	695	446	657	899
RTF	012	010	012	104	097	103	160	456	880	471	787	978
PUR	003	006	009	076	086	097	054	345	851	392	758	975
MIX	002	004	007	065	083	093	029	336	912	359	810	991

^a Tabled values represent the average Type I error and power values across all within-set correlations and numbers-of-variables (9 cases).

Table E4. Average of the Type I Error and Power Values by Distribution and Within-Set Correlation^a

$(\rho_y, \rho_x)^b$	Type I Error						Power					
	$\alpha = .01$			$\alpha = .10$			$\alpha = .01$			$\alpha = .10$		
	1	2	3	1	2	3	1	2	3	1	2	3
[0, 0]												
RAO	010	010	010	100	100	100	283	419	283	577	694	577
RTF	011	010	010	102	100	101	261	399	266	556	676	563
PUR	005	004	005	089	085	087	211	329	214	520	640	530
MIX	005	005	005	085	085	084	215	334	214	526	639	527
[0, -1.12]												
RAO	010	010	010	102	103	102	274	410	271	569	686	570
RTF	010	010	010	102	100	102	244	376	244	534	657	543
PUR	005	005	005	087	086	087	195	309	195	501	620	509
MIX	005	005	004	086	086	089	197	310	197	506	624	510
[.5, 0]												
RAO	011	011	011	101	100	101	281	420	280	573	686	571
RTF	011	011	011	104	100	100	263	401	267	554	671	560
PUR	006	006	005	087	086	085	209	333	211	520	636	526
MIX	006	005	005	085	085	084	211	333	206	524	631	519
[1, .5]												
RAO	011	010	011	099	102	099	288	428	290	575	692	576
RTF	010	010	009	101	102	100	329	466	327	615	727	612
PUR	005	004	005	086	087	085	267	388	265	578	691	577
MIX	004	004	005	082	085	084	260	372	244	559	663	545
[0, 3]												
RAO	012	012	012	100	102	103	297	434	301	590	699	592
RTF	011	011	011	100	099	097	318	451	316	603	711	605
PUR	006	006	006	086	085	084	257	377	258	567	674	569
MIX	005	005	005	082	081	082	278	382	259	585	681	570
[1, 3]												
RAO	011	012	012	104	104	106	299	434	304	585	698	588
RTF	010	010	010	102	100	100	314	448	311	602	710	603
PUR	005	005	005	087	085	085	257	373	255	568	673	569
MIX	004	004	004	085	085	084	265	373	245	569	665	552
[2, 6]												
RAO	018	020	022	110	115	116	322	452	331	587	696	594
RTF	011	011	011	103	104	101	410	522	393	676	765	668
PUR	006	006	005	087	089	087	339	439	323	640	728	633
MIX	005	004	005	084	085	083	328	398	270	609	670	557
[0, 20]												
RAO	032	035	040	135	140	144	362	487	385	634	726	642
RTF	011	011	011	102	101	101	478	568	450	734	793	709
PUR	006	006	005	087	086	086	401	476	372	696	756	674
MIX	004	004	004	081	078	082	436	471	370	730	761	669

^a Tabled values represent the average Type I error and power values across all sample sizes and numbers-of-variables (9 cases).

^b Within-set correlation 1 = (.3, .3), 2 = (.3, .7), 3 = (.7, .7).

Table E5. Average of the Type I Error and Power Values by Distribution and Number-of-Variables^a

V	Type I Error						Power					
	$\alpha = .01$			$\alpha = .10$			$\alpha = .01$			$\alpha = .10$		
	4	6	8	4	6	8	4	6	8	4	6	8
[0, 0]												
RAO	009	011	010	098	101	101	336	318	332	628	608	612
RTF	010	010	010	100	105	098	307	299	321	603	591	600
PUR	005	005	004	093	090	078	267	241	245	590	555	545
MIX	006	005	004	089	089	076	274	246	244	597	555	539
[0, -1.12]												
RAO	011	009	009	103	104	100	322	313	320	624	602	599
RTF	011	010	010	100	104	101	287	284	293	587	572	574
PUR	007	005	004	094	089	078	250	228	221	573	538	519
MIX	006	004	003	095	090	077	252	231	221	577	542	522
[.5, 0]												
RAO	011	011	012	099	102	101	324	326	331	617	607	606
RTF	011	011	010	099	106	100	301	306	323	596	593	596
PUR	007	006	004	092	089	077	261	250	243	582	558	542
MIX	006	005	004	089	088	077	265	249	235	592	555	527
[1, .5]												
RAO	010	010	011	100	099	101	336	328	342	621	609	614
RTF	010	010	009	100	103	100	377	366	378	657	648	648
PUR	006	004	003	092	088	078	333	298	290	644	612	589
MIX	005	004	004	092	085	074	337	285	255	649	584	534
[0, 3]												
RAO	011	012	013	094	108	102	348	341	343	632	626	623
RTF	011	011	011	098	104	094	357	357	371	640	639	640
PUR	007	006	005	092	090	073	315	294	282	626	603	581
MIX	006	005	004	087	084	074	336	303	281	647	611	576
[1, 3]												
RAO	012	012	012	104	103	106	341	343	352	625	625	621
RTF	009	009	011	099	101	102	350	352	370	638	634	642
PUR	006	004	004	092	086	080	308	291	286	624	601	585
MIX	006	004	003	094	081	079	324	290	269	639	594	552
[2, 6]												
RAO	018	020	021	102	115	124	361	365	379	624	626	627
RTF	011	011	011	103	101	104	447	438	440	710	705	694
PUR	006	006	005	095	086	082	397	361	343	697	670	634
MIX	005	005	004	092	081	079	403	325	267	702	614	521
[0, 20]												
RAO	029	038	041	125	139	155	395	405	434	663	661	678
RTF	012	011	011	104	096	105	495	499	502	750	744	742
PUR	008	005	004	098	081	081	441	416	392	737	709	680
MIX	006	004	002	093	077	072	484	422	371	790	718	652

^a Tabled values represent the average Type I error and power values across all sample sizes and within-set correlations (9 cases).

Table E6. Empirical Type I Error Rates And Power Values For
Distribution [0, 0]^a

(ρ_y, ρ_x)	N	V	BAR	Type I Error				Power					
				RAO	RTF	PUR	MIX	BAR	RAO	RTF	PUR	MIX	
$(\alpha = .01)$													
$(.3, .3)$	25	4	007	007	012	003-	003-	067	067	064	029-	021-	
		6	012	011	010	002-	001-	049	048	053	013-	013-	
		8	011	011	012	002-	001-	045	043	049	009-	006-	
	50	4	011	011	012	007	007	238	239	212	166	169	
		6	011	010	010	004-	003-	187	184	164	108-	111-	
		8	011	011	010	004-	004-	176	175	156	080-	083-	
	100	4	009	009	008	006-	006-	606	606	552	525-	543-	
		6	011	011	012	010	011	582	577	537	479	490	
		8	008	008	010	005-	005-	610	610	565	487-	500-	
	$(.3, .7)$	25	4	007	007	010	003-	003-	095	094	090	040-	035-
			6	012	011	009	001-	001-	090	087	093	025-	020-
			8	011	011	011	002-	001-	089	084	091	017-	010-
		50	4	011	011	013	007-	007-	340	340	304	252-	253-
			6	011	010	010	004-	003-	324	323	310	209-	209-
			8	011	011	009	003-	004-	362	359	350	200-	194-
		100	4	009	009	006-	006-	006-	765	765	710-	683-	709-
			6	011	011	010	007	011	830	828	776	733	759
			8	008	008	009	006-	005-	893	893	869	805-	817-
	$(.7, .7)$	25	4	007	007	009	003-	002-	067	067	065	029-	023-
			6	012	011	010	001-	002-	049	048	055	015-	013-
			8	011	011	010	002-	001-	045	043	060	011-	006-
		50	4	011	011	012	009	009	238	239	210	162	168
			6	011	010	008	004-	002-	187	184	171	110-	111-
			8	011	011	009	002-	004-	176	175	170	091-	082-
		100	4	009	009	006-	005-	006-	606	606	553-	521-	541-
			6	011	011	010	008	011	582	577	532	481	489
			8	008	008	011	007	006-	610	610	575	503-	495-
$(\alpha = .05)$													
$(.3, .3)$	25	4	043	043	043	029-	031-	211	211	202	152-	160-	
		6	052	050	054	028-	027-	180	179	179	110-	103-	
		8	053	050	051	021-	017-	170	162	155	078-	072-	
	50	4	057	057	053	046	049	466	466	436	410	421	
		6	053	053	050	040-	040-	400	399	374	321-	324-	
		8	052	051	048	032-	036-	386	381	369	282-	287-	
	100	4	048	048	041	039-	042-	815	814	768	761-	785-	
		6	055	055	055	050	048	792	791	758	738	752	
		8	050	050	054	046	043	815	814	784	750	761	
	$(.3, .7)$	25	4	043	043	043	031-	031-	270	270	248	203-	205-
			6	052	050	058+	033-	027-	254	252	251+	164-	151-
			8	053	050	052	022-	017-	261	254	251	130-	111-
		50	4	057	057	054	048	049	583	584	546	518	531
			6	053	053	049	038-	040-	595	595	563	496-	504-
			8	052	051	044	030-	036-	626	623	608	507-	501-
		100	4	048	048	043	039-	042-	910	910	888	882-	890-
			6	055	055	054	046	048	942	942	922	910	920
			8	050	050	055	045	043	968	967	956	945	947

Table E6 (continued)

(ρ_y, ρ_x)	N	V	BAR	Type I Error				Power				
				RAO	RTF	PUR	MIX	BAR	RAO	RTF	PUR	MIX
(.7, .7)	25	4	043	043	044	027-	029-	211	211	200	156-	154-
		6	052	050	054	032-	032-	180	179	186	113-	099-
		8	053	050	049	023-	020-	170	162	171	087-	075-
	50	4	057	057	057	048	052	466	466	431	405	418
		6	053	053	047	037-	040-	400	399	383	324-	320-
		8	052	051	046	034-	038-	386	381	383	302-	290-
	100	4	048	048	042-	039-	043	815	814	767	759-	783-
		6	055	055	052	045	046	792	791	749	730	747
		8	050	050	054	045	047	815	814	790	758	760
$(\alpha = .10)$ (.3, .3)	25	4	090	090	095	081-	074-	321	321	306	283-	278-
		6	100	098	110	084-	076-	286	283	280	220-	210-
		8	106	103	094	060-	059-	271	266	255	171-	168-
	50	4	104	104	104	100	098	594	594	559	549	561
		6	099	099	102	090	088-	536	536	513	473	483-
		8	101	100	091	077-	080-	540	536	509	448-	455-
	100	4	099	099	101	098	095	888	888	860	856	870
		6	105	106	111+	107	106	875	875	848+	837	852
		8	101	102	109	099	092	890	891	871	846	858
(.3, .7)	25	4	090	090	094	078-	074-	397	397	372	346-	343-
		6	100	098	105	080-	076-	391	388	363	296-	287-
		8	106	103	095	060-	059-	382	378	379	262-	239-
	50	4	104	104	104	098	098	699	698	667	656	667
		6	099	099	101	088-	088-	723	723	698	663-	665-
		8	101	100	089-	072-	080-	761	759	736-	678-	674-
	100	4	099	099	101	096	095	953	953	934	933	939
		6	105	106	105	100	106	970	970	957	954	961
		8	101	102	105	096	092	983	983	976	972	975
(.7, .7)	25	4	090	090	093	081-	074-	321	321	303	280-	281-
		6	100	098	105	076-	069-	286	283	287	234-	211-
		8	106	103	099	060-	056-	271	266	269	194-	169-
	50	4	104	104	107	103	097	594	594	561	548	562
		6	099	099	095	083-	083-	536	536	517	479-	483-
		8	101	100	091	077-	073-	540	536	526	467-	452-
	100	4	099	099	101	097	099	888	888	864	862	874
		6	105	106	108	104	107	875	875	853	839	847
		8	101	102	110	100	096	890	891	882	865	862

^a Tabled values represent the proportion of rejections across 3000 replications at $\alpha = .01, .05$, and $.10$, where N = sample size, NV = no. of variables, BAR = Bartlett, RAO = Rao F, RTF = rank-transform Rao F, PUR = pure-rank, MIX = mixed-rank, "+" indicates a liberal Type I error rate, and a "-" indicates a conservative Type I error rate.

Table E7. Empirical Type I Error Rates And Power Values For
Distribution $[0, -1.12]^a$

(ρ_y, ρ_x)	N	V	BAR	Type I Error				Power					
				RAO	RTF	PUR	MIX	BAR	RAO	RTF	PUR	MIX	
$(\alpha = .01)$													
$(.3, .3)$	25	4	013	013	012	004-	004-	062	062	060	025-	029-	
		6	011	010	012	001-	003-	054	051	051	013-	015-	
		8	011	010	010	001-	001-	049	046	041	007-	010-	
		50	4	012	012	011	008	007	219	219	192	157	154
			6	010	010	010	006-	006-	172	171	157	098-	101
			8	012	012	012	006-	004-	168	167	151	084-	085-
		100	4	009	009	009	007	007	586	586	517	481	488
			6	009	008	008	007	005-	582	580	510	457	460-
			8	008	008	009	006-	004-	580	579	514	432-	430-
	25	4	014+	014+	013	005-	005-	089	089	087	038-	041-	
		6	010	010	009	003-	002-	084	083	084	023-	024-	
		8	010	010	009	001-	001-	079	076	076	015-	012-	
		50	4	010	010	011	008	007	316	317	279	228	224
			6	009	009	011	006-	006-	321	319	289	185-	191-
			8	010	010	011	004-	004-	352	346	317	183-	182-
		100	4	008	008	010	008	007	754	754	675	650	651
			6	009	009	008	006-	005-	829	828	755	706-	710-
			8	007	007	010	007	006-	883	882	823	752	752-
	25	4	014+	014+	014+	005-	003-	063+	063+	065+	030-	030-	
		6	008	008	011	001-	002-	051	049	052	017-	016-	
		8	009	008	009	001-	000-	047	045	047	005-	008-	
		50	4	012	012	012	009	007	222	222	195	155	157
			6	011	011	011	006-	005-	168	167	154	099-	100-
			8	013	013	009	003-	002-	161	159	153	080-	085-
		100	4	009	009	011	008	008	583	584	517	488	491
			6	011	010	009	007	006-	574	571	500	455	458-
			8	006-	006-	009	005-	006-	577-	577-	513	430-	427-
$(\alpha = .05)$													
$(.3, .3)$	25	4	054	054	055	042-	043	200	200	184	146-	146	
		6	048	047	050	029-	027-	170	166	162	098-	097-	
		8	055	050	054	020-	024-	157	151	149	072-	071-	
		50	4	055	055	051	043	045	454	454	410	383	390
			6	051	051	053	044	043	394	394	364	312	310
			8	056	054	050	037-	034-	377	374	345	274-	279-
		100	4	049	048	049	047	045	804	803	749	741	744
			6	057	057	050	046	049	790	789	736	715	719
			8	047	046	049	039-	039-	795	794	732	696	698
	25	4	054	055	057	043	042-	250	251	239	191	196-	
		6	046	045	052	031-	026-	245	242	225	147-	147-	
		8	047	045	052	018-	020-	242	235	225	109-	115-	
		50	4	053	054	048	040-	044	571	571	515	485-	498
			6	050	050	051	039-	039-	585	585	531	468-	475-
			8	057	056	051	036-	035-	615	610	572	471-	474-
		100	4	048	048	045	043	046	909	909	859	854	858
			6	053	053	051	043	047	939	938	904	890	891
			8	046	045	049	042-	040-	968	968	935	919	921

Table E7 (continued)

(ρ_y, ρ_x)	N	V	BAR	Type I Error				Power					
				RAO	RTF	PUR	MIX	BAR	RAO	RTF	PUR	MIX	
(.7, .7)	25	4	060+	060+	061+	047	044	195+	195+	194+	156	155	
		6	050	050	051	028-	027-	172	170	172	105-	105-	
		8	050	047	043	014-	016-	152	147	156	073-	081-	
	50	4	053	054	052	044	044	446	447	406	376	386	
		6	055	055	057	044	047	394	394	358	310	312	
		8	051	049	050	036-	036-	381	377	362	288-	285-	
	100	4	049	049	047	046	046	805	804	750	742	747	
		6	054	053	052	047	048	788	786	729	709	710	
		8	043	043	043	035-	036-	789	788	739	695-	699-	
$(\alpha = .10)$ (.3, .3)	25	4	104	104	103	089-	085-	318	319	294	267-	275-	
		6	098	096	100	071-	073-	280	279	261	210-	207-	
		8	103	101	104	065-	062-	258	250	241	162-	174-	
	50	4	104	104	100	097	098	585	585	536	522	531	
		6	099	099	102	090	088-	535	536	493	460	469-	
		8	104	100	100	083-	083-	513	511	476	424-	429-	
	100	4	100	101	100	096	096	884	884	839	836	838	
		6	110	111+	110	103	104	871	872+	832	819	821	
		8	097	098	101	089-	089-	883	884	836	809-	811-	
	(.3, .7)	25	4	105	105	098	090	085-	383	384	359	331	330-
			6	100	098	100	073-	073-	368	364	348	277-	288-
			8	103	099	101	059-	058-	365	356	342	239-	247-
50		4	105	104	099	095	097	694	694	647	636	641	
		6	105	105	103	092	094	717	717	669	632	639	
		8	108	107	104	088-	085-	754	752	707	639-	646-	
100		4	097	097	097	095	096	953	953	925	923	922	
		6	112+	113+	105	099	101	972+	972+	948	942	947	
		8	094	095	096	087-	087-	986	986	970	963-	961-	
(.7, .7)	25	4	104	104	102	091	096	320	321	298	271	276	
		6	101	099	104	074-	075-	271	268	273	218-	222-	
		8	104	100	096	056-	056-	264	258	253	173-	173-	
	50	4	107	107	103	098	105	586	586	542	532	536	
		6	103	103	107	096	095	531	531	496	460	461	
		8	103	101	102	084-	086-	514	511	503	447-	439-	
	100	4	097	097	097	095	098	888	888	844	839	842	
		6	110	111+	108	099	104	881	881+	831	821	822	
		8	094	095	101	090	087-	883	885	842	819	818-	

^a Tabled values represent the proportion of rejections across 3000 replications at $\alpha = .01, .05$, and $.10$, where N = sample size, NV = no. of variables, BAR = Bartlett, RAO = Rao F, RTF = rank-transform Rao F, PUR = pure-rank, MIX = mixed-rank, "+" indicates a liberal Type I error rate, and a "-" indicates a conservative Type I error rate.

Table E8. Empirical Type I Error Rates And Power Values For
Distribution [.5, 0]^a

(ρ_y, ρ_x)	N	V	BAR	Type I Error				Power					
				RAO	RTF	PUR	MIX	BAR	RAO	RTF	PUR	MIX	
$(\alpha = .01)$													
(.3, .3)	25	4	011	011	012	005-	003-	071	071	065	029-	026-	
		6	012	012	013	003-	003-	059	058	053	013-	012-	
		8	015+	014+	012	001-	001-	050+	048+	052	005-	005-	
	50	4	011	011	012	006-	007	222	222	201	155-	159	
		6	012	011	012	007	007	205	204	190	124	116	
		8	010	010	007	003-	005-	175	173	165	087-	078-	
	100	4	011	011	011	010	009	576	577	546	514	524	
		6	013	012	012	009	009	574	572	534	481	494	
		8	011	010	010	007	008	600	599	557	476	483	
	(.3, .7)	25	4	010	010	013	005-	003-	100	100	089	046-	036-
			6	013	012	009	003-	002-	097	093	092	028-	023-
			8	014+	014+	010	003-	003-	094+	090+	102	014-	008-
		50	4	011	011	011	009	006-	330	330	297	234	243-
			6	010	010	011	005-	006-	353	353	325	238-	230-
			8	010	009	006-	002-	004-	360	359	354-	209-	188-
		100	4	011	011	011	009	009	742	742	697	671	687
			6	011	011	012	009	008	820	818	782	743	762
			8	012	012	011	007	007	893	893	871	813	816
	(.7, .7)	25	4	010	010	011	003-	002-	071	071	070	028-	027-
			6	011	011	012	002-	002-	062	060	055	014-	010-
			8	015+	014+	011	001-	001-	050+	047+	055	007-	003-
		50	4	011	011	011	007	006-	220	221	204	161	152-
			6	011	011	009	005-	004-	202	202	188	124-	104-
			8	010	010	008	003-	004-	176	173	182	096-	073-
		100	4	012	012	011	009	008	576	577	542	513	530
			6	010	010	012	009	007	570	569	537	481	493
			8	012	012	011	008	008	601	600	567	478	459
$(\alpha = .05)$													
(.3, .3)		25	4	051	051	051	037-	037-	199	199	198	160-	152-
			6	050	048	053	028-	024-	183	181	171	100-	101-
	8		060+	057	046	020-	022-	163+	156	157	076-	063-	
	50	4	050	051	053	047	045	451	452	421	392	409	
		6	050	050	050	041-	038-	420	420	401	345-	345-	
		8	047	045	045	028-	031-	383	379	363	291-	280-	
	100	4	052	052	048	046	047	800	799	755	746	769	
		6	055	055	055	051	046	786	785	758	732	753	
		8	048	047	046	042-	042-	817	816	791	754-	753-	
	(.3, .7)	25	4	052	052	055	041-	036-	254	254	245	198-	201-
			6	050	049	054	025-	025-	259	255	254	157-	138-
			8	059+	055	054	022-	021-	255+	247	249	129-	098-
		50	4	051	052	050	043	044	570	571	543	512	526
			6	049	049	054	043	037-	586	586	571	511	512-
			8	046	045	043	031-	030-	622	617	597	502-	485-
		100	4	052	052	047	045	048	897	897	872	864	883
			6	053	052	056	052	045	934	934	919	909	921
			8	050	049	047	043	045	960	965	959	944	947

Table E8 (continued)

(ρ_y, ρ_x)	N	V	Type I Error				Power					
			BAR	RAO	RTF	PUR	MIX	BAR	RAO	RTF	PUR	MIX
(.7, .7)	25	4	051	051	051	037-	035-	205	205	193	151-	152-
		6	053	052	048	028-	029-	180	178	185	118-	100-
		8	059+	055	053	020-	020-	166+	162	168	081-	060-
	50	4	051	051	049	044	043	450	452	430	407	413
		6	049	049	049	039-	033-	422	422	398	347-	340-
		8	049	047	043	031-	028-	385	380	382	316-	272-
	100	4	053	053	047	044	045	797	797	758	747	767
		6	051	050	054	048	049	782	781	761	742	747
		8	051	050	049	043	046	819	817	799	772	754
$(\alpha = .10)$ (.3, .3)	25	4	099	099	105	089-	081-	310	310	295	272-	271-
		6	098	097	106	076-	075-	290	287	278	216-	205-
		8	111+	106	105	061-	062-	267+	262	253	177-	156-
	50	4	097	097	097	091	090	584	584	554	544	561
		6	105	105	104	089-	091	548	548	526	494-	500
		8	102	099	090	074-	078-	527	525	496	442-	445-
	100	4	102	102	101	098	094	874	874	850	846	864
		6	105	106	111+	105	100	867	868	852+	839	849
		8	097	098	112+	099	093	894	895	881+	854	862
(.3, .7)	25	4	099	099	097	088-	082-	382	383	360	331-	327-
		6	099	097	107	078-	074-	387	384	376	309-	286-
		8	111+	107	104	061-	063-	378+	372	372	261-	223-
	50	4	097	097	090	085-	091	689	689	662	652-	669
		6	101	101	101	088-	091	714	714	692	656-	659
		8	099	097	091	075-	076-	740	737	721	660-	649-
	100	4	102	102	101	099	091	947	948	928	926	938
		6	107	107	110	105	101	964	964	954	950	953
		8	098	098	098	091	095	981	981	977	974	974
(.7, .7)	25	4	098	098	107	092	086-	309	309	303	278	275-
		6	100	099	099	071-	072-	289	287	294	234-	209-
		8	113+	108	104	064-	058-	268+	265	272	185-	145-
	50	4	095	095	092	089-	090	581	580	560	545-	561
		6	104	104	099	086-	086-	546	546	522	491-	487-
		8	099	097	090	075-	074-	523	520	514	464-	430-
	100	4	104	104	096	094	095	873	874	848	845	861
		6	104	104	113+	105	103	867	868	844+	832	845
		8	099	100	102	094	093	892	893	880	861	855

^a Tabled values represent the proportion of rejections across 3000 replications at $\alpha = .01, .05$, and $.10$, where N = sample size, NV = no. of variables, BAR = Bartlett, RAO = Rao F, RTF = rank-transform Rao F, PUR = pure-rank, MIX = mixed-rank, "+" indicates a liberal Type I error rate, and a "-" indicates a conservative Type I error rate.

Table E9. Empirical Type I Error Rates And Power Values For
Distribution [1, .5]^a

(ρ_y, ρ_x)		N	V	Type I Error				Power						
				BAR	RAO	RTF	PUR	MIX	BAR	RAO	RTF	PUR	MIX	
$(\alpha = .01)$														
$(.3, .3)$	25	4	011	011	009	003-	002-	079	079	081	032-	021-		
		6	012	012	013	002-	002-	074	073	066	016-	008-		
		8	015+	015+	010	001-	001-	064+	062+	060	009-	004-		
	50	4	012	012	011	007	007	245	245	284	232	218		
		6	010	010	010	005-	004-	191	191	220	148-	122-		
		8	007	007	008	003-	001-	186	183	212	118-	073-		
	100	4	009	009	007	006-	008	595	595	683	650-	693		
		6	009	009	011	008	007	555	554	662	602	611		
		8	013	013	010	007	005-	608	608	690	598	589-		
	25	4	010	010	015+	004-	004-	105	105	126+	065-	038-		
		6	014+	013	007	001-	002-	116+	115	122	032-	018-		
		8	011	009	010	001-	001-	123	120	121	018-	006-		
	50	4	008	008	009	008	004-	339	340	382	324	317-		
		6	010	009	011	004-	004-	359	359	403	287-	244-		
		8	010	010	010	002-	004-	389	383	436	280-	178-		
	100	4	011	011	010	007	009	753	753	816	795	837		
		6	009	008	008	005-	005-	807	806	868	825-	851-		
		8	009	009	007	004-	005-	867	866	917	870-	864-		
	$(.7, .7)$	25	4	010	010	010	002-	003-	078	078	090	043-	031-	
			6	013	013	010	002-	001-	075	074	076	021-	009-	
			8	015+	014+	009	002-	001-	070+	066+	065	013-	003-	
		50	4	010	010	012	008	007	234	234	260	212	199	
			6	009	009	007	003-	003-	201	201	232	157-	107-	
			8	013	012	008	005-	005-	193	191	233	126-	061-	
		100	4	009	009	007	006-	006-	589	589	666	639-	675-	
			6	009	009	009	006-	006-	577	575	648	594-	591-	
			8	008	008	009	007	011	600	599	671	581	517	
$(\alpha = .05)$														
$(.3, .3)$		25	4	042-	042-	045	032-	030-	211-	211-	235	191-	169-	
			6	057	056	055	029-	031-	205	203	206	129-	099-	
	8		052	050	048	022-	018-	187	182	182	083-	050-		
	50	4	056	057	053	048	045	462	463	513	493	502		
		6	042-	042-	048	040-	032-	389-	389-	454	396-	369-		
		8	047	045	051	032-	031-	397	394	443	364-	287-		
	100	4	047	047	047	046	044	803	803	859	851	886		
		6	047	047	050	048	044	781	779	843	824	851		
		8	055	054	054	048	044	815	814	871	842	842		
	25	4	051	051	055	042-	037-	272	272	310	252-	209-		
		6	048	047	048	027-	029-	286	283	297	202-	144-		
		8	055	051	046	019-	017-	285	279	305	152-	087-		
	50	4	051	052	049	041-	046	564	565	628	599-	620		
		6	052	052	056	045	034-	590	590	655	597	566-		
		8	052	049	049	036-	037-	645	640	682	584-	502-		
	100	4	053	053	046	046	050	901	901	937	933	953		
		6	046	046	050	042-	043	926	926	960	954-	968		
		8	045	045	045	039-	038-	957	957	976	966-	966-		

Table E9 (continued)

(ρ_y, ρ_x)	N	V	BAR	Type I Error				Power				
				RAO	RTF	PUR	MIX	BAR	RAO	RTF	PUR	MIX
(.7, .7)	25	4	054	054	057	039-	038-	216	216	233	191-	158-
		6	052	051	051	024-	024-	195	193	213	140-	082-
		8	064+	061+	051	020-	020-	188	187	194	098-	050-
	50	4	051	051	047	042-	046	469	470	502	471-	474
		6	051	051	053	041-	038-	414	414	467	406-	342-
		8	054	054	046	033-	037-	400	397	440	366-	256-
	100	4	047	047	046	044	046	800	799	853	846	877
		6	045	045	057	052	046	783	783	842	826	825
		8	053	052	048	041-	038-	801	799	856	824-	796-
	25	4	091	091	097	084-	078-	330	330	355	327-	303-
		6	103	102	109	079-	078-	299	296	319	257-	209-
		8	101	098	102	063-	053-	295	288	298	197-	131-
(.3, .3)	50	4	103	103	103	096	092	574	574	629	618	639
		6	086-	086-	093	081-	078-	520-	520-	596	554-	523-
		8	099	098	097	081-	075-	537	531	586	524-	458-
	100	4	106	106	100	098	096	883	883	918	916	936
		6	099	100	103	098	099	866	867	908	901	921
		8	107	107	104	097	092	883	883	925	910	910
	25	4	099	100	102	089-	087-	385	385	430	403-	372-
		6	101	099	096	071-	068-	405	403	426	347-	287-
		8	107	104	093	055-	056-	408	401	441	314-	196-
	50	4	099	099	099	094	093	680	679	739	730	752
		6	105	105	105	093	091	713	714	773	748	724
		8	107	106	104	085-	083-	752	750	801	747-	680-
	100	4	105	105	106	103	103	950	950	964	963	978
		6	098	099	103	099	099	966	966	982	979	985
		8	098	098	105	097	081-	984	984	987	985	991-
(.7, .7)	25	4	096	097	104	091	086-	317	318	342	314	291-
		6	103	103	099	074-	068-	299	296	327	261-	190-
		8	112+	108	103	062-	063-	290+	285	313	212-	129-
	50	4	095	095	094	087-	094	591	591	624	614-	630
		6	106	106	107	096	086-	549	549	593	565	505-
		8	097	095	090	075-	076-	536	533	568	513-	422-
	100	4	098	099	093	090	097	874	875	916	914	938
		6	092	092	111+	104	099	867	868	905+	898	910
		8	096	097	097	090	088-	868	869	916	902	887-

^a Tabled values represent the proportion of rejections across 3000 replications at $\alpha = .01, .05$, and $.10$, where N = sample size, NV = no. of variables, BAR = Bartlett, RAO = Rao F, RTF = rank-transform Rao F, PUR = pure-rank, MIX = mixed-rank, "+" indicates a liberal Type I error rate, and a "-" indicates a conservative Type I error rate.

Table E10. Empirical Type I Error Rates And Power Values For
Distribution [0, 3]^a

(ρ_y, ρ_x)			N	V	BAR	Type I Error				Power										
						RAO	RTF	PUR	MIX	BAR	RAO	RTF	PUR	MIX						
$(\alpha = .01)$																				
$(.3, .3)$			25	4	012	012	013	005-	003-	081	081	081	038-	027-						
				6	015+	015+	009	002-	001-	066+	064+	069	018-	012-						
				8	015+	014+	012	003-	003-	060+	058+	061	012-	011-						
			50	4	011	011	012	008	007	251	252	258	206	224						
				6	011	011	010	006-	003-	218	216	227	148-	159-						
				8	012	012	010	004-	003-	194	190	226	123-	114-						
			100	4	009	009	010	008	009	614	614	647	619	688						
				6	013	013	011	008	009	601	598	635	582	651						
				8	013	013	009	008	007	604	603	655	570	618						
			25	4	011	011	011	006-	004-	113	113	111	052-	041-						
				6	013	013	012	003-	001-	116	111	113	036-	018-						
				8	018+	017+	010	002-	003-	112+	105+	119	020-	014-						
			50	4	012	012	010	007	007	357	357	350	295	319						
				6	011	010	011	004-	003-	369	368	390	283-	272-						
				8	012	012	011	004-	003-	395	393	427	262-	231-						
			100	4	009	009	009	007	008	759	760	790	766	817						
				6	013	012	014+	012	009	825	823	854+	826	857						
				8	014+	014+	012	009	006	875+	875+	900	850	872						
$(.7, .7)$			25	4	013	013	011	005-	003-	083	082	078	039-	026-						
				6	013	013	012	002-	001-	068	066	077	024-	011-						
				8	016+	015+	010	002-	002-	066+	062+	072	010-	007-						
			50	4	011	011	011	008	007	256	258	247	204	212						
				6	010	010	010	006-	006-	227	225	226	162-	136-						
				8	012	012	012	006-	004-	201	199	234	132-	093-						
			100	4	009	009	009	007	009	614	615	648	616	668						
				6	015+	015+	013	009	009	601+	600+	623	570	608						
				8	013	013	010	009	008	602	602	643	563	571						
			$(\alpha = .05)$			$(.3, .3)$			25	4	048	048	054	036-	032-	223	224	215	172-	172-
										6	057	057	055	027-	027-	200	197	202	130-	113-
										8	059+	057	048	020-	021-	185+	179	187	090-	086-
						50	4	051	051	056	049	044	479	480	484	459	504			
							6	053	053	054	045	040-	441	441	456	393	414-			
							8	051	050	046	033-	028-	417	412	443	363-	365-			
						100	4	043	043	045	043	043	817	817	838	829	870			
							6	054	054	053	046	046	800	793	836	818	856			
							8	057	056	045	038-	043	806	805	842	815-	839			
$(.3, .7)$						25	4	047	048	049	036-	032-	276	277	274	221-	222-			
							6	055	055	050	028-	025-	285	284	286	185-	160-			
							8	059+	056	047	023-	021-	284+	274	300	157-	112-			
						50	4	051	052	052	044	044	588	589	597	570	613			
							6	052	052	057	045	038-	614	614	638	577	598-			
							8	051	048	048	037-	030-	652	649	660	562-	557-			
						100	4	044	044	043	041-	046	898	898	913	910-	933			
							6	057	056	056	050	042-	935	934	948	942	960-			
							8	056	056	045	040-	044	959	959	970	962-	971			

Table E10 (continued)

(ρ_y, ρ_x)	N	V	Type I Error				Power					
			BAR	RAO	RTF	PUR	MIX	BAR	RAO	RTF	PUR	MIX
(.7, .7)	25	4	048	049	044	033-	028-	227	227	216	172-	168-
		6	058+	057	049	027-	023-	208+	205	208	137-	105-
		8	059+	056	048	022-	025-	190+	184	199	099-	067-
	50	4	052	053	052	047	046	479	480	480	454	491
		6	056	056	056	043	037-	444	444	457	398	389-
		8	050	050	048	030-	030-	427	424	452	369-	331-
	100	4	045	045	044	040-	044	811	811	836	829-	862
		6	057	057	057	052	046	798	797	824	811	841
		8	057	057	052	046	043	802	801	821	800	812
$(\alpha = .10)$ (.3, .3)	25	4	091	091	100	088-	076-	333	334	330	304-	312-
		6	113+	112+	100	075-	070-	320+	317+	308	246-	249-
		8	107	104	101	061-	066-	291	285	304	204-	192-
	50	4	099	099	099	096	092	610	609	614	604	648
		6	105	105	105	090	086-	576	576	591	556	575-
		8	095	093	096	079-	076-	562	557	580	514-	531-
	100	4	090	091	096	093	091	879	879	899	896	919
		6	102	102	111+	105	099	875	875	899+	893	918
		8	105	106	088-	082-	082-	878	880	905-	889-	916-
(.3, .7)	25	4	089-	090	105	094	074-	403-	404	400	371	381-
		6	113+	113+	097	072-	070-	412+	409+	420	341-	330-
		8	108	103	093	056-	061-	409	402	426	304-	254-
	50	4	102	101	096	091	092	705	705	718	706	745
		6	107	107	103	092	085-	729	730	751	721	750-
		8	101	100	094	079-	073-	759	756	768	713-	727-
	100	4	089-	089-	098	095	095	936-	936-	953	951	968
		6	103	104	112+	107	097	966	966	976+	972	981
		8	107	109	089-	082-	083-	978	978	986-	983-	990-
(.7, .7)	25	4	092	092	100	090	077-	334	335	332	305	307-
		6	114+	114+	098	072-	069-	314+	311+	320	257-	232-
		8	104	102	095	061-	064-	305	298	315	216-	176-
	50	4	101	101	092	087-	093	604	604	612	601-	632
		6	108	108	105	094	088-	579	579	592	555	555-
		8	099	098	092	073-	073-	569	567	576	523-	501-
	100	4	094	095	094	093	092	880	880	899	896	914
		6	105	106	106	098	091	871	871	897	886	911
		8	107	108	095	084-	086-	879	879	901	882-	897-

^a Tabled values represent the proportion of rejections across 3000 replications at $\alpha = .01, .05$, and $.10$, where N = sample size, NV = no. of variables, BAR = Bartlett, RAO = Rao F, RTF = rank-transform Rao F, PUR = pure-rank, MIX = mixed-rank, "+" indicates a liberal Type I error rate, and a "-" indicates a conservative Type I error rate.

Table E11. Empirical Type I Error Rates And Power Values For Distribution [1, 3]^a

(ρ_y, ρ_x)	N	V	BAR	Type I Error				BAR	Power				
				RAO	RTF	PUR	MIX		RAO	RTF	PUR	MIX	
$(\alpha = .01)$													
$(.3, .3)$	25	4	011	011	007	003-	003-	089	089	085	038-	030-	
		6	015+	014+	011	003-	003-	073+	071+	066	023-	012-	
		8	012	012	015+	002-	000-	064	060	064+	010-	006-	
		50	4	010	010	008	006-	006-	242	242	251	200-	210-
			6	008	008	008	004-	001-	227	226	229	154-	144-
			8	011	010	010	004-	003-	214	212	219	127-	099-
		100	4	014+	015+	012	009	008	596+	597+	632	607	662
			6	012	011	008	008	007	583	581	619	571	615
			8	012	012	010	007	006-	613	613	656	582	606
	25	4	011	011	008	004-	004-	122	122	110	058-	045-	
		6	016+	015+	010	002-	003-	118+	113+	113	039-	020-	
		8	013	012	013	002-	000-	123	116	129	026-	008-	
		50	4	011	011	007	006-	005-	333	334	348	286-	303-
			6	008	008	009	003-	002-	388	387	385	279-	257-
			8	013	013	009	004-	003-	397	393	419	254-	205-
		100	4	014+	014+	013	010	008	746+	747+	772	751	799
			6	013	013	009	008	006-	821	819	852	811	849
			8	013	013	009	006-	006-	872	872	901	855	871
	25	4	011	011	009	004-	004-	091	091	080	037-	029-	
		6	015+	015+	010	003-	003-	075+	073+	076	025-	011-	
		8	013	013	012	001-	000-	068	064	073	014-	005-	
		50	4	010	010	009	005-	005-	248	249	244	197-	189-
			6	009	009	007	002-	002-	240	239	221	156-	124-
			8	012	011	009	003-	004-	228	226	221	130-	076-
100		4	013	013	011	010	009	598	598	628	600	647	
		6	013	013	008	007	006-	584	582	610	562	578-	
		8	014+	013	011	008	007	615+	615	647	575	544	
$(\alpha = .05)$													
$(.3, .3)$	25	4	056	056	050	036-	032-	224	225	218	177-	167-	
		6	061+	060+	051	030-	028-	199+	198+	200	124-	111-	
		8	058+	055	060+	024-	019-	194+	190	195+	100-	075-	
		50	4	048	048	047	038-	039-	464	464	476	453-	475-
			6	048	048	048	034-	031-	445	445	449	396-	398-
			8	056	055	047	034-	034-	422	417	439	347-	329-
		100	4	061+	060+	058+	055	055	803+	803+	819+	810	860
			6	051	051	047	040-	037-	791	789	833	814-	847-
			8	056	055	052	045	046	810	808	848	814	841
	25	4	054	054	050	036-	033-	283	284	278	226-	207-	
		6	062+	060+	054	029-	028-	293+	290+	284	186-	150-	
		8	057	055	055	024-	020-	294	289	307	160-	101-	
		50	4	049	049	046	038-	040-	579	580	590	560-	601-
			6	049	049	052	039-	032-	614	614	627	570-	575-
			8	057	054	047	030-	034-	640	635	667	574-	520-
		100	4	060+	060+	057	054	053	894+	894+	913	909	936
			6	052	051	047	040-	037-	934	933	949	943-	958-
			8	056	056	049	040-	044	958	957	967	953	966

Table E11 (continued)

(ρ_y, ρ_x)	N	V	BAR	Type I Error				Power				
				RAO	RTF	PUR	MIX	BAR	RAO	RTF	PUR	MIX
$(.7, .7)$	25	4	054	054	050	034-	034-	228	228	224	177-	160-
		6	061+	059+	057	030-	027-	210+	208+	215	138-	095-
		8	059+	056	050	018-	018-	203+	196	214	107-	062-
	50	4	047	047	042-	035-	037-	461	463	481-	451-	460-
		6	048	048	051	041-	034-	451	451	453	397-	378-
		8	054	054	047	032-	038-	428	423	442	361-	291-
	100	4	059+	059+	060+	057	054	794+	793+	820+	812	854
		6	051	051	047	043	039-	787	785	819	804	827-
		8	059+	058+	050	043	048	810	810	830	801	802
$(\alpha = .10)$ $(.3, .3)$	25	4	107	107	096	083-	085-	326	327	330	306-	296-
		6	108	107	108	083-	075-	314	311	304	247-	226-
		8	102	100	106	067-	060-	301	296	299	206-	170-
	50	4	096	096	096	091	092	589	589	607	595	626
		6	099	099	102	091	086-	575	575	587	553	559-
		8	104	103	100	079-	081-	548	545	584	522-	497-
	100	4	108	108	108	104	104	877	877	895	892	925
		6	102	103	094	089-	085-	872	873	900	892-	912-
		8	109	111+	105	096	098	873	875+	912	896	908
$(.3, .7)$	25	4	108	108	093	082-	084-	403	404	398	367-	368-
		6	105	104	110	081-	071-	418	414	412	335-	301-
		8	105	101	102	065-	055-	414	405	425	307-	220-
	50	4	095	095	093	086-	094	692	692	722	711-	734
		6	099	099	101	090	086-	726	726	743	710	719-
		8	106	104	097	077-	085-	758	757	777	728-	698-
	100	4	109	109	108	105	110	940	941	956	954	971
		6	102	103	094	088-	085-	962	962	973	971-	979-
		8	109	109	103	094	094	977	978	982	979	989
$(.7, .7)$	25	4	106	107	093	083-	085-	328	328	329	304-	293-
		6	107	107	103	076-	073-	320	319	313	260-	213-
		8	110	107	104	062-	058-	314	307	320	219-	152-
	50	4	099	099	091	085-	086-	593	593	610	599-	619-
		6	099	100	101	089-	081-	572	573	581	552-	537-
		8	108	107	096	083-	088-	557	554	579	520-	448-
	100	4	111+	111+	110	106	105	874+	874+	894	892	919
		6	105	106	096	090	086-	872	872	897	887	901-
		8	113+	113+	106	095	096	874+	874+	902	888	889

^a Tabled values represent the proportion of rejections across 3000 replications at $\alpha = .01, .05$, and $.10$, where N = sample size, NV = no. of variables, BAR = Bartlett, RAO = Rao F, RTF = rank-transform Rao F, PUR = pure-rank, MIX = mixed-rank, "+" indicates a liberal Type I error rate, and a "-" indicates a conservative Type I error rate.

Table E12. Empirical Type I Error Rates And Power Values For
Distribution [2, 6]^a

(ρ_y, ρ_x)	N	V	Type I Error					Power				
			BAR	RAO	RTF	PUR	MIX	BAR	RAO	RTF	PUR	MIX
$(\alpha = .01)$												
(.3, .3)	25	4	021+	021+	014+	004-	003-	127+	127+	122+	053-	028-
		6	021+	019+	008	003-	001-	117+	115+	104	028-	008-
		8	017+	015+	008	001-	001-	104+	101+	083	013-	001-
	50	4	016+	016+	009	007	005-	275+	275+	365	308	302-
		6	018+	018+	011	006-	004-	250+	248+	311	214-	149-
		8	021+	020+	014+	007	003-	260+	259+	307+	171	069-
	100	4	015+	015+	010	008	008	594+	594+	787	767	858
		6	017+	017+	014+	011	010	582+	580+	798+	755	811
		8	017+	017+	015+	009	007	601+	600+	811+	746	721
(.3, .7)	25	4	020+	020+	012	004-	002-	164+	164+	167	083-	039-
		6	020+	020+	009	003-	001-	173+	171+	166	050-	010-
		8	024+	023+	010	001-	001-	175+	169+	151	030-	003-
	50	4	017+	017+	011	007	004-	361+	361+	461	402	388-
		6	022+	022+	011	006-	005-	399+	398+	485	361-	236-
		8	024+	023+	013	006-	003-	429+	428+	506	330-	121-
	100	4	016+	016+	010	008	006-	728+	728+	885	867	933-
		6	019+	019+	011	009	010	801+	799+	923	901	930
		8	019+	019+	011	008	006-	849+	848+	953	925	919-
(.7, .7)	25	4	021+	021+	011	004-	003-	131+	131+	122	056-	020-
		6	025+	024+	011	002-	002-	135+	132+	105	030-	004-
		8	026+	025+	008	001-	001-	133+	130+	095	015-	001-
	50	4	017+	017+	011	007	006-	277+	277+	348	291	246-
		6	023+	022+	010	006-	005-	266+	264+	297	204-	092-
		8	027+	026+	011	006-	003-	273+	271+	289	164-	031-
	100	4	016+	016+	010	008	006-	590+	591+	769	743	815-
		6	022+	022+	011	008	008	579+	575+	748	703	688
		8	021+	021+	011	006-	009	604+	604+	762	697	535
$(\alpha = .05)$												
(.3, .3)	25	4	067+	067+	056	042-	039-	259+	260+	302	243-	197-
		6	060+	058+	047	023-	020-	254+	252+	263	163-	088-
		8	068+	065+	043	018-	020-	242+	237+	232	113-	037-
	50	4	058+	059+	048	040-	036-	468+	468+	597	569-	605-
		6	065+	065+	055	042-	038-	441+	441+	574	509-	446-
		8	078+	077+	057	039-	035-	436+	433+	540	448-	309-
	100	4	050	049	053	051	045	781	779	919	915	964
		6	066+	065+	052	048	045	775+	775+	923	914	949
		8	057	056	052	044	041-	786	784	930	912	921
(.3, .7)	25	4	066+	067+	057	039-	037-	315+	316+	362	299-	236-
		6	071+	070+	053	027-	019-	334+	331+	356	247-	118-
		8	071+	069+	049	020-	018-	339+	331+	349	184-	041-
	50	4	060+	061+	051	043	038-	565+	567+	696	670-	716-
		6	067+	067+	051	040-	040-	606+	606+	732	676-	616-
		8	083+	080+	054	037-	031-	626+	623+	736	644-	446-
	100	4	053	053	051	049	042-	881	881	964	963	985-
		6	069+	069+	051	046	046	912+	912+	980	976	990
		8	066+	065+	057	048	043	944+	944+	990	986	990

Table E12 (continued)

(ρ_y, ρ_x)	N	V	Type I Error					Power				
			BAR	RAO	RTF	PUR	MIX	BAR	RAO	RTF	PUR	MIX
(.7, .7)	25	4	064+	064+	058+	043	037-	266+	267+	290+	242-	162-
		6	069+	069+	048	028-	019-	263+	260+	266	174-	064-
		8	074+	070+	052	017-	016-	270+	266+	235	117-	023-
	50	4	061+	061+	048	041-	036-	465+	466+	583	556-	558-
		6	065+	065+	048	039-	037-	453+	453+	555	495-	345-
		8	086+	084+	054	038-	034-	452+	448+	515	431-	195-
	100	4	051	051	051	049	044	784	784	908	902	949
		6	068+	067+	053	047	045	769+	767+	902	891	906
		8	068+	068+	059+	050	045	780+	780+	911+	891	830
$(\alpha = .10)$ (.3, .3)	25	4	115+	115+	111+	096	090	353+	354+	410+	382	348
		6	110	110	092	069-	064-	343	345	376	308-	204-
		8	117+	114+	092	050-	059-	338+	334+	339	237-	113-
	50	4	100	100	097	091	090	580	580	710	696	755
		6	108	108	109	097	085-	559	559	695	664	642-
		8	133+	132+	110	091	081-	548+	545+	673	610	489-
	100	4	094	094	101	099	098	861	862	956	954	981
		6	110	110	106	097	093	853	853	957	953	980
		8	108	108	105	095	093	854	854	968	957	969
(.3, .7)	25	4	113+	113+	112+	099	091	413+	414+	483+	456	405
		6	115+	114+	098	069-	064-	445+	443+	489	408-	256-
		8	121+	119+	105	061-	056-	449+	444+	479	345-	123-
	50	4	101	101	098	093	089-	675	674	797	787	836-
		6	117+	117+	105	092	090	706+	706+	831	804	777
		8	142+	141+	101	089-	079-	729+	728+	834	783-	645-
	100	4	092	092	099	096	100	928	928	982	981	994
		6	114+	114+	105	099	094	949+	950+	993	991	998
		8	118+	119+	112+	104	099	972+	972+	997+	995	998
(.7, .7)	25	4	114+	114+	109	098	085-	363+	365+	406	377	305-
		6	126+	124+	095	066-	066-	361+	359+	382	315-	159-
		8	123+	120+	095	062-	056-	360+	356+	361	249-	076-
	50	4	102	102	091	087-	083-	581	580	699	688-	722-
		6	116+	116+	098	084-	083-	567+	568+	679	648-	549-
		8	146+	144+	102	085-	084-	557+	554+	647	591-	348-
	100	4	090	090	104	100	102	860	860	950	947	975
		6	117+	118+	103	099	090	848+	848+	942	938	957
		8	115+	116+	111+	103	103	852+	853+	950+	941	926

^a Tabled values represent the proportion of rejections across 3000 replications at $\alpha = .01, .05$, and $.10$, where N = sample size, NV = no. of variables, BAR = Bartlett, RAO = Rao F, RTF = rank-transform Rao F, PUR = pure-rank, MIX = mixed-rank, "+" indicates a liberal Type I error rate, and a "-" indicates a conservative Type I error rate.

Table E13. Empirical Type I Error Rates And Power Values For Distribution [0, 20]^a

(ρ_y, ρ_x)	N	V	Type I Error					Power					
			BAR	RAO	RTF	PUR	MIX	BAR	RAO	RTF	PUR	MIX	
$(\alpha = .01)$													
$(.3, .3)$	25	4	026+	026+	013	004-	003-	158+	158+	158	083-	060-	
		6	039+	039+	011	003-	001-	144+	141+	137	035-	022-	
		8	034+	032+	013	001-	000-	144+	142+	115	019-	010-	
	50	4	027+	027+	012	007	005-	305+	305+	426	364	470-	
		6	036+	036+	008	006-	002-	295+	294+	422	314-	338-	
		8	037+	037+	010	004-	003-	303+	302+	421	266-	231-	
	100	4	026+	026+	014+	012	009	628+	629+	865+	848	938	
		6	029+	029+	010	008	009	628+	626+	872	842	937	
		8	037+	037+	011	009	004-	664+	663+	886	842	918	
	$(.3, .7)$	25	4	030+	030+	012	005-	004-	203+	203+	195	103-	069-
			6	040+	040+	012	003-	001-	221+	217+	204	064-	022-
			8	039+	037+	008	001-	000-	231+	227+	207	036-	011-
		50	4	027+	028+	010	007	004-	399+	400+	516	451	541-
			6	037+	036+	008	006-	003-	430+	429+	565	442-	422-
			8	043+	043+	011	005-	003-	496+	493+	588	407-	282-
100		4	029+	029+	013	011	010	746+	747+	916	906	961	
		6	035+	035+	012	009	007	808+	806+	954	938	970	
		8	037+	036+	009	006-	004-	856+	856+	964	941	965	
$(.7, .7)$	25	4	032+	032+	012	005-	003-	174+	174+	145	074-	043-	
		6	044+	044+	012	002-	001-	185+	182+	139	046-	016-	
		8	047+	045+	010	001-	001-	189+	186+	136	024-	009-	
	50	4	030+	030+	009	006-	006-	313+	313+	404	336-	378-	
		6	040+	040+	008	004-	004-	323+	322+	381	281-	226-	
		8	055+	054+	011	005-	003-	363+	361+	377	241-	137-	
	100	4	030+	030+	013	011	010	625+	626+	826	809	898	
		6	041+	041+	012	007	007	630+	627+	818	781	844	
		8	048+	048+	010	004-	004-	673+	673+	821	754-	776-	
	$(\alpha = .05)$												
	$(.3, .3)$	25	4	082+	082+	064+	045	033-	321+	321+	353+	295	323-
			6	090+	088+	055	030-	020-	298+	296+	318	210-	178-
8			104+	100+	051	021-	016-	296+	289+	298	139-	096-	
50		4	073+	074+	047	040-	041-	514+	515+	658	636-	762-	
		6	083+	083+	044	034-	035-	489+	489+	666	611-	682-	
		8	096+	095+	051	036-	031-	501+	496+	667	567-	601-	
100		4	076+	076+	058+	056	047	807+	807+	947+	942	982	
		6	079+	078+	046	040-	037-	809+	807+	961	953	988	
		8	082+	081+	045	038-	040-	824+	823+	963	952-	981-	
$(.3, .7)$		25	4	090+	090+	055	043	034-	385+	385+	412	353	358-
			6	095+	093+	052	029-	020-	386+	383+	400	285-	202-
			8	111+	108+	049	016-	015-	408+	401+	399	219-	103-
		50	4	077+	077+	048	039-	038-	600+	601+	740	716-	827-
			6	092+	092+	041-	030-	028-	630+	630+	775-	726-	777-
			8	098+	097+	051	037-	034-	683+	681+	794	713-	685-
		100	4	074+	074+	058+	053	046	877+	877+	973+	971	991
			6	085+	084+	047	043	038-	915+	913+	986	984	994
			8	094+	093+	051	044	035-	937+	936+	992	989	995-

Table E13 (continued)

(ρ_y, ρ_x)	N	V	Type I Error				Power					
			BAR	RAO	RTF	PUR	MIX	BAR	RAO	RTF	PUR	MIX
(.7, .7)	25	4	092+	092+	054	037-	034-	332+	333+	335	282-	264-
		6	103+	102+	050	027-	023-	331+	329+	311	210-	129-
		8	115+	111+	051	020-	015-	344+	337+	296	165-	074-
	50	4	075+	076+	048	041-	037-	508+	509+	626	602-	714-
		6	092+	092+	044	034-	034-	505+	505+	614	558-	559-
		8	111+	110+	052	036-	038-	537+	533+	615	521-	421-
	100	4	074+	073+	057	054	045	795+	795+	932	928	972
		6	093+	093+	048	044	047	799+	798+	934	924	963
		8	103+	102+	049	043	043	824+	822+	932	916	940
$(\alpha = .10)$ (.3, .3)	25	4	128+	128+	107	097	083-	430+	430+	482	457	513-
		6	144+	142+	106	076-	069-	343+	345+	376	308-	204-
		8	161+	158+	104	060-	049-	417+	409+	422	285-	240-
	50	4	117+	117+	088-	085-	090	621+	621+	772-	763-	869
		6	134+	134+	092	080-	077-	598+	598+	766	740-	826-
		8	151+	149+	105	091	078-	611+	608+	781	724	782-
	100	4	123+	123+	119+	115+	104	875+	875+	972+	971+	994
		6	127+	127+	093	088-	086-	872+	873+	981	979-	995-
		8	132+	132+	104	091	091	887+	888+	981	975	994
(.3, .7)	25	4	135+	136+	103	090	080-	484+	484+	538	509	562-
		6	146+	144+	104	077-	060-	484+	481+	530	448-	401-
		8	177+	173+	103	058-	045-	518+	512+	526	379-	251-
	50	4	119+	119+	096	092	092	704+	704+	836	828	902
		6	137+	137+	089-	078-	075-	735+	735+	860-	835-	892-
		8	154+	152+	106	088-	080-	773+	772+	873	828-	845-
	100	4	123+	124+	119+	117+	105	926+	926+	986+	986+	996
		6	127+	128+	089-	083-	086-	955+	955+	993-	992-	998-
		8	146+	146+	101	090	083-	962+	962+	996	996	998-
(.7, .7)	25	4	134+	134+	099	090	083-	435+	436+	458	430	454-
		6	155+	154+	106	076-	065-	422+	419+	425	357-	279-
		8	174+	171+	106	060-	050-	448+	441+	410	298-	172-
	50	4	115+	115+	095	088-	093	621+	621+	742	729-	830
		6	143+	144+	092	082-	078-	616+	616+	728	699-	726-
		8	162+	161+	110	090	086-	636+	635+	725	677	613-
	100	4	129+	129+	109	106	104	868+	868+	961	959	986
		6	138+	138+	091	087-	093	867+	867+	969	964-	983
		8	149+	150+	103	096	087-	878+	878+	962	954	976-

^a Tabled values represent the proportion of rejections across 3000 replications at $\alpha = .01, .05$, and $.10$, where N = sample size, NV = no. of variables, BAR = Bartlett, RAO = Rao F, RTF = rank-transform Rao F, PUR = pure-rank, MIX = mixed-rank, "+" indicates a liberal Type I error rate, and a "-" indicates a conservative Type I error rate.

Table E14. Frequency Distributions of Simulated Data

Interval	[0, 0]	[0, -1.12]	[.5, 0]	[1, .5]	[0, 3]	[1, 3]	[2, 6]	[0,20]
<-8.0	0	0	0	0	0	0	0	1
-8.0 -7.0	0	0	0	0	0	0	0	3
-7.0 -6.0	0	0	0	0	1	0	0	5
-6.0 -5.0	0	1	0	0	3	0	0	7
-5.0 -4.0	0	0	0	0	8	0	0	24
-4.0 -3.0	9	1	0	0	45	4	0	56
-3.0 -2.7	15	0	0	0	29	7	0	28
-2.7 -2.5	32	0	0	0	30	5	0	24
-2.5 -2.3	38	0	0	0	38	15	0	33
-2.3 -2.1	65	0	0	0	51	35	0	37
-2.1 -1.9	108	0	23	0	87	44	0	52
-1.9 -1.7	149	0	153	0	95	77	0	56
-1.7 -1.5	202	720	258	0	158	136	0	77
-1.5 -1.3	317	575	371	0	186	236	0	119
-1.3 -1.1	384	572	531	1197	316	375	0	119
-1.1 -0.9	508	538	644	1287	377	549	809	214
-0.9 -0.7	596	544	694	896	584	798	1940	303
-0.7 -0.5	630	585	718	770	681	861	1256	447
-0.5 -0.3	782	550	772	674	844	984	1002	820
-0.3 -0.1	727	586	751	681	914	964	873	1469
-0.1 0.1	801	580	783	563	1038	890	724	2151
0.1 0.3	751	591	686	549	935	795	587	1455
0.3 0.5	740	577	629	494	837	630	503	797
0.5 0.7	648	572	564	431	689	561	416	431
0.7 0.9	608	540	514	412	529	429	328	313
0.9 1.1	465	575	415	355	385	328	254	210
1.1 1.3	387	528	343	289	300	273	234	150
1.3 1.5	316	562	280	277	215	219	185	102
1.5 1.7	215	802	230	228	143	169	160	86
1.7 1.9	174	0	164	198	124	116	128	77
1.9 2.1	130	0	142	160	91	98	87	49
2.1 2.3	66	0	108	134	68	87	95	55
2.3 2.5	56	0	72	121	44	71	74	34
2.5 2.7	42	0	48	99	33	51	53	25
2.7 3.0	20	0	59	93	38	52	85	33
3.0 4.0	17	0	41	92	62	102	118	76
4.0 5.0	2	0	7	0	14	22	59	32
5.0 6.0	0	0	0	0	2	10	16	11
6.0 7.0	0	0	0	0	3	3	7	8
7.0 8.0	0	0	0	0	1	2	3	4
>8.0	0	0	0	0	2	2	4	7

^a Tabled values represent the frequency distributions of the simulated data using 10,000 deviates.

BIBLIOGRAPHY

THE BIBLIOGRAPHY

- Anderson, T. W. (1958). An Introduction to Multivariate Statistical Analysis. New York: John Wiley & Sons.
- Arnold, H. J. (1964). Permutation support for multivariate techniques. Biometrika, 51, 65-70.
- Bartlett, M. S. (1938). Further aspects of the theory of multiple regression. Proceedings of the Cambridge Philosophical Society, 34, 33-40.
- Belli, G. M. (1983). Robustness and power of multivariate tests for trends in repeated measures data under variance-covariance heterogeneity. Unpublished doctoral dissertation, Michigan State University.
- Bhattacharrya, G. K., Johnson, R. A., & Neave, H. R. (1971). A comparative power study of the bivariate rank sum test and t-test. Technometrics, 13(1), 191-198.
- Box, G. E. P., & Cox, D. R. (1964). An analysis of transformations. Journal of Royal Statistical Society, B26, 211-246.
- Box, G. E. P., & Muller, M. E. (1958). A note on the generation of random normal deviates. Annals of Mathematical Statistics, 29, 610-613.
- Boyer Jr., J. E., Palachek, A. D., & Schucany, W. R. (1983). An empirical study of related correlation coefficients. Journal of Educational Statistics, 8, 75-86.
- Chase, G. R., & Bulgren, W. G. (1971). A Monte Carlo investigation of the robustness of the T^2 . Journal of the American Statistical Association, 66, 499-502.
- Chatterjee, S. K., & Sen, P. K. (1964). Nonparametric tests for the bivariate two-sample location problem. Calcutta Statistical Association Bulletin, 13, 18-58.
- Conover, W. J. (1980). Practical Nonparametric Statistics (2nd ed.), New York: John Wiley & Sons.
- Conover, W. J., & Iman, R. L. (1976). On some alternative procedures using ranks for the analysis of experimental designs. Communications in Statistics, Series A, 5, 1348-1368.

- _____(1980a). Small sample efficiency of Fisher's randomization test when applied to experimental designs. Unpublished manuscript presented at the annual meeting of the American Statistical Association, Houston, August 1980.
- _____(1980b). The rank transformation as a method of discrimination with some examples. Communications in Statistics, Series A, 9, 465-487.
- _____(1981). Rank transformation as a bridge between parametric and nonparametric statistics. The American Statistician, 35, 124-129.
- _____(1982). Analysis of covariance using the rank transformation. Biometrics, 38, 715-724.
- Davis, A. W. (1980). On the effects of moderate nonnormality on Wilks's likelihood ratio criterion. Biometrika, 67, 419-427.
- _____(1982a). On the distribution of Hotelling's one-sample T^2 under moderate nonnormality. Journal of Applied Probability, 19A, 207-216.
- _____(1982b). On the effects of moderate nonnormality on Roy's largest root test. Journal of the American Statistical Association, 77, 896-900.
- Fawcett, R. F., & Salter, K. C. (1987). Distributional studies and the computer: An analysis of Durbin's rank test. The American Statistician, 41, 81-83.
- Feir-Walsh, B. J., & Toothaker, L. E. (1974). An empirical comparison of the ANOVA-F test, normal scores test, and Kruskal-Wallis test under violation of assumptions. Educational and Psychological Measurement, 34, 789-799.
- Fleishman, A. (1978). A method for simulating nonnormal distributions. Psychometrika, 43, 521-532.
- Gaito, J. (1970). Non-parametric methods in psychological research. In Readings in Statistics for the Behavioral Sciences, E. F. Heermann & L. A. Braskamp (eds.), 38-49. Englewood Cliffs, N.J.: Prentice-Hall, Inc.
- Glass, G. V., Peckham, P. D., & Sanders, J. R. (1972). Consequences of failure to meet assumptions underlying the fixed-effects analysis of variance and covariance. Review of Educational Research, 42, 237-288.
- Gibbons, J. D. (1971). Nonparametric Statistical Inference. New York: McGraw-Hill Book Company.
- Gittins, R. (1985). Canonical Analysis: A Review With Applications in Ecology. New York: Springer-Verlag.

- Harris, R. J. (1975). A Primer of Multivariate Statistics. New York: Academic Press.
- Hartley, H. O. (1976). The impact of computers on statistics. In On the History of Statistics and Probability, D. B. Owen (ed.), 421-442. New York: Marcel Dekker, Inc.
- Harwell, M. R., & Serlin, R. C. (1985). A nonparametric test statistic for the general linear model. Paper presented at the Annual Meeting of the American Educational Research Association, Chicago.
- Hogg, R. V., & Randles, R. H. (1975). Adaptive distribution-free regression methods and their applications. Technometrics, 17, 399-407.
- Hollander, M., & Wolfe, D. A. (1973). Nonparametric Statistical Methods. New York: John Wiley & Sons.
- Hopkins, J. W., & Clay, P. P. F. (1963). Some empirical distributions of bivariate T^2 and homoscedasticity criterion M under unequal variance and leptokurtosis. Journal of the American Statistical Association, 58, 1048-1053.
- Iman, R. L. (1974a). Use of a t -statistic as an approximation to the exact distribution of the Wilcoxon signed ranks test statistic. Communications in Statistics, 3, 795-806.
- ____ (1974b). A power study of a rank transform for the two-way classification model when interaction may be present. The Canadian Journal of Statistics, 2(2), 227-239.
- ____ (1976). An approximation to the exact distribution of the Wilcoxon-Mann-Whitney rank sum test statistic. Communications in Statistics, 5, 587-598.
- Iman, R. L. and Conover, W. J. (1976). A comparison of several rank tests for the two-way layout. Technical Report SAND76-0631, Sandia Laboratories, Albuquerque, New Mexico.
- ____ (1978). Approximations of the critical region for Spearman's rho with or without ties present. Communications in Statistics, Series B, 7, 269-282.
- ____ (1979). The use of rank transform in regression. Technometrics, 21, 499-509.
- ____ (1980a). A comparison of distribution free procedures for the analysis of complete blocks. Unpublished manuscript presented at the annual meeting of the American Institute of Decision Sciences, Las Vegas, November, 1980.

- _____. (1980b). Multiple comparisons procedures based on the rank transformation. Unpublished manuscript presented at the annual meeting of the American Statistical Association, Houston, August, 1980.
- International Mathematical and Statistical Libraries. (1983). IMSL Library reference manuals. Houston, Texas.
- Ito, K., & Schull, W. J. (1964). On the robustness of the T^2 test in multivariate analysis of variance when variance-covariance matrices are not equal. Biometrika, 51, 71-82.
- Ito, P. K. (1980). Robustness of ANOVA and MANOVA test procedures. In Handbook of Statistics, P. R. Krishnaiah (ed.), Vol. I, 199-236. New York: North Holland Publishing Company.
- Johnson, N. L., & Kotz, S. (1970). Continuous Univariate Distributions 1. New York: John Wiley & Sons.
- _____. (1970). Continuous Univariate Distributions 2. New York: John Wiley & Sons.
- Kaiser, H. F., & Dickman, K. (1962). Sample and population score matrices and sample correlation matrices from an arbitrary population correlation matrix. Psychometrika, 27, 179-182.
- Kaskey, G., Kolman, B., Krishnaiah, P. R., & Steinberg, L. (1980). Transformations to normality. In Handbook of Statistics, P. R. Krishnaiah (ed.), Vol. I, 321-341. New York: North-Holland Publishing Company.
- Kendall, M. G., & Stuart, A. (1969). The Advanced Theory of Statistics (Vol. I). New York: Hafner Publishing Company.
- _____. (1969). The Advanced Theory of Statistics (Vol. II). New York: Hafner Publishing Company.
- Knapp, T. R. (1978). Canonical correlation analysis: A general parametric significance testing system. Psychological Bulletin, 85, 410-416.
- Kottkamp, R., Mulhern, J. A., & Hoy, W. K. (1985). Secondary School Climate: A Revision of the OCDQ. Rutgers University.
- Learmonth, G. P., & Lewis, P. A. W. (1973). Statistical tests of some widely used and recently proposed uniform random number generators. Naval Postgraduate School, Monterey, California.
- Lehmann, E. L. (1975). Nonparametrics: Statistical Methods Based on Ranks. San Francisco, CA: Holden-Day.

- Lester, P. E. (1983). Development of An Instrument to Measure Teacher Job Satisfaction. Unpublished doctoral dissertation, New York University, New York.
- Marascuilo, L. A., & Levin, J. R. (1983). Multivariate Statistics in the Social Sciences: A Researcher's Guide. Monterey, CA: Brooks/Cole.
- Marascuilo, L. A., & McSweeney, M. (1977). Nonparametric and Distribution-Free Methods For the Social Sciences. Monterey, CA: Brooks/Cole.
- Mardia, K. V. (1970). Measures of multivariate skewness and kurtosis with applications. Biometrika, 57, 519-530.
- _____. (1971). The effect of nonnormality on some multivariate tests and robustness to nonnormality in the linear model. Biometrika, 58, 105-121.
- _____. (1974). Applications of some measures of multivariate skewness and kurtosis to testing normality and to robustness studies. Sankya, Series, B, 36, 115-128.
- Morrison, D. F. (1976). Multivariate Statistical Methods (2nd ed.) New York: McGraw-Hill Book Company.
- Muirhead, R. J., & Waternaux C. M. (1980). Asymptotic distributions in canonical correlation analysis and other multivariate procedures for nonnormal populations. Biometrika, 67, 31-43.
- Muller, K. E., & Peterson, B. L. (1984). Practical methods for computing power in testing the multivariate general linear hypothesis. Computational Statistics and Data Analysis, 2, 143-158.
- Noether, G. E. (1984). Nonparametrics: The early years - impressions and recollections. The American Statistician, 38, 173-178.
- Olson, C. L. (1974). Comparative robustness of six tests in multivariate analysis of variance. Psychological Bulletin, 83, 579-586.
- _____. (1976). On choosing a test statistic in multivariate analysis of variance. Psychological Bulletin, 83, 579-586.
- Pearson, E. S., & Hartley, H. O. (1951). Charts of the power function for analysis of variance tests derived from the non-central F distribution. Biometrika, 38, 112-130.
- Penfield, D. A., & Koffler, S. L. (1985). A power study of selected nonparametric K-sample tests. Paper presented at the Annual Meeting of the American Educational Research Association, Chicago.

- Puri, M. L., & Sen, P. K. (1969). A class of rank order tests for a general linear hypothesis. Annals of Mathematical Statistics, 40, 1325-1343.
- _____. (1971). Nonparametric Methods in Multivariate Analysis. New York: John Wiley & Sons.
- _____. (1985). Nonparametric Methods in General Linear Models. New York: John Wiley & Sons.
- Rao, C. R. (1951). An asymptotic expansion of the distribution of Wilk's criterion. Bulletin of the International Statistics Institute, 33, 177-180.
- Rogosa, D. (1980). Comparing nonparallel regression lines. Psychological Bulletin, 88, 307-321.
- SAS Institute, Inc. (1982). SAS User's Guide: Statistics. Cary, NC: SAS Institute, Inc.
- Scheffe', H. (1959). The Analysis of Variance. New York: John Wiley & Sons.
- Srisukho, D. (1974). Monte Carlo Study of the Power of H-test Compared to F-test When Population Distributions are Different in Form. Unpublished doctoral dissertation, University of California, Berkeley.
- Takeuchi, K., Yanai, H., & Mukherjee, B. N. (1982). The Foundations of Multivariate Analysis. New Delhi, India: Wiley Eastern Limited.
- Tiku, M. L., & Singh, M. (1982). Robust statistics for testing mean vectors of multivariate distributions. Communications in Statistics: Theory and Methods, 11, 985-1001.
- Timm, N. H. (1975). Multivariate Analysis with Applications in Education and Psychology. Monterey, Cal.: Brooks/Cole Publishing Company.
- Tracy, D. S., & Conley, W. C. (1982). Exact sampling distributions of the coefficient of kurtosis using the computer. Journal of Computing Science and Statistics, 14, 262-265.
- Vale, C. D., & Maurelli, V. A. (1983). Simulating multivariate nonnormal distributions. Psychometrika, 48(3), 465-471.
- Williams, E. J. (1959). The comparison of regression variables. Journal of the Royal Statistical Society, Series B, 21, 396-399.
- Zwick, R. J. (1984). A Monte Carlo Investigation of Non-parametric One-Way Multivariate Analysis of Variance in the Two-Group Case. Unpublished doctoral dissertation, University of California, Berkeley.

MICHIGAN STATE UNIV. LIBRARIES



31293006004661