

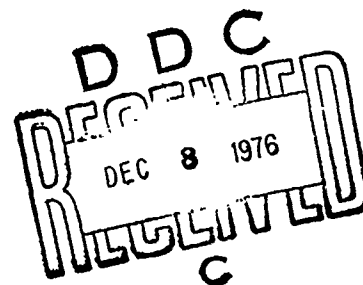
AD A033183

UNIT VERSUS DIFFERENTIAL WEIGHTING SCHEMES FOR DECISION MAKING:

A METHOD OF STUDY AND SOME PRELIMINARY RESULTS

SOCIAL SCIENCE RESEARCH INSTITUTE
UNIVERSITY OF SOUTHERN CALIFORNIA

J. Robert Newman
David Seaver
Ward Edwards



ADVANCED DECISION TECHNOLOGY PROGRAM

CYBERNETICS TECHNOLOGY OFFICE
DEFENSE ADVANCED RESEARCH PROJECTS AGENCY
Office of Naval Research • Engineering Psychology Programs

DISTRIBUTION STATEMENT A

Approved for public release;
Distribution Unlimited

The objective of the Advanced Decision Technology Program is to develop and transfer to users in the Department of Defense advanced management technologies for decision making.

These technologies are based upon research in the areas of decision analysis, the behavioral sciences and interactive computer graphics.

The program is sponsored by the Cybernetics Technology Office of the Defense

Advanced Research Projects Agency and technical progress is monitored by the Office of Naval Research — Engineering Psychology Programs. Participants in the program are:

Decisions and Designs, Incorporated
The Oregon Research Institute
Perceptronics, Incorporated
Stanford University
The University of Southern California

Inquiries and comments with regard to this report should be addressed to:

Dr. Martin A. Tolcott
Director, Engineering Psychology Programs
Office of Naval Research
800 North Quincy Street
Arlington, Virginia 22217

or

LT COL Roy M. Gulick, USMC
Cybernetics Technology Office
Defense Advanced Research Projects Agency
1400 Wilson Boulevard
Arlington, Virginia 22209

TECHNICAL REPORT SSRI 76-5

**UNIT VERSUS DIFFERENTIAL WEIGHTING SCHEMES
FOR DECISION MAKING:
A METHOD OF STUDY AND SOME
PRELIMINARY RESULTS**

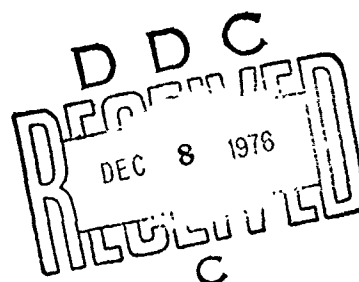
by

J. Robert Newman, David Seaver,
and Ward Edwards

Sponsored by

Defense Advanced Research Projects Agency
ARPA Order No. 3052
Under Subcontract from
Decisions and Designs, Incorporated

July, 1976



EXPRESSION for	
TIS	White Sheet ☒
JRC	Black Sheet ☐
UNANNOUNCED	☐
JUSTIFICATION	☐
BY	
DATE	
AL	
A	

SOCIAL SCIENCE RESEARCH INSTITUTE
University of Southern California
Los Angeles, California 90007

DISTRIBUTION STATEMENT A
Approved for public release;
Distribution Unlimited

Summary

A persistent problem in prediction studies and decision making problems is that of weighting the attributes or dimensions of information assumed relevant to the prediction or decision problem. Intuition and past experience has indicated that the attributes should be differentially weighted with the more important ones receiving higher weights. Recently, however, several empirical and theoretical studies have indicated that there are many situations in which differential weighting may not be necessary and that simple unit weighting, that is, just adding up the attributes of information, may be as good as and in some cases better than differential weighting. The implications of this result, if true, have extraordinary practical and theoretical significance, and the problem of weighting requires very careful study.

In this report, the first of a series, a method of generating realistic data is described, and illustrations of how the method can be used to study the usefulness of different data analysis and prediction are given. The method utilizes a computer simulation which generates an N-by-M data matrix where N is the number of observations and M is the number of variables or measurements taken on each observation. For example, N could be 15 automobiles being considered for possible purchase, and M could be 10 performance and/or quality factors of importance for each of the automobiles. The entries in the data matrix would be simulated measurement values for each factor on each automobile. The method also allows for the simulation of various types of error in the assigned values. The computer program to accomplish this is outlined. Two examples of the use of the method are given. One compares the familiar multiple regression model with simple unit weighting in a prediction problem to predict a well defined criterion variable from a set of predictor variables. The regression model estimates the weights to be assigned to each predictor whereas the unit weighting model merely adds up the predictors and thus does not assign differential weights. The results indicate that multiple regression is superior to unit weighting for prediction purposes, but the differences between the two models are not substantial. The second example compares several ways of forming weighted and unweighted combinations of attributes of dimensions of importance to help persons make practical decisions. Some of the conditions under which differential weighting is important for practical decision making are specified. The conditions under which differential weighting is not important are also specified.

Table of Contents

Summary.....	ii
List of Figures.....	iv
List of Tables.....	v
Acknowledgment.....	vi
Introduction.....	1
Method.....	2
Brief Comment on Rationale.....	4
Simulation 1.....	8
Data Matrices.....	8
The Models and the Basis for Comparison.....	9
Results.....	10
Discussion.....	12
Simulation 2.....	14
Data Generation and Analysis.....	17
Results.....	18
Discussion.....	26
References.....	27
DISTRIBUTION LIST.....	29
DD FORM 1473.....	32

List of Figures

<u>Figure</u>	<u>Page</u>
1 Choices Varying on Two Dimensions	15

List of Tables

	<u>Page</u>
1 Example of a Correlation Matrix Used as an Input to the Simulation	3
2 Means and Standard Deviations for the Computer-Generated Data Matrix Consisting of 25 Observations on Each of Four Variables	5
3 Sample Intercorrelations for the Computer-Generated Data Matrix Consisting of 25 Observations on Each of Four Variables	6
4 Percentage of Times Regression Model Outperforms Unit Weighting Model for Various Sample Sizes and Number of Variables with Measurement Error	11
5 Expected Mean Squared Error (E(MSE)) for the Two Models for Various Sample Sizes and Number of Variables with Measurement Error	13
6 Mean Square Error (MSE) between the Three Models for Various Weighting Schemes and Number of Attributes with Measurement Error	19
7 Mean Square Error (MSE) between the Three Models for Various Weighting Schemes and Number of Attributes with Measurement Error	20
8 Summary Statistics for the Correlation Coefficient between Model 2 (Differential Weights) and Model 3 (Unit Weights)	22
9 Summary Statistics for the Correlation Coefficient between Model 2 (Differential Weights) and Model 3 (Unit Weights)	23
10 Summary Statistics for the Correlation Coefficient Between Model 1 (Differential Weights-Product Term) and Model 3 (Unit Weights)	24
11 Summary Statistics for the Correlation Coefficient between Model 1 (Differential Weights-Product Term) and Model 3 (Unit Weights)	25

Acknowledgment

Support for this research performed by Social Science Research Institute was provided by the Advanced Research Projects Agency of the Department of Defense and was monitored under Contract N00014-76-C-0074 with the Office of Naval Research, under subcontract from Decisions and Designs, Inc.

Unit Versus Differential Weighting Schemes for Decision Making: A Method of Study and Some Preliminary Results

Introduction

A persistent and pervasive problem in decision making research and technology is how to combine presumably relevant information into a composite and then use that composite in making a decision.

In this paper we describe a method to study one aspect of this general problem: how to weight information in forming the composite. With this method, we can explore various differential weighting schemes and for each such scheme a comparison can be made with the simplest possible scheme, namely, just adding the variables or "unit weighting." There is an accumulating body of evidence indicating that such unit weighting may be as good and in some places, "better" than more complicated differential weighting schemes such as multiple regression. This evidence has a theoretical and analytic underpinning as provided by the work of Wilks (1938), Gulliksen (1950, Ch. 20), and more recently Einhorn and Hogarth (1975), Wainer (1976), Wainer and Thissen (1976) and Green (1974). There have also been several empirical studies as represented by the works of Lawshe and Shucker (1959), Wesman and Bennett (1959), and Fischer (1972). There have been at least three computer simulation studies (Schmidt, 1971, 1972; Claudy, 1972), and the approach we take is similar to such simulations. In an important review and analysis, Dawes and Corrigan (1974) argue cogently that simple additive (unit weighting) models are quite appropriate and indeed desirable in many decision making situations.

The implications of these results have extraordinary theoretical and practical significance. Our own interest is in the area of human decision making where the problem facing the decision maker is to choose an action alternative from a set of competing alternatives. The choice is difficult in most practical situations because there are often many attributes or dimensions of importance relevant to the choices available, and these attributes are often in conflict and vary in degree of apparent importance. In several places, Edwards and his associates (Edwards, 1971; Edwards and Guttentag, 1975; Edwards, Guttentag, and Snapper, 1975; Gardiner and Edwards, 1975) have proposed a multi-attribute utility measurement scheme which is very simple to apply in a wide variety of practical decision making contexts. The method involves simple ranking of the attributes and the attachment of importance weights to the attributes by using a modification of ratio scaling. Single attribute utilities are then measured, multiplied by the importance weight of the attribute and summed across attributes to yield an overall utility. Edwards offered this technique as a replacement for the more intricate and complex models advocated by mathematically-oriented decision theorists, such as Raiffa (1968), and Keeney and Raiffa (1976). If unit weighting is equal to or superior to such weighting, then perhaps the Edwards approach should adopt unit weighting and things would become even simpler for the practical decision maker. The entire practice of applied differential psychology as exemplified by Anastasi (1958) and econometrics (Johnson, 1972), just to mention two other important areas, will also become simpler if we can specify the conditions in which unit weighting is appropriate. This is the first of a series of reports to help specify those conditions.

Method

The method proposed is a Monte Carlo simulation of a multivariate process. The simulation generates a random variable vector $X = (x_1, x_2, \dots, x_M)$ from $N(0, \Sigma)$ the multivariate normal distribution with mean zero and variance-covariance matrix:

$$\Sigma = \begin{bmatrix} \sigma_{11} & \dots & \sigma_{1M} \\ \vdots & & \vdots \\ \sigma_{M1} & \dots & \sigma_{MM} \end{bmatrix}$$

The computational procedure for this method is provided by Scheuer and Stoller (1962) which is based on a theorem given in Anderson (1958, p. 19ff). Let Y be a normally distributed random vector with $N(0, I)$ where I is the unit matrix of size M . Anderson demonstrates that there is a matrix C with elements c_{ij} such that if $X = CY$, then $X = (x_1, x_2, \dots, x_n)$ is distributed as $N(0, CC')$, where C' is the transpose of C . The solution to finding C is familiar to psychologists as the "square-root" or "pivot" method of factor analysis and is described in Fruchter (1954) and Harmon (1967). The "square-root" method decomposes the original variance-covariance matrix Σ into a lower triangular factor matrix C such that $CC' = \Sigma$. Once the elements c_{ij} of C have been determined, then the column vector X is obtained from:

$$X = CY \tag{1}$$

where Y is a column vector containing elements $y_1, y_2, \dots, y_M$ which are standardized normal variates, i.e., with zero mean and unit variance. This method can be repeated for any number (N) times and is readily adopted for high speed computation. The end result is to produce an $N \times M$ data matrix with N rows depicting observations, for example subjects or choice alternatives, and M columns depicting measurements such as experimental conditions, psychological tests, or attributes. The elements of the data matrix represent a random sample from a multivariate normal population and thus simulate a "score" for each observation on each of the variables. In order to illustrate and further clarify this method, we will outline the steps of a computer program written to accomplish the above description.

Step 1. An "input" standardized variance-covariance matrix is defined such as that given in Table 1 which presents the intercorrelations between four variables.

The numbers in Table 1 are taken from a study by Sloan and Newman (1955). The first three variables are sub-tests of the Wechsler Intelligence Scale for Children (WISC); the fourth variable is the full-scale weighted score (FSWS), based on all ten sub-tests of the WISC. The investigators were interested in

Table 1
Example of a Correlation Matrix
Used as an Input to the Simulation

Variable				
	1	2	3	4
1	1			
2	.50	1		
3	.43	.45	1	
4	.47	.81	.74	1

developing a short form of the test and variables 1-3 represent the best combination of sub-tests to predict variable 4 (FSWS). Use of the correlation matrix rather than the variance-covariance matrix does not restrict generality and simplifies the computations.

Step 2. The variance-covariance matrix is decomposed by using the square-root method to yield a new matrix C such that $CC' =$ the input matrix. Before proceeding to Step 3, a check is made to see whether this identity holds. For example, if Table 1 is decomposed, we make sure the original table can be reproduced before proceeding.

Step 3. For each member of a given sample size of N observations, a sample of M standardized normal deviates are selected from the univariate normal distribution with zero mean and unit variance, where M stands for the number of variables. For example, with Table 1, $M = 4$ and if the sample size N was equal to 25, the result of this step yields 25 vectors, each containing 4 standardized normal variates. Each of these is treated as a column vector which is pre-multiplied by the matrix C obtained in Step 2. The result is the $N \times M$ data matrix. As an illustration of the end result of this step, Table 2 presents the summary statistics (means and standard deviations) for a particular data matrix based on a sample size 25 created from the input matrix represented by Table 1. Table 3 depicts the sample intercorrelations between the four variables.

Note that in Table 2 the means and standard deviations are not quite zero and one, respectively. This is due to sampling fluctuations. Sampling variability also accounts for the fact that the sample intercorrelations presented in Table 3 are not always equal to the input intercorrelations depicted in Table 1 which formed the basis for the sample. Of course, as the sample size is increased, sampling errors will decrease and the calculated sample statistics will be closer to their true values.

Step 4. The data matrices generated in Step 3 contain sampling error due to the sampling process but are free of measurement error. To make this simulation more realistic, Step 4 adds probabilistic noise to any one or all of the variables. This is accomplished by random sampling either from the uniform probability distribution defined over some interval (a,b) or from the normal probability distribution with parameters $(0,\sigma)$. This step simulates measurement error, and this error can be added to just one, say, the dependent variables, or to any subset of variables in the generated data matrix.

Brief Comment on Rationale.

Before giving examples of how this data-generating method can be used, we would like to comment briefly on two issues that formed the basis for developing the method. These can be stated as questions. The first is: why choose the multi-variate normal probability distribution as the basis for generating data? Our answer is that this distribution has very nice properties. Each of the marginal and conditional probability distributions is also normal. Thus we can alter any of the generated frequency distributions by knowledge of the properties of the normal distribution. For example, suppose we wanted to create a positively skewed distribution on one of the mar-

Table 2
Means and Standard Deviations for the Computer-
Generated Data Matrix Consisting of
25 Observations on Each of Four Variables

	Variable			
	1	2	3	4
Mean	.001	.080	-.128	.05
S.D.	1.256	.955	1.196	1.081

Table 3
Sample Intercorrelations for the Computer-
Generated Data Matrix Consisting of
25 Observations on Each of Four Variables

Variable				
	1	2	3	4
1				
2	.73			
3	.57	.50		
4	.69	.81	.81	

ginal distributions. This could be accomplished by stating what proportion of the cases should be to the right of a certain standard score. As the numbers are generated, if they are at or exceed this standard score value, then a positive constant or a randomly selected positive value could be added to such numbers. As another example, suppose we wanted approximately 2% - 2.5% of the generated numbers to be "outliers," that is, to lie outside an expected range. This could be accomplished by specifying that if a generated number fell at or beyond 2 plus or minus standard score units, then an appropriate constant, or a randomly selected constant, is added to that number. We cannot think of any type of realistic (or, for that matter, unrealistic) data that cannot be simulated by this method. Unfortunately, we do not have any convenient theoretical or empirical reason, except that provided by previous studies, to enable us to state what kind of data would be of interest to generate.

The second question is: how does the method lend itself to the case where there are no well-defined independent (predictor) or dependent (criterion) variables? Actually, this is the case of most interest to us since many applied decision problems are of this kind. Perhaps the best way to answer this question is to give a concrete example. Suppose a large police department is faced with the problem of selecting a new motor vehicle to serve as the official patrol car in the city for the next three years. You are faced with the responsibility for guidance in selecting the "best" vehicle.

The attributes or dimensions of importance for any set of competing motor vehicles might be such things as performance, expense, safety, reliability, and vehicle comfort. Each of these might have sub-dimensions. For example, performance would most certainly include acceleration speed, and top speed in this context. Any subset of these dimensions may be correlated or uncorrelated. (One way to estimate what these correlations might be would be to take the ratings on a large number of motor vehicles given by the magazine Consumer Reports on each of the dimensions mentioned above and obtain an intercorrelation matrix from these ratings.) As soon as these correlations are specified, then data which would represent a vector of the numerical values on each of the dimensions for each candidate vehicle can be generated. Various rank orderings of the dimensions on a scale of importance could also be simulated. In practice, this ordering and attachment of importance weights are done by experts concerned about the final decision. Also, these same experts might well have different preferences or utilities of how each candidate motor vehicle rates on each of the importance dimensions. These utility measurement numbers can also be simulated by this method. We now have two sets of numbers, the importance weights of the dimensions and the utility measurement of each candidate vehicle on the dimensions. Various ways of forming composites and aggregating these two numbers, including that of just unit weighting, for example, adding up each vehicle's score on each dimension and choosing that vehicle with the highest sum, can now be specified and compared.

We now turn to two applications of the data generating method described above. The first concerns a straightforward prediction problem typical of that found in many disciplines. The second focuses on various ways of aggregating information in a decision analysis situation.

Simulation 1

In this study, we compared the familiar multiple regression model applied to a well-defined prediction problem, with the simplest model possible; for example, unit weighting. By a well-defined prediction problem, we mean one in which there is a criterion variable and several predictor variables, which are correlated with that criterion as expressed by validity coefficients. The predictor variables are correlated with each other as well. One reason for studying this problem first is that it is probably the most common problem in applied psychology and other fields. Correlation matrices of the type depicted in Table 1 abound in psychology. The three validity coefficients are moderate to moderately high, and the intercorrelation between the three predictor variables are moderate. The multiple regression model is the most recommended to apply to such matrices for the purposes of obtaining the optimal prediction of the criterion.

A second reason for this study is that the regression model is the one that has been most frequently compared with unit weighting in the literature, and previous investigators have often reported that unit weighting outperformed the regression model. This is such an important finding that we wanted to study it carefully. Incidentally, we purposely chose a prediction problem the characteristics of which would lead one to expect the regression model to do very well when compared to unit weighting. The correlations in Table 1, which will be the input to the data generating procedures, are all positive. There are only three predictor variables and the validity coefficients and predictor intercorrelations are all positive. If there were no sampling or measurement error, then it would be possible to calculate the true standardized regression coefficients and these would be quite different from one another and will be presented later in this paper. Also, the generated data are all sampled from a multivariate normal distribution and there was little chance for any "outliers" to appear in the data.

Data Matrices.

Using as input the standardized variance-covariance (correlation) matrix given in Table 1, a set of $N \times M$ data matrices were generated where N is the sample size and $M = 4$, in this case, is the number of variables. Variable 4 was designed as the criterion or dependent variable. The sample size N took on four values, 25, 50, 75, and 100. After each matrix was generated, it was "copied" into a second matrix and random noise (measurement error) was added to either one, two, three, or four variables. The error was sampled from the uniform probability distribution defined over the unit interval (0,1). Thus, for each run, a data matrix was generated that was free of measurement error for each of the sample sizes, and this matrix was then contaminated with measurement error in from one to four of the variables. In one set, error was added only to the criterion variable. In other sets, error was added to one or more of the independent or predictor variables up to 3, but not to the criterion; finally, error was added to all four variables. Thus, there were 24 types of data matrices to be used in comparing the models described in the next section. For each type, 100 data matrices were sampled to form the basis of model comparison.

The Models and the Basis for Comparison.

The two models to be compared were simple unit weighting and multiple regression. The first, of course, is formed by simply summing the first three independent variables (1-3). This was added as variable 5 to each data matrix. The second one was generated by the familiar least squares estimates for the regression coefficients, and this formed variable 6 for each data matrix. Incidentally, before each model was formed, all data values for all four original variables were standardized. The squared multiple correlation coefficient (R^2) was calculated between variable 4 (criterion) and variable 5 (unit weighting prediction of criteria) and between variable 4 and variable 6 (multiple regression prediction of the criterion). In comparing the two models, we followed the scheme used by Einhorn and Hogarth (1975) in their analytic study of the efficacy of unit weighting. The expected mean squared errors of the regression and unit weighting models respectively are given by two well-known formulas:

$$E(\text{MSER}) = \frac{\sum_{i=1}^n e_i^2}{\text{DFR} = N-k-1} \quad (2)$$

where $\sum_{i=1}^n e_i^2$ is the sum of the squared residuals of the regression model; and

$$E(\text{MSEU}) = \frac{\sum_{i=1}^n u_i^2}{\text{DFU} = N-2} \quad (3)$$

where $\sum_{i=1}^n u_i^2$ is the sum of the squared residuals from the equal weighting

model. The denominators DFR and DFU of (2) and (3) are the degrees of freedom for the respective models. For the regression model (2), k plus 1 degrees of freedom are lost for the k regression coefficients and the additive constant that must be estimated from the data. For the unit weighting model (3), the degrees of freedom are actually that for the equal weighting model where all the regression coefficients of (2) are set equal to some constant. Thus, there are only two values, the one regression coefficient and the additive constant, that need to be estimated from the data. Strictly speaking, there are no degrees of freedom lost for the unit weighting model since nothing is estimated from the data. However, as Einhorn and Hogarth point out, the unit weighting model will correlate perfectly (1.0) with the equal weighting model and, thus, equation (3) allows us to use standard statistical theory in comparing the two prediction models.

The relative predictive efficiency of the two models may be assessed by comparing the ratios of their mean squared errors:

$$\frac{E(\text{MSER})}{E(\text{MSEU})} = \frac{(N-2) \sum_{i=1}^N e_i^2}{(N-k-1) \sum_{i=1}^N u_i^2} \quad (4)$$

which, since the criterion variable is the same in both models and $1-R^2$ equals the sum of the squared residuals in (4), is re-expressed in the convenient form:

$$\frac{E(MSER)}{E(MSEU)} = \frac{(N-2)}{(N-k-1)} \frac{(1-R_R^2)}{(1-R_U^2)} \quad (5)$$

where R_R^2 and R_U^2 are the squared multiple correlations using the regression and unit weighting models, respectively. When the ratio in (5) is less than one, the regression model performs more accurately than unit weighting. When the ratio is greater than one, unit weighting will perform better than the regression model. For $k \geq 2$, the ratio $(N-2) / (N-k-1)$ of (5) will always be greater than one and thus favor the unit model. The ratio $(1-R_R^2) / (1-R_U^2)$ of (5), however, will always be less than or equal to one on initial fit thus favoring the regression model. For this reason, (5) was used to compare the two models on cross-validated regression models. This cross validation was accomplished sequentially in the sampling process; that is, the coefficients estimated in sample 1 were used to predict the actual values in sample 2; those estimated in sample 2 were used to predict the values in sample 3, and so on.

Results.

Table 4 presents the percentage of times the regression model outperformed the unit weighting model according to the criterion provided by (5), for the four sample sizes and the number of variables with measurement error. Note that except for the case of the smallest sample size ($N=25$), the regression model almost always does better than the unit weighting model. With $N=25$, the regression model goes from a better performance percentage of 92 down to 69 as the number of variables infested with measurement error increases from 0 to 4. For small sample sizes, the initial estimates of the regression coefficients can be quite unstable and often do not hold up well under cross validation especially when the number of variables with measurement error increases. Thus, while on a percentage of better performance basis the regression model appears superior to unit weighting, the use of the model under these conditions is questionable, a familiar finding.

The results in Table 4 would seem to indicate that the regression model should be preferred over unit weighting for sample sizes of 50 or more, at least for the type of data matrices studies here. However, consider Table 5, which presents the expected mean squared error for the two models for the various sample sizes and the number of variables that have measurement error. Each value in Table 5 represents an average of 100 samples and can be considered very stable.

The results in Table 5 indicate that while the regression model is in all cases superior to unit weighting, the difference between them is not all that dramatic, the regression model reducing the absolute error somewhere between 7 and 13 percent in comparison to unit weighting.

Table 4
Percentage of Times Regression Model Outperforms
Unit Weighting Model for Various Sample
Sizes and Number of Variables with Measurement Error

Sample Size	Number of Variables with Measurement Error					
	0	1(c) ^a	1	2	3	4
25	92	82	91	85	75	69
50	98	94	96	96	95	92
75	99	98	99	99	98	96
100	99	99	99	98	99	98

Note: Each percentage is based on 100 replications of each sample size.

^aerror added to the criterion variable only.

Discussion.

These results are in general agreement with those reported by Einhorn and Hogarth (1975), Schmidt (1971) and Claudy (1972). These investigators found regression weights to be superior to unit weights for sample sizes of 50 or greater, and our simulation supports this finding. For small sample sizes (e.g., $N < 25$), none of the above investigators recommend the use of regression weights over unit weights. We concur, although our results are not as pessimistic; that is, the regression weights do show a superiority on the average, but this advantage dissipates rapidly when measurement error is added to the variables, and we would have to conclude that the estimation of regression weights is not worth the effort for small sample sizes.

The most striking feature of our results, however, is that while regression weights may be superior to unit weights, in many situations the superiority is certainly not very great. It should be remembered that we purposely chose a type of situation in which one might expect a priori the regression model to do quite well; that is, the sample was from a multivariate normal population, the validity coefficients between the predictor and criterion variables were moderately strong, and the intercorrelations between the predictors were moderate. There were only three predictors, and all previous findings have shown that with few predictors, regression models are often superior to unit weighting models, but this advantage disappears as the number of predictors increases. Also, the input correlation matrix given in Table 1 actually presents the true population correlations for the data-generating model. Sampling and measurement error omitted, the true standardized regression weights for the three predictor variables in this study are: -.0452, .6155, and .4825 respectively. Assigning unit weights would seem to be ridiculously high and result in considerable error. Even when conditions are quite favorable to the regression model, however, our results show that this advantage is never greater than 13% in absolute value comparison between the two models. In light of the complexity of the regression model, as contrasted to the unit weighting model, this is hardly a command performance.

Of course, in terms of relative improvement, the regression model does look quite good compared to the unit model. If one forms the ratio: $E(MSER)/E(MSEU)$ of the values given in Table 5, then the relative improvement of the regression model over the unit model goes from a low of 21% to a high of 67%, depending on what condition is considered. There are many practical situations in which such relative improvement could have a high "payoff."

Our results do not agree with those of Dawes and Corrigan (1974) or Wainer (1976), who reported that unit weighting was actually superior to regression weights. The explanation of this lies in the way the regression coefficients are estimated in the initial sample. If they are poorly estimated, they will not hold up well on a validating sample, and equal weights will be as good or superior. The conditions under which regression weights will be wrongly estimated in the initial sample are: (a) very small sample size relative to the number of predictors; (b) a large number of predictors; (c) the presence of "outliers" which are outside the expected range of sample values; (d) non-normality in the original sample. Equal or unit weights are insensitive to all such conditions and thus will often be superior to regres-

Table 5
Expected Mean Squared Error [E(MSE)] for the Two Models for
Various Sample Sizes and Number of Variables with Measurement Error

	N							
	25		50		75		100	
No. of Variables With Error	E(MSEU)	E(MSER)	E(MSEU)	E(MSER)	E(MSEU)	E(MSER)	E(MSEU)	E(MSER)
0	.30	.20	.30	.19	.32	.20	.29	.18
1(c) ^a	.35	.26	.36	.26	.36	.26	.35	.24
1	.30	.19	.31	.19	.33	.20	.30	.18
2	.31	.23	.32	.23	.32	.22	.32	.21
3	.35	.27	.34	.24	.33	.23	.33	.23
4	.40	.33	.37	.29	.38	.29	.37	.29

^a error in the criterion variable only.

sion weights when such conditions exist in the initial sample. By design, none of these conditions existed in our data-generating process. Thus when conditions are "just right" for the regression model, that model will behave as the statistical theory says it should behave. However, in practical situations, one cannot expect everything to be "just right," and any investigator should be very cautious in applying the regression model when the conditions for the model are not "just right."

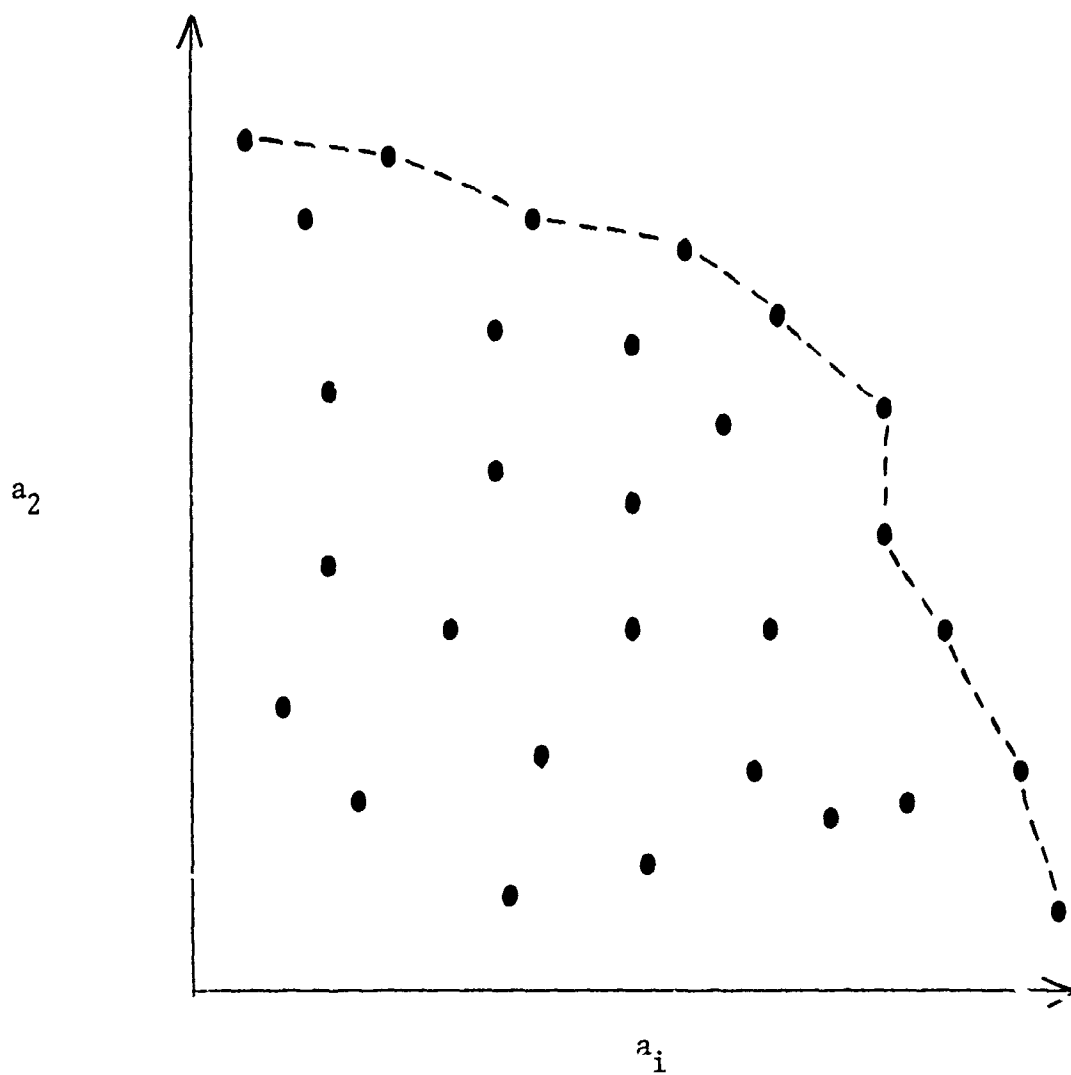
It should be pointed out that there are ways to improve the estimation of regression weights in the initial sample. Mosteller and Tukey (1968) have described the "jackknife" procedure; Hoerl (1964) recommends ridge estimation of regression coefficients; and Lindley and Smith (1972) describe a Bayesian procedure to regression and demonstrate its superiority to conventional regression estimation. Wainer and Thissen (1976) recommend smoothing the initial estimates, especially when there is a theoretical structure for the regression coefficients.

We intend to explore such techniques and compare them with unit weighting in subsequent studies.

Simulation 2

Our second simulation focuses on a common problem in multi-attribute decision analysis situations. This problem has two parts: how to assign importance weights to the attributes considered relevant to the choice alternatives and how to aggregate or combine these weighted attributes into a composite to aid in the final decision. The choice of attributes to be included and the assignment of the importance weight to each attribute is usually done by expert human judgment. Unlike the problem studied in the previous section, there is no clear-cut dependent or criterion variable to which the attributes may be related, and thus we are not trying to predict any criterion from knowledge of the attributes. Also, there is no clear cut theory such as a least squares regression model which stipulates how the weighted attributes are to be combined or aggregated. For example, consider the decision problem alluded to earlier, that of trying to decide which of many automobiles available should be chosen as the "official car." Some of the attributes considered important in making this decision might be such things as fuel economy, small exterior size, large interior size, passing/acceleration ability, low interior noise, ride quality, cost, and the like. These characteristics interact and trade-offs are often necessary. For example, in order to obtain excellent fuel economy, it might be necessary to sacrifice acceleration. This could be accomplished by considering lighter cars, but these, in turn, could adversely affect ride quality, interior size, and so forth. These interactions are reflected in both positive and negative correlations between the attributes. In addition, we must consider that in many choice situations we would expect the average intercorrelation among the attributes of alternatives being actively considered to be negative.

To explain this, consider choices among alternatives that have two relevant attributes, a_1 and a_2 . Figure 1 illustrates the possible values of a_1 and a_2 that the choice alternatives might assume, with each attribute oriented



Choices Varying on Two Dimensions

Figure 1

so more is better. Each point represents a single choice alternative with its two-dimensional values. In situations where a single alternative is to be chosen, any rational decision maker should consider only alternatives that are not dominated. In the two-attribute case, this means that the decision maker would not consider any alternative for which there is another alternative that is at least as good on both attributes and better on one. In Figure 1 only alternatives on the northeast boundary would be considered. Thus, for all alternatives being considered, being better on one dimension must mean being worse on the other. This leads to a negative correlation between the attributes. In the two-attribute case, this negative correlation could be as high as -1.0. Obviously, as the number of attributes increases, the lower bound on the average intercorrelation increases. In the particular situation where all intercorrelations are assumed equal, the lower bound is $-1/(k-1)$. In the case where multiple alternatives are to be chosen, the situation is slightly more complicated. If j alternatives are to be chosen, with $j > 1$, an alternative must be dominated by at least j other alternatives to be eliminated from consideration. Thus, it appears that in most multi-attribute choice situations there are likely to be negative inter-attribute correlations, a circumstance not previously investigated in comparing unit and differentially weighted models.

The second reason for extending this analysis is that in the multi-attribute choice situation, the combination rule for aggregating the single attribute utilities is not always strictly an additive rule. When the single attribute utilities are combined additively, the model becomes mathematically similar to the regression model. However, interaction terms are often called for in the multi-attribute utility combination rule. Thus, a general formula for the total utility U_j of any alternative j could be:

$$U_j = \sum_i w_i u_i(a_{ij}) + f(u_1(a_{1j}), u_2(a_{2j}), \dots, u_k(a_{kj})) \quad (6)$$

where $u_i(a_{ij})$ is the utility assigned to alternative j on attribute i , the w_i 's are importance weights, and f is a general term to be used in forming the composite U_j that allows for required interaction terms such as cross product terms. If no interactions are needed in this model, the f term vanishes, and the model becomes simple additive. If such additivity can be expected to hold for the case of positive correlation between the attributes, we know what to expect if we use (6) as the basis for making a final decision. Just about any version of (6) is as good as any other. Thus, if you set $f(u(a_{ij}))=0$ and the $w_i=1$ for all i , you have the unit weighting model which should perform as well as any other version of (6). Since this version is the simplest, any rational decision maker should certainly adopt it. The analytic proof of this rule is due to Wilks (1938), who demonstrated under reasonable conditions that the average value of the correlation between any two linear combination of attributes differs from unity by terms of order $1/k$ and the variance of the correlation is of order $1/k^2$, where k is the number of attributes. The larger the mean value of positive correlation between pairs of attributes, the more rapidly the mean value of the correlation between linear combinations approaches 1.

Because of the complexity the combination rule for U_j can assume due to

the flexibility of the f term, we have restricted this particular simulation to the simplest possible case, the two-attribute case. For this study the correlation between the two attributes was also restricted at moderately high positive (.8) and moderately high negative (-.8).

Assuming that $u_i(a_{ij}) = a_{ij}$ for all i which does not restrict our results since this can occur either by assuming u_i 's are linear or the a_{ij} 's are already measures of utility, we studied three versions of equation (6):

$$\text{Model 1} = U_1 = w_1 a_{1j} + w_2 a_{2j} + w_3 a_{1j} a_{2j} ;$$

$$\text{Model 2} = U_2 = w_1 a_{1j} + w_2 a_{2j} ;$$

$$\text{Model 3} = U_3 = a_{1j} + a_{2j} .$$

The three weights w_i , $i = 1, 2, 3$, were allowed to vary in increments of .1 on w_1 , the weight assigned to attribute 1 under the following schema:

$$w_1 = .1$$

$$w_2 = w_1 + .1$$

$$w_3 = (1 - w_1 - w_2) .$$

Thus, for model 1, which has a cross-product term, the first pass through attribute 1 received a weight of .1, attribute 2 a weight of .2 and the product attribute 2 a weight of .3 and the product term a weight of .5. Since the weights were constrained to sum to one, the last pass-through gave a weight of .5 to attributes 1 and 2, respectively and a 0 weight to the product term.

For model 2, the first pass through attribute 2 received twice the weight of attribute 1, the second time the weights assigned to attributes 1 and 2 were .2 and .3, respectively, and the final pass resulted in equal weights (.50) being assigned to the two attributes. (There was no w_3 for this model.)

Model 3, of course, is simple unit weighting. All three models are identical when models 1 and 2 are in the equal weighting version.

Data Generation and Analysis.

The data generation was similar to that described in the previous section. The input "matrix," however, was the simplest possible, consisting either of a .8, or -.8 correlation between the two attributes. With either input, we chose to generate two sample sizes (N) of data, 25 and 50. Thus the output of the data generation process was a $N \times 2$ table of numbers with each row of the table representing a utility on each of the two attributes.

The three aggregation models were then formed and the correlation between each of the models obtained. Thus, the correlation coefficient between the models is the dependent variable and was used to form the basis for assessing

the similarity or dissimilarity of the models. To make the analysis somewhat akin to that presented in the previous section, we used $1 - r_{ij}^2$, the non-overlapping or error variance, as the basic measure to compare the dissimilarity of the models where r_{ij}^2 is the square of the correlation coefficient between model i and j . The larger this expression, the more dissimilar the models are. Each correlation coefficient was based on either 25 or 50 sample values, and we repeated each sample size 50 times; thus, the results of the model comparison can be based on an average of 50 repetitions. Also, as in the previous simulation, we considered 3 conditions of measurement error: none, measurement error added to one attribute, and measurement error added to both attributes.

Before presenting the results we would like to clarify in specific terms just what the simulation consisted of. Consider the decision problem of trying to decide which automobile would be the best choice, and only two attributes influence the decision: fuel economy and small exterior size. These attributes could be scaled to correlate positively. Other things being equal, the smaller the exterior size of a car the less fuel it consumes. If now the sample size is 25, then the data generation program generates a numerical utility for each of 25 competing automobiles on each of the two attributes. The three models are then formed for each of these 25 cars. If this were a real decision problem, then the car that has the highest composite utility for whatever model is being used, would be chosen. We then compared the similarity-dissimilarity of the three models by repeating the entire simulation 50 times and calculating the correlations between the three models each time. For the case of negative correlation between the attributes, then, the two attributes considered relevant might be fuel economy and large interior size. These can be expected to be negatively correlated in practice. For this case we performed the entire simulation again.

Results.

Tables 6 and 7 present the mean square error (MSE) between the three models for the various weighting schemes and the number of attributes with measurement error. Table 6 is for the sample size being equal to 25, and Table 7 is for the sample size 50. The numbers outside the parentheses are the mean square error for the positively correlated attributes, and the values inside the parentheses are the mean square error when the attributes were negatively correlated. Because there is very little difference between Tables 6 and 7, indicating that sample size does not influence the main results, we will focus our discussion around Table 6. Note that the mean square error between model 2 (differential weighting) and model 3 (equal weighting), labeled MSE23, is in all cases about 0 indicating almost perfect similarity between unit and differential weighting when the attributes are positively correlated. This was the expected finding. With negative correlation between the attributes, as indicated by the MSE23 within the parentheses, there is dissimilarity between the two models, and this dissimilarity is greatest for the case of the weights being most dissimilar. As the weights get closer together in relative size, the models become more similar. It is somewhat instructive to compare the summary statistics on the correlation coefficients, for the -.8 attribute intercorrelation case, calculated between models 2 and 3 for the 50 repetitions of the sim-

Table 6

Mean Square Error (MSE) between the Three Models for Various
Weighting Schemes and Number of Attributes with Measurement Error^{a,b}

No. of Attributes With Error	MSF12			MSE13			MSE23		
	0	1	2	0	1	2	0	1	2
<u>Weights</u>									
W_1 W_2 W_3									
.1 .2 .7	.78 (.88)	.57 (.69)	.43 (.81)	.78 (.93)	.59 (.85)	.43 (.79)	.01 (.48)	.02 (.46)	.02 (.43)
.2 .3 .5	.52 (.86)	.37 (.73)	.28 (.71)	.52 (.91)	.38 (.86)	.28 (.73)	.00 (.25)	.00 (.24)	.00 (.23)
.3 .4 .3	.20 (.67)	.16 (.58)	.12 (.45)	.20 (.72)	.16 (.72)	.13 (.43)	.00 (.14)	.00 (.15)	.00 (.12)
.4 .5 .1	.02 (.13)	.02 (.11)	.01 (.11)	.02 (.20)	.02 (.23)	.01 (.20)	.00 (.10)	.00 (.09)	.00 (.09)
.5 .5 0	.01 (.12)	.01 (.14)	.01 (.11)	.01 (.18)	.01 (.14)	.01 (.16)	.00 (.08)	.00 (.07)	.00 (.06)

^a The values in the parentheses are for attributes which are negatively correlated (-.8).

^b Each MSE is based on 50 repetitions of a sample size equal to 25.

Table 7

Mean Square Error (MSE) between the Three Models for Various

Weighting Schemes and Number of Attributes with Measurement Error^{a,b}

No. of Attributes With Error	MSE12			MSE13			MSE23		
	0	1	2	0	1	2	0	1	2
<u>Weights</u>									
W_1									
W_2									
W_3									
.1	.80 (.90)	.60 (.73)	.45 (.86)	.80 (.96)	.61 (.92)	.45 (.84)	.01 (.50)	.02 (.49)	.02 (.45)
.2	.58 (.86)	.41 (.71)	.27 (.72)	.58 (.91)	.42 (.85)	.27 (.72)	.00 (.25)	.00 (.23)	.00 (.21)
.3	.24 (.70)	.18 (.56)	.14 (.52)	.24 (.73)	.19 (.70)	.14 (.53)	.00 (.16)	.00 (.14)	.00 (.13)
.4	.02 (.15)	.02 (.13)	.02 (.11)	.02 (.24)	.02 (.27)	.02 (.16)	.00 (.10)	.00 (.09)	.00 (.08)
.5	.01 (.11)	.01 (.12)	.02 (.10)	.01 (.17)	.02 (.12)	.02 (.16)	.00 (.07)	.00 (.06)	.00 (.05)

^a The values in the parentheses are for attributes which are negatively correlated (-.8).

^b Each MSE is based on 50 repetitions of a sample size equal to 50.

ulation. Some of these are presented in Table 8, for the condition of no measurement error in the attributes, and Table 9 for the condition of measurement error added to both utilities.

Note in Table 8 when attribute 1 receives twice the weight of attribute 2, the mean correlation between the two models is .71, and there is a wide range in the computed correlations going from a low of .32 to a high of about .9. There is a slight tendency for the adding of measurement error to change things, but, for all practical purposes, adding error does not seem to change the results much. (Table 9). Tables 8 and 9 also indicate that as the differential weights become more similar, as indicated by part (b) and (c) of the tables, then the similarity between the two models becomes greater (Mean = .95), and the range of calculated values is reduced substantially.

Now consider the mean square error comparing model 1 (weighted with a cross product term) and model 3 (unit weighting), labeled as MSE13 in Table 6 and 7. Here we see that there are substantial differences between the models for both positively and negatively correlated attributes. The inclusion of a product term in model 1, which in turn gets a high weight relative to the other terms, gives quite different results than simple unit weighting. For the case of positively correlated attributes, this dissimilarity is diminished as the differential weights become more alike and is virtually eliminated when the weights assigned to the attributes are almost alike (.4, .5) and the cross product term gets a low relative weight (.1). For the case of negatively correlated attributes, while the similarity between the models increases as the weights become more alike, there remains a substantial dissimilarity between the models. Incidentally, unlike the comparison between models 2 and 3 (MSE23), adding measurement error to the utilities on the attributes does make a difference in comparing models 1 and 3 (MSE13). Adding error seems to make the models more similar, as indicated by decreases in the mean square error.

To further clarify this finding, we present the summary statistics for the calculated correlation coefficients, for the most extreme and least extreme differential weights; Table 10 presents these for the case of no measurement error and Table 11 for the case of measurement error added to both utilities. Table 10 indicates that when the differential weights are extreme, with the product term of model 1 receiving a large relative weight, the average correlation between this model and model 3 (unit weight) is quite small, especially for the case of negative correlation between the attributes, and the range of the calculated coefficients is quite wide. When the weights for attributes are closer to being equal (.4, .5), and the product term receives a low relative weight, then the difference between the two models is virtually eliminated for the positively correlated attributes but is maintained for the negatively correlated attributes. Table 11 indicates that when measurement error is added to the attributes, then the differences between the two models are not as great.

These results clearly indicate that a model that includes a cross-product term as the weighted utility composite (such as model 1 of this simulation) will yield substantially different results from a simple unit weighting model, especially when the product term receives a relatively high weight.

Table 8

Summary Statistics for the Correlation Coefficient
between Model 2 (Differential Weights)
and Model 3 (Unit Weights)

	(a) $w_1=.2, w_2=.1$	(b) $w_1=.4, w_2=.5$	(c) $w_1=.5, w_2=.5$
Mean	.71	.95	.96
Std. Dev.	.12	.02	.02
Min.	.32	.90	.92
Max.	.89	.99	.99
Range	.57	.09	.07

Note. N=50 for each condition. No measurement error.

Table 9

Summary Statistics for the Correlation Coefficient
between Model 2 (Differential Weights)
and Model 3 (Unit Weights)

	(a) $w_1=.2, w_2=.1$	(b) $w_1=.4, w_2=.5$	(c) $w_1=.5, w_2=.5$
Mean	.75	.95	.97
Std. Dev.	.10	.02	.01
Min.	.40	.89	.93
Max.	.91	.98	.99
Range	.51	.09	.06

Note. N=50 for each condition. Measurement error added to both attributes.

Table 10
Summary Statistics for the Correlation Coefficient
between Model 1 (Differential Weights-Product Term)
and Model 3 (Unit Weights)

Corr. Between Attributes	Weights		
	$w_1 = .1, w_2 = .2, w_3 = .7$	$w_1 = .4, w_2 = .5, w_3 = .1$	
	+	-	+
Mean	.36	.16	.99
Std. Dev.	.31	.22	.004
Min.	-.41	-.34	.98
Max	.89	.58	.99
Range	1.30	.92	.01

Note. Based on 50 repetitions of a sample size equal to 25. No measurement error.

Table 11
Summary Statistics for the Correlation Coefficient
between Model 1 (Differential Weights-Product Term)
and Model 3 (Unit Weights)

Weights			
$w_1 = .1, w_2 = .2, w_3 = .7$ $w_1 = .4, w_2 = .5, w_3 = .1$			
Corr. Between Attributes	+	-	+
Mean	.73	.31	.99
Std. Dev.	.19	.23	.003
Min.	-.01	-.16	.98
Max.	.95	.77	.99
Range	.96	.93	.01

Note. Based on 50 repetitions of a sample size equal to 25 measurement error added to both attributes.

When the attributes are positively correlated, this difference will disappear when all the main terms receive similar weights and the product term receives considerably less relative weight. However, for the case of negatively correlated attributes, the differences between the two models are never eliminated. The practical implication of this finding is that when this condition holds, and for the type of decision problem considered (for example, selection of the "best automobile"), the final selection could be quite different, depending upon what model was being used.

Tables 6 and 7 also compare the mean square error between model 1 and 2 (MSE12), but the results are virtually the same as those just discussed for models 1 and 3 and will not be repeated.

Discussion

We will confine our discussion to the case of negatively correlated attributes, since, if the attributes are positively correlated, we know that differential weighting is not necessary if additivity holds.

For the simple case of just two attributes that are negatively correlated, then, differential weighting does make a difference, and in practical situations this difference can be very important. This results is in line with a more realistic study, which investigated the case of a mixture of positive and negative intercorrelations among eleven attributes (Newman, 1976). Of course, negative correlations can often be and should be changed in sign by rescaling the measured attributes. However, the Newman paper demonstrates that even when all the attributes are oriented in the "right" direction, the negative intercorrelations can still appear; and this condition can result in differential weighting affecting what the final choice in a decision situation might be.

This raises the intriguing question of just what weighting scheme should be used since the choice can critically affect the final decision. We have no answer to this question and, indeed, there may not be one. There is a strong need to develop a theoretical rationale for differential weighting when the attributes of importance are a mixture of positive and negative intercorrelations. We are pursuing this problem theoretically, but for the time being, we will continue with empirical studies.

REFERENCES

- Anastasi, A. Differential Psychology. 3rd ed. New York: Macmillan, 1958.
- Anderson, T.W. An Introduction to Multivariate Statistical Analysis. New York: John Wiley and Sons, 1958.
- Claudy, J.G. A comparison of five variable weighting procedures. Educational and Psychological Measurement, 1973, 32, 311-322.
- Dawes, R.M. and Corrigan, B. Linear models in decision making. Psychological Bulletin, 1974, 81, 95-106.
- Edwards, W. Social utilities. The Engineering Economist, 1972, 6, 119-129.
- Edwards, W. and Guttentag, M. Experiments and evaluations: A re-examination. Chapter in Bennett, C. and Lumsdaine, A. (Eds.), Experiments and Evaluations. New York: Academic Press, 1975.
- Edwards, W. Guttentag, M., and Snapper, K. Effective evaluation: A decision theoretic approach. In Streuning, E.L. and Guttentag, M. (Eds.), Handbook of Evaluation Research, Vol. 1. Beverly Hills, Ca.: Sage Publications, 1975.
- Einhorn, H.J. and Hogarth, R.M. Unit weighting schemes for decision making. Organizational Behavior and Human Performance, 1975, 13, 171-192.
- Fischer, G.W. Four methods for assessing multi-attribute utilities: An experimental validation. Engineering Psychology Laboratory, University of Michigan, 1972.
- Fruchter, B. An Introduction to Factor Analysis. Princeton, N.J.: Van Nostrand, 1954.
- Gardiner, P. and Edwards, W. Public values: Multi-attribute utility measurement for social decision making. In Kaplan, J., and Schwartz, S. (Eds.), Human Judgment and Decision Processes. New York: Academic Press, 1975.
- Green, B.J. Parameter sensitivity in multivariate methods. Mimeographed manuscript. Baltimore: Department of Psychology, Johns Hopkins University, 1974.
- Gulliksen, H. Theory of Mental Tests. New York: Wiley, 1960.
- Harmon, H. Modern Factor Analysis. 2nd ed. Chicago: University of Chicago Press, 1967.
- Hoerl, A.E. Ridge analysis. Chemical Engineering Progress Symposium Series, 1964, 60, 67-77.

- Johnston, J. Econometric Methods. 2nd ed. New York: McGraw-Hill, 1972.
- Keeney, R.L. and Raiffa, H. Decisions with Multiple Objectives. New York: Wiley, to be published (1976).
- Lawshe, C.H. and Schucker, R.E. The relative efficiency of four test weighting methods in multiple regression. Educational and Psychological Measurement, 1959, 19, 103-144.
- Lindley, P.V. and Smith, A.F.M. Bayes estimates for the linear model. Journal of the Royal Statistical Society, Series B, 1972, 34, 1-41.
- Mosteller, F. and Tukey, J. Data analysis, including statistics. In Ludzey, G. and Akonson, E. (Eds.), The Handbook of Social Psychology. 2nd ed. Reading, Ma.: Addison-Wesley, 1968.
- Newman, J.R. Differential weighting in multi-attribute utility measurement: When it should not and when it does make a difference. Social Science Research Institute Technical Report, University of Southern California, July, 1976.
- Raiffa, H. Decision Analysis: Introductory Lectures on Choices Under Uncertainty. Reading, Ma: Addison-Wesley, 1968.
- Scheuer, E.M. and Stoller, D.S. On the generation of normal random vectors. Technometrics, 1962, 4, 278-281.
- Schmidt, R.L. The relative efficiency of regression and simple unit predictor weights in applied differential psychology. Educational and Psychological Measurement. 1971, 31, 699-714.
- Schmidt, R.L. The reliability of differences between linear regression weights in applied differential psychology. Educational and Psychological Measurement, 1972, 32, 879-886.
- Sloan, W. and Newman, J.R. The development of a Wechsler-Bellevue II short form. Personnel Psychology, 1955, 8, 347-353.
- Wainer, H. Estimating coefficients in linear models: It don't make no nevermind. Psychological Bulletin, 1976, 83, 213-217.
- Wainer, H. and Thissen, D. Three steps towards robust regression. Psychometrika, 1976, 41, 9-33.
- Wesman, A.G. and Bennett, G.K. Multiple regression vs. simple addition of scores in prediction of college grades. Educational and Psychological Measurement, 1959, 19, 243-246.
- Wilks, S.S. Weighting systems for linear functions of correlated variables when there is no dependent variable. Psychometrika, 1938, 3, 23-40.

Research Distribution List

Department of Defense

Assistant Director (Environment and Life Sciences)

Office of the Deputy Director of Defense Research and Engineering (Research and Advanced Technology)

Attention: Lt. Col. Henry L. Taylor
The Pentagon, Room 3D129
Washington, DC 20301

Office of the Assistant Secretary of Defense (Intelligence)

Attention: CDR Richard Schlaff
The Pentagon, Room 3E279
Washington, DC 20301

Director, Defense Advanced Research Projects Agency

1400 Wilson Boulevard
Arlington, VA 22209

Director, Cybernetics Technology Office

Defense Advanced Research Projects Agency
1400 Wilson Boulevard
Arlington, VA 22209

Director, Program Management Office

Defense Advanced Research Projects Agency
1400 Wilson Boulevard
Arlington, VA 22209
(two copies)

Administrator, Defense Documentation Center

Attention: DDC-TC
Cameron Station
Alexandria, VA 22314
(12 copies)

Department of the Navy

Office of the Chief of Naval Operations (OP-987)

Attention: Dr. Robert G. Smith
Washington, DC 20350

Director, Engineering Psychology Programs (Code 455)

Office of Naval Research
800 North Quincy Street
Arlington, VA 22217
(three copies)

Assistant Chief for Technology (Code 200)

Office of Naval Research
800 N. Quincy Street
Arlington, VA 22217

Office of Naval Research (Code 230)

800 North Quincy Street
Arlington, VA 22217

Office of Naval Research

Naval Analysis Programs (Code 431)
800 North Quincy Street
Arlington, VA 22217

Office of Naval Research

Operations Research Programs (Code 434)
800 North Quincy Street
Arlington, VA 22217

Office of Naval Research (Code 436)

Attention: Dr. Bruce McDonald
800 North Quincy Street
Arlington, VA 22217

Office of Naval Research

Information Systems Program (Code 437)
800 North Quincy Street
Arlington, VA 22217

Office of Naval Research (ONR)

International Programs (Code 1021P)
800 North Quincy Street
Arlington, VA 22217

Director, ONR Branch Office

Attention: Dr. Charles Davis
536 South Clark Street
Chicago, IL 60605

Director, ONR Branch Office

Attention: Dr. J. Lester
495 Summer Street
Boston, MA 02210

Director, ONR Branch Office

Attention: Dr. E. Glove and Mr. R. Lawson
1030 East Green Street
Pasadena, CA 91106
(two copies)

Dr. M. Bertin

Office of Naval Research
Scientific Liaison Group
American Embassy - Room A-407
APO San Francisco 96503

Director, Naval Research Laboratory

Technical Information Division (Code 2627)
Washington, DC 20375
(six copies)

Director, Naval Research Laboratory (Code 2029)

Washington, DC 20375
(six copies)

Scientific Advisor
Office of the Deputy Chief of Staff
for Research, Development and Studies
Headquarters, U.S. Marine Corps
Arlington Annex, Columbia Pike
Arlington, VA 20380

Headquarters, Naval Material Command
(Code 0331)
Attention: Dr. Heber G. Moore
Washington, DC 20360

Headquarters, Naval Material Command
(Code 0344)
Attention: Mr. Arnold Rubinstein
Washington, DC 20360

Naval Medical Research and Development
Command (Code 44)
Naval Medical Center
Attention: CDR Paul Nelson
Bethesda, MD 20014

Head, Human Factors Division
Naval Electronics Laboratory Center
Attention: Mr. Richard Coburn
San Diego, CA 92152

Dean of Research Administration
Naval Postgraduate School
Monterey, CA 93940

Naval Personnel Research and Development
Center
Management Support Department (Code 210)
San Diego, CA 92152

Naval Personnel Research and Development
Center (Code 305)
Attention: Dr. Charles Gettys
San Diego, CA 92152

Dr. Fred Muckler
Manned Systems Design, Code 311
Navy Personnel Research and Development
Center
San Diego, CA 92152

Human Factors Department (Code N215)
Naval Training Equipment Center
Orlando, FL 32813

Training Analysis and Evaluation Group
Naval Training Equipment Center
(Code N-00T)
Attention: Dr. Alfred F. Smode
Orlando, FL 32813

Department of the Army

Technical Director, U.S. Army Institute for the
Behavioral and Social Sciences
Attention: Dr. J.E. Uhlaner
1300 Wilson Boulevard
Arlington, VA 22209

Director, Individual Training and Performance
Research Laboratory
U.S. Army Institute for the Behavioral and
and Social Sciences
1300 Wilson Boulevard
Arlington, VA 22209

Director, Organization and Systems Research
Laboratory
U.S. Army Institute for the Behavioral and
Social Sciences
1300 Wilson Boulevard
Arlington, VA 22209

Department of the Air Force

Air Force Office of Scientific Research
Life Sciences Directorate
Building 410, Bolling AFB
Washington, DC 20332

Robert G. Gough, Major, USAF
Associate Professor
Department of Economics, Geography and
Management
USAF Academy, CO 80840

Chief, Systems Effectiveness Branch
Human Engineering Division
Attention: Dr. Donald A. Topmiller
Wright-Patterson AFB, OH 45433

Aerospace Medical Division (Code RDH)
Attention: Lt. Col. John Courtright
Brooks AFB, TX 78235

Other Institutions

The Johns Hopkins University
Department of Psychology
Attention: Dr. Alphonse Chapanis
Charles and 34th Streets
Baltimore, MD 21218

Institute for Defense Analyses
Attention: Dr. Jesse Orlansky
400 Army Navy Drive
Arlington, VA 22202

Director, Social Science Research Institute
University of Southern California
Attention: Dr. Ward Edwards
Los Angeles, CA 90007

Perceptronics, Incorporated
Attention: Dr. Amos Freedy
6271 Variel Avenue
Woodland Hills, CA 91364

Director, Human Factors Wing
Defense and Civil Institute of
Environmental Medicine
P.O. Box 2000
Downsville, Toronto
Ontario, Canada

Stanford University
Attention: Dr. R.A. Howard
Stanford, CA 94305

Montgomery College
Department of Psychology
Attention: Dr. Victor Fields
Rockville, MD 20850

General Research Corporation
Attention: Mr. George Pugh
7655 Old Springhouse Road
McLean, VA 22101

Oceanautics, Incorporated
Attention: Dr. W.S. Vaughan
3308 Dodge Park Road
Landover, MD 20785

Director, Applied Psychology Unit
Medical Research Council
Attention: Dr. A.D. Baddeley
15 Chaucer Road
Cambridge, CB 2EF
England

Department of Psychology
Catholic University
Attention: Dr. Bruce M. Ross
Washington, DC 20017

Stanford Research Institute
Decision Analysis Group
Attention: Dr. Allan C. Miller III
Menlo Park, CA 94025

Human Factors Research, Incorporated
Santa Barbara Research Park
Attention: Dr. Robert R. Mackie
6780 Cortona Drive
Goleta, CA 93017

University of Washington
Department of Psychology
Attention: Dr. Lee Roy Beach
Seattle, WA 98195

Eclectech Associates, Incorporated
Post Office Box 179
Attention: Mr. Alan J. Pesch
North Stonington, CT 06359

Hebrew University
Department of Psychology
Attention: Dr. Amos Tversky
Jerusalem, Israel

Dr. T. Owen Jacobs
Post Office Box 3122
Ft. Leavenworth, KS 66027

Unclassified

SECURITY CLASSIFICATION OF THIS PAGE (When Data Entered)

REPORT DOCUMENTATION PAGE		READ INSTRUCTIONS BEFORE COMPLETING FORM
1. REPORT NUMBER	2. GOVT ACCESSION NO.	3. RECIPIENT'S CATALOG NUMBER
9 Technical Rept.		Jul 75 - Sep 76
4. TITLE (and Subtitle)		5. TYPE OF REPORT & PERIOD COVERED
Unit Versus Differential Weighting Schemes for Decision Making: A Method of Study and Some Preliminary Results		Technical 7/75 - 9/76
6. AUTHOR		6. PERFORMING ORG. REPORT NUMBER
J. Robert Newman, David Seaver Ward Edwards		SSRI 76-5
7. PERFORMING ORGANIZATION NAME AND ADDRESS		8. CONTRACT OR GRANT NUMBER(s)
Social Science Research Institute University of Southern California Los Angeles, California 90007		Prime Contract # N00014-76-C0074 Subcontract # 75-030-0711
9. CONTROLLING OFFICE NAME AND ADDRESS		10. PROGRAM ELEMENT, PROJECT, TASK AREA & WORK UNIT NUMBERS
Advanced Research Projects Agency 1400 Wilson Blvd. Arlington, Virginia 22201		
11. MONITORING AGENCY NAME & ADDRESS (if different from Controlling Office)		12. REPORT DATE
Office of Naval Research 800 North Quincy Street Arlington, VA 22217		July 1976
13. DISTRIBUTION STATEMENT (of this Report)		13. NUMBER OF PAGES
Approval for public release, distribution unlimited		39
15. DISTRIBUTION STATEMENT (of the abstract entered in Block 20, if different from Report)		15. SECURITY CLASS. (of this report)
N00014-76-C-0074		UNCLASSIFIED
16. SUPPLEMENTARY NOTES		15a. DECLASSIFICATION/DOWNGRADING SCHEDULE
Support for this research performed by Social Science Research Institute was provided by the Advanced Research Projects Agency of the Department of Defense and was monitored under Contract N00014-76-C-0074 with the Office of Naval Research, under subcontract from Decisions and Designs, Inc.		
17. KEY WORDS (Continue on reverse side if necessary and identify by block number)		
decision analysis multivariate prediction computer simulation multi-attribute utility analysis		data aggregation and analysis weighting schemes
18. ABSTRACT (Continue on reverse side if necessary and identify by block number)		
A method for generating realistic data is described, and illustrations of how the method can be used to study the efficacy of different data analysis models are given. The method is a Monte Carlo computer simulation of a multivariate process and generates a N by M data matrix where N is the number of observations and M is the number of variables or measurements. The computer program to accomplish this is outlined. Two examples of the use of the method are given. One compares the familiar multiple regression		

DD FORM 1 JAN 73 1473

EDITION OF 1 NOV 65 IS OBSOLETE
S/N 0102-LF 014-6601

Unclassified

SECURITY CLASSIFICATION OF THIS PAGE (When Data Entered)

390664

Unclassified

SECURITY CLASSIFICATION OF THIS PAGE(When Data Entered)

model with simple unit weighting in a well defined prediction problem. The results indicate that multiple regression is superior to unit weighting for prediction purposes, but the differences between the two models are not great. The second example compares several ways of forming weighted and unweighted composites in a multi-attribute decision making context. Some of the conditions in which differential weighting is important in such contexts are specified.



Unclassified

SECURITY CLASSIFICATION OF THIS PAGE(When Data Entered)