

AD735129

Reproduced by
**NATIONAL TECHNICAL
INFORMATION SERVICE**
Springfield, Va. 22151

Security Classification

DOCUMENT CONTROL DATA - R & D

Security classification of title, body of abstract and indexing annotation must be entered when the overall report is classified

1. ORIGINATING ACTIVITY (Corporate author)		2a. REPORT SECURITY CLASSIFICATION	
University of Georgia		Unclassified	
		2b. GROUP	
		Unclassified	
3. REPORT TITLE			
"Configuration and Classification of Clusters in n-Dimensions"			
4. DESCRIPTIVE NOTES (Type of report and, inclusive dates)			
5. AUTHOR(S) (First name, middle initial, last name)			
Surendra J. Trivedi			
6. REPORT DATE		7a. TOTAL NO. OF PAGES	7b. NO. OF REFS
December 1971		156	38
8a. CONTRACT OR GRANT NO.		9a. ORIGINATOR'S REPORT NUMBER(S)	
N00014-69-A-0423		Technical Report No. 75	
b. PROJECT NO.		Dept. of Statistics and Computer Science	
NR 042-261			
c.		9b. OTHER REPORT NO(S) (Any other numbers that may be assigned this report)	
d.		No. 17 THEMIS	
10. DISTRIBUTION STATEMENT			
Reproduction in whole or part is permitted for any purpose of the United States Government.			
11. SUPPLEMENTARY NOTES		12. SPONSORING MILITARY ACTIVITY	
		ONR - Washington, D.C.	
13. ABSTRACT			
Many experiments involve measuring a number of response variables simultaneously. As a result of giving one stimulus to an experimental unit, what we obtain is not just one response but several responses. In statistical language, we deal with a multivariate situation as opposed to univariate situations. Usually, many stimuli, called factors, are considered at many levels in the same experiment. Many statistical techniques are available to analyse this type of data and to draw conclusions therefrom. The present work considers one such technique, the identification of subgroups of individuals on the basis of responses, i.e., a special case of cluster analysis. Many different algorithms proposed for detecting clusters have been reviewed. These fall into two classes--(i) those which detect clusters of variables--factor analysis--and (ii) those which detect clusters of experimental units--cluster analysis. What we have done in the present work lies in-between these two techniques. We first perform cluster analysis on p responses for each unit, and sort the experimental units into groups. After assigning experimental units to groups, we look for those individuals whose total response could be expressed in (p-1) combinations of original responses, irrespective of the groups or clusters to which they belong. Those individuals whose total response could be described in terms of (p-1) combinations of original responses are said to lie on a "simple structure plane." A geometric probability approach has been used to determine whether a given point lies on a simple structure plane. The orthogonal distance of a point from the plane is used as a criterion. If many points lie on a given simple structure plane, the plane is said to be overdetermined. Probability expression has been derived to determine whether a simple structure plane can be regarded as being overdetermined. Those individuals which lie on a (p-1)-dimensional hyperplane in a p-dimensional space, need one variable less in their description.			

LD FORM 1473 (PAGE 1)
1 NOV 65
S/N 0101-807-6811

Unclassified

Security Classification

Security Classification

DOCUMENT CONTROL DATA - R & D

Security classification of title, body of abstract and indexing annotation must be entered when the overall report is classified

1. ORIGINATING ACTIVITY (Corporate author)		20. REPORT SECURITY CLASSIFICATION	
		20. GROUP	
3. REPORT TITLE			
4. DESCRIPTIVE NOTES (Type of report and, inclusive dates)			
5. AUTHOR(S) (First name, middle initial, last name)			
6. REPORT DATE		7a. TOTAL NO. OF PAGES	7b. NO. OF REFS
8a. CONTRACT OR GRANT NO.		9a. ORIGINATOR'S REPORT NUMBER(S)	
b. PROJECT NO.			
c.		9b. OTHER REPORT NO(S) (Any other numbers that may be assigned this report)	
d.			
10. DISTRIBUTION STATEMENT			
11. SUPPLEMENTARY NOTES		12. SPONSORING MILITARY ACTIVITY	
13. ABSTRACT			
Page 2. In this case, (p-1) linear combinations of the original p variables, could be used for describing these points. For real life interpretation, this procedure would be useless as we would be left with (p-1) artificial variables instead of p observable variables. It would seem reasonable, then, to eliminate, for the data points close to one of the subspaces, that observable variable which contributes the least to the description of this selection of data points. A correlation approach has been used to determine the variable to be eliminated, and it turns out that the variable which correlates most strongly in absolute value with the artificial variable should be discarded. Two computer programs have been developed to assist the user in analysing his data using this approach. The first one, named CLUSTR, identifies clusters and simple structure planes. The second one, an interactive graphics program, named ELLIPSE, can be used to visualize the configuration of clusters and the underlying simple structures. The clusters are projected onto many 2-dimensional spaces and displayed on the IBM 2250 Graphics terminal. If the plane of projection selected is orthogonal to a simple structure plane, the points lying on the particular simple structure plane, will make a band of narrow width more or less resembling a straight line. For other planes of projection this will not hold. The clustering of points is, however, unaffected by the choice of a particular plane of projection.			

Unclassified

Security Classification

KEY WORDS	LINK A		LINK B		LINK C	
	ROLE	WT	ROLE	WT	ROLE	WT
CLASSIFICATION CLUSTER ANALYSIS FACTOR ANALYSIS MULTIVARIATE ANALYSIS STATISTICAL COMPUTATION CONVERSATIONAL COMPUTATION GRAPHICS TERMINALS						

DD FORM 1473 (BACK)

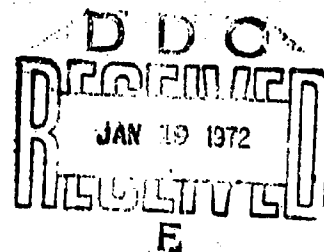
Unclassified

THEMIS REPORT NUMBER 17
TECHNICAL REPORT NUMBER 75
CONFIGURATION AND CLASSIFICATION OF CLUSTERS IN
n-DIMENSIONS

SURENDRA J. TRIVEDI

Reproduction in whole or in part is permitted for any purpose of the United States Government. This Research was supported, in part, by the Office of Naval Research, Contract Number N00014-69-A-0423, NR 042-261.

ROLF BARGMANN
PRINCIPAL INVESTIGATOR



THE UNIVERSITY OF GEORGIA
DEPARTMENT OF STATISTICS AND COMPUTER SCIENCE
ATHENS, GEORGIA 30601

DECEMBER 1971

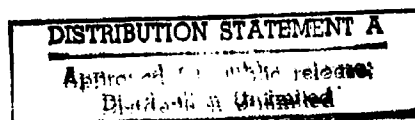


TABLE OF CONTENTS

CHAPTER		Page
I.	INTRODUCTION	1
	1.1 Multivariate Experiments	1
	1.2 Review of the Related Literature	2
	1.3 Real and Virtual Clusters	4
	1.4 d-clusters and k-clusters	5
	1.5 Identification of Mixed Samples	6
	1.6 Definition of the Problem	9
II.	PRINCIPAL COMPONENT AND OTHER PROJECTIONS	14
	2.1 Introduction	14
	2.2 Principal Component Projections	14
	2.3 Principal Component Displays	17
	2.4 Other Projections	18
III.	CLUSTERS IN SUBSPACES--THE SIMPLE STRUCTURE SUBSPACE .	19
	3.1 The Problem	19
	3.2 Mathematical Formulation	20
	3.3 A Basic Probability	21
	3.4 Determination of Simple Structure Configuration	26
IV.	CLUSTERS IN SUBSPACES--IDENTIFICATION	30
	4.1 Introduction	30
	4.2 Overdetermined Subspaces	30
	4.3 Identification of Subspaces	32
	4.4 Elimination of Variables	38

TABLE OF CONTENTS (continued)

CHAPTER		Page
V.	DESCRIPTION OF COMPUTER PROGRAMS	43
	5.1 The Computer Programs	43
	5.2 An Algorithm for Identification of Points Lying on a Subspace	43
	5.3 Loading the Data Set (Utility Program (IEBGENER)	46
	5.4 The Display Program and the Conversational System	46
VI.	USER'S GUIDE	68
	6.1 Introduction	68
	6.2 The CLUSTER Program	68
	6.3 IEBGENER Utility Routine	70
	6.4 The ELLIPSE Program	71
	6.5 An Illustration	78
	BIBLIOGRAPHY	114
	APPENDIX	117
	A	117
	B	140

CHAPTER I

INTRODUCTION

1.1 Multivariate Experiments

Many experiments involve measuring a number of response variables simultaneously. Thus, for example, in determining the effectiveness of a new drug, a person's systolic and diastolic blood pressure may be observed, before and after administration of the drug; also his pulse rate, temperature and other physiological data may be recorded. As a result of giving one stimulus to an experimental unit, what we obtain is not just one response but several responses. In statistical language, we deal with a multivariate situation (many responses) as opposed to univariate situations (only one response). An experiment is rarely so simple as described above. Usually, many stimuli, which will be called factors, are considered at many levels in the same experiment. Ultimately, the experimenter has a large collection of data before him. The problem is how to interpret the data. Depending upon the experimenter's objective, many statistical techniques are available to analyse the data and to draw conclusions therefrom. In the present work, we consider one such technique, the identification of subgroups of individuals on the basis of responses, i.e., a special case of cluster analysis. In the following sections, we review some of the problems of the subject and then outline the special problem in the present dissertation.

1.2 Review of the related literature

Consider an experiment where many responses are recorded on each experimental unit. We shall assume that there are n experimental units, and, on each unit, p responses are measured. The resulting data can be put in the form of an $n \times p$ matrix. Each row of this matrix corresponds to one experimental unit. Each experimental unit can be represented by a point in a p -dimensional space. It is possible that some of these points will be so close to one another that they form a "cluster." The problem of detecting clusters has been considered by many authors during the last 30 years. Tryon [37] in 1939 gave many algorithms based on the correlation matrix of variables, for the related problem of assigning variables to groups. The technique, much similar to the concept of the "coefficient of belonging" described by Harman [14], was developed on the assumption that correlations among variables belonging to the same group should be much higher than correlations between these variables and those not belonging to the group. Holzinger and Harman [15] defined their coefficient of belonging or B-coefficient as "100 times the ratio of the average of the intercorrelations among the variables of a group to their average correlation with all the remaining variables."

Sokal and Michener proposed the weighted mean-pair method to identify clusters. This method was originally applied to an entomological problem [33]. A more recent description of this method has been given by Sokal and Sneath [34] who recommend it as "the best of a class of commonly used methods of cluster analysis." The method operates on an $N \times N$ similarity matrix. Those two individuals, i and j say, which have the highest similarity are paired together, i.e., put in the same cluster.

Any appropriate measure of similarity e.g., product moment correlation coefficient, coefficient of association, etc., can be used, although by far the most commonly used measure is the product moment correlation coefficient. After the individuals i and j have been paired together, columns (and rows) i and j of the similarity matrix are replaced by a single column consisting of the means of the elements in rows (and columns) i and j . The process is then repeated on the new matrix of order $(N-1)$, when either two new individuals have the highest similarity and form a new pair, or the existing pair combines with a further individual to make a cluster of three. The process continues and at each stage a new cluster consisting of a pair of individuals is formed or the new individual is assigned to one of the already existing clusters of previously combined individuals.

Edwards and Cavalli-Sforza [7] suggest dividing the points into two sets such that the sum of squares of distances between sets is a maximum. Thus according to this method, one can find only two clusters, no more and no less. Since the total sum of squares is a constant for a given sample, maximizing between-groups sum of squares is equivalent to minimizing the within-groups sum of squares. The method consists of examining all the $2^{N-1} - 1$ two-set partitions of N individuals and selecting the one which gives the minimum within-set sum of squares. The method is not suitable for a large value of N as the time required on a computer to examine all the two-set partitions is enormous. It was estimated that with $N = 21$, the time required to examine all the partitions, on a computer with 5 micro-second access time would be 100 hours and that for $N = 41$, it would be 54,000 years. Thus even with the help of the fastest computers, the method would be impracticable. One more algorithm

based on ecological applications was proposed by Williams and Lambert [38]. Gower [12] gives an excellent comparison of the last three mentioned algorithms together with some of his own modifications proposed for these algorithms.

Another important study was made by Neyman and Scott [27]. They extensively studied the clustering of galaxies in the universe and presented the theories of "simple clustering" and "multiple clustering." "Simple clustering" was based on the assumption that galaxies occur in clusters and that the cluster centers are uniformly distributed throughout the universe. In "multiple clustering" it was assumed that the cluster centers radiate from super clusters.

From the above discussion, it is clear that in the statistical literature, we come across two types of clusters--(i) the clusters of variables and (ii) the clusters of individuals or points or experimental units. Without going into the details of the confusion that these two concepts have created and their uses and misuses, we only note that given a multivariate sample, it is possible to identify the underlying clusters by applying any of the suitable techniques available.

We now describe in detail an algorithm proposed by Bargmann and Graney [5] to determine clusters with the object of identifying mixed samples of multivariate normal distributions. With the help of methods to be developed in the present work, we may then study the configuration of such subgroups.

1.3 Real and Virtual Clusters

Real clusters are defined to be clusters of points in the original space. Virtual clusters on the other hand are clusters of projections of

points in a space of dimension lower than that of the original space. To borrow an example from Graney [10], if one were to look for clusters of stars as observed from the earth, one would be dealing with virtual clustering. The observer perceives the stars as projections onto the surface of the celestial sphere. To determine the real clusters of stars, it would be necessary to measure the distance of each star from the observer. It is also obvious from the above example that, virtual clusters may not necessarily be real clusters, and vice versa. One may tend to think that real clusters would necessarily be virtual clusters also. This, however, is not so. If one were standing in the midst of a real cluster, he may not find any cluster at all. In this connection, it should also be noted that in the algorithm proposed by Graney [10], if a cluster is centered at the center of the entire system, it would not be detected. This is similar to identification of galaxies. Being a member of our own galaxy, we do not obtain, by direct observation, a description of the configuration of our galaxy. For that purpose, we will have to make calculations based upon distance measurements (or observe from a different galaxy).

1.4 d-clusters and k-clusters

As experimental units are assigned to clusters, a decision has to be made whether two units are close enough to justify their inclusion in the same cluster. One may look at this problem in two directions: The so-called d-clusters and k-clusters. Consider a region S of fixed radius d . This is said to be overdetermined at the α -level of significance if the number of points in the region exceeds a value k_0 such that, under the null hypothesis of uniform distribution of points,

$$P [k_0 \text{ or more points in } S] < \alpha$$

This type of cluster is known as d-cluster since it results from a region of fixed radius d.

The other type of clustering results when a fixed number of points fall into a region with sufficiently small radius. Let there be a fixed number of points, say, k. Let d_0 be the radius of a sphere (or hypersphere) just sufficient to enclose all these k points within the sphere. It is apparent that d, the radius of the sphere is a random variable. If

$$P [d < d_0] < \alpha$$

where α is the predetermined level of significance, then these k points are said to form an overdetermined cluster. Such a cluster is known as the k-cluster since it results from a fixed number of points falling within a sufficiently small sphere.

1.5 Identification of Mixed Samples

Consider the case of k-cluster discussed above. As an example of this type of cluster, consider an experiment in which the number of people given a drug from a certain class of drugs (which produce similar effects) is limited. Then one would like to know if the symptoms among some individuals are more closely alike than those of other individuals. In other words, it would be appropriate to see if there are some individuals whose symptoms are so close that they form a cluster. But this also implies that we are looking for a principle of classification which distinguishes this group of individuals. Such a situation can arise in linear analysis. For the sake of simplicity, we shall describe the situation in terms of univariate analysis, but it applies equally well to multivariate analysis. In a two-factor experiment one being applied at r levels and the other being

applied at s levels, we will have $r \times s$ cells; let us assume that there are n_{ij} observations in cell (i, j) . Analysis of these data on the basis of a linear statistical model assumes homogeneity of cell variances. The hypothesis of the equality of cell variances can be tested by application of Bartlett's test. If this hypothesis should be rejected, three possible causes can be responsible: (i) The cell variances are not constant, σ^2 , but proportional to some known v_{ij} (a different one for every cell). A variance-stabilization transformation, or weighted regression, or both, can be employed to correct this situation. (ii) There may be "mavericks" -- misclassified or improperly recorded items. They can be omitted before the analysis of the data. (iii) There may be a third factor of classification present. If this is the case, what we regard as "error sum of squares" is in fact not the error sum of squares but the variance component sum of squares error, plus sum of squares due to the third factor which we have not taken into account. In multivariate analysis of variance, this quantity would be the $(H + E)$ matrix instead of the E matrix alone where E stands for the "error" SSP matrix and H stands for the "hypothesis" SSP matrix.¹ Hence the problem reduces to that of detecting the third hidden factor. It is clear that the sample within each cell comes from two or more populations instead of from just one. It will be necessary to "unmix" the samples within each cell before a valid linear statistical analysis of the data can be performed. Instead of describing the technique of unmixing the mixed samples, we refer to Graney [10], for a full description of the technique, which also contains a number of illustrations.

¹ SSP sum of squares and products, sometimes also called "Wishart" matrix.

There has been, in recent years, a considerable resurgence of interest in cluster analysis. Ling [23] has discussed many techniques and he has tried to classify these. It is apparent that there is little purpose in inventing yet new similarity indices, distance measures, or search algorithms. The present dissertation deals with a point intermediate to the two problems which cluster analysis has attempted: (a) classification of individuals, on the basis of responses and (b) classification of response variables assuming a homogeneous group of individuals (really factor analysis). Cattell [6], and Stephenson [35] view the entire complex as "factor analysis" and call the first problem the "Q technique," the second problem the "R technique," and the combined problem, the "P technique." Unfortunately, their techniques do not lend themselves to a study of configurations because in their attempt to regard every problem as a correlational one, the authors perform analyses which become self-contradictory. For example, sums of squares and products can be used as terms in estimates of correlation between variables only if the individuals are independent, and conversely, "correlations" between individuals can be estimated by a product-moment approach only if the variables are uncorrelated. Quadratic forms would be needed otherwise, and the sums of squares and products can be quite meaningless. The present study avoids this confusion by separating, at each stage, the clustering problem (cluster of individuals) from the problem of structures of variables within clusters.

Guttman [11] has proposed a yet another technique to reduce the dimensionality of data. The technique, "nonmetric" in nature, works on an $n \times n$ symmetric matrix R , and determines those transformations which yield

an Euclidean coordinate system X ($X : n \times m$), such that $XX' = F$ for m a minimum and, furthermore, satisfy all inequalities that whenever $r_{ij} > r_{kl}$ then $f_{ij} > f_{kl}$ for the non-diagonal elements of R and F , ($i \neq j, k \neq l$). This avoids the problem of communalities and "when some lawful structure or pattern is present in the data, e.g., a simplex, a circumplex, or a radex, a nonmetric analysis will reveal the configuration, whereas a metric approach will obscure the lawfulness." For detailed discussion of this approach and the algorithms developed in this connection, refer to Guttman [11] and Lingoes and Guttman [24]. It is clear that this technique will reduce the dimensionality of all the data points. Again, as in factor analysis, it will not be affected by mean shifts. It is thus an alternative to structural and factor analyses and not an "in-between" solution as proposed by us in section 1.6.

1.6 Definition of the Problem:

In the above sections, we have given a brief outline of traditional approaches to cluster analysis. The discussion reveals that, given a multivariate sample, it is always possible to detect some underlying cluster structure and to assign points (experimental units) to the clusters to which they belong. This is not the only kind of analysis that can be performed. The same data could also be subjected to factor analysis, which assigns response variables to classes. Cluster analysis will group the individuals bringing out the mean effects and leaving the variable structure unaltered. Factor analysis will describe the data in terms of artificial variables without telling anything about the group means involved. What we propose to do in the present work lies in-between these two extremes. We first perform cluster analysis on p responses for

each unit, and sort the experimental units into groups. After assigning experimental units into groups, we look for those individuals whose total response could be expressed in $(p-1)$ combinations of original responses, irrespective of the groups or clusters to which they belong. Thus the ultimate purpose of our analysis is to elicit more information from a set of multivariate data than is possible by cluster or factor analysis alone. Frequently both of these extremes produce trivial results. In medical applications, cluster analysis tends to producing clusters of patients who are healthy, slightly ill, and very ill. By contrast, factor analysis identifies a collective set of symptoms such as "fever," "pain," and "chills."

The algorithm developed in the present work begins by identifying clusters in the sense of estimating parameters of mixed multivariate normal distributions with equal variance-covariance matrices. If this were not done the unit ellipsoids around the grand mean would be affected, uncontrollably, by mean shifts. Within each of these ellipsoids, we look for well overdetermined subspaces of lower dimensionality, in the sense that we want a very significant number of individuals to fall into a region close to these highly overdetermined subspaces. These will be called "Simple Structure Planes."

We use this term because of its similarity with one criterion which Thurstone [36] used in describing a "simple structure" in the common-factor space of factor analysis. If all principal components of the within-cluster matrix were calculated and "rotated," the results of our study could also be produced. This, however, would be a tremendously wasteful procedure.

Those individuals which have a large distance from the subspace, yet belong to the original cluster, would differ from the others in that they require at least one additional diagnostic.

Geometrically, the Euclidean distance on a metric given by the unit ellipsoid, is obtained as follows: A vector is passed from the center of the ellipsoid (O) to the point in question (A). This vector intercepts the unit ellipsoid at point (S). The Euclidean distance is then $(\text{length OA})/(\text{length OS})$. For this reason, the ellipsoid is called "unit ellipsoid." It generalizes the concept of a "unit interval" in one dimension (hence metric). The computer output reports these distances as a vector $\underline{y}$.

Now, the largest projected distance of a point, onto a 2-dimensional subspace, from the simple structure unit ellipsoid, is at most equal to the actual distance in p dimensions. Hence a point showing an appreciable distance from the simple structure ellipsoid, on any projection, would be representative of a point requiring one variable more, for adequate description, than the points lying in the simple structure region. This property is not related to the clustering of the points. Points in the same simple structure subspace may be far removed from each other; they may belong to different clusters.

The points having a large (positive or negative) distance from each of these planes can now be scrutinized. They share some characteristics, which makes them different from the other points on this plane. It is to be noted here that this procedure was not intended to find sub-clusters--one could do that by tightening the control constants in the

original program--but to study subspaces of lower dimensionality.² Thus the multivariate data are viewed from a new point of reference, and one may identify principles permitting a different taxonomy of experimental units (patients, plants, etc.).

There is a certain analogy of techniques between the subspace solution and the "simple structure rotation" in factor analysis. This is expected in virtual clusters. The cosine of the angular distance between two vectors can be regarded as a correlation if the vectors represent variables. It is in this sense that our n data points correspond to n correlated variables of factor analysis and our variables themselves correspond to the factors. With this understanding we can apply the rotational techniques to the original data matrix in order to obtain points lying on simple structure planes.

In the present dissertation, we have synthesized these two techniques--cluster analysis and simple structure identification--into a single program. When the user subjects his data to this program, named CLUSTER, he gets clusters and simple structure planes as the output. Next, we have developed an interactive graphics program, named ELLIPSE, which can be used to visualize the configuration of clusters and the underlying simple structures. Configuration of multidimensional clusters can be determined only if their dimensionality is reduced. For this purpose, it is necessary to project the clusters onto several 2 or 3 dimensional subspaces. The emphasis in the present work is on projecting the clusters

²Those readers who assume that this is analogous to factor analysis should be reminded that the latter increases the dimensionality from p correlated to $p + k$ (at least partially) uncorrelated variables. There is, of course, no relationship to an incomplete "component analysis" which produces singular solutions.

onto many 2-dimensional spaces and displaying them on the IBM 2250 Graphics terminal. If the plane of projection selected is orthogonal to a simple structure plane, the points lying on the particular simple structure plane will make a band of narrow width more or less resembling a straight line. The display program can also be used in an exploratory manner. The user can supply many different vectors. If, by using some vector of projection, he sees narrow bands as described above, the corresponding vector is orthogonal to a simple structure plane. Such visual determination of simple structure planes, however, is rather difficult especially if the user is required to make inferences on configurations in four or more dimensions, on the basis of 2-dimensional displays. It is implicit from the above discussion that the determination of simple structure is equivalent to the fact that the points lying on a simple structure plane need at least one variable less in their description. Thus if p variables have been measured on all the experimental units of the sample $(p-1)$ variables are adequate in the case of those experimental units which lie on a simple structure plane. In the case of these experimental units, one of the p variables can then be expressed by a linear combination of the remaining $(p-1)$ variables. The statistical interpretation of this phenomenon is considered in Chapter 3.

CHAPTER II

PRINCIPAL COMPONENT AND OTHER PROJECTIONS

2.1 Introduction

As stated in the previous chapter, one of the goals of the present work is to display many different projections of multidimensional clusters. In this chapter, we give an account of projections along eigenvectors or principal components of the metric ellipsoids onto 2-dimensional subspaces. The principal component projections are necessarily "orthogonal." It is worth noting at the outset that very little information was gained by these principal component projections. Not that we were surprised by this finding, but some social scientists seem to attribute a lot more to this particular mathematical reference frame than it deserves.

2.2 Principal Component Projections

Let $Z = (Z : n \times p)$ be a data matrix of the multivariate observations. The n data points can be represented in a p -dimensional space. It is assumed that the clusters formed by these n points have been identified as well as the points belonging to the clusters. It is further assumed that points, which cannot be assigned to one of these clusters, have been eliminated. Now, if the data come from a p -variate normal distribution (or a mixture of p -variate normal distributions with different mean vectors but the same variance-covariance matrix), the clusters would be elliptical in shape, and the ellipsoids would have the same orientation. If there are k clusters, and if we denote by $\mu_1, \mu_2, \dots, \mu_k$ the

cluster centers, the equations of the ellipsoids enclosing these clusters can be written as

$$(\underline{x} - \underline{\mu}_i)' \Sigma^{-1} (\underline{x} - \underline{\mu}_i) = \text{constant}$$

for $i = 1, 2, \dots, k$, where $\underline{x}$ stands for the running coordinates. The points belonging to a particular cluster will be enclosed within the respective ellipsoids. In practice, since the population (or populations) from which the sample was drawn, will not be exactly normal, we will not expect all the data points belonging to a particular cluster to lie within the corresponding ellipsoid. However, unless the population is far from normal, the ellipsoidal fit will be quite good. To visualize how the points are distributed in p -dimensional space and how they look in relation to their enclosing ellipsoids, we need to project the points onto several 2-dimensional spaces. The technique is simple and classic and is spelled out here for the sake of completeness. Let us take the equation of the i th unit ellipsoid¹

$$(\underline{x} - \underline{\mu}_i)' \Sigma^{-1} (\underline{x} - \underline{\mu}_i) = 1 \quad (2.2.1)$$

Since Σ is a positive definite (or at least positive semi-definite) matrix, there exists an orthogonal matrix Q such that

$$Q' \Sigma Q = D_\lambda \quad (2.2.2)$$

where D_λ is a diagonal matrix with diagonal elements equal to the characteristic roots of Σ , and Q is the matrix of the eigenvectors of Σ . (2.2.2)

¹This generalizes the "standard unit interval" in univariate analysis.

can be written as

$$\Sigma = Q D_{\lambda} Q' \quad (2.2.3)$$

and hence

$$\Sigma^{-1} = Q D_{1/\lambda} Q \quad (2.2.4)$$

where $D_{1/\lambda}$ is the inverse of D_{λ} . (The reciprocals of the characteristic roots of Σ are the diagonal elements of $D_{1/\lambda}$). Substituting (2.2.4) into (2.2.1), we have

$$(\underline{x} - \underline{\mu}_1)' Q D_{1/\lambda} Q' (\underline{x} - \underline{\mu}_1) = 1 \quad (2.2.5)$$

or, if we let $\underline{y}'_1 = (\underline{x} - \underline{\mu}_1)' Q$, this reduces to

$$\underline{y}'_1 D_{1/\lambda} \underline{y}_1 = 1 \quad (2.2.6)$$

(2.2.6) is the standard principal axes reduction of the conic (2.2.1). The transformation $\underline{y}'_1 = (\underline{x} - \underline{\mu}_1)' Q$ is the orthogonal rotation of the original reference axes in the direction of the principal axes of the ellipsoids. The direction cosines of the principal axes of all the k ellipsoids are identical because of the assumption of equal metric (homogeneity of dispersion matrices). If we write the matrix Q as

$$Q = [\underline{q}_1, \underline{q}_2, \dots, \underline{q}_p]$$

where $\underline{q}_1, \underline{q}_2, \dots, \underline{q}_p$ are the eigenvectors of Σ , the transformation

$\underline{y}'_1 = (\underline{x} - \underline{\mu}_1)' Q$ can be written as

$$\underline{y}'_1 = (\underline{x} - \underline{\mu}_1)' [\underline{q}_1, \underline{q}_2, \dots, \underline{q}_p].$$

Since Q is orthogonal, $\underline{z}'_1 Q$ will be coordinates of an original data point $\underline{z}'_1$, which is the i th row of the data matrix Z , with reference to the new coordinate axes. In particular $\underline{z}'_1 \underline{q}_1, \underline{z}'_1 \underline{q}_2$ will represent the

orthogonal projection of the original data point $\underline{z}_1$ onto the 2-dimensional plane determined by the eigenvectors $\underline{q}_1$ and $\underline{q}_2$. Let us assume that the eigenvectors are arranged in descending order, i.e., $\underline{q}_1$ is the eigenvector corresponding to the largest root, $\underline{q}_2$ that corresponding to the second largest root, etc. Then the 2-dimensional plane determined by $\underline{q}_1$ and $\underline{q}_2$ is the plane containing the two largest principal axes of the ellipsoids and we would be projecting the data points onto this plane. We can take all the $\binom{p}{2} = p(p-1)/2$ pairs of eigenvectors and project the data points on these planes.

2.3 Principal Component Displays

The IBM 2250 Graphics terminal was used to display projections on the $p(p-1)/2$ 2-dimensional planes formed by each pair of the eigenvectors. The data used by Graney [10] were taken, and the three clusters identified by him were projected on all the three 2-dimensional pairs formed by the 3 variables. The unit ellipses (i.e., projections of the unit ellipsoid) around the clusters were also displayed. The orientation of these ellipses, is, in the case of principal components, of course parallel to the axes of reference. If, however, one of the axes is not a principal component, the inclination of the principal axes of the ellipses with the reference axes can be displayed. This inclination may present some evidence regarding the nature of the data points, which was obscured in the principal component plot. Because of the disappointing lack of information contained in the principal component plots, we did not even bother to document the computer program. Instead, the program which has been documented permits projection around arbitrary pairs of reference axes. These computer programs are described in detail in Chapter V.

2.4 Other Projections

Let the equation of the n-dimensional unit ellipsoid be

$$(\underline{x} - \underline{\mu}_i)' \Sigma^{-1} (\underline{x} - \underline{\mu}_i) = 1$$

where Σ is the $p \times p$ variance-covariance matrix, $\underline{x}$ are the running coordinates and $\underline{\mu}_i$ is the cluster center of the i th cluster. There will be as many ellipsoids as the number of clusters identified. For the sake of notational convenience, we will drop the subscript i from $\underline{\mu}_i$. As all the ellipsoids are referred to the same metric Σ^{-1} , they differ only with respect to their location. In the discussion that follows, we are concerned with the shape rather than the location. The matrix Σ , as a population parameter, is unknown and we will replace it by its unbiased estimate, s , the matrix of mean squares and mean products within groups. The projection on the 2-dimensional plane formed by the variables i and j can be written as

$$s^{ii}(x_i - \mu_i)^2 + s^{jj}(x_j - \mu_j)^2 + 2s^{ij}(x_i - \mu_i)(x_j - \mu_j) = 1$$

where s^{ij} is the (i, j) th element of s^{-1} . These equations for various values of i and j are employed in displaying the projections. For the purposes of computer programming, this equation is further simplified into

$$s^{ii}y_i^2 + s^{jj}y_j^2 + 2s^{ij}y_iy_j = 1$$

where $x_i - \mu_i = y_i$ and $x_j - \mu_j = y_j$. By employing the standard reduction techniques, the coordinates to plot the ellipses can be easily calculated. Further aspects on this phase of the computer program are treated in Chapter V.

CHAPTER III

CLUSTERS IN SUBSPACES--THE SIMPLE STRUCTURE SUBSPACE

3.1 The Problem

The previous two chapters considered the problem of cluster identification and cluster configuration. We now come to the second part of our inquiry--identification of experimental units or points which could be described by measuring a fewer number of variables on them. Before we proceed further with the identification of the points lying on a subspace, we want to consider the situations where such a problem can arise. A familiar example would be one of medical diagnosis. A number of patients are measured on a number of medical symptoms. Sometimes these measurements are repeated on the same patients for a number of days consecutively. Here the symptoms measured are our variables and the patients are subjects or experimental units. If the measurements are taken on different days, it would add a factor of classification; let us assume that this factor (i.e., days) has not been recorded. Of course, from these data one can construct a variance-covariance matrix and from that obtain a correlation matrix. This could be subjected to factor analysis. Factor analysis would reveal groupings of the symptoms in these data. Application of Thurstone's [36] principle of simple structure could reveal such groupings. Disappointing examples of this kind of analysis resulting in the symptoms like "chill," "fever," and "pain" are not so uncommon.

The data could also be subjected to cluster analysis to form groups of patients. In this type of analysis, the patients would be classified into groups depending upon the absence or severity of symptoms they have in common with respect to other patients belonging to the same group. Thus, for the patients belonging to the same group, almost all the patients would be measuring equally on the average on different symptoms; they may either measure high on the same symptom compared to other patients belonging to different groups or measure low, etc. To summarize, factor analysis would tell us about the symptoms, and cluster analysis would help us to group patients according to the "degree of severity," say, of the symptoms. According to cluster analysis, we may have two patients belonging to different groups perhaps because one measured low on one symptom and the other measured high on the same symptom. It may happen that the particular symptom would have left the final diagnosis unchanged. To this extent, this symptom could be discarded so far as these two individuals are concerned and then they will belong to the same "group." But neither the factor analysis nor cluster analysis would bring out this fact; nor would it bring out the fact that "days" is a hidden factor. We must extend both techniques before we can identify such configurations. The problem is formulated in mathematical terms in the next section where we present, in detail, a specific approach.

3.2 Mathematical Formulation

In previous chapters we have dealt with the techniques of cluster analysis. It was pointed out that given an $n \times p$ data matrix, it is possible to identify the underlying clusters. We now ask the next

question--are there any data points which instead of lying in the p -dimensional space, lie on the $(p-1)$ -dimensional hyperplane? This question is important as an answer to it, among other things, will reveal the following things: (i) The points that lie on the $(p-1)$ -dimensional hyperplane. The determination of the points lying on the $(p-1)$ -dimensional hyperplane will help us to describe these points in terms of $(p-1)$ variables instead of p variables. (ii) We will essentially devise a "discriminant function" which helps us to split the original sample into two groups of points--one group which needs all the p original variables to describe the points belonging to it and another group which needs $(p-1)$ instead of p variables to describe the points belonging to it. In the second case, we have also to consider the question of which variable to discard from the original p variables. Note that it is also implied in the second case that the $p \times p$ variance-covariance matrix of the points lying on the $(p-1)$ -dimensional hyperplane will be singular, as its rank will be $(p-1)$ and not p .

3.3 A Basic Probability

At the outset, let us make clear what we mean by points lying "approximately" in a subspace. Our attempt will be to look for those points which lie close to a $(p-1)$ -dimensional hyperplane in a p -dimensional space. Clearly, any $(p-1)$ vectors always lie exactly on a $(p-1)$ -dimensional hyperplane, whereas a minimum of p vectors is necessary to overdetermine a $(p-1)$ -dimensional hyperplane. We will, therefore, look for those $(p-1)$ -dimensional hyperplanes which have p or more vectors lying close to them. The singular case where p or more vectors lie, exactly, on a plane will be ignored, since in this instance, those points lying on the hyperplane will

satisfy a linear relationship between the p-variables; this cannot happen unless there is a deterministic linear relationship between the p-variables in the population, and thus, any sample will reflect this relationship and the $p \times p$ estimated variance-covariance matrix of any sample from such a population will be singular. Since we require inverses of dispersion matrices, we must exclude these redundancies. To determine, whether a point overdetermines a hyperplane, we shall consider its Euclidean distance from the hyperplane and examine whether this distance could be considered negligible in a probabilistic sense. We need an expression for this probability. The derivation follows the reasoning given by Bargmann [2]. Let P be a point in 2-dimensional space. We are interested in the orthogonal distance of this point from a line, which is a hyperplane in 2-dimensional case. Without loss of generality, we can assume the X-axis to be this line (Figure 3.3.1). Let the vector OP subtend an angle θ at the origin with the X-axis. We must now assume that the reference axes span a Cartesian frame (orthogonal, equal units along each axis). This implies that a Gram-Schmidt transformation of the original observations (using S, the within estimated dispersion matrix as the original metric) must precede the calculation of probabilities of overdetermination. With this assumption we can now draw a circle with OP as radius. Let M be the foot of the perpendicular drawn from P on the X-axis. Then

$$PM = OP \cdot \sin \theta$$

$$\text{Therefore, } \theta = \sin^{-1} PM/OP = \sin^{-1} 2PM/2OP = \sin^{-1} a/2h$$

Where $a = PP'$ and h is the radius of the circle. Thus, the probability that a point falls within the angle θ on the circumference of the circle is given by

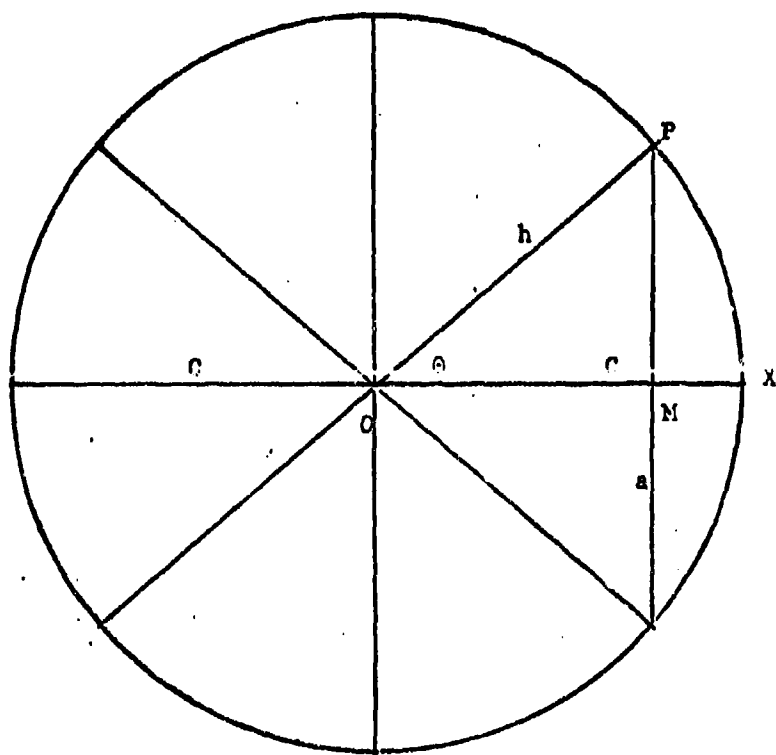


Figure 3.3.1

$$\begin{aligned}
 P_2 &= \frac{2 \times \text{length of arc generated by angle } 2\theta}{\text{circumference of the circle}} \\
 &= 2 \cdot OP \cdot 2\theta / 2\pi \cdot OP \\
 &= 2\theta / \pi \\
 &= (2/\pi) \arcsin(a/2h) \qquad (3.3.1)
 \end{aligned}$$

We can view the above probability either way--

- (i) the probability of a point falling within an angle θ as stated above
- (ii) the probability of a point falling within the orthogonal distance of $+a/2$ to $-a/2$.

The technique involved in generalizing the above probability to higher dimensions, say k , is simple. We have to obtain the surface element of a sphere of radius h in k dimensions and the surface element cut off by the k -dimensional arc which subtends an angle θ at the center of the sphere. We shall illustrate the procedure for 3-dimensions before generalizing it to k -dimensions. The 3-dimensional sphere can be obtained by revolving a semicircle around its diameter. The surface element generated by arc between $y = -a/2$ and $y = a/2$ is twice the element generated by arc between $y = 0$ and $y = a/2$ and hence the 3-dimensional surface element between $y = -a/2$ and $y = +a/2$ can be expressed as

$$C_3 = 2 \cdot 2\pi \int_0^{a/2} x ds$$

where $x = \sqrt{h^2 - y^2}$ and $ds = h dy / \sqrt{h^2 - y^2}$

$$\begin{aligned}
 \text{Hence } C_3 &= 2 \cdot 2\pi \int_0^{a/2} (h\sqrt{h^2 - y^2}) / (\sqrt{h^2 - y^2}) dy \\
 &= 2\pi ah
 \end{aligned}$$

and S_3 , the total surface element will be

$$2 \cdot 2\pi \int_0^h (h \cdot \sqrt{h^2 - y^2}) / (\sqrt{h^2 - y^2}) dy$$

$$= 4\pi h^2$$

Hence, $P_3 = C_3/S_3 = a/2h$

(3.3.2)

For k dimensions, the $(k-1)$ -dimensional semisphere is required to be revolved around a $(k-2)$ -dimensional hyperplane. In the above notation, the probability P_k can be expressed as

$$P_k = \frac{2 \cdot \frac{S_{k-1}}{h^{k-2}} \cdot h \cdot \int_0^{a/2} \frac{(\sqrt{h^2 - y^2})^{k-2}}{\sqrt{h^2 - y^2}} dy}{2 \cdot \frac{S_{k-1}}{h^{k-2}} \cdot h \cdot \int_0^h \frac{(\sqrt{h^2 - y^2})^{k-2}}{\sqrt{h^2 - y^2}} dy}$$

$$= \frac{\int_0^{a/2} (\sqrt{h^2 - y^2})^{k-3} dy}{\int_0^h (\sqrt{h^2 - y^2})^{k-3} dy}$$

In the numerator of the above expression, substitute $y = h \sin \phi$. Then, $dy = h \cos \phi d\phi$ and the integral reduces to

$$h^{k-2} \int_0^{\alpha/2} \cos^{k-2} \phi d\phi$$

where $\alpha/2 = \arcsin a/2h$. By the same substitution, the denominator could be reduced to

$$h^{k-2} \int_0^{\pi/2} \cos^{k-2} \phi d\phi$$

If we let $\sin^2 \phi = z$, this integral can be reduced to the complete Beta integral

$$\frac{1}{2} B\left(\frac{1}{2}, \frac{k-1}{2}\right) = \frac{1}{2} \int_0^1 z^{-1/2} (1-z)^{(k-3)/2} dz$$

and the denominator could be reduced to

$$\frac{1}{2} \int_0^{\sin^2 \alpha/2} z^{-1/2} (1-z)^{(k-3)/2} dz$$

Therefore, P_k , the ratio can be expressed as

$$P_k = B(\sin^2 \alpha/2; 1/2, (k-1)/2) \quad (3.3.3)$$

where the right side of (3.3.3) is the incomplete Beta function. This is the probability of a point falling within $(-a/2, +a/2)$ in k dimensions.

We will now make use of this probability expression in studying the simple structure configuration.

3.4 Determination of Simple Structure Configuration

Let us assume that we have a total of n points in p -dimensional space. As stated in the previous section, the probability, in p -dimensional space, of a point falling within $(-a/2, a/2)$ is given by

$$P_p = B(\sin^2 \alpha/2; \frac{1}{2}, (p-1)/2) \quad (3.4.1)$$

where $\alpha/2 = \arcsin a/2h$. Without loss of generality, h could be taken to be

1. We are interested in determining whether these are points lying on a

subspace. In a p -dimensional space, $(p-1)$ points always lie on a $(p-1)$ -dimensional subspace. Thus out of a total of n points, we have available only $(n-p+1)$ free points. Likewise, if we are to make any statement about r points lying on a subspace, we can only look to $(r-p+1)$ points as $(p-1)$ of them will always lie on a subspace. If we denote by P_p the probability of a single point lying within a fixed distance $\pm a/2$ from a given $(p-1)$ -dimensional hypersphere, in p -space, the probability that $(r-p+1)$ points will fall in that region out of $(n-p+1)$ is given by

$$\binom{n-p+1}{r-p+1} P_p^{r-p+1} (1-P_p)^{n-r}$$

Thus, if a region is to contain more than r points, the probability would be given by the cumulative binomial distribution

$$\sum_{i=r}^n \binom{n-p+1}{i-p+1} P_p^{i-p+1} (1-P_p)^{n-i} \quad (3.4.2)$$

Also on the basis of the $(n-p+1)$ free points available, the expected number of points (in excess of $(p-1)$) falling in a $(p-1)$ -dimensional region is

$$(n-p+1)P_p \quad (3.4.3)$$

The sum of binomial terms (3.4.2) can be evaluated by the incomplete Beta function. The result can be summarized as

$$\begin{aligned} &P[r \text{ or more of } n \text{ points lie within } \pm a/2] \\ &= B(P_p; r-p+1, n-r+1) \end{aligned} \quad (3.4.4)$$

where P_p itself is an incomplete Beta function. If this probability is very small we shall say that the simple structure subspace determined by

the points lying on it is overdetermined. However, the probability of a subspace being overdetermined itself depends on the probability P_p . We, therefore, need to consider the interval width and the probabilities, so that the "probability of overdetermination" may be a meaningful concept. Bargmann [2], in studying the overdetermined subspaces in relation to factor analysis fixed the ratio a/h to be ± 0.10 . These criteria, which reflect common usage in factor analysis, did not suit our requirements.

The object of this test, as will be discussed in Chapter IV, is to suggest to the viewer of a graphics display, some vectors which are normal to overdetermined hyperplanes. If no such concentration were present, the expected number of points, would be, (according to (3.4.3) and the discussion on free points)

$$(n-p+1)P_p + (p-1) \quad (3.4.5)$$

After considerable experimentation, with n between 20 and 50, and p between 2 and 5, we found that taking P_p such that the expected number of points is $(p+5)$, i.e.,

$$P_p = 6/(n-p+1) \quad (3.4.6)$$

gave satisfactory results in the identification of overdetermined hyperplanes by the Beta test (3.4.4). A value of P_p smaller than this was too stringent so that a considerable number of points would have to lie on a subspace before it could be considered well determined. The suggested value of P_p was fairly moderate. We compensated for this value by tightening up the probability level for declaring a subspace to be overdetermined. We set this probability at 0.01.

The arbitrariness of choice of P_p and levels of significance may be disquieting to some readers. They may be reminded though, that we are dealing with a phase of data analysis which is exploratory. A sample from

a p -variate normal distribution, with a dispersion matrix of rank $(p-1)$, will always be on a $(p-1)$ -dimensional subspace, exactly. There is no test for a hypothesis in the population. Rather, we deal with an instance where, after some linear transformations of the original variables, one of them has negligible variance, after some points have been deleted from the sample. Consequently, a decision as to what is negligible, and what is "close to a plane" is really as arbitrary as declaring that, viewed from some point in the universe, most (but not all) of the stars of a galaxy lie close to a plane. In the final analysis, only the graphic display of certain projections will reveal such intuitive configurations.

In our examples, the subspaces were overdetermined at much smaller values than 0.01 which points to the fact that (3.4.6) was quite useful.

For the rest, the choice of critical values is as arbitrary as "0.05 level" (because we have five fingers?) and the 0.01 level of significance. As a guide for displaying configurations, our two levels of P_p and α were "useful."

CHAPTER IV

CLUSTERS IN SUBSPACES--IDENTIFICATION

4.1 Introduction

In the previous chapter, we addressed ourselves to the statistical aspects of well determined subspaces and derived a few pertinent geometric probability expressions which may guide us to find such spaces. No mention was made of techniques for finding such configurations. In the present chapter, we consider (a) the technique of determining points lying on a subspace; and (b) the problem of which variable could be discarded for the points which describe an overdetermined subspace.

4.2 Overdetermined Subspaces

The problem of determining subspaces in our case is rather similar to the determination of simple structures in factor analysis. The essential difference, however, lies in the fact that a factor analyst looks for simple structure among variables. From data points, he constructs a correlation matrix and gets a factor matrix by applying one of the many suitable techniques available for this purpose. If he so desires, after obtaining an initial solution of the factor matrix, he may obtain a "preferred" representation, e.g., Lawley's form [14], etc. For a factor analyst such a solution may not serve his purpose if he is interested in relating the artificial variables to observable ones. In that case he will have to resort to some other forms of representation, such as the simple structure technique. The number of artificial variables required

to explain an observable variable is known as the complexity of the observable variable. For the purpose of interpretation of artificial variables, it is desirable that complexities of observable variables be low. Both analytic and geometrical techniques are available to a factor analyst to express the final solution in a form suitable for interpretation.

We have a slightly different problem. First of all, we do not look for any artificial variables to represent the observable variables. Thus whereas a factor analyst works on a factor matrix, we work on the data matrix itself. The starting point for a factor analyst is the correlation matrix obtained from the data matrix. A somewhat similar standardization is employed in our case, except that we standardize the data matrix on the basis of cluster means and the "within" matrix and then normalize the points to unit length. After displaying these points, and the unit ellipses around the cluster means, on all 2-dimensional directions, we proceed to single out those points which could be described in terms of fewer variables. In this connection, it does not concern us how far apart these points are, as long as they lie on a subspace of lower dimensionality. Thus this technique has an advantage that it can identify points lying on a subspace even though they may be belonging to different populations, or clusters. This is precisely what we had in mind. The technique of cluster analysis assigns points to the populations to which they belong. Leaving this structure intact, our new technique determines points which lie on a subspace.

A discussion may be in order regarding the number of subspaces one can find. If we can determine one overdetermined plane, the chances are that there are many planes in the vicinity of a plane already found.

The reason is that, by "tilting" the plane already found, we can still retain many of the points belonging to the original subspace found, pick up a few new points and obtain another overdetermined plane. However, we can reduce the multiplicity of subspaces by ignoring these additional planes found which lie within a certain range of the original plane. This can be done by requiring a minimum angle between the normals to two distinct planes. What angle should be maintained between two planes before declaring them as distinct is a matter of choice. The "orthogonality" preferred by some factor analysts is, at least for our problem, quite useless.

4.3 Identification of Subspaces

In this section, the general identification technique will be described. Let there be n experimental units, each with p measurements. The data matrix of order $n \times p$ will be designated as X . This matrix is subjected to a cluster analysis program. If the data are normally distributed, we need to use virtual clusters; hence we used the program developed by Bargmann and Graney [5] for this purpose. After NG (computer program notation) such clusters have been found, we may define a matrix A , with elements a_{ij} , of order $n \times NG$ such that

$$\begin{aligned} a_{ij} &= 1 \quad \text{if unit } i \text{ belongs to cluster } j \\ &= 0 \quad \text{otherwise} \end{aligned} \quad (4.3.1)$$

Let D_k (of order $NG \times NG$) denote a diagonal matrix with elements k_j , $j = 1, 2, \dots, NG$ where k_j is the number of points assigned to cluster j . Subtraction of cluster centers or cluster means from each point produces a matrix Y , formally given by,

$$Y = X - AD_k^{-1}A'X \quad (4.3.2)$$

$$= (I - AD_k^{-1}A')X \quad (4.3.3)$$

Now we can obtain an estimate of the common dispersion matrix Σ , using the within-cluster sample dispersion matrix

$$S = (1/n_e)Y'Y \quad (4.3.4)$$

where $n_e = n - NG$. For the determination of subspaces, we must first transform our reference frame to a Cartesian metric (which is, of course, merely a computational device, and never actually displayed). As a convenient technique, we used the Gram-Schmidt ("Forward Doolittle") reduction,

$$S = TT' \quad (4.3.5)$$

where T is a lower triangular matrix with positive diagonal elements, hence unique. In terms of this Cartesian reference frame, the data matrix is now transformed into

$$Z = Y(T')^{-1} \quad (4.3.6)$$

To find the "simple structure" subspaces, we look for unit vectors $\underline{t}$, such that

$$Z\underline{t} = \underline{v} \quad (4.3.7)$$

and $\underline{v}$ has the property that as many elements as possible are close to zero in the following sense:

Let $\underline{z}_i'$ denote the i th row of Z and v_i denote the i th element of $\underline{v}$. Then it is clear that

$$\underline{z}_i' \underline{t} = v_i \quad (4.3.8)$$

The i th element in $\underline{v}$, namely v_i , is considered close to zero if

$$v_i / \sqrt{\underline{z}_i' \underline{z}_i} < \sin \alpha/2 \quad (4.3.9)$$

where $\sin \alpha/2$ is determined by (3.4.1).

Let us now consider the significance of the expressions (4.3.8) and (4.3.9). $\underline{t}$ is a unit vector and so is $\underline{z}_1'/\sqrt{\underline{z}_1'\underline{z}_1}$. Their "inner product," $\underline{z}_1'\underline{t}/\sqrt{\underline{z}_1'\underline{z}_1}$, therefore is the cosine of the angle between the vectors $\underline{t}$ and $\underline{z}_1'$. But

$$\underline{z}_1'\underline{t}/\sqrt{\underline{z}_1'\underline{z}_1} = v_1/\sqrt{\underline{z}_1'\underline{z}_1} \quad (4.3.10)$$

and hence $v_1/\sqrt{\underline{z}_1'\underline{z}_1}$ is the cosine of the angle between the vectors $\underline{t}$ and $\underline{z}_1'$. Since $\underline{t}$ is unit normal to the subspace (4.3.9) is established.

For a given number of experimental units, n , and p measurements on each of them, we can determine P_p using (3.4.6) and hence $\sin \alpha/2$ using (3.3.4). (4.3.9) is then a test to determine if the vector corresponding to a given data point (reduced in terms of z 's) is close to the subspace to which $\underline{t}$ is orthogonal. If the vector (and hence the data point) is close to the subspace, we regard the point as falling in the overdetermined region, and treat the corresponding element of $\underline{y}$ as "zero." We can examine each element of $\underline{y}$ in this manner and determine how many "zero" elements are there. It is the count of these "zero elements" which we subject to the Beta test (3.4.4). If this test is significant, we say that the subspace is overdetermined and report it as a solution, provided it is not "close" to any subspace already found.

The transformation vectors $\underline{t}$ are found in a manner analogous to Thurstone's "Analytical Method" combined with his earlier "Single Plane Method" [36]. According to this method, each row of the reduced data matrix Z is used as a point of departure to find a vector $\underline{t}$. Let us start with the i th row vector $\underline{z}_i'$. We shall assume that $z_{i1}, z_{i2}, \dots, z_{ip}$ are elements of $\underline{z}_i'$. Then $\underline{t}_0$, the initial approximation to $\underline{t}$ is obtained by normalizing $\underline{z}_i'$, i.e.,

$$\underline{t}_0 = \underline{z}_1 / \sqrt{\underline{z}_1' \underline{z}_1} = (t_{01}, t_{02}, \dots, t_{0p})$$

with this trial, the projections

$$\underline{v}_0 = \underline{Z} \underline{t}_0$$

are calculated and the elements v_{0i} are tested for closeness to zero.

If a value is very close, a large weight is assigned to the point, for the subsequent weighted regression technique. If it is large, the weight may be zero. Following Thurstone, we use discrete (step) weights, as follows:

$$\text{If } 0 < v_{0i} / \sqrt{\underline{z}_1' \underline{z}_1} < \sin \alpha/2$$

$$w_i = \text{BND} \quad (\text{BND, bound, initially 8}).$$

$$\text{If } \sin \alpha/2 < v_{0i} / \sqrt{\underline{z}_1' \underline{z}_1} < 2 \sin \alpha/2$$

$$w_i = \text{BND} - 1$$

$$\text{If } v_{0i} / \sqrt{\underline{z}_1' \underline{z}_1} > \text{BND} \times \sin \alpha/2$$

$$w_i = 0.$$

With these weights we can follow a weighted regression scheme to obtain an improved vector $\underline{t}_1$. This can be further simplified by a relation (due to Thurstone) which expresses

$$t_{1j}^* = \frac{t_{0j} - u_j}{1 - t_{0j} u_j} \quad (4.3.11)$$

$$\text{where } u_j = \frac{\sum_{i=1}^n z_{ij} v_{0i} / w_i}{\sum_{i=1}^n z_{ij}^2 / w_i} \dots \quad (4.3.12)$$

$$t_{1j} = t_{1j}^* / \sqrt{t_{11}^{*2} + t_{12}^{*2}} \quad \text{are then}$$

elements of the new trial vector $\underline{t}_1$. We calculate

$$\underline{v}_1 = \underline{Z} \underline{t}_1$$

and once again the weights are assigned following the scheme detailed above.

However, before assigning these new weights the BND value is reduced to 3, in the next cycle to 2 and then kept at 1 for the remaining iterations. Beginning with $\underline{t}_0$, seven iterations are performed which give rise to $\underline{t}_1, \underline{t}_2, \dots, \underline{t}_7$. The last one, $\underline{t}_7$, is retained as $\underline{t}$ and $Z\underline{t} = \underline{v}$ is formed. If the number of "zero" entries in this $\underline{v}$ is significantly large as explained earlier, we will have determined a well defined subspace, to which $\underline{t}$ is normal. The entire process is repeated, with each row taken as a trial value. A given row may or may not identify a subspace. If it leads to an overdetermined subspace, the solution, except for the first one, is checked as explained in section 4.2 to ensure that the subspace is "different" from any of the ones already found from previous rows. For this purpose, we require that the cosine of the angle between two planes should not exceed 0.7 meaning that the planes were apart by at least 45° . Once again we would like to mention that this is an arbitrary requirement; we found this useful in our experimental studies.

Now we must retransform to the original data frame. The vectors $\underline{t}$ that we obtain above, are with reference to the matrix Z . Our original objective was to remove mean shifts and then look for experimental units which lie on subspaces. However, in working with Z , we have removed the mean shifts and also reduced the metric to an orthogonal Cartesian frame. This was a matter of convenience. To restore the original metric, we must obtain the solution in terms of Y . This can be easily obtained as under. (4.3.6) is the transformation of Y into Z and (4.3.7) is the simple structure solution in terms of Z . Substituting (4.3.6) into (4.3.7) we obtain

$$Y(T')^{-1}\underline{t} = \underline{v} \quad (4.3.13)$$

from which we conclude that $(T')^{-1}\underline{t}$ is the transformation in terms of Y .

The computer program reports both $\underline{t}$ and $(T')^{-1}\underline{t}$. The vector $\underline{t}$ is reported under the heading "Vector which transforms original factor matrix into the above plane no. $v_{\underline{t}}$ " and $(T')^{-1}\underline{t}$ is reported under the heading "Transformation vector to transfer raw data to simple structure." The corresponding elements of $\underline{y}$ are also reported. We shall also refer to elements of $\underline{y}$ as the "loadings" or "scores" of original points with reference to the simple structure plane determined.

For each $\underline{t}$, only the vectors $(T')^{-1}\underline{t}$ are conveyed to the display program, since we plot the original data points, in terms of the observed variables. As well be explained in Chapter 5, the display program takes one of the observed variables as one axis and some specified vector as another axis. For a given choice of a variable and a vector, the vector is reduced in such a manner that it forms an orthogonal frame of reference with the chosen observed variable. A question may be raised as to why one of the axes always corresponds to an observed variable. In this connection, it should be pointed out that a vector specified by the user, is equivalent to an artificial variable, a linear combination of the observable ones. The elements of the vector given by the user serve as weights in forming this linear combination. When we view the displays with reference to an orthogonal frame of reference consisting of an observable variable against an artificial variable, we can make an inference as to how an observable variable compares with an artificial variable. If both reference axes were to correspond to artificial variables, the displays would be hard to interpret in a realistic sense. This is our main consideration in insisting that one of the axes correspond to an observable variable. Further, if we should permit a user to select any two vectors, these could be translated

into an orthogonal frame of reference in many different ways leading to utter confusion.

4.4 Elimination of Variables

One notable difference between the rotational problem in factor analysis, and the problem presented in this dissertation, is the fact that the former focuses attention on those points (variables, in that case) which are far removed from the subspaces. By contrast, the latter pays attention to only those points which are so close to a subspace that they define it. It is these data points only which require one variable less for adequate description. It should be noted, again, that such a subspace is not a single region within the p -dimensional space. Rather, for each simple structure plane, there are as many (parallel) $(p-1)$ -dimensional hyperplanes as there are clusters. Points which, in this sense, fall into the same subspace may be far removed from each other, inasmuch as they may be in different clusters. But even with their distinct neighbors, they share the property that the same $(p-1)$ variables are sufficient to explain their characteristics. It is for this reason that we expect entirely new principles of classification of data points, different from what could be expected by varying or refining cluster analysis or factor analysis techniques.

Mathematically, we could use for description, $(p-1)$ linear combinations of the original p variables, with the combinations chosen within the hyperplane orthogonal to the $(T^{-1})' \underline{t}$ vector. For real life interpretation, this procedure would be useless. Surely we are better off with p observable variables than with $(p-1)$ artificial ones. The principle of parsimony is not just a principle restricted to dimensionality. It would

seem reasonable, then, to eliminate, for the data points close to one of the subspaces, that observable variable which contributes the least to the description of this selection of data points.

As a measure of proximity of each variable to the artificial variable which defines the subspace, we propose to use the correlation between the original variables and the artificial variable. This can be obtained as follows. Recall that Y is obtained from the original data matrix after subtraction of the appropriate cluster means. Hence

$$\mathbf{1}'Y = \mathbf{0}' \quad (4.4.1)$$

where $\mathbf{1}'$ is a row vector consisting of all 1's and $\mathbf{0}'$ is a null row vector.

Also,

$$\begin{aligned} \mathbf{1}'\underline{y} &= \mathbf{1}'Z\underline{t} \\ &= \mathbf{1}'Y(T')^{-1}\underline{t} && \text{(By (4.3.6))} \\ &= \mathbf{0}' && (4.4.2) \end{aligned}$$

and

$$\begin{aligned} \underline{y}'\underline{y} &= \underline{t}'Z'Z\underline{t} \\ &= n_e \underline{t}'I_t \\ &= n_e && (4.4.3) \end{aligned}$$

since $\underline{t}$ is a unit vector. By virtue of the fact that $Z'Z = n_e I$, and because of (4.4.1) $(1/n_e)\underline{y}'\underline{y}$ will be an unbiased estimate of the covariances between the original variables and an artificial one on which the data points have scores which are the elements of $\underline{y}$. There will be as many $\underline{y}$ vectors as there are overdetermined subspaces (each corresponding to a different, but possibly overlapping, selection of data points). For each of these, we must determine its correlations with the original variables. Thus, if there are 4 variables and 3 different solutions (overdetermined subspaces), we will have a total of $4 \times 3 = 12$ correlations. Now, for a given $\underline{y}$ (a given subspace),

$$\begin{aligned}
Y'Y &= Y'Z\underline{t} \\
&= Y'Y(T')^{-1}\underline{t} && \text{(By (4.3.6))} \\
&= n_e T T' (T')^{-1} \underline{t} && \text{(By (4.3.4) and (4.3.5))} \\
&= n_e T \underline{t} && (4.4.4)
\end{aligned}$$

T is the same lower triangular matrix that was used in reducing the metric underlying the matrix Y to a Cartesian reference frame. From (4.4.4), we conclude that

$$(1/n_e)Y'Y = T\underline{t} \quad (4.4.5)$$

Hence the estimated covariances between the original variables and a single artificial variable, are elements of $T\underline{t}$. To obtain the correlations, we have to divide each of these elements by the square root of the estimate of the variance of the artificial variable and the square root of the estimate of the variance of the observable variable. Using (4.4.2) and (4.4.3), we conclude that the estimate of the variance of the artificial variable is

$$(1/n_e)\underline{v}'\underline{v} = 1 \quad (4.4.6)$$

since a sum of squares ($\underline{v}'\underline{v}$) must be divided by degrees of freedom to produce a variance estimate. Hence to obtain the correlations, we have to divide the elements of $T\underline{t}$ by the square root of the estimate of the variance of the observable variable only. In the computer program, after the vector $\underline{t}$ is obtained, the product $T\underline{t}$ is formed and the correlations are then calculated by division of each of the elements of $T\underline{t}$ by the square root of the estimate of the variance of the corresponding observable variable. The estimates of the variances of the observables are still available in the final step of the cluster analysis part of the program and these values are stored for use at this time. T is also stored.

Once the correlations between the original variables and their scores with reference to a subspace (i.e., vector $\underline{y}$) are available, it is easy to determine which variable should be discarded. In discriminant analysis, we come across the concept of correlations between original variables and the discriminant function. The discriminant function is nothing but a linear combination of the original variables which discriminates best between experimental units in a specified sense. The purpose for which we want to discriminate is important as the discriminant functions for different purposes are usually different. Given these things, the variable which correlates most strongly in absolute value with the discriminant function is the most important in discriminating. In our study the vector $\underline{t}$ which transforms the original observations into $\underline{y}$ plays a similar role in the sense that the elements of a well defined subspace represented by $\underline{y}$ has most elements near zero (see section 4.3) and a few far removed from zero. In analogy with discriminant analysis, $\underline{y}$ is the vector that produces the two groups of data points. Thus, the variable which correlates most strongly with this artificial variable, contributes most to the discrimination process. It is this maximally correlated observed variable which should be eliminated for, after it has been discarded the other variables contribute far less to the discrimination between these two sets of data points. If only one variable is to be sought which would explain the difference between these data points which fall into the subspace and those that do not, it would be this one which is closest (has highest correlation in absolute value) to the expendable artificial variable. Note the exact opposite of this technique to discriminant analysis, where we seek the best discriminator. Here we identify the "most expendable" artificial variable. By our correlational technique, we have identified,

for each subspace, that observable variable which is closest to the artificial variable not needed in the description of the selected data subset. Note again, the need for this correlational interpretation. If it were argued that the artificial variable itself ought to be discarded, we would be left with $(p-1)$ artificial variables, a rather unsatisfactory situation. After the correlational approach, we have $(p-1)$ observed variables left.

CHAPTER V

DESCRIPTIONS OF COMPUTER PROGRAMS.

5.1 The Computer Programs

The algorithms explained in previous chapters have been synthesized into two computer programs which are available at the University of Georgia. The first program, named CLUSTER, and described in the next section, identifies the clusters and overdetermined subspaces. The second one is available as a conversational system for the IBM 2250 Graphics unit. In this chapter, we describe these two computer programs. The following chapter will contain instructions regarding the use of these programs, and the interpretation of graphical displays.

5.2 An Algorithm for Identification of Points Living on a Subspace

In this program, beginning with the data matrix, we first identify the clusters. This is essentially the algorithm proposed by Bargmann and Graney [5]. However, the algorithm proposed by Bargmann and Graney stops at the identification process. Since we also need to identify points lying on an overdetermined subspace, we extend the algorithm further. A complete listing of the program is contained in Appendix A. This program can be logically divided into 2 parts. Up to statement number 811, it is essentially the reproduction of the program developed by Bargmann and Graney [5], where a complete documentation of this part can be found. We have made a small change in the program to suit our needs. As explained in Bargmann and Graney [5], their program makes three passes to identify

clusters. The passes made in their program are controlled by the statement immediately following statement number 444. In our program, CLUSTER, this has been replaced by the transition to the subspace-identification program. Further, the program developed by Bargmann and Graney [5] does not calculate cluster means or the "within" matrix after three passes. Their program computes these quantities at the beginning of the first pass, and after the first and second passes. We require these quantities for our search for the points lying on subspaces. These quantities are calculated in statements between number 811 and number 20001. At this stage, we also punch cards containing the means for each cluster and each variable. These will be needed prior to the execution of ELLIPSE described below. In statements between number 11020 and number 10001, we standardize the original points and reduce them to zero mean and a Cartesian metric. This, then, is the beginning of the second logical part of the program. Here, we determine the points lying on subspaces, using the method of weighted least squares together with Thurstone's "Analytical Method" and the "Single Plane Method." The algorithms (including the change of the BND variable) have been presented in sections 4.2 and 4.3

The output of this program consists of two parts--a print out and a punched deck. The printed output contains the results of three passes made to find clusters. The results listed after the third pass are the final results relating to cluster analysis. It shows which point belongs to which cluster. It also shows at what level the points got included in the cluster. The second part of the printed output gives various simple structure solutions. It gives vectors which transform the original observations into simple structure planes. For each of these vectors, the correlations between the original variables and the "scores" of points

with reference to this vector, the number of points falling into the simple structure plane corresponding to this vector, and the probability of these many points falling into this simple structure plane, are given. A simple structure plane is not included as a solution if the probability of the number of points falling into this plane is greater than 0.01 or if it is within 45° of a plane already found.

Apart from the printed output, a punched deck is also produced. These are data cards which are later loaded into data sets, as described below. The first few cards in this deck--equal in number to the number of clusters formed--are the cards containing cluster means. The next set of cards contains vectors which transform the original data points into simple structure solutions. The number of clusters is designated by NG and the number of simple structure solutions is designated by NSOL. This program also reproduces the cards for each data point with the following additional information. For each data point, the card contains a serial number in columns 1-3, the number of the cluster to which it belongs in column 4, and the simple structure planes in which it is included in columns 5-14. In column 4, a '1' is punched if the point belongs to cluster number 1, '2' if the point belongs to cluster number 2, etc. '0' is punched in column 4 if the point was not assignable to any of the clusters. The simple structure planes to which the point belongs is indicated in columns 5-14 as follows: A '1' in column 5 indicates that the point falls into simple structure number 1, a '1' in column 6 indicates that the point belongs to simple structure number 2, etc. (with provision for up to 10 solutions). This is certainly more than adequate capability. Zeros or blanks in columns 5 to 14 (the computer program punches zeros) indicate absence of this point in the corresponding simple structure plane. Columns 15-70

contain the original coordinates of each data point. The entire punched deck output is also printed out. The cluster means are printed at the end of the third pass, and the vectors which transform the raw data into simple structure solutions are printed immediately after the printout of the corresponding solution. The information punched into the last NP cards of the punched deck is also printed as a final summary. The user may find it helpful to keep this summary with him while he studies the displays.

5.3 Loading the Data Set (Utility Program IEBGENER)

After the user has subjected his data to the CLUSTER program, he should next run the IEBGENER program. This program supplies the output of CLUSTER program as input to the ELLIPSE program. A header card, containing the number of points, the number of variables, the number of groups, and the number of overdetermined subspaces identified by the CLUSTER program, is put before the punched deck produced by the CLUSTER program, and the IEBGENER program is executed. A sample deck set-up for the IEBGENER program is given in Chapter VI; this is a utility routine which transfers the cards to disk.

5.4 The Display Program and the Conversational System

The second program of the package serves to display projections of the clusters and subspaces, on an IBM 2250 Graphics Display Unit. It enables a user, at the console, to communicate with the system and to manipulate displays appearing on the scope. The program, named ELLIPSE, is capable of handling up to 8 variables, 50 data points, 10 simple structure solutions and 5 groups or clusters. The numbering of the

variables is implicit. The first coordinate, for each data point, is regarded as variable number 1, the second coordinate as variable number 2, etc.

The program is an interactive one in the sense that the user, at the console, can decide on variations for later displays on the basis of what he saw in the earlier ones. The input of the data is so formatted and programmed that, if a user wishes to reassign a point from one cluster to another, or if he wishes to change cluster centers, etc., he only needs to make a change to this effect in the corresponding data cards. This capability gives the user an opportunity to redefine clusters and subspaces on the basis of the displays generated. The interaction between the user and the system is achieved through the use of the programmed function keys and the alphameric keyboard which is a part of the 2250 Graphics Display Unit.

The program includes a main program and 3 subroutines. The main program calls two subroutines CALC and EXIBIT. CALC reads the entire input into the ELLIPSE program. The input consists of a header card containing the number of points, the number of variables, the number of clusters and the number of simple structure solutions; cluster means for each variable; vectors which transform the original data points (raw data) into various simple structure solutions; and the data points, together with information regarding the cluster to which a data point belongs and whether or not it lies on a given overdetermined subspace. As explained earlier, this input is given to the program through the execution of the IEBGENER Utility routine. The first (executable) statement of the CALC subroutine reads the header card. Following this, the DO 14 loop reads cluster means, the DO 114 loop reads the transformation vectors and the

DO 10 loop reads the data points. The array IG is used to store information regarding the cluster to which a data point belongs. Recall that for each data point, columns 5-14 contained '1' to indicate the corresponding simple structure plane to which this point belongs. This configuration of '1's and '0's is read as a 10 digit integer number and stored in the first column of the two-dimensional array INNSOL. The serial number of a data point is stored in the second column of INNSOL. In the DO 422 loop, the array IG is examined to determine the size of each cluster. The DO 429 loop, then, determines the total number of points assigned to clusters. The subroutine COR2 is now called which yields the within sum of squares and products matrix based on all the points belonging to clusters. The upper triangular part of this symmetric matrix is stored, columnwise, in the array DSPROD. In the DO 434 loop, each element of the array DSPROD is divided by n_c , the degrees of freedom, to yield an estimate S of the common variance-covariance matrix, Σ . Immediately following the statement number 434, SINV, a sub-routine from the IBM Scientific Subroutine Package, is called to invert the matrix S. The inverse of this matrix S, stored as the first column of the two-dimensional array A, is the metric (based on all points belonging to clusters) employed for unit ellipsoids around the clusters. The subroutine CALC then calculates metrics corresponding to overdetermined subspaces. For each of the overdetermined subspaces, the DO 426 loop calculates a metric corresponding to each of the overdetermined subspaces on the basis of the points belonging to it. The inner DO 427 loop examines the 10 digit integer numbers stored in the first column of the array INNSOL to determine the points lying on a particular subspace. For each subspace, the subroutine

CORIS is utilized to calculate a new within-cluster sum of squares and products matrix. In the DO 442 loop, each element of this matrix is divided by the number of points belonging to the overdetermined subspace, to yield an estimate of the variance-covariance matrix based on the points belonging to the subspace only. The subroutine SINV is then called to invert this matrix. The inverse matrix is used as the metric for ellipsoids corresponding to the overdetermined subspace. The metric corresponding to the overdetermined subspace number 1 is stored as the second column of the two-dimensional array A, the metric corresponding to the overdetermined subspace number 2 is stored as the third column of the array A, etc. This is accomplished in the DO 433 loop. The DO 15 loop, beginning at statement number 450, calculates the maximum and minimum values for each variable. These maximum and minimum values are required later for scaling purposes to accommodate each of the projected data points within the screen limits. A flow chart for the subroutine CALC is given in figure 5.4.1.

Subroutine EXIBIT

The subroutine EXIBIT begins with a call to the DISPLA subroutine of the GRAF (Graphics Addition To Fortran) package [16]. This results in setting up GDSX, GDSY, GTEXT, GPOINT, GDSE, GDSER and GINPUT as display variables. The subroutine LIGHTS of the GRAF package is next called to turn on the lights corresponding to the programmed function keys numbered 1 up to the number of variables, keys 27 to 29 and 31. After these preliminaries, the subroutine MESSGE is called to display an informative message about the program, for the benefit of the user. The subroutine MESSGE returns the control to the subroutine EXIBIT as soon as the user

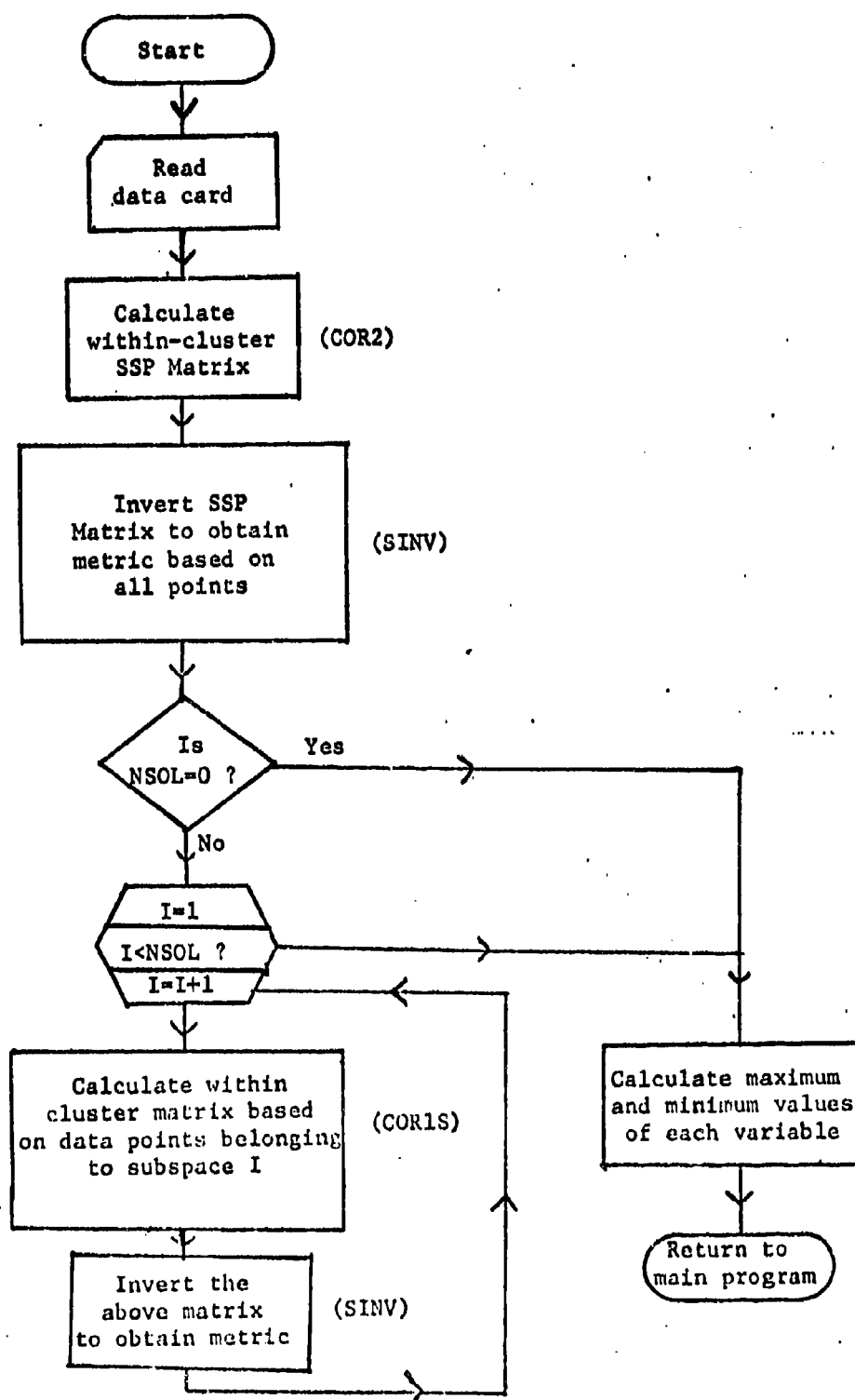
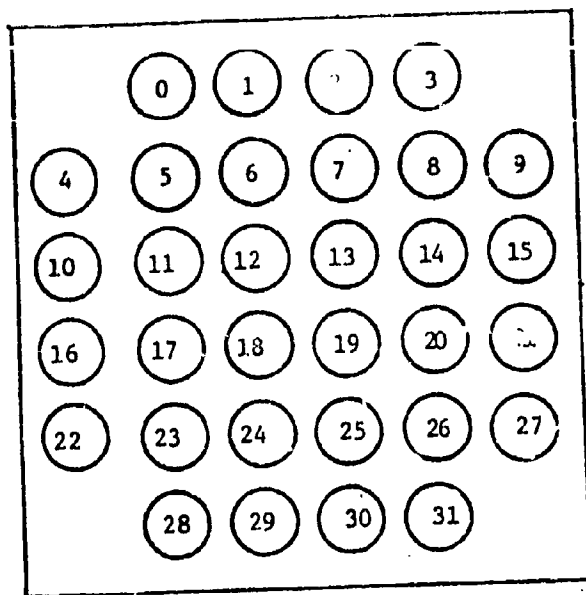
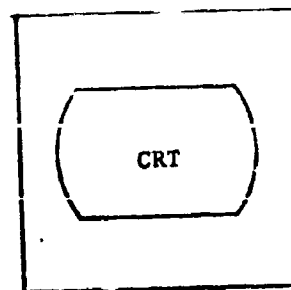


Figure 5.4.1

presses any programmed function key. The message appearing on the screen is erased, a variable NTEST (used later to register the depressed key) is set initially equal to 1 and a variable NGG (a flag which is used to indicate change of vectors) is set equal to 0 and the control goes to statement number 20. This statement is a call to the subroutine KEYIN, with NTEST, the input argument having been set equal to 1. The subroutine KEYIN accepts from the user a vector and the number of a variable, in order to form a 2-dimensional plane. The user indicates his choice of the number of the variable by pressing the corresponding programmed function key. The variable NAXIS is used as an output argument of the subroutine KEYIN and on return contains the number of the programmed function key pressed by the user. As will be seen later, on return from KEYIN, the variable NAXIS must either have a value equal to the number of the variable the user wishes to utilize, or 30. If NAXIS equals 30, it is implied that the user wishes to stop and the display program comes to an end. Otherwise, the subroutine ELLPSE is called with the current value of NAXIS (the number of the variable) as input argument. The subroutine ELLPSE displays the projections of the original data points, and the unit ellipsoids having the metric based on all the points belonging to clusters, onto the 2-dimensional plane formed by the vector and the variable supplied by the user, provided there is no singularity or redundancy involved (see ELLPSE below). After the above displays appear, the user is expected to press a programmed function key. If he presses key 29 or 31, the control comes to statement number 80. This results in erasing the current displays and then a call to the subroutine KEYIN, with NTEST, the input argument, being 29 or 31. As before, the call to the subroutine KEYIN is followed by a call to the subroutine ELLPSE and the cycle



Programmed Function Keyboard



Alphameric Keyboard

IBM 2250 Display Unit

continues. If a key corresponding to the number of an overdetermined subspace was pressed, the control comes to statement number 61, and if key 30 was pressed, the display program comes to an end. If the control comes to statement number 61, the subroutine RELPSE is called to display the projections of the overdetermined subspace. If none of the above mentioned keys is pressed, an error message, as contained in the format statement number 62, appears on the screen. The error message continues to appear until the user presses a proper key or, in case of singular or redundant situations, rectifies the situation. If the subroutine RELPSE is called, after the projections of the overdetermined subspace are displayed, the user is expected to press a programmed function key. Once again, if key 29 or 31 is pressed, the current displays are erased, the control comes to statement number 20 and the subroutine KEYIN is called. The program terminates if key 30 was pressed. If the key corresponding to the number of the overdetermined subspace was pressed, that part of the current display pertaining to the projection of the overdetermined subspace is erased, and the subroutine RELPSE is called to display the projections of the overdetermined subspace that is now being requested. An error message appears if none of the abovementioned keys is pressed, and the cycle continues in this manner. A flowchart of the subroutine is given in figure 5.4.2.

Subroutine KEYIN

The subroutine KEYIN accepts a vector and the number of the variable from the user, to form a 2-dimensional plane. It has one input argument, NTEST, and an output argument, NAXIS. The first statement in the subroutine tests the value of NTEST. If it equals 31, that part of

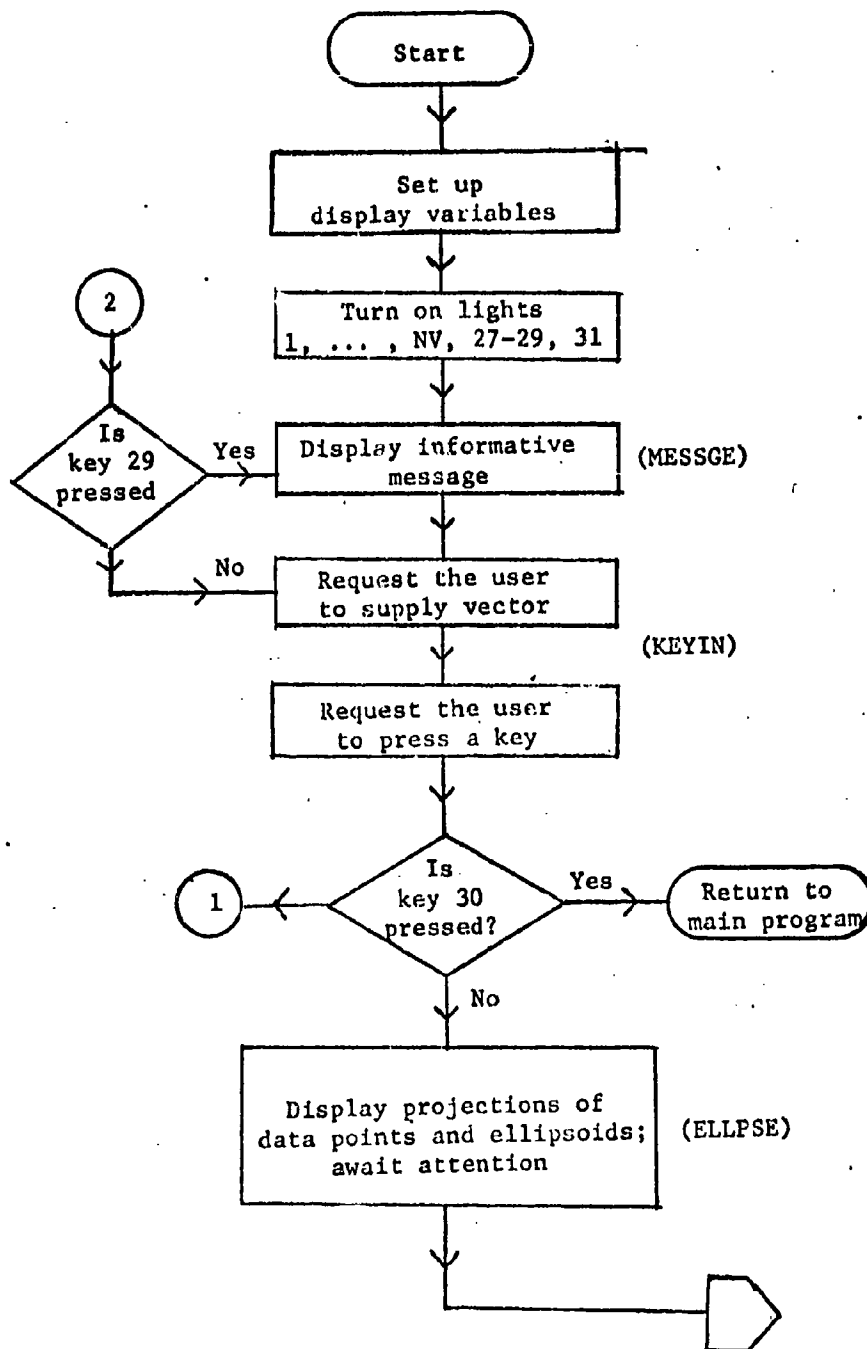


Figure 5.4.2a

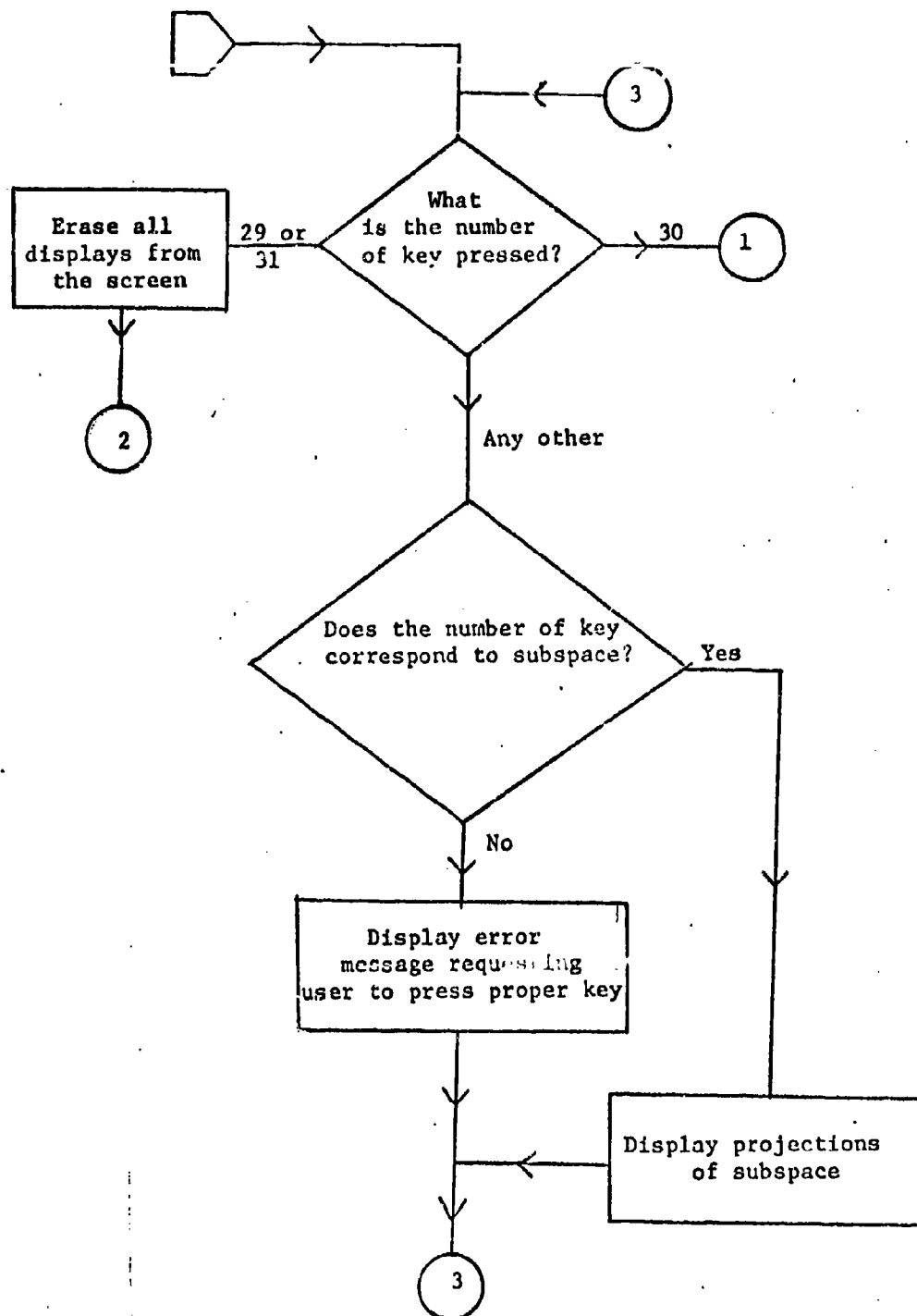


Figure 5.4.2b

the subroutine which accepts a vector from the user is skipped, and the control comes to statement number 101. Otherwise, the control comes to statement number 99, and whatever appears on the screen is erased to prepare to accept a vector from the user. On the first call to the subroutine KEYIN, the input argument, NTEST, is set equal to 1 so that the control invariably comes to statement 99. On subsequent occasions, however, NTEST will have a value of 29 or 31. The DO 110 loop displays the transformation vectors suggested by the rotation (CLUSTER) program for the information of the user. The message contained in the format statement number 3000, requesting the user to supply a vector, then appears on the screen. The program now awaits the user to supply the vector. The calls to SCTDV and DVTDM subroutines of the GRAF package transfer the vector supplied by the user from the screen to the display variable table and from there to the dummy unit 4. The vector is read from the dummy unit 4 into the array RINPUT. Before, however, the vector is read, the DO211 loop transfers the current values stored in the RINPUT array to the RWRKNG array. This is done to insure that the vector previously supplied by the user is not destroyed in case he wants to continue with that vector. The DO 213 loop checks if the user supplied a null vector. If so (as is the case when he just wants to continue with the previous vector), this is always replaced by the previous non-null vector as would be apparent from the DO 216 loop. The only exception, as will be seen from the statement following the statement number 213, is the first call to the subroutine. This would result in singularity and the user will receive an error message instead of the displays.

After the vector is accepted, the control comes to statement number 101. The message contained in the format statement number 1658 appears on the screen. The loop beginning with the statement number 60 insures that no value other than 30, or the number corresponding to the variable which the user wishes to utilize, will be returned as the value of the output argument NAXIS. The user can, of course, go back to statement 99, the beginning of the program, if he presses key 28. This gives him a chance to amend the vector already supplied. A flowchart of the subroutine is given in figure 5.4.3.

Subroutine ELLPSE

This subroutine is used to display projections of original data points, and the unit ellipsoids having the metric based on all points belonging to clusters, onto the 2-dimensional plane formed by the vector and the variable supplied by the user. The DO 10 and DO 11 loops set up a matrix R formally given by

$$R = \begin{bmatrix} 0 & a_1 \\ 0 & a_2 \\ \vdots & \vdots \\ \vdots & \vdots \\ 1 & 0 \\ \vdots & \vdots \\ \vdots & \vdots \\ 0 & a_p \end{bmatrix} \quad (5.4.1)$$

where 1 in the first column appears in the row corresponding to the number of the variable chosen by the user, and $(a_1, a_2, \dots, 0, \dots, a_p)$ is the vector supplied by the user and modified to form an orthogonal axis

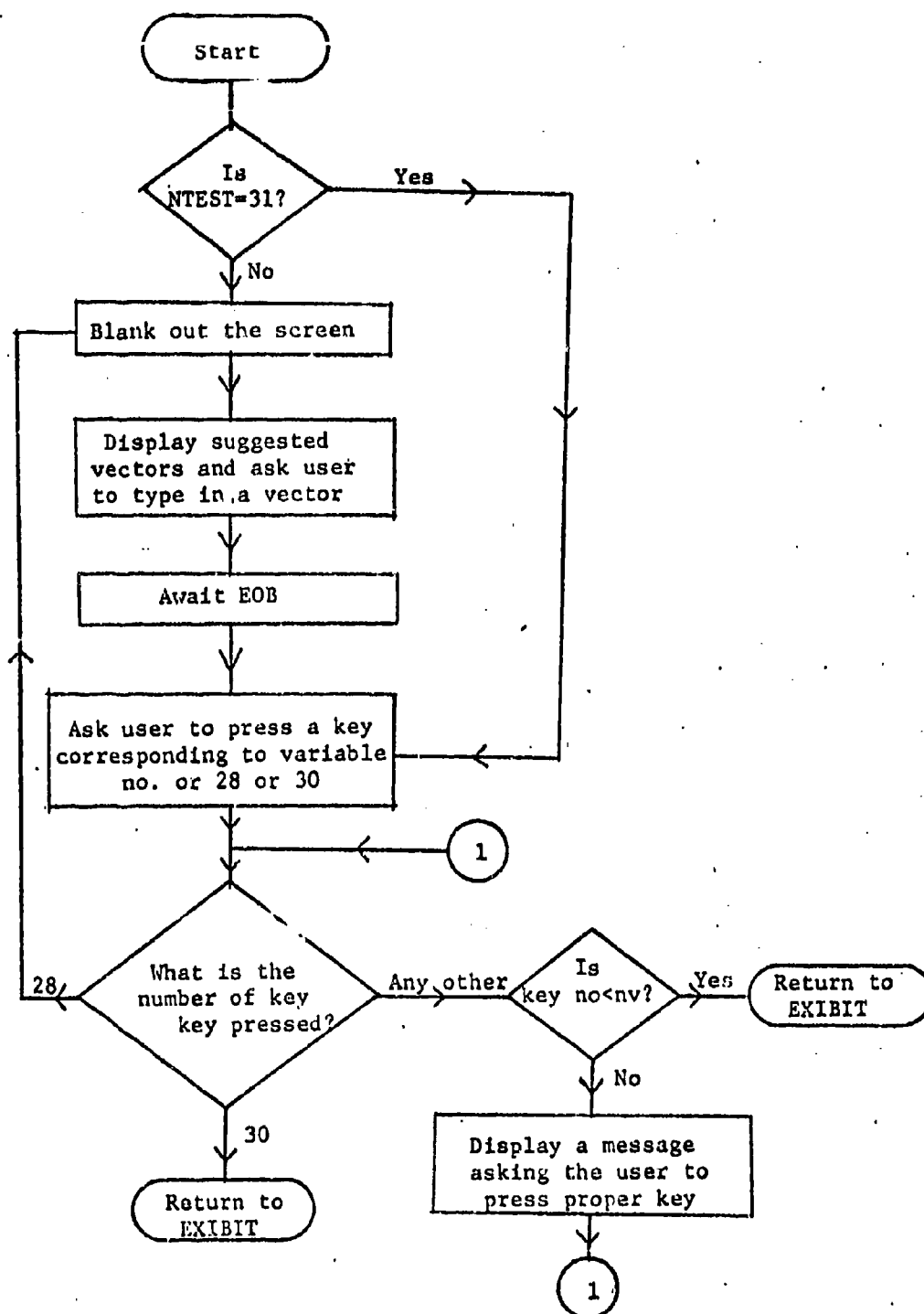


Figure 5.4.3

with the desired variable. If the user supplies a vector collinear with his choice of the observed variable, the second column of the matrix R will have all zero elements. The DO 20 loop, and the statement immediately following this loop, check for the above-mentioned possibility of singularity. If no singularity is present, the control comes to statement number 211; otherwise, the error message, as contained in the format statement number 1659, is displayed on the screen, and control is returned to the subroutine EXIBIT. At statement 211, the DO 21 loop is set up to normalize the second column of the matrix R . The DO 12 loop picks up the metric S^{-1} based on all points belonging to clusters and the subroutines MPRD and GTPRD of the IBM Scientific Subroutine Package are called, to form the matrix product $R'S^{-1}R$. The DO 13 loop, and the statement immediately following it, calculate the matrix product XR , where X is the matrix of original data points. For a projected data point, the element of the first column of the matrix is treated as its x-coordinate and the element of the second column is treated as y-coordinate. The DO 110 loop creates orders to plot points with these sets of x and y coordinates. The actual plotting is done by displaying, on the screen, at the place where the point should appear, its serial number, so that the user may know which data point projects into what region of the 2-dimensional display.

In the statements immediately following the DO 110 loop, the semi-axes of the ellipses to be displayed are calculated. They are based on the metric $R'S^{-1}R$ (of the projected ellipsoids). The cluster centers (centers of ellipses) are likewise transformed into $\mu_1'R$. The points describing the circumference of the ellipses

$$(\underline{x}' - \underline{\mu}_1 R) R' S^{-1} R (\underline{x} - R' \underline{\mu}_1) = 1 \quad (5.4.2)$$

where $\underline{x}' = (x_1, x_2)$ are the running coordinates, are constructed by angular sweep. Beginning with an initial angle of 5° the DO 150 loop calculates 72 different points describing the circumference of the ellipses. The DO 100 loop creates orders to plot these points and the ellipses appear on the screen when the statement immediately following the DO 100 loop is executed. A flowchart of the subroutine is given in figure 5.4.4

Subroutine RELPSE

This subroutine is used to display projections of ellipsoids having metric based on the points belonging to the overdetermined subspace being superimposed. If the metric for the i th subspace is denoted by S_i^{-1} , the subroutine calculates the points describing the circumference of ellipses

$$(\underline{x}' - \underline{\mu}_1 R) R' S_i^{-1} R (\underline{x} - R' \underline{\mu}_1) = 1 \quad (5.4.3)$$

where R is as defined in ELLPSE. The technique employed to calculate these points is similar to the one employed before. Like ELLPSE, RELPSE also checks for singularity. If it is present, no displays of ellipses appear. It is to be noted that, if the user chooses one of the vectors suggested to him in the display (the vectors identified in the CLUSTER program), the ellipses constructed in RELPSE, if the user depresses the corresponding key, is quite flat, as intended. It is conceivable that, in such an instance, singularities could occur (though we did not see any in our examples) and hence the program checks for them.

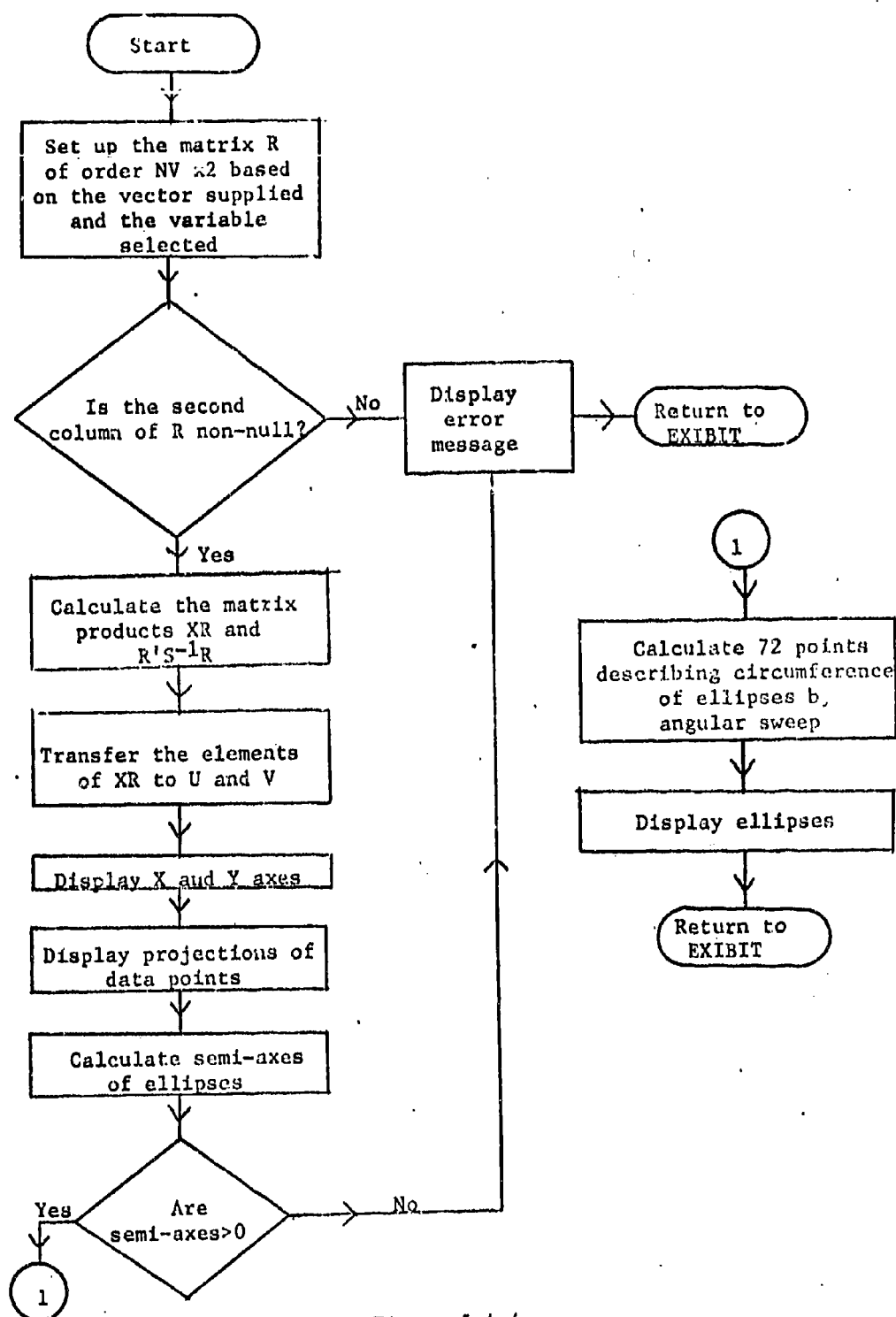
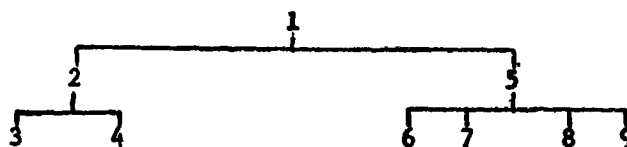


Figure 5.4.4

Overlay Structure

To reduce storage requirements, an overlay structure was designed as follows:



1. (Root) contains the main program and Function KD.
2. Contains subroutine CALC
3. Contains subroutine COR1S
4. Contains subroutine COR2
5. Contains subroutine EXIBIT
6. Contains subroutine MESSGE
7. Contains subroutine KEYIN
8. Contains subroutine RELPSE
9. Contains subroutine ELLPSE

Segment 1, along with the main program and function KD, also contains the system support routines IBCOM#, ARIH#, FIOCS#, ADCON#, and system utilities IHCUATBL, IHCUOPT and IHCTRCH. Segment 5 contains all of the GRAF routines required except BUFRS, CUR\$\$, RCUR\$, READSC, SCNDVDK and SCTDV, which are included in segment 7. With the help of this overlay structure and equivalencing of a few arrays in the ELLPSE subroutine, it was possible to reduce the storage requirements to 64K bytes.

Deck Layout for ELLIPSE

```

//STEP1 EXEC FORTGC
//FORT.SYSLIN DD DSN=NAME=&&CHAIN(ROOT),SPACE=(TRK,(150,10,5)),      C
                UNIT=SYSDA,DISP=(NEW,PASS)

//FORT.SYSIN DD *
```

Main program here

/=

//STEP5 EXEC FORTGC

//FORT.SYSLIN DD DSNAME=&&CHAIN(LINKA3),DISP=(MOD,PASS),UNIT=SYSDA

//FORT.SYSIN DD *

Function KD Source Deck

/*

//STEP2 EXEC FORTGC

//FORT.SYSLIN DD DSNAME=&&CHAIN(LINKA),DISP=(MOD,PASS),UNIT=SYSDA

//FORT.SYSIN DD *

Subroutine CALC Source Deck

/*

//STEP7 EXEC FORTGC

//FORT.SYSLIN DD DSNAME=&&CHAIN(LINKA5),DISP=(MOD,PASS),UNIT=SYSDA

//FORT.SYSIN DD *

Subroutine COR1S Source Deck

/*

//STEP6 EXEC FORTGC

//FORT.SYSLIN DD DSNAME=&&CHAIN(LINKA4),DISP=(MOD,PASS),UNIT=SYSDA

//FORT.SYSIN DD *

Subroutine COR2 Source Deck

/*

//STEP4 EXEC FORTGC

//FORT.SYSLIN DD DSNAME=&&CHAIN(LINKA2),DISP=(MOD,PASS),UNIT=SYSDA

//FORT.SYSIN DD *

Subroutine EXIBIT Source Deck

Subroutine EXIBIT Source Deck

/*

//STEP0 EXEC FORTGC

//FORT.SYSLIN DD DSN=LINKA8,DISP=(MOD,PASS),UNIT=SYSDA

//FORT.SYSIN DD *

Subroutine MESSGE Source Deck

/*

//STEP10 EXEC FORTGC

//FORT.SYSLIN DD DSN=LINKA9,DISP=(MOD,PASS),UNIT=SYSDA

//FORT.SYSIN DD *

Subroutine KEYIN Source Deck

/*

//STEP3 EXEC FORTGC

//FORT.SYSLIN DD DSN=LINKA1,DISP=(MOD,PASS),UNIT=SYSDA

//FORT.SYSIN DD *

Subroutine ELLPSE Source Deck

/*

//STEP8 EXEC FORTGC

//FORT.SYSLIN DD DSN=LINKA6,DISP=(MOD,PASS),UNIT=SYSDA

//FORT.SYSIN DD *

Subroutine RELPSE Source Deck

/*

//S10 EXEC LKED,PARM=(LET,LIST,OVLY,XREF)

//LKED.SYSLMOD DD DSN=SYS1.GRAFLIB(ELLIPSE),DISP=SHR,

C

//SPACE=(TRK,(0,0))

//LKED.SYSLIB DD DSN=SYS1.GRAFLIB,DISP=SHR

```
// DD DSN=SYS1.UGALIB,DISP=SHR
// DD DSN=SYS1.SSPLIB,DISP=SHR
// DD DSN=SYS1.FORTLIB,DISP=SHR
// DD DSN=SYS1.LINKLIB,DISP=SHR
// DD DSN=SYS1.GRAPHLIB,DISP=SHR
//LKED.MODULE DD DSN=&&CHAIN,DISP=OLD
//LKED.SYSIN DD *
  INCLUDE MODULE(ROOT)
  INCLUDE MODULE(LINKA3)
  INCLUDE SYSLIB(IBCOM#)
  INCLUDE SYSLIB(ARITH#)
  INCLUDE SYSLIB(FIOCS#)
  INCLUDE SYSLIB(ADCON#)
  INCLUDE SYSLIB(IHCUATBL)
  INCLUDE SYLIB(IHCUOPT)
  INCLUDE SYSLIB(ERRMON)
  INCLUDE SYSLIB(IHCTRCH)
  OVERLAY ONE
  INCLUDE MODULE(LINKA)
  INCLUDE SYSLIB(SINV)
  INCLUDE SYSLIB(MFSO)
  OVERLAY TWO
  INCLUDE MODULE (LINKA5)
  OVERLAY TWO
  INCLUDE MODULE (LINKA4)
  OVERLAY ONE
  INCLUDE MODULE (LINKA2)
```



```
INCLUDE SYSLIB(GAFERR)
INCLUDE SYSLIB(LIGHTS)
INCLUDE SYSLIB(WRFMT$)
INCLUDE SYSLIB(DETEKT)
INCLUDE SYSLIB(PLOT)
INCLUDE SYSLIB(DETAIN)
INCLUDE SYSLIB(DISPLA)
INCLUDE SYSLIB($$OVER)
INCLUDE SYSLIB(CHAR)
INCLUDE SYSLIB(POINT)
INCLUDE SYSLIB(LINE)
INCLUDE SYSLIB(PLACE)
INCLUDE SYSLIB($$$$BT)
INCLUDE SYSLIB($$INIT)
INCLUDE SYSLIB(DUMMY$)
INCLUDE SYSLIB($VOVER)
INCLUDE SYSLIB(CLOSE)
INCLUDE SYSLIB(LINE$$)
INCLUDE SYSLIB(UNPLOT)
INCLUDE SYSLIB(PLACE$)
INCLUDE SYSLIB(POINT$)
INCLUDE SYSLIB(ERASE)
INCLUDE SYSLIB(BLANK)
INCLUDE SYSLIB(RESET)
OVERLAY TWOA
INCLUDE MODULE (LINKAS)
OVERLAY TWOA
```

INCLUDE MODULE (LINKA9)

INCLUDE SYSLIB(SCNDVDK)

INCLUDE SYSLIB(CUR\$\$)

INCLUDE SYSLIB(BUFRS)

INCLUDE SYSLIB(SCTDV)

INCLUDE SYSLIB(RCUR\$)

INCLUDE SYSLIB(READSC)

OVERLAY TWOA

INCLUDE MODULE (LINKA6)

OVERLAY TWOA

INCLUDE MODULE (LINKA1)

CHAPTER VI

USER'S GUIDE

6.1 Introduction

The user who is interested in using the programs described in the previous chapter would find himself in one of the following situations:

(i) He may not have analysed his multivariate data and may not yet have identified clusters and subspaces. If so, he should first subject his data to the CLUSTER program. The next section contains instructions on how to use (execute) this program.

(ii) He may have analysed his data using the CLUSTER program but has not loaded the data set for the ELLIPSE program. If so, he should execute the IEBGENER Utility routine and load the data set. This is described in section 6.3.

(iii) Finally, the user may have gone through the steps (i) and (ii) above and may want to see the displays of projected clusters and subspaces. The use of the ELLIPSE program including an indication of what to look for in the displays is described in section 6.4.

Section 6.5 contains an illustration.

6.2 The CLUSTER program

The analysis of the user's multivariate data begins with the identification of clusters and overdetermined subspaces, if any. For this purpose, the user must first analyse his data using the CLUSTER program.

This is a batch program. The following data cards need to be supplied:

Data Card 1

Columns 1-3, number of points (individuals or experimental units)

Columns 6-7, number of variables (responses) measured on each experimental unit

Columns 8-11, alpha level for cluster core on first pass
(suggested value 0.90)

Columns 13-16, alpha level for cluster extension on first pass
(suggested value 0.50)

Columns 18-21, alpha level for cluster core on second pass
(suggested value 0.90)

Columns 23-26, alpha level for cluster extension on second pass
(suggested value 0.50)

Columns 28-31, alpha level for cluster core on third pass
(suggested value 0.90)

Columns 33-36, alpha level for cluster extension on third pass
(suggested value 0.50)

Data Card 2

This is a variable format card and should contain the FORMAT by which each experimental unit is to be read.

The remaining data cards contain the observations, one card (or record which may consist of several cards) contains the coordinates of one point.

The numbering of these points is implicit, according to the sequential order of these cards.

Our suggestion above that 0.90 should be used as alpha level for cluster core and 0.50 as alpha level for cluster extension is empirical. Of course, he can use any other set of values. For a detailed discussion of this matter the reader is directed to Graney [10].

The output of this program has been discussed at length in section 5.2. The punched deck produced by this program is required for the ELLIPSE program.

6.3 IEBGENER Utility routine

As explained in section 5.3, it is necessary to execute the IEBGENER Utility routine to load the output of the CLUSTER program into a data set required as input to the ELLIPSE program. The deck set-up for the execution of this utility routine is as follows:

- (i) JOB card
- (ii) //STEP EXEC PGM=IEBGENER
- (iii) //SYSPRINT DD SYSOUT=A
- (iv) //SYSIN DD DUMMY
- (v) //SYSUT2 DD DSN=SYS1.R2250,VOL=SER=UGA231,DISP=SHR,UNIT=2314
- (vi) //SYSUT1 DD DATA,DCB=(RECFM=FB,LRECL=80,BLKSIZE=320)
- (vii) Data Cards
- (viii) /*

The data cards consist of a header card followed by the punched deck produced by the CLUSTER program (of course, the user could produce his own data cards, and any assignment of points to clusters or subspaces which he desires. In this respect the CLUSTER program is merely intended to give him guidance--but a very strong one indeed). The header card is made up as follows (all numbers right justified);

Columns 1-4, number of points (individuals or experimental units)

Columns 5-8, number of variables (responses) measured on each experimental unit

Columns 9-12, number of clusters identified by the CLUSTER program

Columns 13-16, number of overdetermined subspaces (simple structure solutions) identified by CLUSTER program

After the IEBGENER Utility routine is executed, one is ready for the ELLIPSE program.

6.4 The ELLIPSE program

This program works under the control of GMS (Graphics Monitor System). For greater detail on the operation of GMS see Penn [31]. In order to be operational under the conversational GMS, the load module of the program has to be a member of the partitioned data set GRAPHLIB. The user should verify, by typing the command \$NAMES on the console typewriter, that the program, in fact, is a member of the GRAPHLIB data set. If not, the user will first have to compile and link edit the program using the deck set-up given in section 5.4. The user can then execute the program using the command \$LINK ELLIPSE to link to it.

Photographs made during the use of the program are reproduced here. The user will find it helpful to refer to them while studying the rest of the section. Some of the photographs will be specifically discussed in the next section.

The execution of ELLIPSE begins with the display of an informative message. The user should carefully read the message. (Note especially the use of the programmed function key 30. This is to be pressed only when it is desired to stop the execution of the program.) The user should then press any key other than 0 or 30. He is now asked to type the coordinates of a vector which he wishes to utilize in order to form a

THIS PROGRAM DISPLAYS CLUSTERS BY PROJECTING THEM ON
VARIOUS 2-DIMENSIONAL SUBSPACES. SIMPLE STRUCTURE
SOLUTIONS CAN ALSO BE INDICATED. REFER TO YOUR
INSTRUCTION CHART. HAVE YOU ENTERED ALL NECESSARY
DATA VIA IEDGENER?

THE BOTTOM ROW OF THE PROGRAM
FUNCTION KEYS IS LIT UP. THEY WILL BE USED
AS DIRECTED. THE DARK ONE, KEY NO. 30, IS
THE PANIC BUTTON. IT WILL RETURN YOU TO THE
MONITOR.

IN CASE OF PANIC PRESS KEY 30
NOW PRESS ANY KEY TO GET STARTED

ACCORDING TO THE ROTATION PROGRAM THE
FOLLOWING VECTORS TRANSFORM RAW DATA INTO
SIMPLE STRUCTURE SOLUTIONS NO. 1 TO 3

-0.309-0.275 0.256-0.309
0.012 0.843-0.155 0.516
-0.306-0.259-0.316 0.860

ENTER A TRANSFORMATION VECTOR.
NO MORE THAN 7 CHARACTERS, DECIMAL POINT
MUST BE TYPED
DEPRESS JUMP KEY AFTER EACH COORDINATE

X(1)=

X(2)=

X(3)=

X(4)=

X(5)=

X(6)=

X(7)=

X(8)=

IF PREVIOUS VECTOR OK, JUST PRESS EOS

ACCORDING TO THE ROTATION PROGRAM THE
FOLLOWING VECTORS TRANSFORM RAW DATA INTO
SIMPLE STRUCTURE SOLUTIONS NO.1 TO 3

-0.309-0.275 0.856-0.309

0.012 0.843-0.155 0.516

-0.306-0.259-0.316 0.860

ENTER A TRANSFORMATION VECTOR.
NO MORE THAN 7 CHARACTERS, DECIMAL POINT
MUST BE TYPED
DEPRESS JUMP KEY AFTER EACH COORDINATE

X(1)=-0.306

X(2)=-0.259

X(3)=-0.316

X(4)=0.860

X(5)=

X(6)=

X(7)=

X(8)=

IF PREVIOUS VECTOR OK, JUST PRESS EOB

PRESS ONE KEY CORRESPONDING TO THE AXIS
- VARIABLE YOU WISH TO SELECT, OR 30 IF
YOU WISH TO STOP. IF YOU WISH TO MAKE
CHANGES IN YOUR VECTOR PRESS KEY 28.

2-dimensional plane, the other axis being an observable variable. Notice that vectors which transform the original data points into overdetermined subspaces appear in this display for the user's ready reference. The user should try one or more of these vectors but he may also supply any number of other vectors, if he feels that this would help him interpret the structure of his data. The coordinates of the vector intended to be used should be typed in through the alphanumeric keyboard. The place where the digit (or any character for that matter) typed will appear on the screen is indicated by a cursor. The use of the "JUMP" key after one coordinate is entered, will cause the cursor to move over to the place for the next coordinate. After all the coordinates of a vector are entered, the user should press EOB. This is done by pressing both the 'ALT' and the '5' key of the alphanumeric keyboard. The first time the user is asked to supply a vector, he must give a non-null vector. (Later on, null vectors will be acceptable and simply mean that there is no change. At that time the user would just press EOB when this display appears and thus indicate that he does not wish to change the previous vector.)

After the vector is supplied, the user is asked to choose a variable as the other axis of a 2-dimensional plane. The choice is made by pressing the programmed function key corresponding to the number of the variable; if variable number 1 is desired, press key 1, etc.

After a vector and a variable have thus been chosen, the desired 2-dimensional plane will appear on the screen. The abscissa corresponds to the variable, the ordinate corresponds to the chosen vector. The legend (numbers given alongside the coordinate axes) indicates minimum and maximum values. The serial number of each data point is projected at the appropriate coordinates. Ellipses, i.e., projections of unit ellipsoids,

based upon the within-cluster metric of all points, are drawn around each cluster center. These ellipses generalize the concept of a standard unit interval in one dimension.

One aspect to be observed on the screen at this time is whether the points which supposedly belong to the same cluster, are close together (regardless which axes or vectors are used), and whether one could distinguish them visually from points belonging to a different cluster. There may be some overlaps but the clusters should be visually distinct. If certain points exhibit the above phenomenon in all displays, then they can be regarded as forming a cluster.

The user has now a choice. He may wish to superimpose one of the overdetermined subspaces, or he may wish to change the vector, or the variable, or both. To change the vector, he presses key 29, to change the number of the variable, he presses key 31, and to superimpose an overdetermined subspace, he presses the key corresponding to the number of that subspace.

Superimposition of an overdetermined subspace results in the display of projections of unit ellipsoids having metric based on only those points which belong to the overdetermined subspace. If the plane of projection selected is orthogonal to the overdetermined subspace, the projections of unit ellipsoids under reference will be flat and elongated. Further, the projections of the points lying on the subspace will make a narrow band (almost resembling a straight line). The plane of projection would be orthogonal to the overdetermined subspace, if the vector which transforms the original data points into this overdetermined subspace is supplied as the desired vector. As many different planes as there are

dimensions can be constructed by taking this normal vector against each of the variable axes.

The plane of projection, in a way, is the direction from which we look at the ellipsoids embedded in the p -dimensional space. The ellipsoids having metric based on the points belonging to an overdetermined subspace must necessarily be flat and elongated. However, they must be viewed from the proper direction. Otherwise, this elongation may not appear or, in some instances, they will appear as very small circles.

Once again, the user can press key 29 to supply a new vector, key 31 to change the number of the variable and the key corresponding to the number of an overdetermined subspace to superimpose the subspace. The program continues in this manner. It can be terminated at any time by pressing key 30.

It should be noted that the projections onto a plane formed by any pair of observable variables is a special case. If the user desires to have projections onto the plane formed by variables 1 and 2, say, all he has to do is supply $(1.0, 0.0, \dots, 0.0)$ as the vector and press key 2. In fact, since the program does not necessarily require normalized vectors as input, any vector of the form $(a, 0, 0, \dots, 0)$ where $a \neq 0$, results in the selection of variable 1 as one of the axes.

6.5 An Illustration

The programs described above were used in the analysis of artificial data. The data consisted of 45 points and 4 variables. These data were generated in the following manner. First, 180 normal deviates with zero mean and unit variance were generated; they made up the 180 elements of a 45×4 matrix, numbered column-wise. The first 25 measurements on the 4th

variable were then redefined by the relation

$$z(I) = z(I)/10 + (z(I - 135) + z(I - 90) + z(I - 45))/3$$

$$I = 136, 137, \dots, 160$$

i.e., the first 25 observations on variable 4 were replaced by one tenth of the original observation plus the mean of the first three variables. Similarly, the last 25 observations on variable 3 were replaced by the relation

$$z(I) = z(I)/10 + (z(I - 90) + z(I - 45) + z(I - 45))/3$$

$$I = 111, 112, \dots, 135$$

The modification of the data matrix by these two operations built in 2 subspaces. In effect, the first 25 observations on variable 4 became almost a linear combination of the remaining 3 variables and the last 25 observations on variable 3 similarly became almost a linear combination of the remaining 3 variables. Still, however, all the variables had zero mean, and to introduce mean shifts, the vectors (3, 9, 5, 9), (5, 9, 7, 3) and (7, 3, 7, 5) were added to the first 15 observations, second 15 observations and the last 15 observations respectively. Thus the entire data matrix became a simulated sample drawn from normal populations with mean vectors (3, 9, 5, 9), (5, 9, 7, 3) and (7, 3, 7, 5) and two built in subspaces. This data matrix was then subjected to the first program of cluster identification and subspace determination. This program correctly assigned the first 15 points to one cluster, the second 15 points to another cluster and the last 15 points to a third cluster. The subspace identification part of the program gave 3 overdetermined subspaces. This was not at all surprising since the two planes were already built in and, as frequently happens in such instances (see section 4.2), a plane was found somewhat in-between the two constructed planes. The cosines between

the normals to these three planes were as follows:

	1	2	3
1	1.00000	-0.29684	0.00673
2	-0.29684	1.00000	0.02303
3	0.00673	0.02303	1.00000

The planes identified were as under (serial numbers of points belonging to the planes are given).

Plane 1: --12, 21, 22, 23, 24, 25, 27, 28, 29, 30, 31, 32, 33, 34, 35, 37, 38, 39, 40, 41, 42, 43, 44, 45, (24 points)

Plane 2: --2, 9, 14, 15, 17, 18, 29, 30, 32, 34, 37, 39, (12 points)

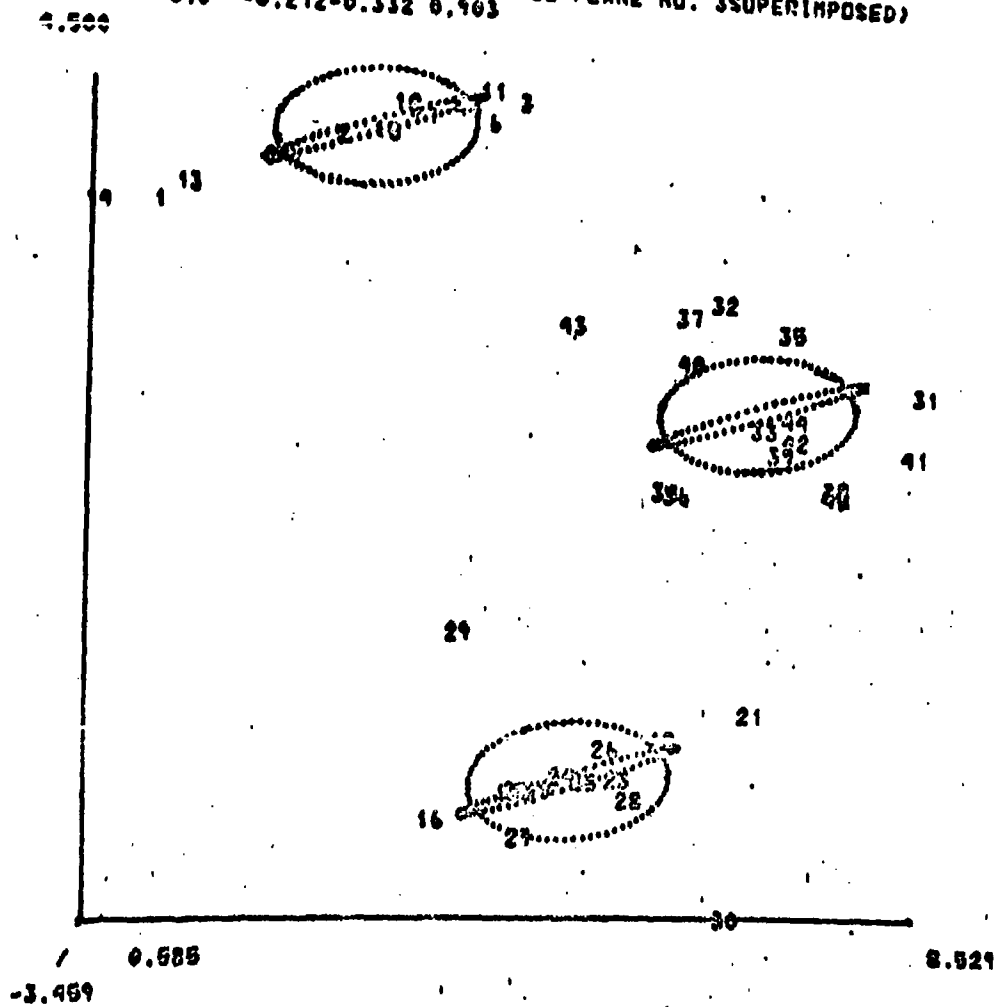
Plane 3: --1, 2, 4, 7, 8, 9, 10, 11, 13, 14, 15, 16, 17, 18, 19, 20, 21, 22, 24, 25, (20 points)

According to the built-in subspaces, one subspace should have contained the last 25 points, i.e., points 21-45 and the other subspace should have contained the first 25 points. Plane 1 given above did pick up 23 of the last 25 points and an extra point number 12. Likewise, plane 3 picked up 20 of the first 25 points. Plane 2 picked up 6 of the points belonging to plane 1 and 6 of the points belonging to plane 3. Thus plane 2 is somewhat of a mixture of the planes 1 and 3.

The results of the CLUSTER program were then displayed using the ELLIPSE program. Notice especially the following displays in which the 2-dimensional planes were formed by selecting a suggested vector (-0.306, -0.259, -0.316, 0.860) and an observable variable. The vector under consideration is the vector 3, which transformed the original data points into simple structure plane 3.

(a) Variable 1 against vector 3--After normalizing, the vector reduced to (0.0, -0.272, -0.332, 0.903). Plane 3 was superimposed. Points number 1, 2, 3, 4, 5, 7, 8, 9, 10, 11, 12, 13, 14, 15, 16, 17, 18, 19, 20, 21, 22, 23, 24, 25, 26 can be seen to be lying on the simple structure plane.

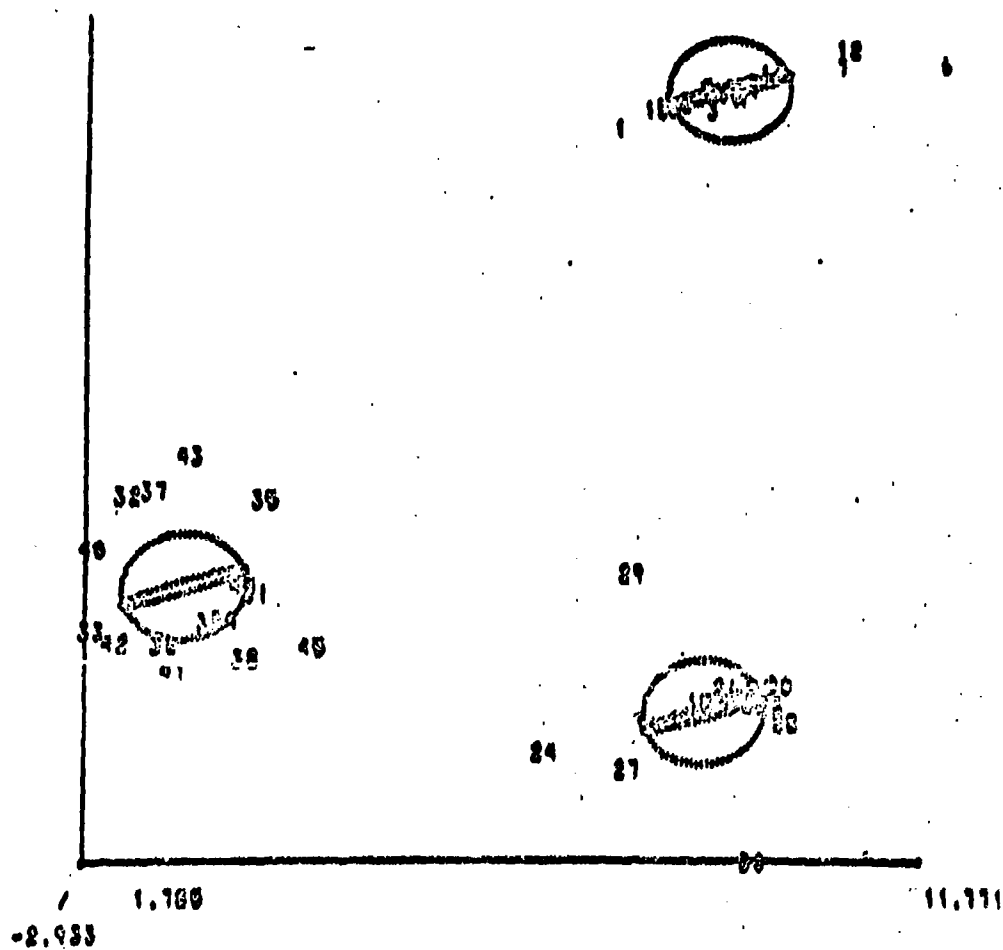
VARIABLE 1 AGAINST VECTOR BELOW (SIMPLE PLANE NO. 3 SUPERIMPOSED)
 0.0 -0.212 -0.332 0.403



(b) Variable 2 against vector 3--After normalizing, the vector reduced to (0.317, 0.0, 0.327, 0.390). Plane 3 was superimposed. Points number 1, 3, 5, 6, 7, 8, 9, 10, 11, 12, 13, 14, 15, 16, 17, 18, 19, 20, 21, 22, 23, 24, 25, 26 can be seen lying on the simple structure plane.

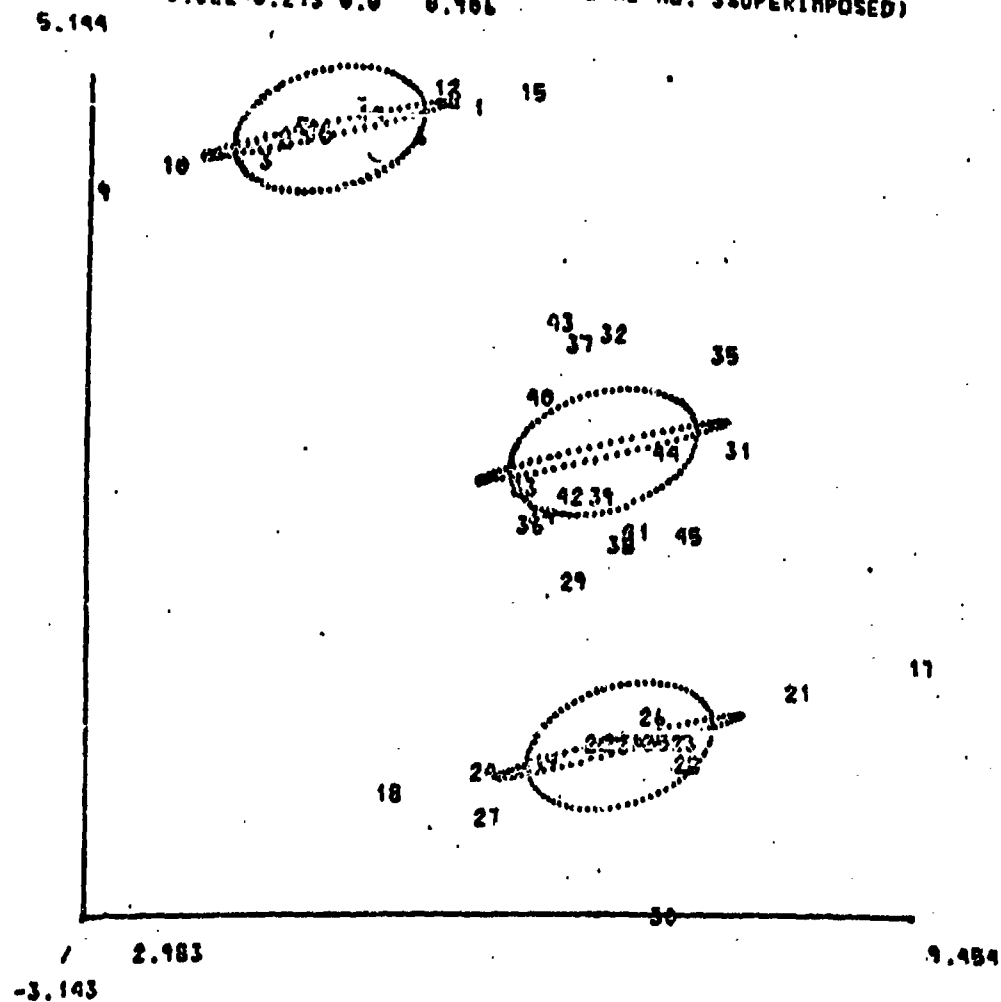
VARIABLE Z AGAINST VECTOR DELTA (SIMPLE PLANE NO. 3 SUPERIMPOSED)

6.058



(c) Variable 3 against vector 3--After normalizing, the vector reduced to (-0.322, -0.273, 0.0, 0.906). Plane 3 was superimposed. Points number 1-26 can be seen lying on the simple structure plane.

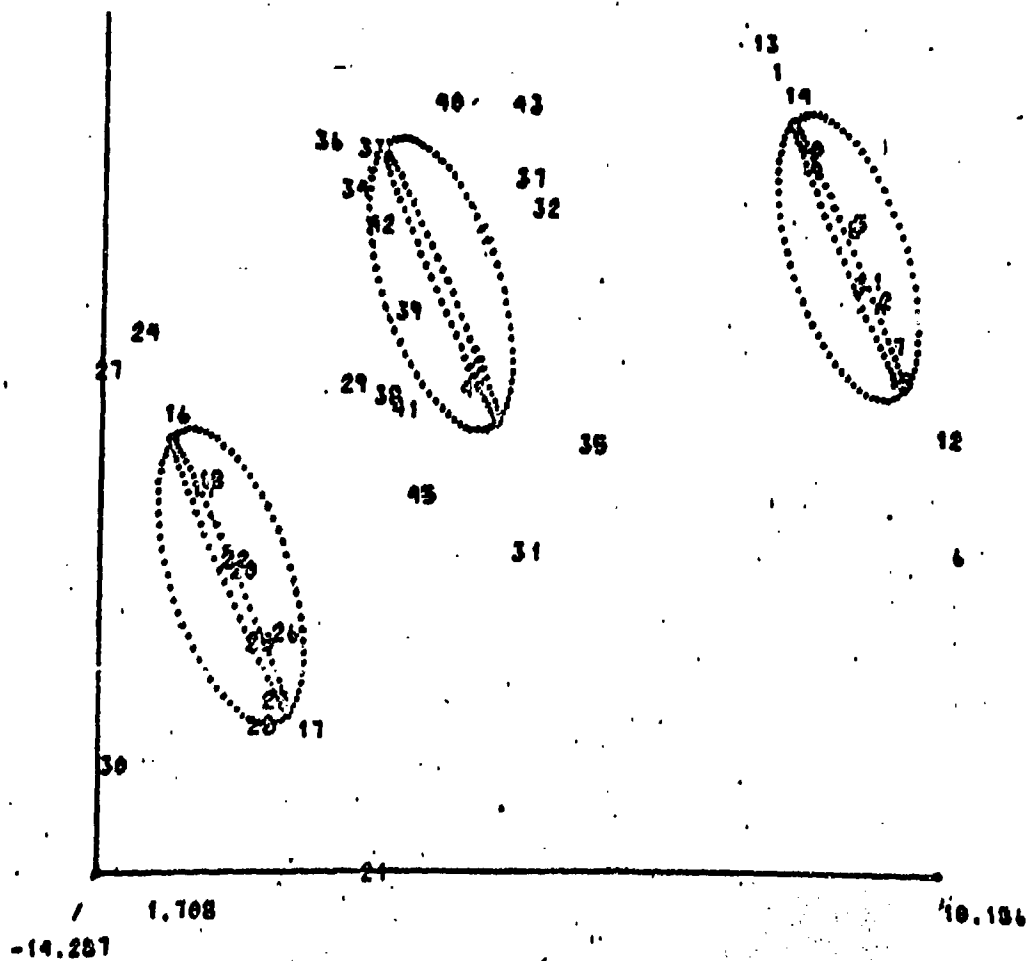
VARIABLE AGAINST VECTOR DELOW, (SIMPLE PLANE NO. 3 SUPERIMPOSED)
 -0.322 -0.213 0.0 0.906



(d) Variable 4 against vector 3--The vector, after normalizing, reduced to $(-0.599, -0.507, -0.619, 0.0)$. Plane 3 was superimposed. Points number 1-27 with the exception of number 5 can be seen lying on the simple structure plane.

VARIABLE 4 AGAINST VECTOR BELOW (SIMPLE PLANE NO. 3 SUPERIMPOSED)
 -0.549 -0.507 -0.619 0.0
 -8.062

88

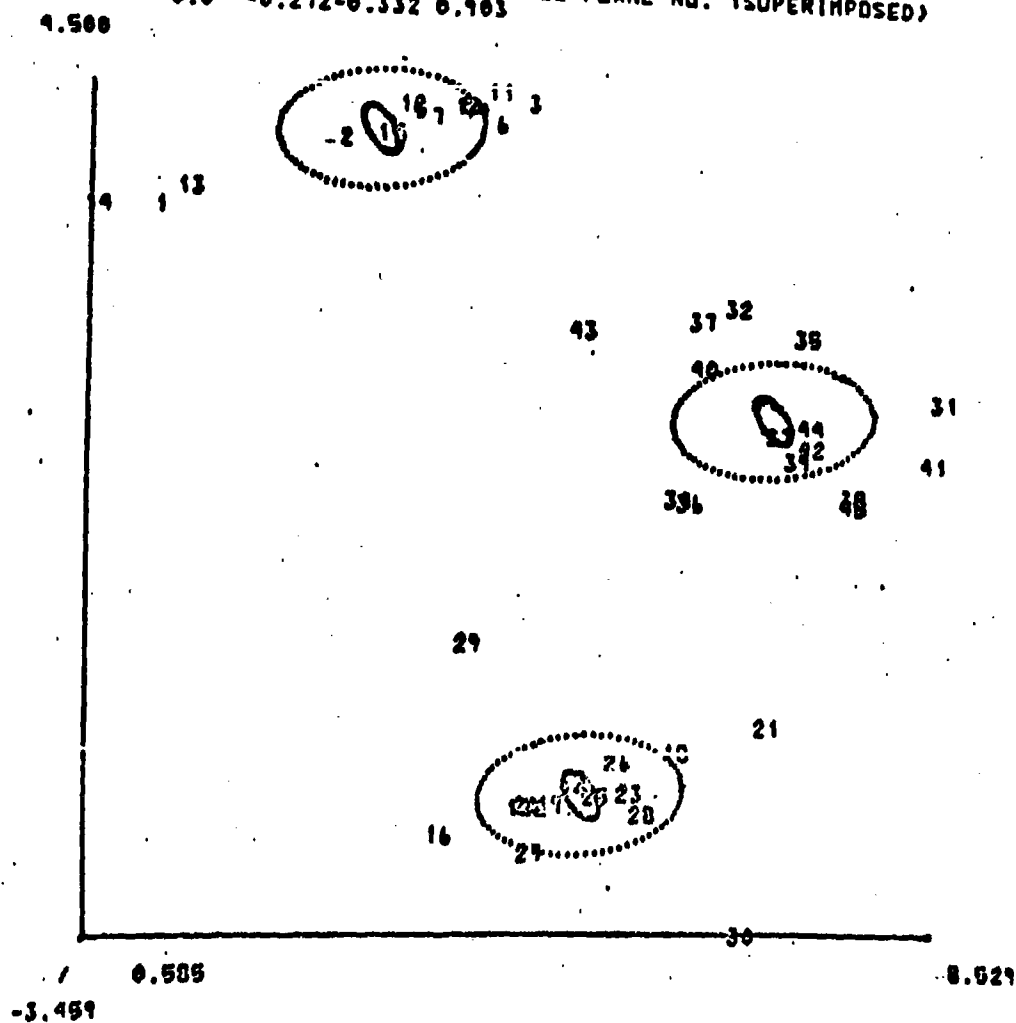


The above 4 displays pertain to each of the 4 variables against vector 3 and superimposition of simple structure plane 3. Vector 3 is the vector which transforms the raw data into simple structure plane 3. Points 1-25 with the exception of 3, 5, 6, 12 and 23 lie on this plane. Vector 3 is the normal to this simple structure plane 3. Hence any plane passing through vector 3 is orthogonal to the simple structure plane 3. When projections onto this orthogonal plane are taken, the points lying on the simple structure plane should fall within a narrow band (almost resembling a straight line) and the above 4 displays clearly bring out this fact. Variables 1, 2, 3, and 4 make 4 different planes, respectively, passing through vector 3 and all orthogonal to the simple structure plane 3. This can also be thought of as rotating a plane passing through vector 3 around the vector 3. Projections are taken when this rotating plane passes through the axis corresponding to each of the variables. Attention should also be drawn to these displays.

(e) Variable 1 against vector 3--The vector, after normalizing, reduced to (0.0, -0.272, -0.332, 0.903). Simple structure plane 1 was superimposed. Since the plane passing through vector 3 and the axis corresponding to variable 1 is not orthogonal to simple structure plane 1, we do not see points lying on a narrow band resembling a straight line in this display.

VARIABLE 1 AGAINST VECTOR BELOW (SIMPLE PLANE NO. 1 SUPERIMPOSED)
 0.0 -0.212 -0.332 0.103

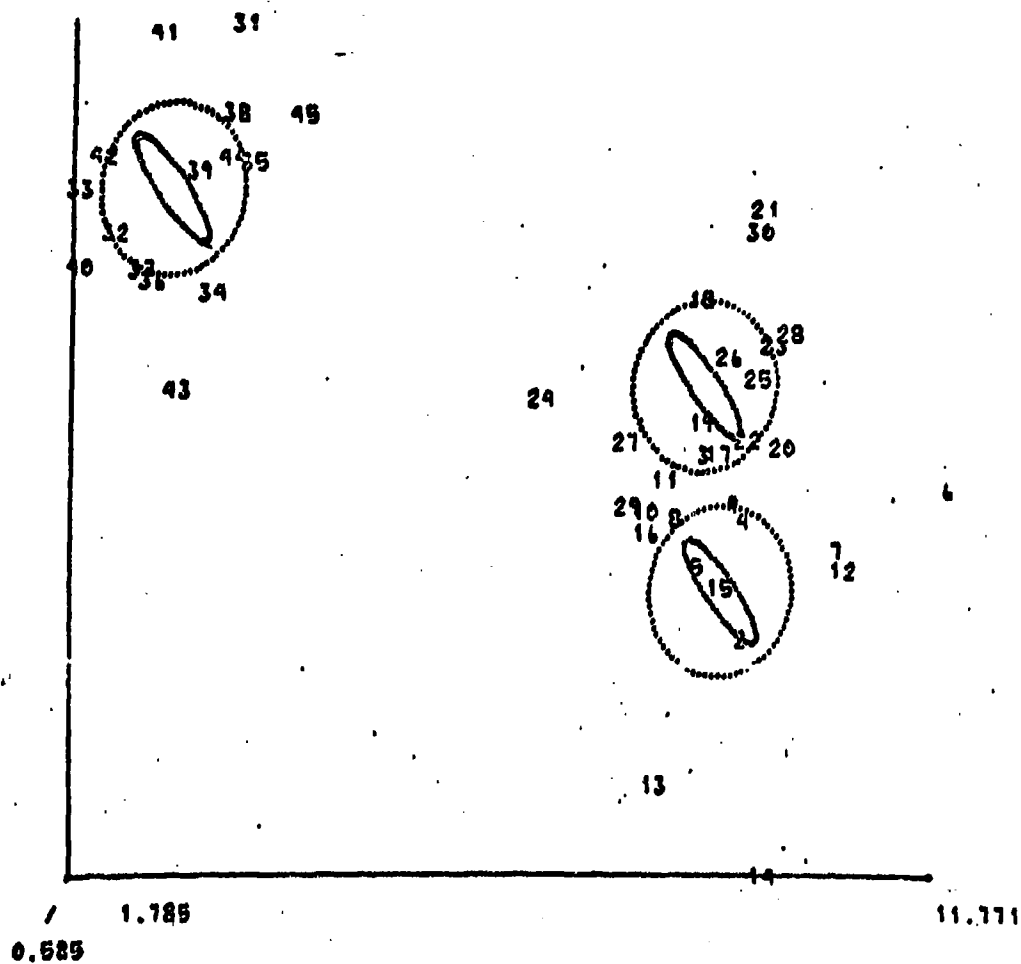
91



(f) Variable 2 against variable 1--Plane 3 was superimposed. Once again, for the reasons mentioned in (e) above, we do not see a good simple structure in this display. We simply are not looking at the ellipsoids from the proper position.

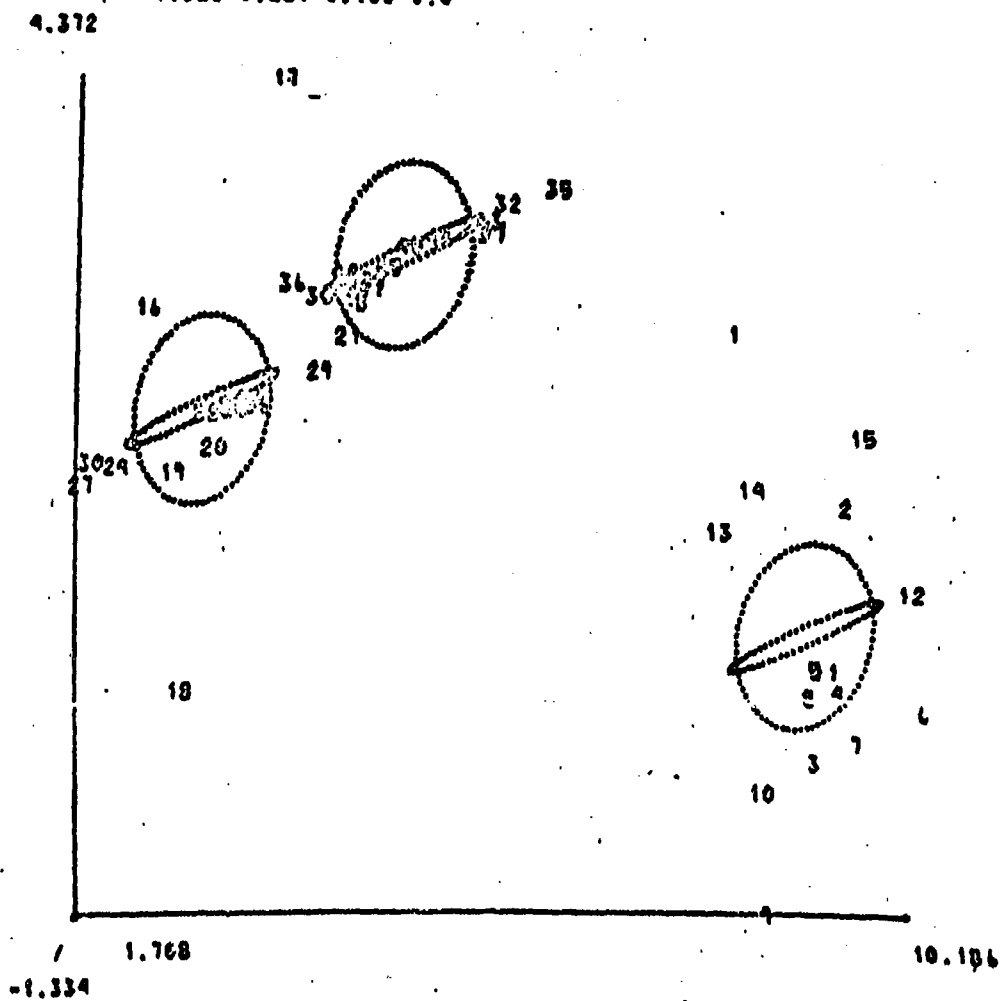
VARIABLE 2 AGAINST VECTOR BELOW. (SIMPLE PLANE NO. 3 SUPERIMPOSED)

8.529



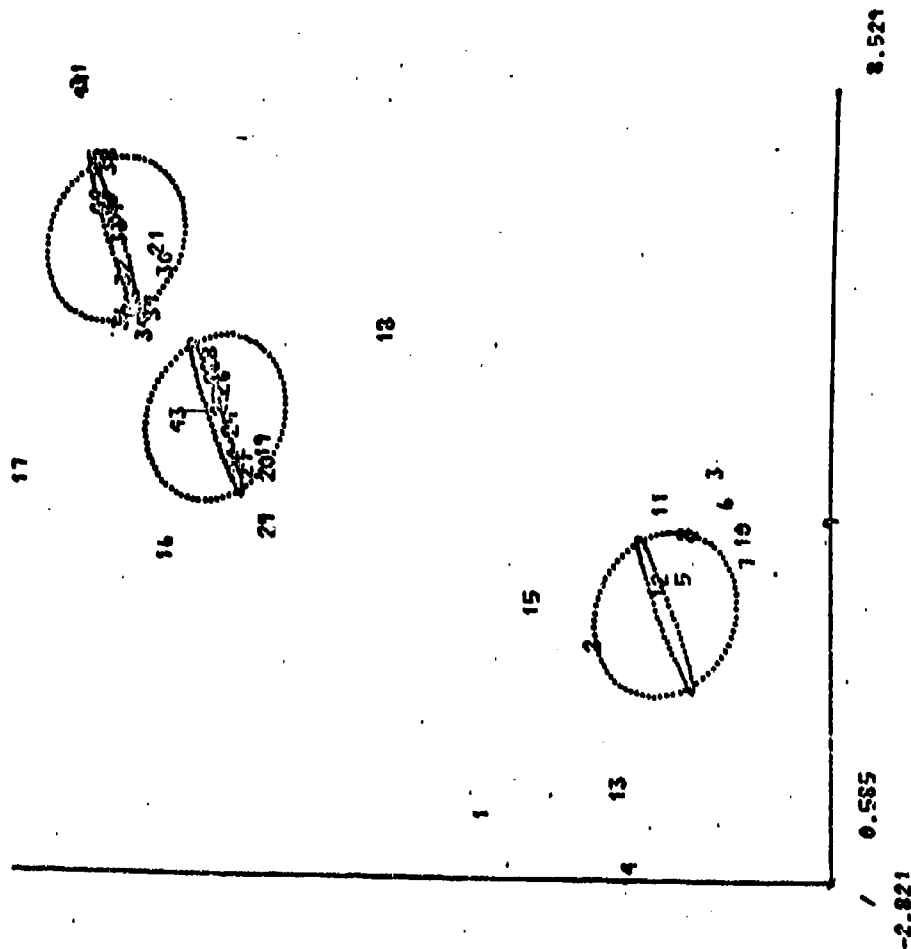
(g) Variable 4 against vector 1--The vector, after normalizing, reduced to $(-0.325, -0.289, 0.900, 0.0)$. Plane 1 was superimposed. Since any plane passing through vector 1 is orthogonal to simple structure plane 1, we expect to see a good simple structure in this display and we do. Points number 21, 22, 23, 25, 26, 28, 30, 31, 32, 33, 34, 35, 37, 38, 39, 40, 41, 42, 43, 44, 45, lie within a narrow band resembling a straight line.

VARIABLE AGAINST VECTOR RELON. (SIMPLE PLANE NO. 1 SUPERIMPOSED)
 -0.325 -0.289 0.400 0.0



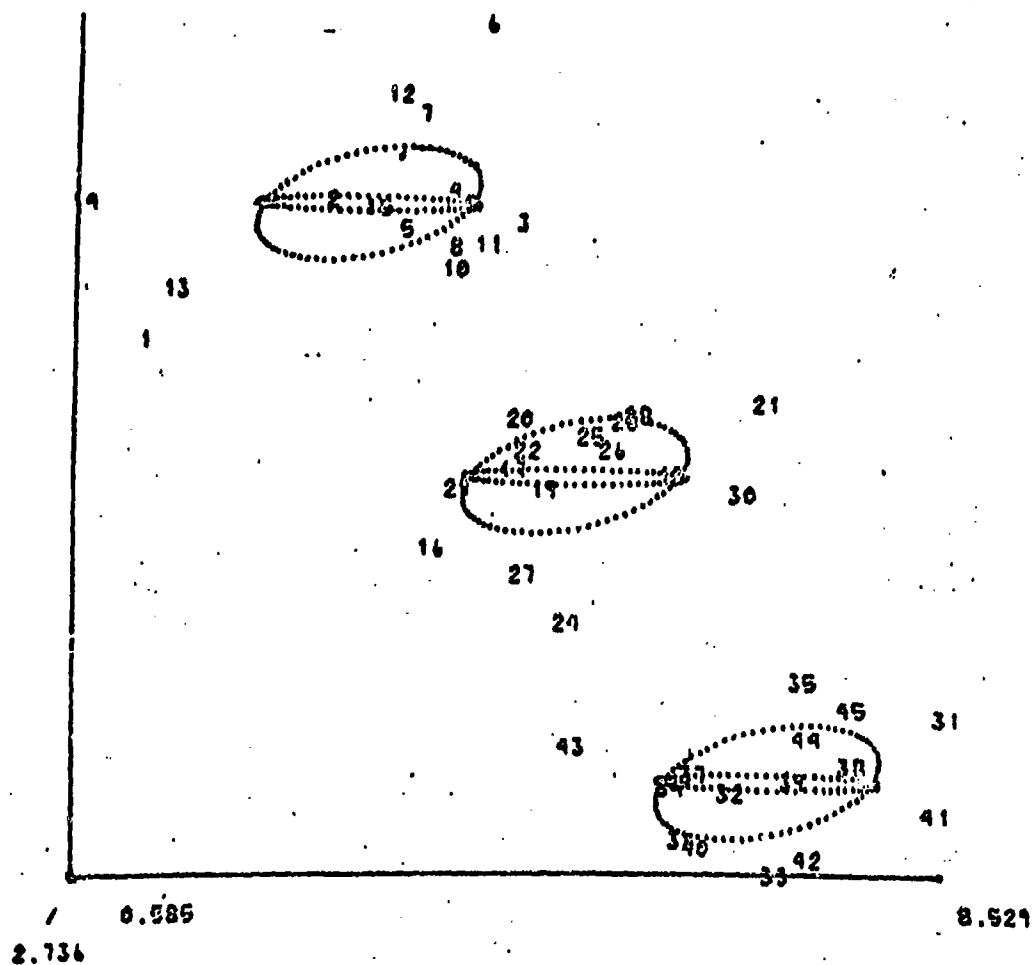
The user should find the other displays easy to understand. A good simple structure is seen only when a plane is selected which passes through a vector that transforms the raw data into a simple structure solution, and the corresponding simple structure plane is superimposed. It should, however, be noted that the clustering of points is not affected by this principle and hence, in all displays, the clusters can be easily identified.

VARIABLE AGAINST VECTOR BELOW: (SIMPLE PLANE NO. 1 SUPERIMPOSED)
 0.6 -0.284 0.100-0.325
 4.632



VARIABLE 1 AGAINST VECTOR BELOW, (SIMPLE PLANE NO. 23 SUPERIMPOSED)
 0.0 0.843-0.195 0.516

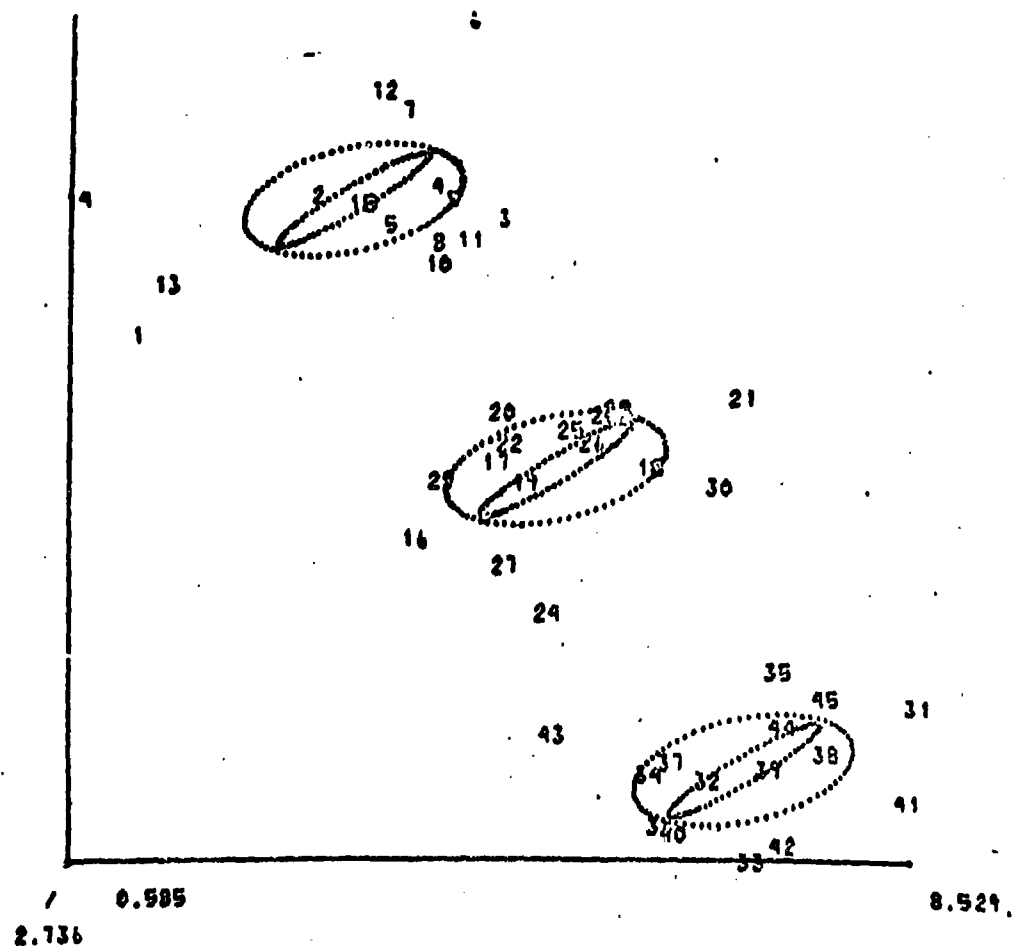
14.334



VARIABLE AGAINST VECTOR BELOW, (SIMPLE PLANE NO. 3 SUPERIMPOSED)

0.0 0.843-0.155 0.516

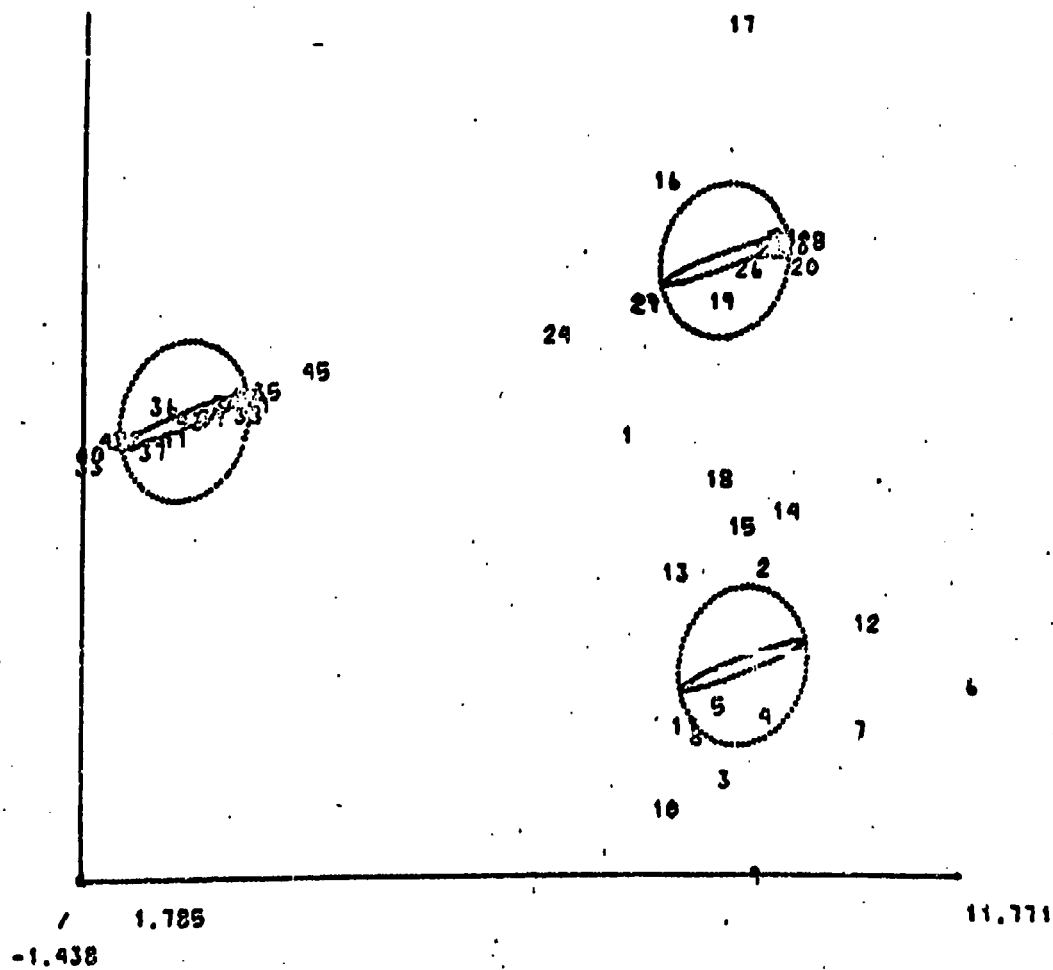
14.334



VARIABLE 2 AGAINST VECTOR BELOW, (SIMPLE PLANE NO. 1 SUPERIMPOSED)

-0.322 0.0 0.891-0.322

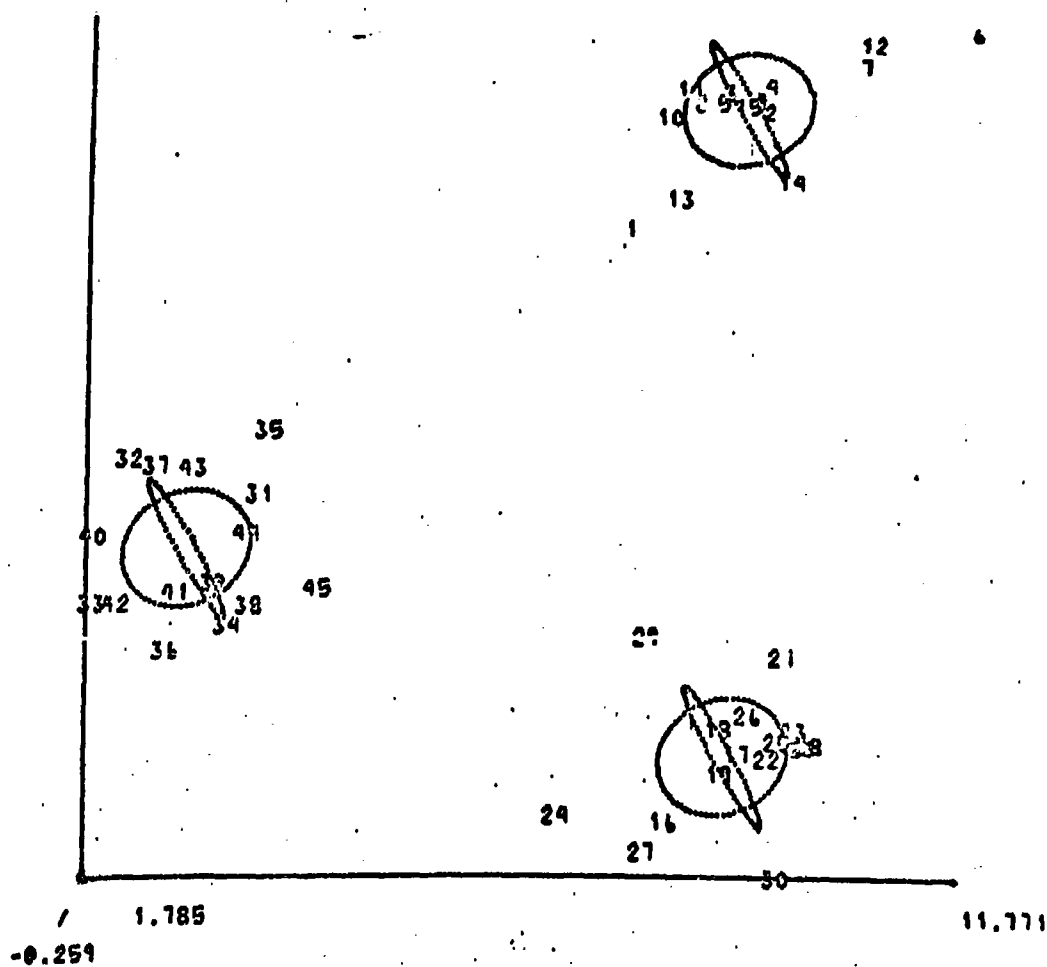
5.751



VARIABLE 2 AGAINST VECTOR BELOW (SIMPLE PLANE NO. 2 SUPERIMPOSED)

0.022 0.0 -0.208 0.151

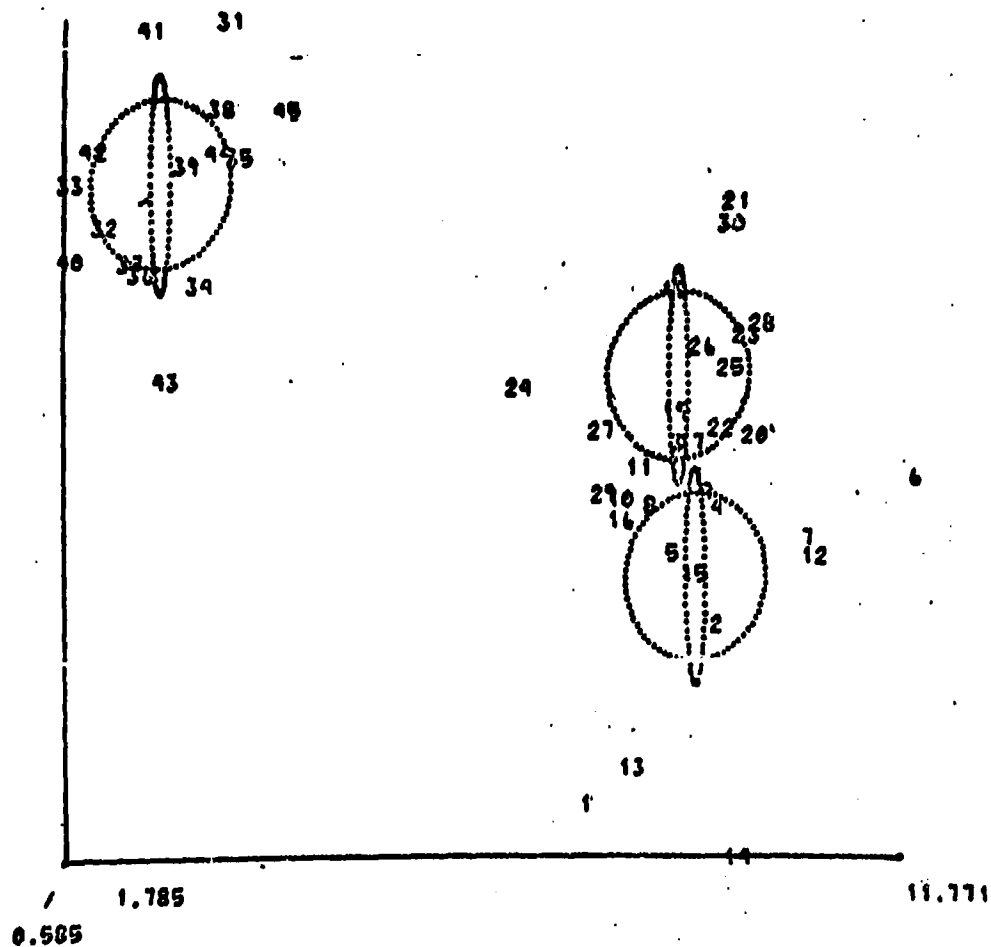
8.293



VARIABLE 2 AGAINST VECTOR BELOW. (SIMPLE PLANE NO. 2 SUPERIMPOSED)

1.000 0.0 0.0 0.0

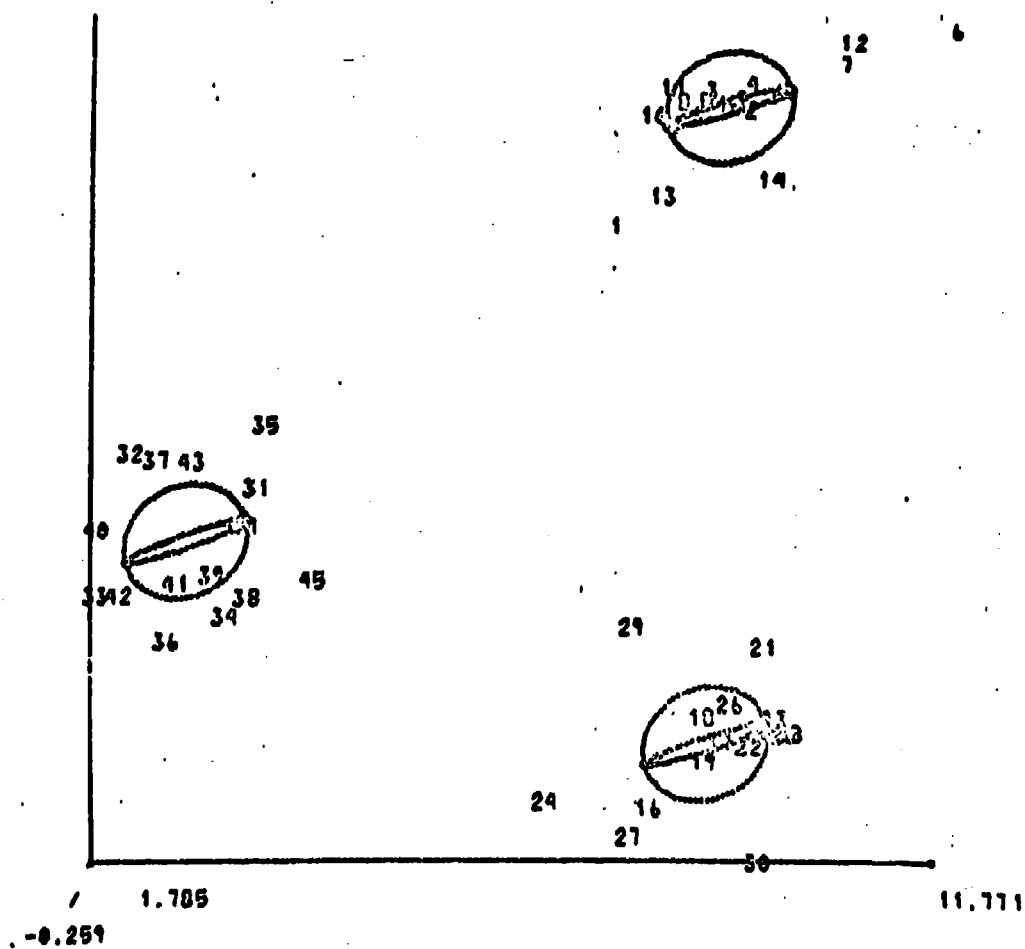
8.529



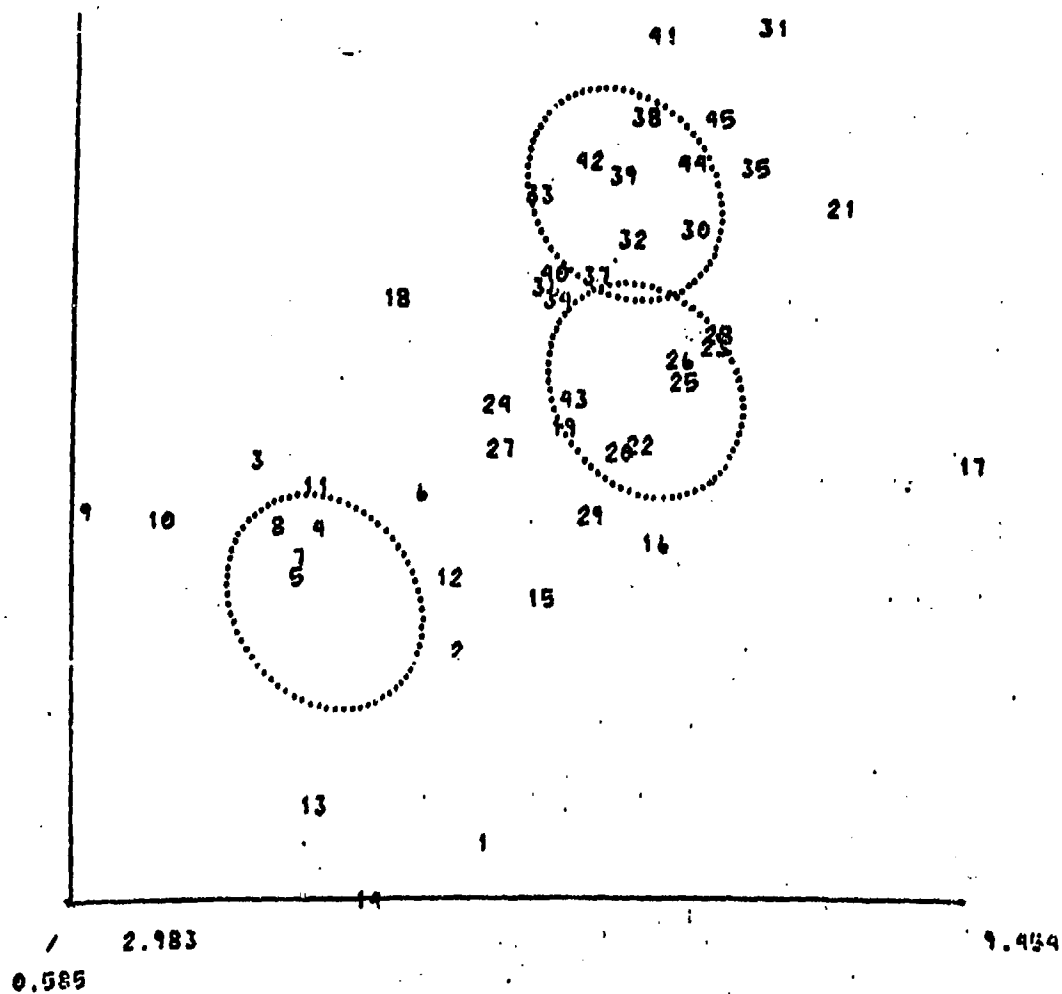
VARIABLE AGAINST VECTOR BELOW (SIMPLE PLANE NO. SUPERIMPOSED)

0.022 0.0 -0.250 0.451

8.293

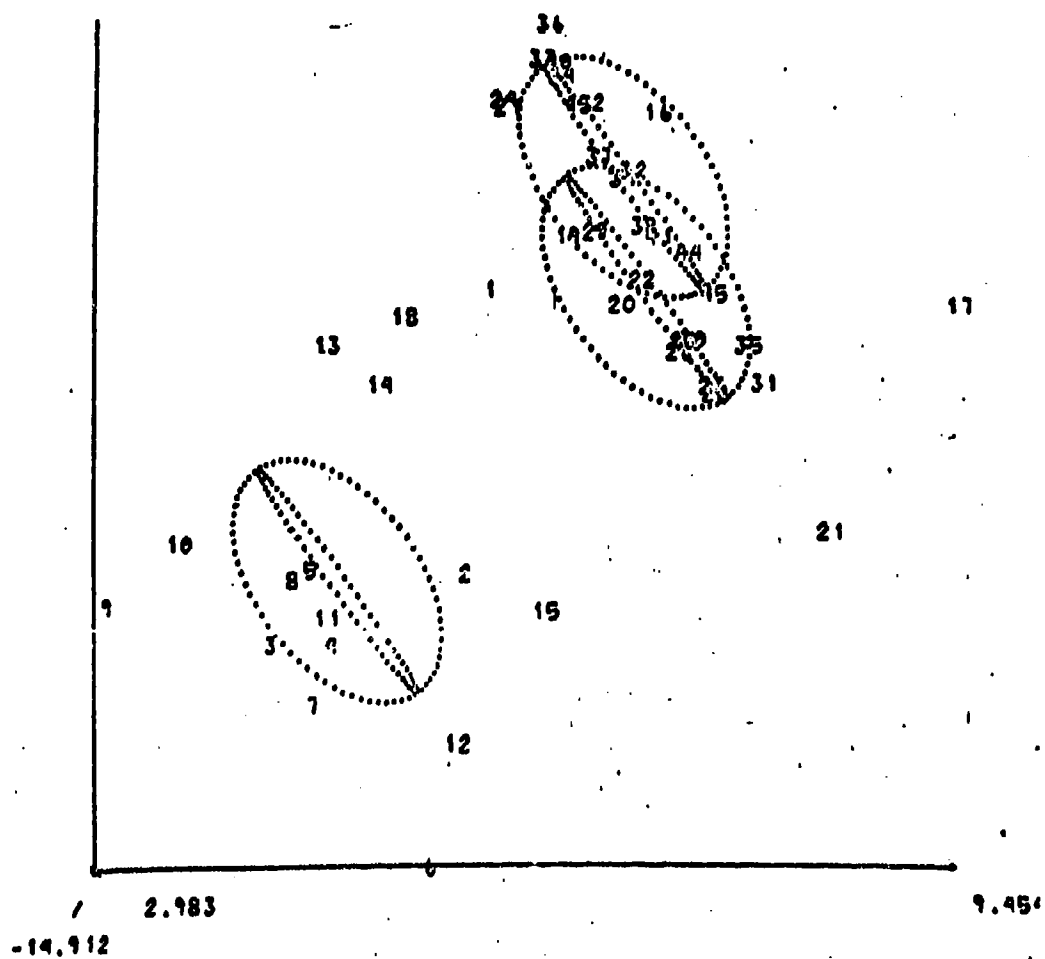


VARIABLE 3 AGAINST VECTOR
1.000 0.0 0.0 0.0
8.529



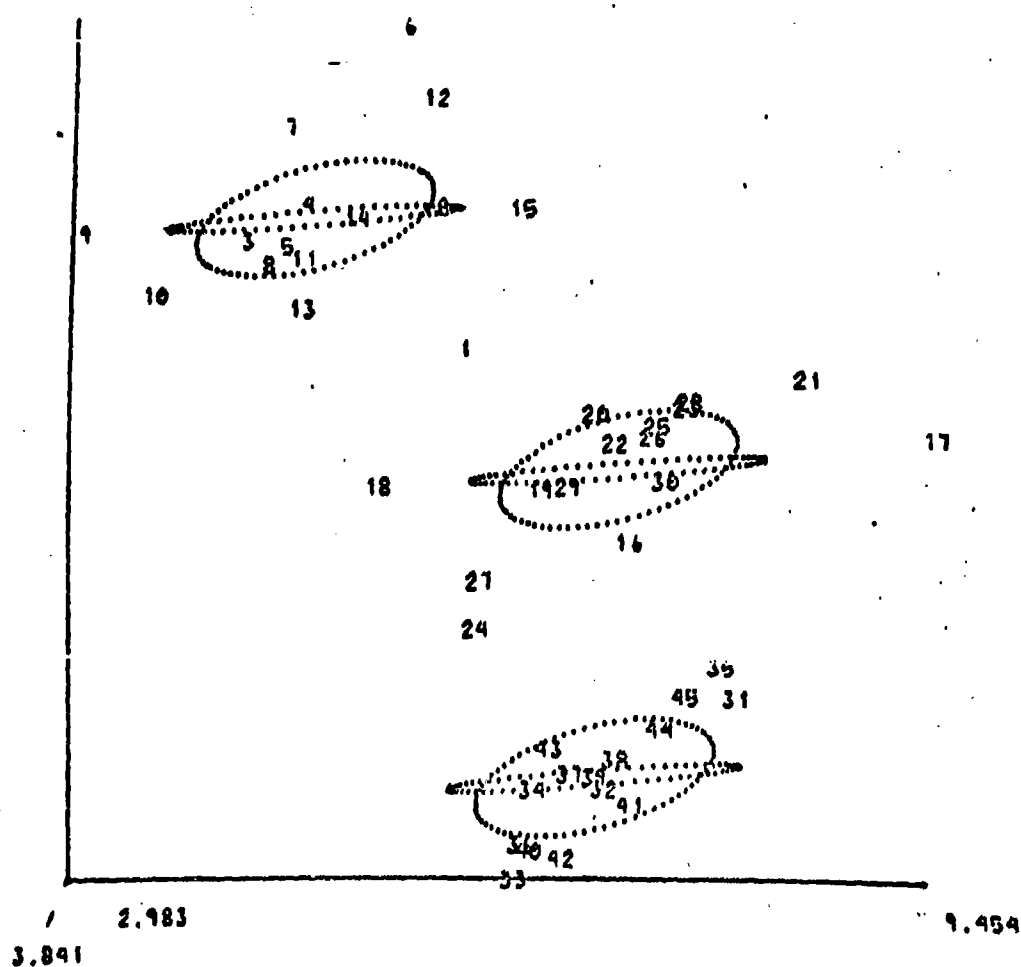
VARIABLE AGAINST VECTOR BELOW (SIMPLE PLANE NO. 1 SUPERIMPOSED)
 -0.518 -0.533 0.0 -0.518

-7.359



VARIABLE AGAINST VECTOR BELOW. (SIMPLE PLANE NO. 2 SUPERIMPOSED)
 0.012 0.853 0.0 0.522

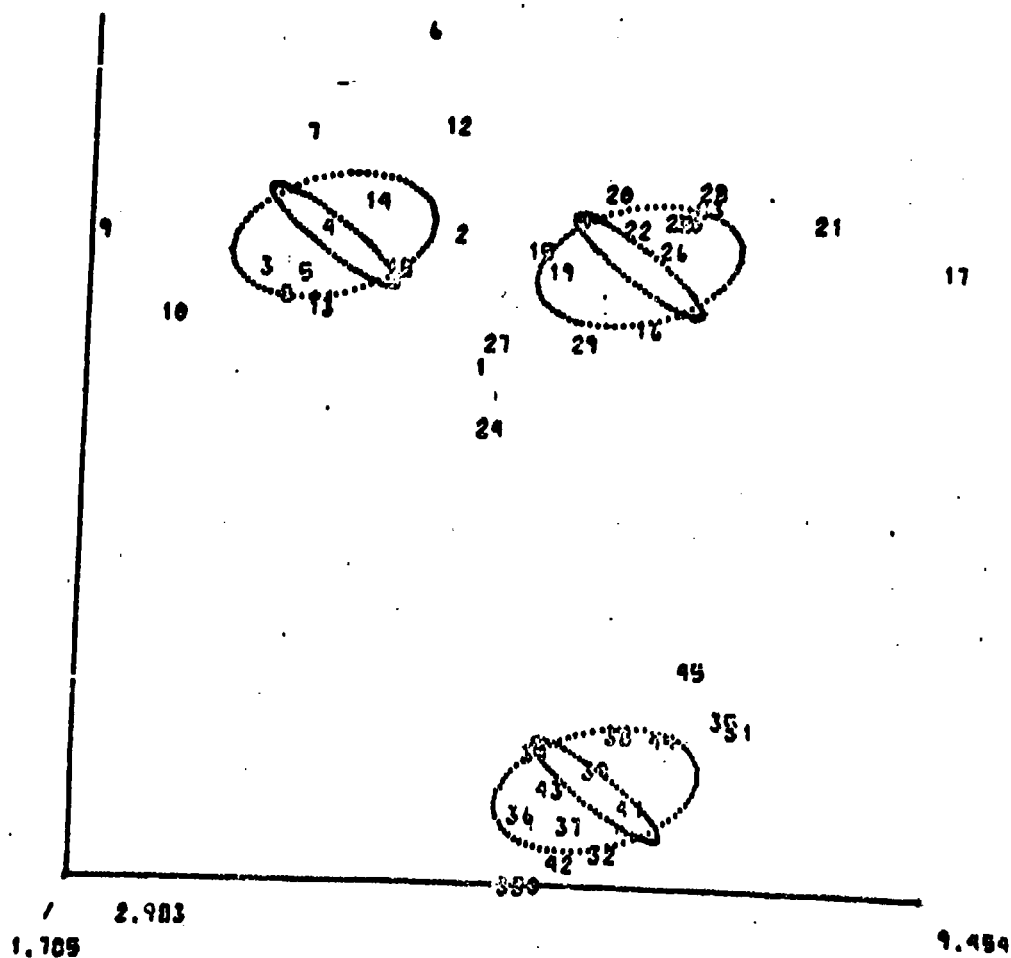
15.408



VARIABLE 3 AGAINST VECTOR BELOW, (SIMPLE PLANE NO. 35 SUPERIMPOSED)

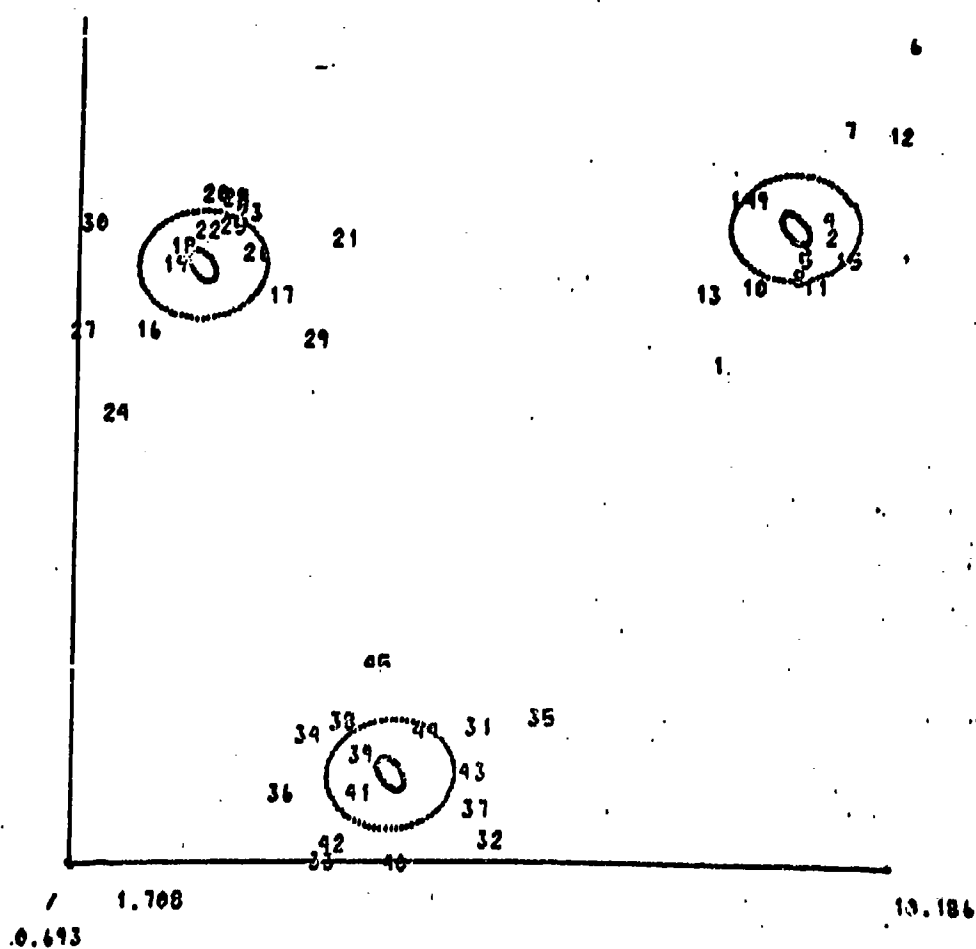
107

11.771

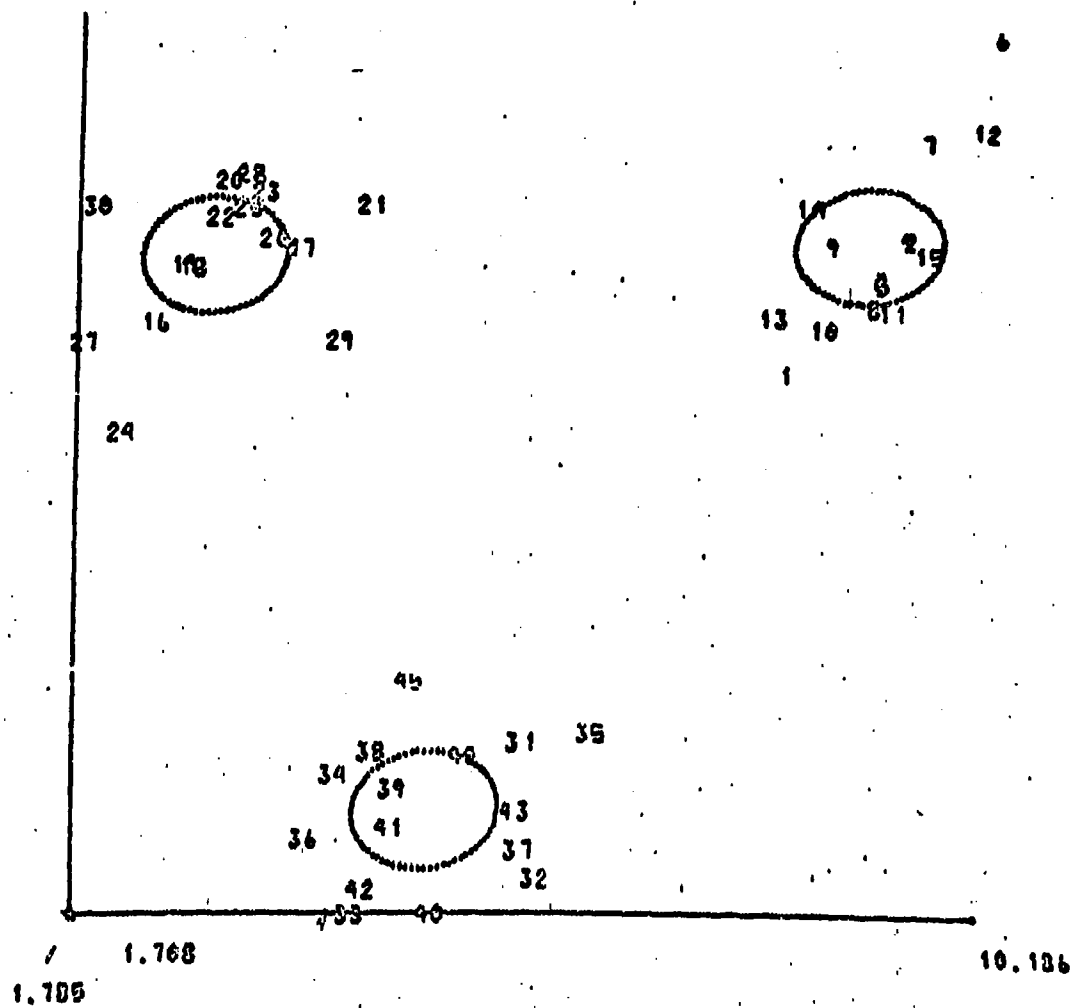


VARIABLE 4 AGAINST VECTOR BELOW, (SIMPLE PLANE NO. 2 SUPERIMPOSED)
 0.019 0.183-0.181 0.0

10.658

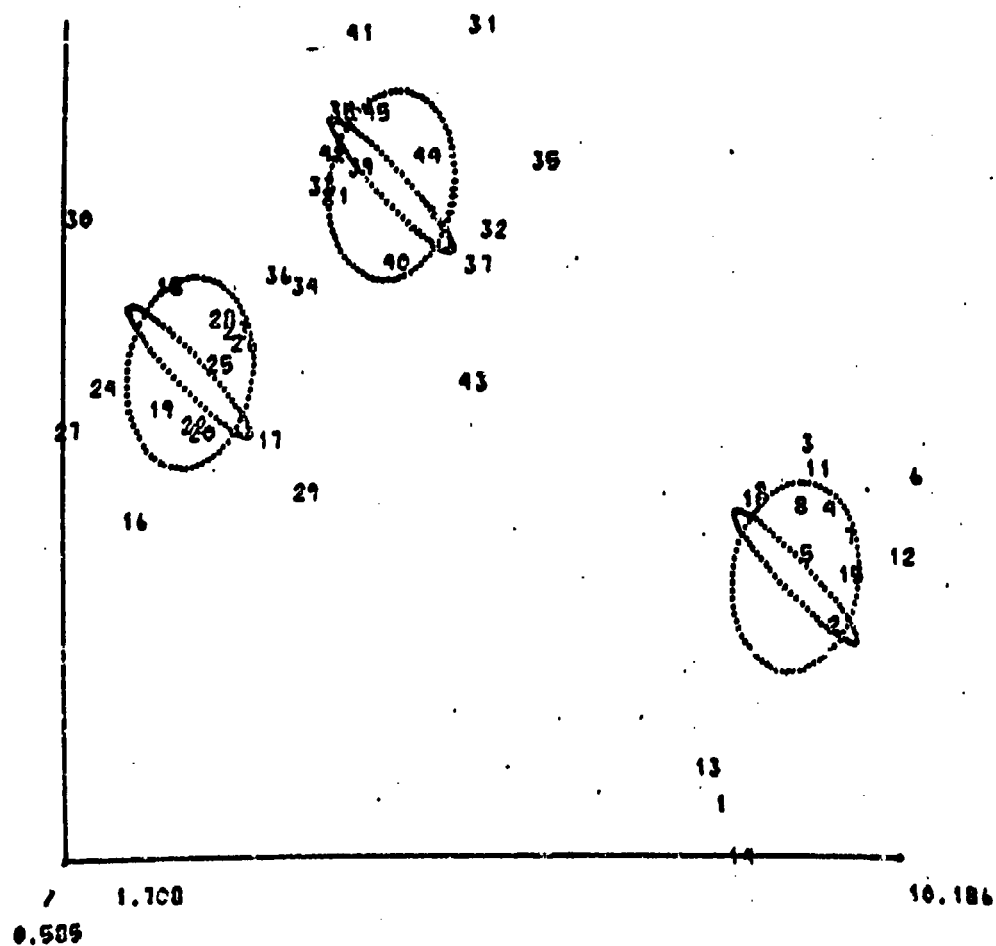


VARIABLE AGAINST VECTOR
 0.0 1.000 0.0 0.0
 11.711



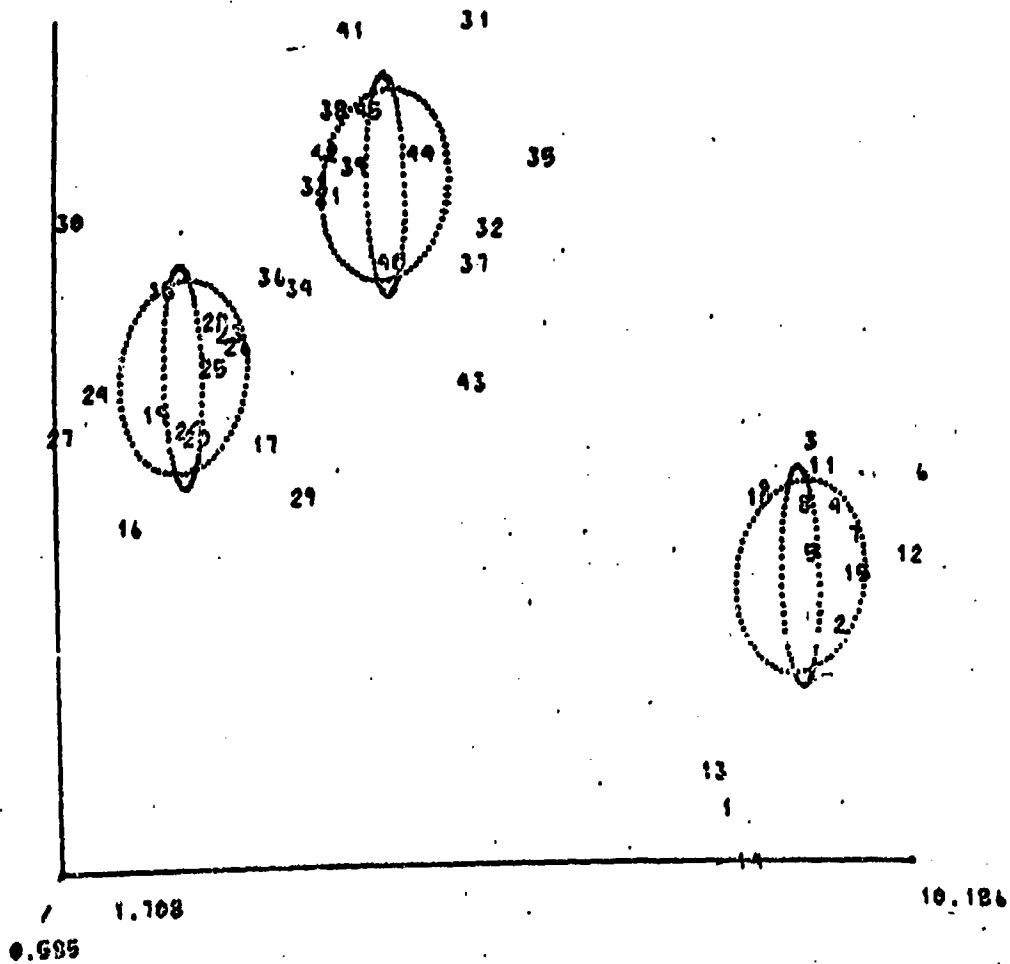
VARIABLE 4 AGAINST VECTOR BELOW, (SIMPLE PLANE NO. 1 SUPERIMPOSED)

8.529



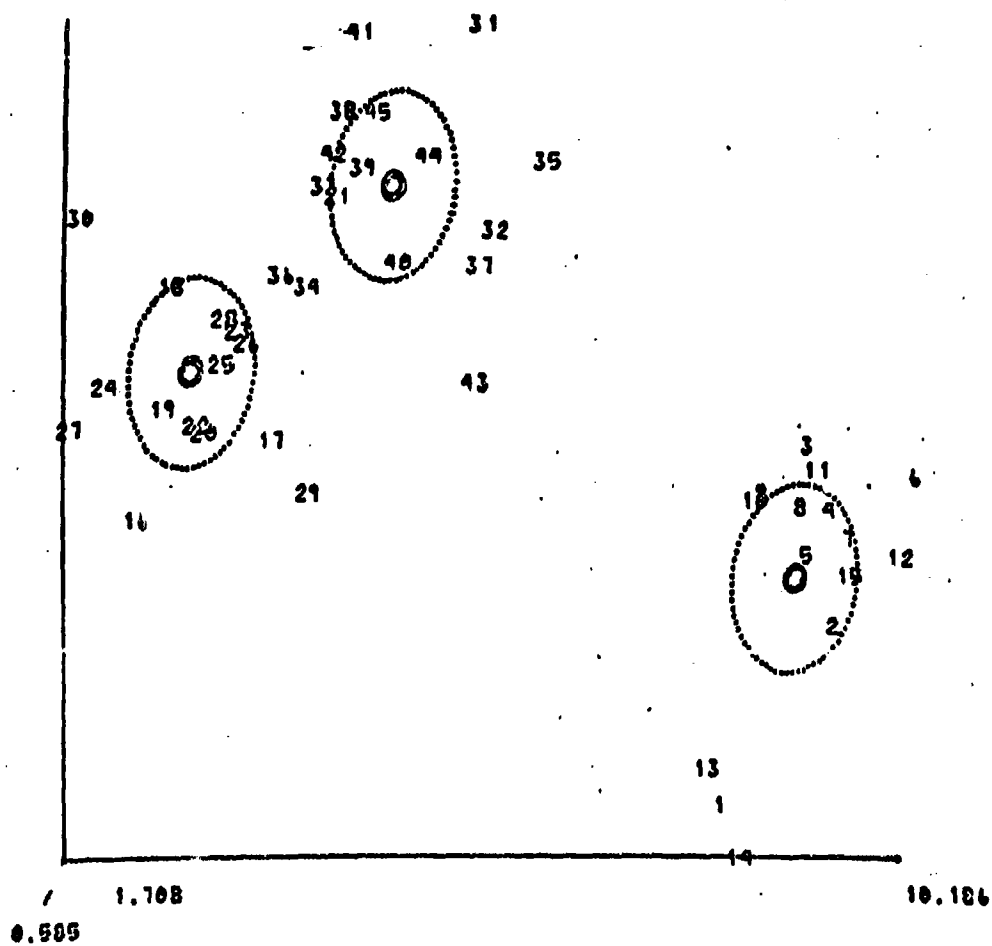
VARIABLE 4 AGAINST VECTOR BELOW, (SIMPLE PLANE NO. 2 SUPERIMPOSED)
 1.000 0.0 0.0 0.0

8.529



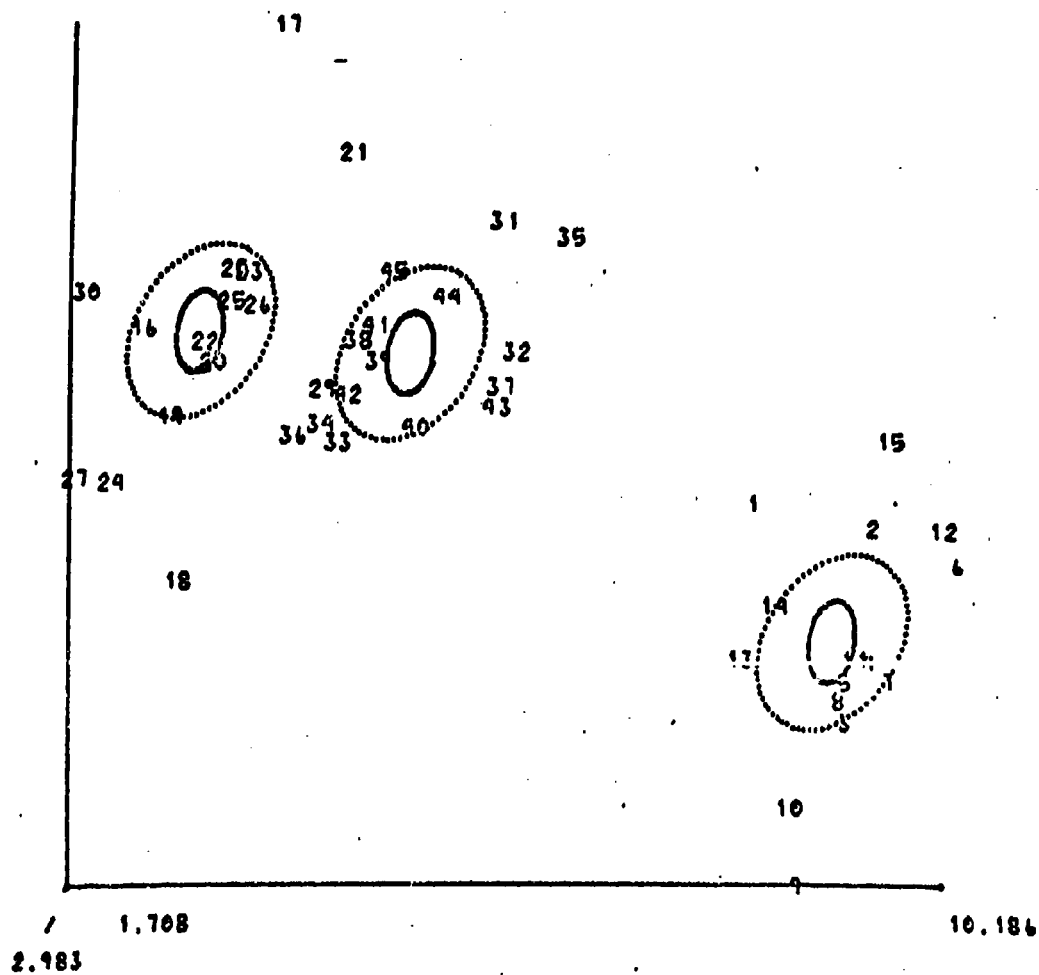
VARIABLE 4 AGAINST VECTOR BELOW, (SIMPLE PLANE NO. 3 SUPERIMPOSED)

8.521



VARIABLE 4 AGAINST VECTOR BELOM. (SIMPLE PLANE NO. 2 SUPERIMPOSED)

7.454



BIBLIOGRAPHY

- (1) Abramowitz, M., and Stegun, I., Handbook of Mathematical Functions, National Bureau of Standards, Washington, D. C., 1964.
- (2) Bargmann, R., "Signifikanzuntersuchungen der Einfachen Struktur in der Faktoren-Analyse," Mitteilungsblatt Fur Mathematische Statistik, 1954.
- (3) Bargmann, R., "Factor Analysis," unpublished address presented at the 2nd annual Goddard Computer Science Symposium, Dallas, 1966.
- (4) Bargmann, R., and Richard Graney, Tables for Significance Tests for Virtual Clusters, Technical Report No. 36, University of Georgia, Dept. of Statistics.
- (5) Bargmann, R., and Richard Graney, An Algorithm for Identifying and Testing Virtual Clusters, Technical Report No. 42, University of Georgia, Department of Statistics.
- (6) Cattell, R. B., Factor Analysis, Harper & Bros., New York, 1957.
- (7) Edwards, A. W. F. and Cavalli-Sforza, L. L., "A method for Cluster Analysis," Biometrics, 21, 1965.
- (8) Feller, W., An Introduction to Probability Theory and its Applications, Nos. I and II, John Wiley & Sons, New York, 1966.
- (9) Fortier, J. and Solomon, H., "Clustering Procedures," Multivariate Analysis, edited by P. R. Krishnaiah, Academic Press, New York, 1966.
- (10) Graney, Richard W., Tests of Concentration and Identification of Mixed Samples, unpublished doctoral dissertation submitted to University of Georgia, 1967.
- (11) Cuttman, L. "A general nonmetric technique for finding the smallest coordinate space for a configuration of points," Psychometrika, vol. 33, pp. 469-506.
- (12) Gower, J. C., "A comparison of some methods of cluster analysis," Biometrics, 1967.
- (13) Gower, J. C., "Some distance properties of latent roots and vector methods used in multivariate analysis," Biometrika, 1966.
- (14) Harman, H. H., Modern Factor Analysis, University of Chicago Press, Chicago, 1967.

- (15) Holzinger, K. and Harman, H. H., Factor Analysis, University of Chicago Press, Chicago, 1941.
- (16) Hurwitz, A., Yeaton, J. and Schaffer, R., Graphics Additions to FORTRAN (GRAF) 1968.
- (17) IBM, System/360 Scientific Subroutine Package, White Plains, N. Y., 1968.
- (18) Kendall, M. G., The Geometry of n Dimensions, C. Griffin and Co., London, 1961.
- (19) Kendall, M. G. and Stuart, A., The Advanced Theory of Statistics, Nos. I, II, and III, C. Griffin & Co., London, 1967.
- (20) King, B., "Stepwise Clustering Procedures," Journal of the American Statistical Association, Vol. 62, 1967.
- (21) Kruskal, J. B., "Multidimensional Scaling by optimizing goodness of fit to a nonmetric hypothesis," Psychometrika, 29, 1964..
- (22) Kruskal, J. B., "Nonmetric multidimensional scaling: a numerical method," Psychometrika, 29, 1964.
- (23) Ling, R. F., "Cluster Analysis," Technical Report No. 18, Yale University, Department of Statistics, 1971.
- (24) Lingoes, J. C. and Guttman, L., "Nonmetric Factor Analysis: A rank reducing alternative to linear factor analysis," Multivariate Behavioral Research, 2, 1967, p. 485-505.
- (25) Miller, R. L. and Kahn, J. S., Statistical Analysis in the Geological Sciences, John Wiley and Sons, New York, 1962.
- (26) Morrison, D., Multivariate Statistical Methods, McGraw Hill, New York, 1967.
- (27) Neyman, J. and Scott, E., "Clustering of Galaxies," Third Berkeley Symposium on Mathematical Statistics and Probability, University of California Press, 1956.
- (28) Pearson, K., Tables of the Incomplete Beta Function, Cambridge University Press, Cambridge, 1968.
- (29) Pearson, K., Tables of the Incomplete Gamma Function, Cambridge University Press, Cambridge, 1946.
- (30) Pearson, K., "On Lines and Planes of Closest Fit," Phil. Mag., 6, 1901.
- (31) Penn, L., An On-Line Statistical Computer System for Lay Usage, Technical Report No. 68, University of Georgia, Department of Statistics, 1971.

- (32) Roy, S. N., Some Aspects of Multivariate Analysis, John Wiley & Sons, New York, 1957.
- (33) Sokal, R. R. and Michener, C. D., "A Statistical Method for evaluating Systematic relationships," University of Kansas Sci. Bull., 38, 1958.
- (34) Sokal, R. R. and Sneath, P. H., Principles of Numerical Taxonomy, W. H. Freeman, San Francisco and London, 1963.
- (35) Stephenson, W., "Inverted Factor Technique," British Journal of Psychology, 1936.
- (36) Thurstone, L. L., Multiple Factor Analysis, University of Chicago Press, 1947.
- (37) Tryon, R. C., Cluster Analysis, Edwards Bros., Ann Arbor, 1939.
- (38) Williams, W. T. and Lambert, J. M., "Multivariate Methods in Plant Ecology," I. J. Ecology, 47, 1959.

APPENDIX A

Source Listings for CLUSTER Program

44-38861-50

19 JUL 1963

[illegible]


```

15N 0119 LPT=LPT
15N 0120 ENSC=L-0
15N 0121 KX=1
15N 0122 35 KX=KX+1
15N 0123 WSUM=0.0
15A 0124 KX=KX-1
15A 0125 DC 26 J=1,KX
15N 0126 WSUM=WSUM+KX(ICOR(J),LPT)
15N 0127 24 ENSC=ENSC+L-0+2.0*WSUM
15N 0128 ICR(KX)=LPT
15N 0129 36 SUMMAX=30-C
15N 0130 CC 29 I=1,NF
15N 0131 CC 25 J=1,NF
15N 0132 IF(I-EQ-ICOR(J)) GO TO 29
15N 0133 28 CONTINUE
15N 0134 DC 38 J=1,NP
15A 0135 IF(I-EQ-TEXT(J)) GO TO 29
15N 0136 38 CONTINUE
15N 0137 CASUM=0.0
15N 0138 DC 31 KX=1,KR
15N 0139 31 CASUM=CASUM+KX(ICOR(KX),I)
15N 0140 IF(CASUM-ALT-SUMMAX) GO TO 22
15N 0141 SUMMAX=CASUM
15N 0142 LPT=I
15N 0143 32 CONTINUE
15N 0144 29 CONTINUE
15A 0145 IF(SUMMAX-LE-0-G1 G3 TO 56
15N 0146 56 FORMAT('POINT ',I3,' IN CLUSTER IS POINT ',I3,'. CORR IS',F8.4)
15N 0147 ALF=TEST(RMAX,KKX)
15N 0148 57 WRITE(6,57) ALF
15N 0149 57 FORMAT(' ALF= ',F8.4)
15N 0150 IF(ALF-LE-ALFCC) GO TO 33
15N 0151 IF(R=EC-2) ICL(LPT)=2
15N 0152 IF(IALF-GT-ALFCOR)AND-(ALF-LE-ALFEXT)) GO TO 34
15N 0153 CC TC 95
15N 0154 33 ICL(LPT)=1
15N 0155 CL(LPT,KCL)=ALF
15N 0156 CC TO 35
15N 0157 34 IF(ICL(LPT)-EC-1) GO TO 37
15N 0158 ICL(LPT)=2
15N 0159 37 CONTINUE
15N 0160 CL(LPT,KCL)=ALF
15N 0161 IE=IE+1
15N 0162 TEXT(IE)=LPT
15N 0163 CC TC 36
15N 0164 96 KCL=KCL+1
15N 0165 IF(KCL-EG-21) GO TO 95
15N 0166 CC TO 97
15N 0167 95 CONTINUE
15N 0168 DC 02 KX=1,2
15N 0169 JA=10*(KX-1)+1
15N 0170 JC=JA+9
15N 0171 WRITE(6,45)
15N 0172
15N 0173
15N 0174
15N 0175
15N 0176
15N 0177
15N 0178
15N 0179
15N 0180
15N 0181
15N 0182
15N 0183

```

```

15N 0184      45 FORMAT(1,'PI NO.',I2X,'CLUSTERS')
15N 0185      IFML(EQ.2) GO TO 113
15N 0186      WRITE(6,112)
15N 0187      112 FORMAT(13X,'NC. 1',I4X,'NC. 2',I4X,'NC. 3',I4X,'NO. 4',I4X,
15N 0188      1 'NC. 5',I4X,'NC. 6',I4X,'NO. 7',I4X,'NO. 8',I4X,
2 'NC. 9',I4X,'NO. 10',I4X)
      GO TO 115
15N 0189      113 WRITE(6,114)
15N 0190      114 FORMAT(13X,'NC. 11',I4X,'NC. 12',I4X,'NC. 13',I4X,'NO. 14',I4X,
15N 0191      1 'NC. 15',I4X,'NC. 16',I4X,'NO. 17',I4X,'NO. 18',I4X,
2 'NC. 19',I4X,'NC. 20',I4X)
      GO TO 115
15N 0192      115 CONTINUE
15N 0193      WRITE(6,146)
15N 0194      146 FGS=1/(1-Q)
15N 0195      CG=47/(1-1P)
15N 0196      47 WRITE(6,48) 1, (CL(I,J),J=JA,JB)
15N 0197      48 FCG=1/(1-Q) 1X,E3.6X -1012X(F8.6)
15N 0198      48 CONTINUE
15N 0199      DC 401 I=1,200
15N 0200      401 ICF(I)=0
15N 0201      DC 402 I=1,20
15N 0202      402 NCU(I)=0
15N 0203      CC 403 J=1,20
15N 0204      DC 404 I=1,1P
15N 0205      IF(1CL(I,J).GT.0.0) AND.(CL(I,J).LE.ALFCOR) GO TO 405
15N 0206      GO TO 404
15N 0207      405 NCU(I)=NCC(I)+1
15N 0208      404 CONTINUE
15N 0209      403 CONTINUE
15N 0210      DO 407 I=1,20
15N 0211      407 NCU(I)=0
15N 0212      DO 409 K=1,20
15N 0213      409 NCU(I)=NAX-1
15N 0214      DC 406 I=1,20
15N 0215      DO 408 J=1,20
15N 0216      408 IF(1-EC.NC(I,J)) GO TO 406
15N 0217      406 C=1/AGE
15N 0218      IF(NC(I)) LE.NMAX GO TO 406
15N 0219      406 CONTINUE
15N 0220      406 NCU(I)=NAX-NC(I)
15N 0221      IZ=1
15N 0222      406 CONTINUE
15N 0223      NCU(I)=IZ
15N 0224      409 CONTINUE
15N 0225      411 CONTINUE
15N 0226      J=0
15N 0227      KPT=0
15N 0228      419 J=J+1
15N 0229      KPT=KPT+1
15N 0230      IF(NC(NH(I)) LE.2) GO TO 99
15N 0231      419 CONTINUE
15N 0232      DO 413 M=1,1P
15N 0233      413 IF(1CL(M,NH(I)).GT.0.0) GO TO 414
15N 0234      GO TO 413
15N 0235      414 ICF=ICF+1
15N 0236      413 CONTINUE
15N 0237      413 CONTINUE
15N 0238      413 CONTINUE
15N 0239      413 CONTINUE
15N 0240      413 CONTINUE
15N 0241      413 CONTINUE
15N 0242      413 CONTINUE

```



```

15M 0243 IF((J,J).EQ.2).OR.((J,J).EQ.3)) GO TO 501
15M 0244 DC 415 JX=1.20
15M 0245 IA TC 3
15M 0246 IF((CRJ,J).EQ.NOR(J,J)) GO TO 415
15M 0247 DC 416 J=1.4P
15M 0248 IF((C(L).NOT(J,J)).NE.0.0).AND.((C(L).NOR(J,J)).NE.0.0))IMATC=IMATC+1
15M 0249
15M 0250 416 CONTINUE
15M 0251 ICTI=ICTI-2
15M 0252 IF((IPATC.GT.2).AND.((IMATC.EE.0))) GO TO 417
15M 0253 IF((IPATC.GT.2).AND.((IMATC.EE.0))) GO TO 417
15M 0254 GO TO 415
15M 0255 417 GO 415 IAX=1.4P
15M 0256 IF((C(L).NOR(CR,C(J,J)).LE.24FCOR).AND.((C(L).NOR(J,J).GT.0.0)).AND.
15M 0257 IIG(IAX).EQ.0)) IC(IX)=NPT
15M 0258 IIG(IAX).EQ.0)) IC(IX)=NPT
15M 0259
15M 0260 418 CONTINUE
15M 0261 418 CONTINUE
15M 0262 415 CONTINUE
15M 0263 502 CONTINUE
15M 0264 NL=0
15M 0265 DO 420 IQ=1,NP
15M 0266 IF((C(IQ).EQ.NPT) NC=NO+1
15M 0267 420 CONTINUE
15M 0268 IF(NC.GT.2) GO TO 999
15M 0269 DO 421 IQ=1,NP
15M 0270 IF((C(IQ).EQ.NPT) IC(IQ)=0
15M 0271 421 CONTINUE
15M 0272 555 CONTINUE
15M 0273 GO TO 429
15M 0274 99 CONTINUE
15M 0275 DO 425 I=1,20
15M 0276 DC 425 I=1,20
15M 0277 425 NIG(I)=0
15M 0278 DO 422 I=1,20
15M 0279 DC 422 I=1,20
15M 0280 DO 423 J=1,NP
15M 0281 IF((C(J).EQ.1) NIG(I)=NIG(I)+1
15M 0282 423 CONTINUE
15M 0283 423 CONTINUE
15M 0284 422 CONTINUE
15M 0285 WRT(16,4301) (NIG(I),I=1,20)
15M 0286 430 FORMAT(' ',20(I2,2X))
15M 0287 AG=0
15M 0288 DO 424 I=1,NP
15M 0289 IF((C(I).LT.MC) GO TO 424
15M 0290 MC=IG(I)
15M 0291 424 CONTINUE
15M 0292 NAI=0
15M 0293 DO 501 I=1,20
15M 0294 DC 501 I=1,20
15M 0295 501 NAI=NIG(NIG(I))
15M 0296 WRT(16,552) (NAI,
15M 0297 502 FORMAT('NO. OF PTS ASSIGNED TO GRUPO
15M 0298 I55',533)
15M 0299 DO 551 I=1,NP
15M 0300 551 REACT(I,IC5) (X(I),J=1,NV)
15M 0301 REACT(I,
15M 0302 GO TO 444
15M 0303 444 CONTINUE
15M 0304 REACT(I,5
15M 0305 CALL COR2
15M 0306 NV = (NV+(NV+1))/2
15M 0307 DO 10000 I=1,NV
15M 0308
15M 0309
15M 0310
15M 0311
15M 0312
15M 0313
15M 0314
15M 0315
15M 0316
15M 0317
15M 0318
15M 0319
15M 0320
15M 0321
15M 0322
15M 0323
15M 0324
15M 0325
15M 0326
15M 0327
15M 0328
15M 0329
15M 0330
15M 0331
15M 0332
15M 0333
15M 0334
15M 0335
15M 0336
15M 0337
15M 0338
15M 0339
15M 0340
15M 0341
15M 0342
15M 0343
15M 0344
15M 0345
15M 0346
15M 0347
15M 0348
15M 0349
15M 0350
15M 0351
15M 0352
15M 0353
15M 0354
15M 0355
15M 0356
15M 0357
15M 0358
15M 0359
15M 0360
15M 0361
15M 0362
15M 0363
15M 0364
15M 0365
15M 0366
15M 0367
15M 0368
15M 0369
15M 0370
15M 0371
15M 0372
15M 0373
15M 0374
15M 0375
15M 0376
15M 0377
15M 0378
15M 0379
15M 0380
15M 0381
15M 0382
15M 0383
15M 0384
15M 0385
15M 0386
15M 0387
15M 0388
15M 0389
15M 0390
15M 0391
15M 0392
15M 0393
15M 0394
15M 0395
15M 0396
15M 0397
15M 0398
15M 0399
15M 0400
15M 0401
15M 0402
15M 0403
15M 0404
15M 0405
15M 0406
15M 0407
15M 0408
15M 0409
15M 0410
15M 0411
15M 0412
15M 0413
15M 0414
15M 0415
15M 0416
15M 0417
15M 0418
15M 0419
15M 0420
15M 0421
15M 0422
15M 0423
15M 0424
15M 0425
15M 0426
15M 0427
15M 0428
15M 0429
15M 0430
15M 0431
15M 0432
15M 0433
15M 0434
15M 0435
15M 0436
15M 0437
15M 0438
15M 0439
15M 0440
15M 0441
15M 0442
15M 0443
15M 0444
15M 0445
15M 0446
15M 0447
15M 0448
15M 0449
15M 0450
15M 0451
15M 0452
15M 0453
15M 0454
15M 0455
15M 0456
15M 0457
15M 0458
15M 0459
15M 0460
15M 0461
15M 0462
15M 0463
15M 0464
15M 0465
15M 0466
15M 0467
15M 0468
15M 0469
15M 0470
15M 0471
15M 0472
15M 0473
15M 0474
15M 0475
15M 0476
15M 0477
15M 0478
15M 0479
15M 0480
15M 0481
15M 0482
15M 0483
15M 0484
15M 0485
15M 0486
15M 0487
15M 0488
15M 0489
15M 0490
15M 0491
15M 0492
15M 0493
15M 0494
15M 0495
15M 0496
15M 0497
15M 0498
15M 0499
15M 0500
15M 0501
15M 0502
15M 0503
15M 0504
15M 0505
15M 0506
15M 0507
15M 0508
15M 0509
15M 0510
15M 0511
15M 0512
15M 0513
15M 0514
15M 0515
15M 0516
15M 0517
15M 0518
15M 0519
15M 0520
15M 0521
15M 0522
15M 0523
15M 0524
15M 0525
15M 0526
15M 0527
15M 0528
15M 0529
15M 0530
15M 0531
15M 0532
15M 0533
15M 0534
15M 0535
15M 0536
15M 0537
15M 0538
15M 0539
15M 0540
15M 0541
15M 0542
15M 0543
15M 0544
15M 0545
15M 0546
15M 0547
15M 0548
15M 0549
15M 0550
15M 0551
15M 0552
15M 0553
15M 0554
15M 0555
15M 0556
15M 0557
15M 0558
15M 0559
15M 0560
15M 0561
15M 0562
15M 0563
15M 0564
15M 0565
15M 0566
15M 0567
15M 0568
15M 0569
15M 0570
15M 0571
15M 0572
15M 0573
15M 0574
15M 0575
15M 0576
15M 0577
15M 0578
15M 0579
15M 0580
15M 0581
15M 0582
15M 0583
15M 0584
15M 0585
15M 0586
15M 0587
15M 0588
15M 0589
15M 0590
15M 0591
15M 0592
15M 0593
15M 0594
15M 0595
15M 0596
15M 0597
15M 0598
15M 0599
15M 0600
15M 0601
15M 0602
15M 0603
15M 0604
15M 0605
15M 0606
15M 0607
15M 0608
15M 0609
15M 0610
15M 0611
15M 0612
15M 0613
15M 0614
15M 0615
15M 0616
15M 0617
```

```

15N 0307      DSPRGC(I) = CSPRGC(I)/(ANIC-NC)
15N 0308      WRITE (4,3061) I,DSPRGC(I)
15N 0309      10000 CONTINUE
15N 0310      DC 11010 I=1,NV
15N 0311      11010 VARI(I) = CSPEC(I)
15N 0312      DC 20000 I=1,NC
15N 0313      IF (NIC(I)-EC(I)) GO TO 20000
15N 0314      DC 20000 J=1,NV
15N 0315      CGSUM(I,J) = CGSUM(I,J)/NIC(I)
15N 0316      20000 CONTINUE
15N 0317      DC 20001 I=1,NC
15N 0318      WRITE(6,11020) (CGSUM(I,J),J=1,NV)
15N 0319      WRITE(7,11020) (CGSUM(I,J),J=1,NV)
15N 0320      20001 CONTINUE
15N 0321      11020 FCPRAT(I,X,E7,3)
15N 0322      CALL CNFSG(CSPRGC,NV,I,E-10,IER)
15N 0323      DC 10001 I=1,NP
15N 0324      IF (IC(I)-EC(I)) GO TO 10003
15N 0325      KK = IC(I)
15N 0326      DC 10002 J=1,NV
15N 0327      CCA = XI(I,J)
15N 0328      10002 OV(I,J) = CCF - CGSUM(KK,J)
15N 0329      GC TO 10005
15N 0330      10003 DC 10004 J=1,NV
15N 0331      COK = X(I,I,J)
15N 0332      10004 GVL(J) = COK - CGSUM(J)
15N 0333      10005 CALL ENIOS(EV,I,NV,CSPRGC,-1,IER)
15N 0334      DC 10007 J=1,NV
15N 0335      XI(I,J) = GVL(J)
15N 0336      10007 CONTINUE
15N 0337      WRITE(15,2) I,X(I,J),J=1,NV)
15N 0338      10001 CONTINUE
15N 0339      NV = INVRNV(I)/2
15N 0340      DC 11000 I=1,NV
15N 0341      SPROD(I) = ESPROD(I)
15N 0342      NV = NV+2
15N 0343      REWIND 15
15N 0344      KFACT = NV
15N 0345      USE WITH SUBROUTINES FACT,WRIT,UTERN,IRSLD
15N 0346      C
15N 0347      NGUT06
15N 0348      NIN=5
15N 0349      C
15N 0350      C
15N 0351      C
15N 0352      C
15N 0353      C
15N 0354      C
15N 0355      C
15N 0356      C
15N 0357      C
15N 0358      C
15N 0359      C
15N 0360      C
15N 0361      C

```

```

FCAN 000
FCANAD00
FCAN3000
FCAN 009
FCAN 011

```

```

15N 0362      DC 1002 J01,AP
15N 0363      IF ERTJCC 1003,1003,1004
15N 0364      1003 GJJC # 0.
15N 0365      GO TO 1002
15N 0366      1004 DJJC # S0M1 ERTJCC
15N 0367      1002 CONTINUE
15N 0368      ASOL # 0
15N 0369      CM = OFLOAT(KFACT-11/2,0)
15N 0370      SATMTA = GSCRT(SETAP16,CO/DFLOAT(MP-KFACT*11,0-500,DN1))
15N 0371      IF (SATMTA-1,0-1) SATMTA = 0.1
15N 0372      DC 1001 J01,AP
15N 0373      IF ERTJCC 1001,1001,1005
15N 0374      1005 ICG # 0
15N 0375      DC 1006 # 1,KFACT
15N 0376      IALP # 1K-1,MP # 1
15N 0377      1004 FRIMK # FRIALP</O1K
15N 0378      BND # 0.
15N 0379      1021 DC 1005 J01,AP
15N 0380      SUM # 0.
15N 0381      DO 1007 KAL,KFACT
15N 0382      1007 SUM # SUM # FRIALP</FRIMK
15N 0383      KAL # 1,MP # 1
15N 0384      FCJJC # SUM
15N 0385      J01 # MP # 1
15N 0386      IF ERTJCC 1008,1008,1001
15N 0387      CAT = ABS(AINT(SMGL10,FCJJC/O1J1)*G-1/SATMTA)
15N 0388      IF CAT # 0.
15N 0389      1011 FCJJC # 0.
15N 0390      GO TO 1006
15N 0391      1010 FCJJC # SCRT BND - CAT
15N 0392      1008 CONTINUE
15N 0393      1012 ICGC 1000,1000,2008
15N 0394      100C NCZE # 0
15N 0395      DC 1001 J01,AP
15N 0396      J01 # MP # 1
15N 0397      IF ERTJCC 1001,1002,1001
15N 0398      1002 NCZE # NCZE # 1
15N 0399      1001 CONTINUE
15N 0400      IF NCZE = 2< 1003,1003,2008
15N 0401      1003 SNO # SNO - 1.
15N 0402      GO TO 1021
15N 0403      2008 IF SNO = 7< 2013,1020,1020
15N 0404      2013 INOR # 0.
15N 0405      DC 1013 J01,KFACT
15N 0406      AV # 0.
15N 0407      ASO # 0.
15N 0408      DC 1014 J01,MP
15N 0409      J01 # 1K-1,MP # 1
15N 0410      J01 # MP # 1
15N 0411      AV # AV # FRIALP</FCJJC</O1K
15N 0412      1014 ASO # ASO # FRIALP</FRIALP</FCJJC
15N 0413      IF ASO # 0. G< GO TO 1013
15N 0414      E # AV/ASO
15N 0415      FRIMK # FRIMK # E</21. - E</FRIMK
15N 0416      1013 INOR # INOR # FRIMK</FRIMK
15N 0417      INOR # SNO # 21NOR<
15N 0418
15N 0419
15N 0420

```

FCAN 195
FCAN 196
FCAN 197
FCAN 198
FCAN 199
FCAN 200

FCAN 201
FCAN 202
FCAN 203
FCAN 204
FCAN 205
FCAN 206
FCAN 207
FCAN 208
FCAN 209
FCAN 210
FCAN 211
FCAN 212
FCAN 213
FCAN 214

FCAN 216
FCAN 217
FCAN 218
FCAN 219
FCAN 220
FCAN 221
FCAN 222
FCAN 223
FCAN 224
FCAN 225
FCAN 226
FCAN 227
FCAN 228
FCAN 229
FCAN 230
FCAN 231
FCAN 232
FCAN 233
FCAN 234
FCAN 235
FCAN 236
FCAN 237
FCAN 238
FCAN 239
FCAN 240

FCAN 241
FCAN 242
FCAN 243
FCAN 244

[illegible]


```

15A 0533      1618 SETUP(KO) = SUM
15A 0534      WRITE(OUT,1624) KC,3SETUP3JC,J01,X6EC
15A 0535      1625 FORMAT(1H 13.6F10.4)
15A 0536      1619 CONTINUE
15A 0537      1602 CONTINUE
15A 0538      DO 1644 I=1,NP
15A 0539      READ (14,105)(X(I,J),J=1,NV)
15A 0540      1644 CONTINUE
15A 0541      1644 REMIND 14
15A 0542      IF (NSOL-GE-10) GO TO 1646
15A 0543      L = NSOL - 1
15A 0544      DO 1645 I=L,10
15A 0545      DC 1645 K8=1,NP
15A 0546      1645 INNSOL(NK)=0
15A 0547      1645 CONTINUE
15A 0548      1645 DO 1647 I=1,NP
15A 0549      WRITE (7,1650) I,IG(I),(INNSOL(I,J),J=1,10),(X(I,J),J=1,NV)
15A 0550      1647 CONTINUE
15A 0551      1650 FORMAT(1X,12,11,1011,EP7.3)
15A 0552      1648 WRITE (6,1760)
15A 0553      1760 FORMAT(1H1)
15A 0554      1760 WRITE (6,1790)
15A 0555      1790 FORMAT(1X,1790)
15A 0556      1790 WRITE (6,1791)
15A 0557      1791 FCORAT(10X,PCINT 8',1X,'CLUSTER 8',5X,1',1X,'2',1X,'3',1X,'4',1X
15A 0558      1,5',1X,'5',1X,'6',1X,'7',1X,'8',1X,'9',1X,'10',10X,'COORDINATES(OBSERVAT
15A 0559      3(CN51)')
15A 0560      DO 1710 I=1,NP
15A 0561      WRITE (6,1752) I,IG(I),(INNSOL(I,J),J=1,10),(X(I,J),J=1,NV)
15A 0562      1710 CONTINUE
15A 0563      1742 FORMAT(14X,12,6X,11,10X,10(11,1X),5X,8(F7.3,2X))
15A 0564      1843 STOP
15A 0565      END
FCAN 322
FCAN 323
FCAN 324
FCAN 325
FCAN 349

```

LEVEL 15 (JUNE 70)

OS/360 FORTRAN M

DATE 71-257/21.42.06

```
COMPILER OPTIONS - NAME= MAIN,OPT=00,LINE=NT=58,SIZE=0000K,  
SOURCE,BCD,NOLIST,NODECK,LOAD=AP,NOECIT,NOID,NOXREF  
ISN CCC2 FUNCTION KD(I,J)  
ISN CCC3 IF(I.LT.J) GC TC 1  
ISN CCC5 KC=(I*(I-1))/2+J  
ISN CCC6 GC TC 95  
ISN CCC7 1 KC=(J*(J-1))/2+I  
ISN CCC8 99 RETURN  
ISN CCC9 END
```


LEVEL 19 (JUNE 70)

OS/360 FOR RAN H

DATE 71.21.7/21.42.16

```

COMPILER OPTIONS - NAME= MAIN,OPT=0C,LINECNT=58 SIZE=0000K,
SOURCE=BCD,NOLIST,NODECK,LOAD=AP,NODEBIT,NOLIO,NOREF
SUBROUTINE CCG1
IMPLICIT REAL*8(D)
COMMON UCSUM(20,20),DXSUM(20),OSPROD(1,00),X(50,20),NV,NP,NG,
11C(100),MIG(20),ACQ(20),NCR(20)
AC=(NV+NP+1)/2
DO 1 I=1,NV
1 DXSUM(I)=0.00
DO 2 I=1,AC
2 DSPRCC(I)=0.00
DO 3 I=1,NV
DO 4 J=1,NP
DX=X(I,J)
4 DXSUM(I)=DXSUM(I)+DX
DXSUM(I)=DXSUM(I)/DCFLCAT(NP)
WRITE(6,200) I,DXSUM(I)
200 FORMAT(' MEAN OF VAR(',I2,')= ',F9.4)
3 CONTINUE
WRITE(6,300)
300 FORMAT(' OSPROD IS UPPER TRI PART OF MATRIX OF CORRECTED S.S. AND
IS.P. FOR NV VARS. STORED IN COLS.')
DO 5 J=1,NV
DO 6 I=1,J
KI=KCI(I,J)
DO 7 L=1,NP
DXL=X(L,J)
DX2=X(L,I)
7 DSPRCC(KI)=DSPRCC(KI)+DXL+DX2
OSPROD(KI)=OSPROD(KI)-DXSUM(I)*DXSUM(L)+DFLOAT(NP)
WRITE(6,100) KI,DSPRCC(KI)
100 FORMAT(' DSPRCC(',I3,')= ',F13.4)
6 CONTINUE
5 CONTINUE
RETURN
END

```

DATE 71-257/21-42-20

DS/360 FORTRAN M

LEVEL 19 (JUNE 70)

```

COMPILER OPTIONS - NAME= MAIN,OPT=CO,LINECNT=5,SIZE=0000K,
SOURCE,BCD,NOLIST,MODECK,LOA,1,M2,MODEBT,NOID,NOMREF
SUBROUTINE COR2
IMPLICIT REAL*8(C)
COMMON CGSUM(20,20),OXSUM(20),DSPROC(400),X(50,20),NV,NP,NS,
1IC(100),NIC(20),ACQ(20),NCR(20)
NC=(NV*(NV+1))/2
CC 2 I=1,AC
2 CSPRCC(I)=O,DC
CC 11 I=1,NG
CC 12 J=1,NV
12 OGSUM(I,J)=C,DO
11 CONTINUE
CC 13 J=1,NG
CC 3 I=1,NV
CC 4 J=1,NP
DO 5 I=1,NV
DO 6 J=1,NG
DO 7 L=1,NP
1F(IG(I),NE,K) GO TO 4
OGSUM(K,I)=OGSUM(K,I)+O
4 CONTINUE
3 CONTINUE
13 CONTINUE
CC 57 I=1,NV
CC 58 K=1,NG
WRITE(6,155) I,P,OGSUM(I,I)
155 FORMAT(1X,5P,1X)
58 CONTINUE
57 CONTINUE
WRITE(6,300)
300 FORMAT(' DSPROC IS UPPER TRI PART OF MATRIX OF CORRECTED S.S. AND
15.P. FOR NV VARS, STORED BY COLS.')
CC 5 J=1,NV
CC 6 I=1,J
KI=K(I,I)
CC 7 L=1,NP
1F(IG(I),EC,0) GO TO 7
OXL=X(I,L)
OXL=X(I,L)
OXL=X(I,L)
DSPRCC(I)=CSPRCC(I)+O/I/OXZ
7 CONTINUE
CC 14 I=1,NG
1F(IG(I),EC,0) GO TO 23
CCOR=CCOR+OGSUM(I,J)*OGSUM(L,I)/NIC(I)
CC TO 14
23 NC=NC+1
14 CONTINUE
DSPRCC(I)=CSPRCC(I)-CCOR
WRITE(6,100) I,I,DSPRCC(I)
100 FORMAT(' DSPRCC('',I3,'')= ',F13.4)
6 CONTINUE
5 CONTINUE
RETURN
END

```

NOT REPRODUCIBLE

DATE 71.257/21.42.25

OS/360 FORTAN M

LEVEL 19 (JUNE 70)

COMPILER OPTIONS - NAME= MAIN,OPT=00,LINECHT=58,SIZE=200CK,

SOURCE,BCD,NCLIST,NODECK,LOAD,MIP,NCECIT,NOLD,NORREF

FUNCTION TEST(RPAX,K)

IMPLICIT REAL*8(D)

DOUBLE PRECISION BETAX

DOUBLE PRECISION YORPA

DOUBLE PRECISION DLGSM

COMMON: DGSDP(20,20),DKSUM(20),DSPROO(40),X(50,20),NV,NP,NG,

11(1,00),NIG(20),MCO(20),NCR(20)

CRMAX=RNAX

CRMAX=1.00-ERRAX**2

OK=K

DE=NV

DN=(CV-1.00)**.500

DN=NP

DN=DN-DK*1.00

DN=..500*BETAX(CRMAX,CV..500)

CRK=CRK-1.00

OTEST=BETAX(CDN,OKK,CON)

TEST=OTEST

RETURN

END

15N C002

15M C003

15N C004

15N C005

15N C006

15N C007

15N C008

15N C009

15N C010

15N C011

15N C012

15N C013

15N C014

15N C015

15N C016

15N C017

15N C018

15N C019

15N C020

NOT REPRODUCIBLE


```

15N C068      OPU # DMP
15N C069      DMP # 5.00*10P-OL< C DL
15N C070      OFUNE # BETAXCMP,DM,ENK - DP
15N C071      IFDFUNE < 123.1.40
15N C072      DU # CMP
15N C073      CUX # DFUNE
15N C074      CL # CMFU
15N C075      DLX # DFUAZ3<
15N C076      IFJJ-10< 86.85.195
15N C077      IF=DECR-GE.C-100*DU-DL<< GO TO 184
15N C079      OPU # CMP
15N C080      CAP # CU - 5.00*DEC.
15N C081      CFUNE # BETAXCMP,DM,ENK - DP
15N C082      IFDFUNE < 41.1.124
15N C083      CU # CMFU
15N C084      DLX # DFUAZ3<
15N C085      DL # CMP
15N C086      DLX # CFUNE
15N C087      IFJJ-10< 86.85.195
15N C088      DL # CMP
15N C089      DLX # DFUNE
15N C090      IFJJ-10< 86.85.195
15N C091      DL # CMP
15N C092      DLX # DFUNE
15N C093      IFJJ-10< 86.85.195
15N C094      BETAP # DMP
15N C095      RETURN
15N C096      DRES # CFUNE & DP
15N C097      WRITESOUT,156< CP,DM,DM,DRES
15N C098      FCMAT:10.5X,43MNO CONVERGENCE IN BETAP IN DOUBLE PRECISION /
15N C099      11X,SHINPU P = 016.10.4H M = 018.10.4H N = 018.10.9H LAST X =
15N C100      2016.10.13P FRCUCES P = 018.10<
15N C101      GC TO 1
15N C102      DMP # DP
15N C103      CC TO 1
15N C104      DMP # DZ
15N C105      WRITE 3NOUT,101< CP,DM,DM
15N C106      14M N = 012.4, 28M RESULT HAS BEEN SET TO ZERO <
15N C107      GC TO 1
15N C108      DMP # DU - 30L-CP<+3DU/DNC
15N C109      GC TO 1
15N C110      DMP # DP+3CU/DNC
15N C111      GC TO 1
15N C112      END

```

DATE 71-257/21-42.31

135

OS/360 FORTRAN M

LEVEL 19 (JUNE 70)

COMPILER OPTIONS - NAME= MAIN,OPT=CC,LINECNT=56,S ZE=000CK,

SOURCE,BCD,NOLIST,NOCHECK,LOAD,N,P,NCEdit,NOID,NOXREF

DCURLE PRECISION FUNCTION YCRMX(DZ)

EXPLICIT REAL*8(C)

DCURLE PRECISION YCRMX

CPI = .39894228041433

DX = CABS(DZ)

IF(CX.GT.3.D0) GO TO 10

DAL = 0.00

DEL = 1.00

CEM = CX

CEM = 1.00

CAN = 0.00

DAN = CFN * 1.00

CAI = -(2.00*CAN - 1.00) * DX * DX

CEI = 4.00*CAN - 1.00

DAL = DBI*CAN + CAI*CAL

DEL = CBI*CEM + CAI * DEL

CAI = DX*DX - CAI

CEI = 2.00 + CEI

CEM = CBI*CAL + CAI*CAN

CEM = CBI*CEI + DAIC*CEM

CFA = CAL/DEL

CFB = CAF/CEM

IF(CFB.EQ.0.00) GO TO 20

IF(CABS((CFB-CFA)/(CFA).LE.1.D-14) GO TO 20

GO TO 5

10

CAL = 0.00

CEL = 1.00

CAN = 1.00

CEM = UX

CEI = DX

DAN = 1.00

CFA = 1.00/UX

DAN = CAN + 1.00

CEI = CAN - 1.00

CAC = CBI*CAN + CAI*CAL

OPC = CEI*CEM + CAI*CEL

CFB = CAC/CEC

CAL = CAN

DEL = CEM

CAN = CAC

CEM = CEC

IF(CFB.EQ.0.00) GO TO 20

IF(CABS((CFB-CFA)/(CFA).LE.1.D-14) GO TO 10

DFA = CFB

GO TO 15

YCRMX = CPI*CFB*DEXP(-DX*CX/2.D0)

IF(CX.LE.3.D0) YCRMX = 0.500 - YCRMX

IF(DZ.GT.0.00) YCRMX = 1.00 - YCRMX

RETURN

END

20

15N 0002

15N 0003

15N 0004

15N 0005

15N 0006

15N 0007

15N 0008

15N 0009

15N 0010

15N 0011

15N 0012

15N 0013

15N 0014

15N 0015

15N 0016

15N 0017

15N 0018

15N 0019

15N 0020

15N 0021

15N 0022

15N 0023

15N 0024

15N 0025

15N 0026

15N 0027

15N 0028

15N 0029

15N 0030

15N 0031

15N 0032

15N 0033

15N 0034

15N 0035

15N 0036

15N 0037

15N 0038

15N 0039

15N 0040

15N 0041

15N 0042

15N 0043

15N 0044

15N 0045

15N 0046

15N 0047

15N 0048

15N 0049

15N 0050

15N 0051

15N 0052

15N 0053

15N 0054

15N 0055

15N 0056

15N 0057

15N 0058

LEVEL 19 (JUNE 7C)

OS/360 FORTRAN M

DATE 71.257/21.42.43

```

CCMPILER OPTIONS - NAME= MAIN,OPT=CG,LINECNT=53,SIZE=0000K,
SOURCE,BCC,NOLIST,MODECK,LOAD,MAP,NOEDIT,NOIO,NOXREF
COUPLE PRECISION FUNCTION DLGCM(DX)
IMPLICIT REAL *8(D)
DY=DX
CTERM=1.00
IF (DX) 1,1.2
DLGCM=0.00
1 RETURN
2 IF (DY-12.00) 3,3.4
3 CTERM=CTERM*DY
CY=CY+1.00
COT2
4 DLGCM= (CY-.500)* DLGCM(CY)-DY+1.00/(12.00*CY)-1.00/(136.00*DY**3)
1* 1.00/ (126.00*DY**5) -1.00/ (1680.00*DY**7 )*.91893853320467309
2 -DLGCM(CTERM)
RETURN
END

```

NOT REPRODUCIBLE


```

15N 0131 4 DA = DNT
15A 0132 GC TO 7
15N 0133 1 WRITE(6,LOG) CM,CM,CX
15N 0134 100 FORMAT(///5X, 4CHERROR 1A INPUT PARAMETER DETAX SET TO 0. ,
15N 0135 3 1 /5X, 2H=, D2G.8, 2HN=, D2G.8, 2HX=.D20.8)
15N 0136 3 DETAX=CZ
15N 0137 2 RETURN
15N 0138 DETAX=DE
15N 0139 2 RETURN
15A 0139 ENC

```

APPENDIX B

Source Listings for ELLIPSE Program

14/31/15

DATE = 71257

MAIN

FORTRAN IV G LEVEL 19

C THIS IS THE MAIN PROGRAM

```
0001      COMMON GDSZ,GDSY,GTEXT,CPOINT,GDSE,GDSER,X(5,8),NV,MP,NGG,NSQL,  
          1NO,CLMEAN(5,8),A(36,11),XMAX(8),XMIN(8),XLO,XLO,XHI,YHI,NDET(5),  
          2R(16),KWRKNG(16),GINPUT,RINPUT(8),SUC(10,8),INNSQL(50,2),IGWORK(50  
          3),NIG(5),JDSPROD(36),IG(50),DCSUM(5,8),DXSUM(1),ARKNG(40)  
0002      CALL CALC  
0003      CALL EXIBIT  
0004      STOP  
0005      END
```

NOT REPRODUCIBLE

14/31/24

DATE = 71257

KD

FORTRAN IV G LEVEL 19

```
0001 FUNCTION KD(I,J)
0002 IF(I.LT.J) GO TO 1
0003 KD=(1*(I-1))/2+J
0004 GO TO 99
0005 1 KD=(J*(J-1))/2+I
0006 99 RETURN
0007 END
```

142

NOT REPRODUCIBLE

FORTRAN IV 5 LEVEL 19

CALC

DATE = 71257

4/31/45

```

0054 DO 430 K=1,NG
0055 DO 431 J=1,NP
0056 IF (IGNORK(J) .EQ. K) NIG(K) = NIG(K) + 1
0057 431 CONTINUE
0058 430 CONTINUE
0059 NNIG = 0
0060 DO 441 K=1,NG
0061 NIG = NNIG + NIG(K)
0062 CALL CORIS
0063 DO 442 K=1,NVV
0064 DSPROD(K) = DSPROD(K)/NNIG
0065 CALL SINVDSPROD,NV,1.E-7,IER)
0066 WRITE (6,1661) IER
0067 DO 433 J=1,NVV
0068 A(J,I+1) = DSPROD(J)
0069 433 CONTINUE
0070 426 CONTINUE
0071 450 DO 15 I=1,NV
0072 XMAX(I) = X(1,I)
0073 XMIN(I) = X(1,I)
0074 DO 15 J=1,NP
0075 IF (X(J,I).GT.XMAX(I)) XMAX(I) = X(J,I)
0076 IF (X(J,I).LT.XMIN(I)) XMIN(I) = X(J,I)
0077 15 CONTINUE
0078 WRITE (6,1651) XMAX(1)
0079 1651 FORMAT(1X,'XMAX(1)=',F7.3)
0080 RETURN
0081 END

```

NOT REPRODUCIBLE

FORTRAN IV G LEVEL 19

CORIS

DATE = 71257

14/32/73

```

0001 SUBROUTINE CORIS
0002 COMMON GCSX,GDSY,GTEXT,GPOINT,GDSE,GDSER,X(50,8),NV,NP,NGG,NSQL,
      1NG,CLMEAN(5,8),A(36,11),XMAX(8),XMIN(8),XLO,YLO,XH-1,YHI,NDET(5),
      2R(16),RWRKNG(16),GINPUT,RINPUT(8),SUG(10,8),INNSCL(50,2),IGWOK(50
      3),NIG(5),DSPROD(36),IG(50),DGSUM(5,8),DXSUM(8),AWRKNG(40)
0003 DO 10 I=1,NV
0004 DO 11 J=1,I
0005 KI = K0(I,J)
0006 DSPROD(KI) = 0.0
0007 DO 12 L=1,NP
0008 IF (IGWOK(L).EQ.0) GO TO 12
0009 IF (IG(L).EQ.0) GO TO 7
0010 K = IG(L)
0011 DX1 = X(L,J) - CLMEAN(K,J)
0012 DX2 = X(L,I) - CLMEAN(K,I)
0013 DSPROD(KI) = DSPROD(KI) + DX1*DX2
0014 GO TO 12
0015 7 DX1 = X(L,J) - DXSUM(J)
0016 DX2 = X(L,I) - DXSUM(I)
0017 DSPROD(KI) = DSPROD(KI) + DX1*DX2
0018 12 CONTINUE
0019 WRITE (6,100) KI,DSPROD(KI)
0020 100 FORMAT(' DSPROD(',I3,')=',F13.4)
0021 11 CONTINUE
0022 10 CONTINUE
0023 RETURN
0024 END

```

NOT REPRODUCIBLE


```

FORTRAN IV C LEVEL 19          CCM:          DATE = 7/25/77          14/32/11

SUBROUTINE COR2
COMMON CCM, CDSY, CTENT, ZPOINT, CDSE, CDSER, X(50,8), NV, NP, NCC, NSOL,
1NG, CLMEAN(5,8), XMAX(1), XMIN(1), XLO, XLO, XHI, XHI, XDET(5),
2X(10), XAKNG(15), GINPOT, INPUT(8), SUC(10,8), INNSCL(50,2), ICDORR(50
3), NIG(5), DSPROD(136), ICI(50), DCSUM(5,8), DCSUM(8), AARRANG(40)
NC=(NV*(NV+1))/2
DO 2 I=1,NC
2 DSPROD(I)=0.00
DO 11 I=1,NCC
DO 12 J=1,NV
12 DCSUM(I,J)=0.00
11 CONTINUE
DO 13 K=1,NCC
DO 3 I=1,NV
DO 4 J=1,NP
DO 5 J=1,NV
DO 6 I=1,JJ
KI=KCI(I,J)
DO 7 L=1,NP
IF (ICORR(KI,EG,0)) GO TO 7
OAI=KIL(J)
DI2=XIL(I)
DSPROD(KI)=DSPROD(KI)+CAL*DMZ
7 CONTINUE
DO 14 L=1,NCC
DO 15 I=1,NV
IF (NIG(I,EG,0)) GO TO 23
DCL=UCOR+DCSUM(I,J)*DCSUM(L,1)/NIG(I)
GO TO 14
23 NCC = NCC-1
14 CONTINUE
DSPROD(KI)=DSPROD(KI)+CCM
OAI=OAI+100
AI=DSPROD(KI)
100 FORMATT= DSPROD(1,13,0)=0.0, F13.4)
6 CONTINUE
5 CONTINUE
RETURN
END

```

NOT REPRODUCIBLE

```

FORTRAN IV C LEVEL 10          EXHIBIT          DATE = 71257          14/32/21

0001      SUMOUTLINE EXHIBIT
0002      COMMON CCSX,CCSY,CTEXT,CPOINT,CSE,COSER,K(50,3),NV,MP,KCC,NSOL,
          ING,CLMEANS,BI,AC(36,11),KX(19),XMIN(8),XLO,YLO,XHI,YHI,NDET(15),
          ZR(10),RWKNG(16),GINPUT,INPUT(8),SUC(10,3),INKSOL(50,2),ICORR(50
          3),NIG(5),DSPRNG(36),IC(5),DCSUM(5,8),DESUM(8),APRNG(40)
          CALL DISPLA(CCSX,CCSY,CTEXT,CPOINT,CSE,COSER,GINPUT)
          WRITE (6,1654)
          1654 FORMAT(1X,'DISPLA SUCCESSFUL')
          I = LIGHTS(1,1)-NV,27-27,31
          CALL MESSE(CTEXT,CCSX,CCSY,CPOINT,CSE,COSER,GINPUT)
          CALL MESSE
          WRITE (4,1655)
          1655 FORMAT('FOLLOWING ANY DISPLAY OF PROJECTIONS PRESS')
          1-KEY 25 TO CHANGE THE VECTOR FORMING ONE OF THE AXES.
          2-KEY 31 TO CHANGE THE VARIABLE FORMING ANOTHER AXIS.
          3-KEY ON PRESS ANY KEY TO DO ON.
          CALL MSET(CTEXT)
          LP = PLOT(CTEXT)
          LP = DETAIN(NDET)
          LP = UNPLOT(CTEXT)
          CALL MSET(CTEXT,CCSX,CCSY,CPOINT,CSE,COSER,GINPUT)
          NDET = 1
          KCC = 0
          26 CALL KEYIN (NAXIS,NTEST)
          IF (NAXIS.EQ.3) GO TO 40
          CALL ECLPSE (NAXIS)
          16 KA = DETAIN(NDET)
          KA = 30,114
          17 IF (NAXIS.EQ.30) GO TO 40
          NTEST = KA
          IF (NAXIS.EQ.29) GO TO 80
          IF (NAXIS.EQ.31) GO TO 80
          IF (NAXIS.EQ.31) GO TO 61
          LP = UNPLOT(CTEXT)
          CALL MSET(CTEXT,CCSX,CCSY,CPOINT,CSE,COSER,GINPUT)
          82 FORNAT(1X,'CAROR OR SINGULAR SOLUTION - PRESS ANOTHER KEY')
          1-IF TRAPPED,PRESS KEY 20
          CALL MSET(CTEXT)
          LP = PLOT(CTEXT)
          GO TO 16
          80 LP = UNPLOT(CCSX,CCSY,CPOINT,CTEXT)
          CALL MSET(CTEXT,CCSX,CCSY,CPOINT,CSE,COSER,GINPUT)
          LP = UNPLOT(COSER)
          CALL MSET(CTEXT,CCSX,CCSY,CPOINT,CSE,COSER,GINPUT)
          GO TO 20
          61 CALL ECLPSE (KA,NAXIS)
          IF (NAXIS.EQ.15) GO TO 103
          KA = 15
          GO TO 17
          103 KA = DETAIN(NDET)
          KA = NDET(4)
          LP = UNPLOT(COSER)
          CALL MSET(CTEXT,CCSX,CCSY,CPOINT,CSE,COSER,GINPUT)
          GO TO 17
          40 CALL BLANK

```

NOT REPRODUCIBLE

14/32/21

DATE = 71257

EXIBIT

FORTRAN IV G LEVEL 19

```
0050      CALL RESET(GDSX,GDSY,GDSE,GDSER,GTEXT,GPOINT,GINPUT)  
0051      RETURN  
0052      END
```

FORTRAN IV G LEVEL 19

HFSSGE

DATE = 71257

14/12/29

0001
0002

SUBROUTINE MESSAGE

COMMON GDSX,GDSY,GTEXT,GPOINT,GDSE,GDSER,X(50,8),NV,NP,NGG,NSOL,
ING,CLMEAN(5,8),A(36,11),XMAX(8),XMIN(8),XLO,XHI,Y-I,NDIET(5),
2X(15),ZKXNG(16),GINPUT,RINPUT(S),SUS(10,8),INNSCL(50,2),IGWORN(50
3),NIG(5),DSPROD(36),IG(5),DGSUM(5,8),DXSUM(8),AARKNG(40)
WRITE (4,460)

0003
0004

460 FORMAT('PLT'S PROGRAM DISPLAYS CLUSTERS BY PROJECTING THEM ON')/
1'POVARIOUS 2-DIMENSIONAL SUBSPACES.SIMPLE STRUCTURE'/
2'PO SOLUTIONS CAN ALSO BE INDICATED.REFER TO YOUR'/
3'PO INSTRUCTION CHART.HAVE YOU ENTERED ALL NECESSARY'/
4'PO DATA VIA IERGENER?'/
5'PO',10X,'THE BOTTOM ROW OF THE PROGRAM'/'/
6'PO FUNCTION KEYS IS LIT UP.THEY WILL BE USED'/'/
7'PO AS DIRECTED.THE DARK ONE,KEY NO. 30,IS'/'/
8'PO THE PANIC BUTTON.IT WILL RETURN YOU TO THE'/'/
9'PO MONITOR.'/'PO IN CASE OF PANIC PRESS KEY 30'/'/
1'PO NOW PRESS ANY KEY TO GET STARTED'/'

0005
0006
0007
0008
0009
0010
0011
0012
0013

CALL XRFMT(GTEXT)
LP = PLCT(GTEXT)
WRITE (6,1653) LP
1653 FORMAT(IX,'LP AFTER FIRST MESSAGE IS',I4)
I = DETAIN(NDIET)
I = UNPLCT(GTEXT)
CALL RESET(GTEXT,GCSX,GDSY,GPOINT,GDSE,GDSER,GINPUT)
RETURN
END

NOT REPRODUCIBLE

```

FORTRAN IV C LEVEL 19          F-VM          DATE = 71257          14/32/38

0001      SUBROUTINE KEVIN (MAX S,TEST)
0002      COMMON CDSX,CDSY,CDEX,CDEY,CDSR,CDSI,X(50),Y(50),VX,VY,MCG,N,SOL,
          INCL,PEANUS(8),ALB(2),I,AKA(19),AMIN(3),ALQ,YLG,XML,YML,MDIT(5),
          Z(10),RANKG(10),GINS(10),RINP(10),SUG(10,8),JNSOL(50,2),ICAD(150)
          31,NIC(1),JSPR(136),I,ISOL,OSUMS(8),DASUM(8),ABSNAG(40)
          IF (TEST.EQ.31) GO TO 101
          99 CALL BLANK
          CALL READS
          CALL RESET(IGTEXT,CDSX,CDSY,CDEY,CDSR,CDSI,CDEX,CDEY,CDSR,CDSI)
          WRITE (4,190) NSOL
          150 FORMAT('P1')/90 ACCORDING TO THE ROTATION PROGRAM THE/
          100 FOLLOWING VECTORS WERE FOUND: NEW DATA INTO:/
          200 SIMPLE STRUCTURE SOLUTIONS NO.1 TO/13)
          CALL WRITE(IGINPUT)
          03 110 J=1,NSOL
          04 111 DO 112 J=1,NSOL
          05 112 DO 113 I=1,19) (SUG(I)-I)*X(I),VX)
          06 113 DO 114 I=1,19) (SUG(I)-I)*Y(I),VY)
          07 114 DO 115 I=1,19) (SUG(I)-I)*Z(I),VZ)
          08 115 DO 116 I=1,19) (SUG(I)-I)*X(I),VX)
          09 116 DO 117 I=1,19) (SUG(I)-I)*Y(I),VY)
          10 117 DO 118 I=1,19) (SUG(I)-I)*Z(I),VZ)
          11 118 DO 119 I=1,19) (SUG(I)-I)*X(I),VX)
          12 119 DO 120 I=1,19) (SUG(I)-I)*Y(I),VY)
          13 120 DO 121 I=1,19) (SUG(I)-I)*Z(I),VZ)
          14 121 DO 122 I=1,19) (SUG(I)-I)*X(I),VX)
          15 122 DO 123 I=1,19) (SUG(I)-I)*Y(I),VY)
          16 123 DO 124 I=1,19) (SUG(I)-I)*Z(I),VZ)
          17 124 DO 125 I=1,19) (SUG(I)-I)*X(I),VX)
          18 125 DO 126 I=1,19) (SUG(I)-I)*Y(I),VY)
          19 126 DO 127 I=1,19) (SUG(I)-I)*Z(I),VZ)
          20 127 DO 128 I=1,19) (SUG(I)-I)*X(I),VX)
          21 128 DO 129 I=1,19) (SUG(I)-I)*Y(I),VY)
          22 129 DO 130 I=1,19) (SUG(I)-I)*Z(I),VZ)
          23 130 DO 131 I=1,19) (SUG(I)-I)*X(I),VX)
          24 131 DO 132 I=1,19) (SUG(I)-I)*Y(I),VY)
          25 132 DO 133 I=1,19) (SUG(I)-I)*Z(I),VZ)
          26 133 DO 134 I=1,19) (SUG(I)-I)*X(I),VX)
          27 134 DO 135 I=1,19) (SUG(I)-I)*Y(I),VY)
          28 135 DO 136 I=1,19) (SUG(I)-I)*Z(I),VZ)
          29 136 DO 137 I=1,19) (SUG(I)-I)*X(I),VX)
          30 137 DO 138 I=1,19) (SUG(I)-I)*Y(I),VY)
          31 138 DO 139 I=1,19) (SUG(I)-I)*Z(I),VZ)
          32 139 DO 140 I=1,19) (SUG(I)-I)*X(I),VX)
          33 140 DO 141 I=1,19) (SUG(I)-I)*Y(I),VY)
          34 141 DO 142 I=1,19) (SUG(I)-I)*Z(I),VZ)
          35 142 DO 143 I=1,19) (SUG(I)-I)*X(I),VX)
          36 143 DO 144 I=1,19) (SUG(I)-I)*Y(I),VY)
          37 144 DO 145 I=1,19) (SUG(I)-I)*Z(I),VZ)
          38 145 DO 146 I=1,19) (SUG(I)-I)*X(I),VX)
          39 146 DO 147 I=1,19) (SUG(I)-I)*Y(I),VY)
          40 147 DO 148 I=1,19) (SUG(I)-I)*Z(I),VZ)
          41 148 DO 149 I=1,19) (SUG(I)-I)*X(I),VX)
          42 149 DO 150 I=1,19) (SUG(I)-I)*Y(I),VY)
          43 150 DO 151 I=1,19) (SUG(I)-I)*Z(I),VZ)
          44 151 DO 152 I=1,19) (SUG(I)-I)*X(I),VX)
          45 152 DO 153 I=1,19) (SUG(I)-I)*Y(I),VY)
          46 153 DO 154 I=1,19) (SUG(I)-I)*Z(I),VZ)
          47 154 DO 155 I=1,19) (SUG(I)-I)*X(I),VX)
          48 155 DO 156 I=1,19) (SUG(I)-I)*Y(I),VY)
          49 156 DO 157 I=1,19) (SUG(I)-I)*Z(I),VZ)
          50 157 DO 158 I=1,19) (SUG(I)-I)*X(I),VX)
          51 158 DO 159 I=1,19) (SUG(I)-I)*Y(I),VY)
          52 159 DO 160 I=1,19) (SUG(I)-I)*Z(I),VZ)
          53 160 DO 161 I=1,19) (SUG(I)-I)*X(I),VX)
          54 161 DO 162 I=1,19) (SUG(I)-I)*Y(I),VY)
          55 162 DO 163 I=1,19) (SUG(I)-I)*Z(I),VZ)
          56 163 DO 164 I=1,19) (SUG(I)-I)*X(I),VX)
          57 164 DO 165 I=1,19) (SUG(I)-I)*Y(I),VY)
          58 165 DO 166 I=1,19) (SUG(I)-I)*Z(I),VZ)
          59 166 DO 167 I=1,19) (SUG(I)-I)*X(I),VX)
          60 167 DO 168 I=1,19) (SUG(I)-I)*Y(I),VY)
          61 168 DO 169 I=1,19) (SUG(I)-I)*Z(I),VZ)
          62 169 DO 170 I=1,19) (SUG(I)-I)*X(I),VX)
          63 170 DO 171 I=1,19) (SUG(I)-I)*Y(I),VY)
          64 171 DO 172 I=1,19) (SUG(I)-I)*Z(I),VZ)
          65 172 DO 173 I=1,19) (SUG(I)-I)*X(I),VX)
          66 173 DO 174 I=1,19) (SUG(I)-I)*Y(I),VY)
          67 174 DO 175 I=1,19) (SUG(I)-I)*Z(I),VZ)
          68 175 DO 176 I=1,19) (SUG(I)-I)*X(I),VX)
          69 176 DO 177 I=1,19) (SUG(I)-I)*Y(I),VY)
          70 177 DO 178 I=1,19) (SUG(I)-I)*Z(I),VZ)
          71 178 DO 179 I=1,19) (SUG(I)-I)*X(I),VX)
          72 179 DO 180 I=1,19) (SUG(I)-I)*Y(I),VY)
          73 180 DO 181 I=1,19) (SUG(I)-I)*Z(I),VZ)
          74 181 DO 182 I=1,19) (SUG(I)-I)*X(I),VX)
          75 182 DO 183 I=1,19) (SUG(I)-I)*Y(I),VY)
          76 183 DO 184 I=1,19) (SUG(I)-I)*Z(I),VZ)
          77 184 DO 185 I=1,19) (SUG(I)-I)*X(I),VX)
          78 185 DO 186 I=1,19) (SUG(I)-I)*Y(I),VY)
          79 186 DO 187 I=1,19) (SUG(I)-I)*Z(I),VZ)
          80 187 DO 188 I=1,19) (SUG(I)-I)*X(I),VX)
          81 188 DO 189 I=1,19) (SUG(I)-I)*Y(I),VY)
          82 189 DO 190 I=1,19) (SUG(I)-I)*Z(I),VZ)
          83 190 DO 191 I=1,19) (SUG(I)-I)*X(I),VX)
          84 191 DO 192 I=1,19) (SUG(I)-I)*Y(I),VY)
          85 192 DO 193 I=1,19) (SUG(I)-I)*Z(I),VZ)
          86 193 DO 194 I=1,19) (SUG(I)-I)*X(I),VX)
          87 194 DO 195 I=1,19) (SUG(I)-I)*Y(I),VY)
          88 195 DO 196 I=1,19) (SUG(I)-I)*Z(I),VZ)
          89 196 DO 197 I=1,19) (SUG(I)-I)*X(I),VX)
          90 197 DO 198 I=1,19) (SUG(I)-I)*Y(I),VY)
          91 198 DO 199 I=1,19) (SUG(I)-I)*Z(I),VZ)
          92 199 DO 200 I=1,19) (SUG(I)-I)*X(I),VX)
          93 200 DO 201 I=1,19) (SUG(I)-I)*Y(I),VY)
          94 201 DO 202 I=1,19) (SUG(I)-I)*Z(I),VZ)
          95 202 DO 203 I=1,19) (SUG(I)-I)*X(I),VX)
          96 203 DO 204 I=1,19) (SUG(I)-I)*Y(I),VY)
          97 204 DO 205 I=1,19) (SUG(I)-I)*Z(I),VZ)
          98 205 DO 206 I=1,19) (SUG(I)-I)*X(I),VX)
          99 206 DO 207 I=1,19) (SUG(I)-I)*Y(I),VY)
          100 207 DO 208 I=1,19) (SUG(I)-I)*Z(I),VZ)
          101 208 DO 209 I=1,19) (SUG(I)-I)*X(I),VX)
          102 209 DO 210 I=1,19) (SUG(I)-I)*Y(I),VY)
          103 210 DO 211 I=1,19) (SUG(I)-I)*Z(I),VZ)
          104 211 DO 212 I=1,19) (SUG(I)-I)*X(I),VX)
          105 212 DO 213 I=1,19) (SUG(I)-I)*Y(I),VY)
          106 213 DO 214 I=1,19) (SUG(I)-I)*Z(I),VZ)
          107 214 DO 215 I=1,19) (SUG(I)-I)*X(I),VX)
          108 215 DO 216 I=1,19) (SUG(I)-I)*Y(I),VY)
          109 216 DO 217 I=1,19) (SUG(I)-I)*Z(I),VZ)
          110 217 DO 218 I=1,19) (SUG(I)-I)*X(I),VX)
          111 218 DO 219 I=1,19) (SUG(I)-I)*Y(I),VY)
          112 219 DO 220 I=1,19) (SUG(I)-I)*Z(I),VZ)
          113 220 DO 221 I=1,19) (SUG(I)-I)*X(I),VX)
          114 221 DO 222 I=1,19) (SUG(I)-I)*Y(I),VY)
          115 222 DO 223 I=1,19) (SUG(I)-I)*Z(I),VZ)
          116 223 DO 224 I=1,19) (SUG(I)-I)*X(I),VX)
          117 224 DO 225 I=1,19) (SUG(I)-I)*Y(I),VY)
          118 225 DO 226 I=1,19) (SUG(I)-I)*Z(I),VZ)
          119 226 DO 227 I=1,19) (SUG(I)-I)*X(I),VX)
          120 227 DO 228 I=1,19) (SUG(I)-I)*Y(I),VY)
          121 228 DO 229 I=1,19) (SUG(I)-I)*Z(I),VZ)
          122 229 DO 230 I=1,19) (SUG(I)-I)*X(I),VX)
          123 230 DO 231 I=1,19) (SUG(I)-I)*Y(I),VY)
          124 231 DO 232 I=1,19) (SUG(I)-I)*Z(I),VZ)
          125 232 DO 233 I=1,19) (SUG(I)-I)*X(I),VX)
          126 233 DO 234 I=1,19) (SUG(I)-I)*Y(I),VY)
          127 234 DO 235 I=1,19) (SUG(I)-I)*Z(I),VZ)
          128 235 DO 236 I=1,19) (SUG(I)-I)*X(I),VX)
          129 236 DO 237 I=1,19) (SUG(I)-I)*Y(I),VY)
          130 237 DO 238 I=1,19) (SUG(I)-I)*Z(I),VZ)
          131 238 DO 239 I=1,19) (SUG(I)-I)*X(I),VX)
          132 239 DO 240 I=1,19) (SUG(I)-I)*Y(I),VY)
          133 240 DO 241 I=1,19) (SUG(I)-I)*Z(I),VZ)
          134 241 DO 242 I=1,19) (SUG(I)-I)*X(I),VX)
          135 242 DO 243 I=1,19) (SUG(I)-I)*Y(I),VY)
          136 243 DO 244 I=1,19) (SUG(I)-I)*Z(I),VZ)
          137 244 DO 245 I=1,19) (SUG(I)-I)*X(I),VX)
          138 245 DO 246 I=1,19) (SUG(I)-I)*Y(I),VY)
          139 246 DO 247 I=1,19) (SUG(I)-I)*Z(I),VZ)
          140 247 DO 248 I=1,19) (SUG(I)-I)*X(I),VX)
          141 248 DO 249 I=1,19) (SUG(I)-I)*Y(I),VY)
          142 249 DO 250 I=1,19) (SUG(I)-I)*Z(I),VZ)
          143 250 DO 251 I=1,19) (SUG(I)-I)*X(I),VX)
          144 251 DO 252 I=1,19) (SUG(I)-I)*Y(I),VY)
          145 252 DO 253 I=1,19) (SUG(I)-I)*Z(I),VZ)
          146 253 DO 254 I=1,19) (SUG(I)-I)*X(I),VX)
          147 254 DO 255 I=1,19) (SUG(I)-I)*Y(I),VY)
          148 255 DO 256 I=1,19) (SUG(I)-I)*Z(I),VZ)
          149 256 DO 257 I=1,19) (SUG(I)-I)*X(I),VX)
          150 257 DO 258 I=1,19) (SUG(I)-I)*Y(I),VY)
          151 258 DO 259 I=1,19) (SUG(I)-I)*Z(I),VZ)
          152 259 DO 260 I=1,19) (SUG(I)-I)*X(I),VX)
          153 260 DO 261 I=1,19) (SUG(I)-I)*Y(I),VY)
          154 261 DO 262 I=1,19) (SUG(I)-I)*Z(I),VZ)
          155 262 DO 263 I=1,19) (SUG(I)-I)*X(I),VX)
          156 263 DO 264 I=1,19) (SUG(I)-I)*Y(I),VY)
          157 264 DO 265 I=1,19) (SUG(I)-I)*Z(I),VZ)
          158 265 DO 266 I=1,19) (SUG(I)-I)*X(I),VX)
          159 266 DO 267 I=1,19) (SUG(I)-I)*Y(I),VY)
          160 267 DO 268 I=1,19) (SUG(I)-I)*Z(I),VZ)
          161 268 DO 269 I=1,19) (SUG(I)-I)*X(I),VX)
          162 269 DO 270 I=1,19) (SUG(I)-I)*Y(I),VY)
          163 270 DO 271 I=1,19) (SUG(I)-I)*Z(I),VZ)
          164 271 DO 272 I=1,19) (SUG(I)-I)*X(I),VX)
          165 272 DO 273 I=1,19) (SUG(I)-I)*Y(I),VY)
          166 273 DO 274 I=1,19) (SUG(I)-I)*Z(I),VZ)
          167 274 DO 275 I=1,19) (SUG(I)-I)*X(I),VX)
          168 275 DO 276 I=1,19) (SUG(I)-I)*Y(I),VY)
          169 276 DO 277 I=1,19) (SUG(I)-I)*Z(I),VZ)
          170 277 DO 278 I=1,19) (SUG(I)-I)*X(I),VX)
          171 278 DO 279 I=1,19) (SUG(I)-I)*Y(I),VY)
          172 279 DO 280 I=1,19) (SUG(I)-I)*Z(I),VZ)
          173 280 DO 281 I=1,19) (SUG(I)-I)*X(I),VX)
          174 281 DO 282 I=1,19) (SUG(I)-I)*Y(I),VY)
          175 282 DO 283 I=1,19) (SUG(I)-I)*Z(I),VZ)
          176 283 DO 284 I=1,19) (SUG(I)-I)*X(I),VX)
          177 284 DO 285 I=1,19) (SUG(I)-I)*Y(I),VY)
          178 285 DO 286 I=1,19) (SUG(I)-I)*Z(I),VZ)
          179 286 DO 287 I=1,19) (SUG(I)-I)*X(I),VX)
          180 287 DO 288 I=1,19) (SUG(I)-I)*Y(I),VY)
          181 288 DO 289 I=1,19) (SUG(I)-I)*Z(I),VZ)
          182 289 DO 290 I=1,19) (SUG(I)-I)*X(I),VX)
          183 290 DO 291 I=1,19) (SUG(I)-I)*Y(I),VY)
          184 291 DO 292 I=1,19) (SUG(I)-I)*Z(I),VZ)
          185 292 DO 293 I=1,19) (SUG(I)-I)*X(I),VX)
          186 293 DO 294 I=1,19) (SUG(I)-I)*Y(I),VY)
          187 294 DO 295 I=1,19) (SUG(I)-I)*Z(I),VZ)
          188 295 DO 296 I=1,19) (SUG(I)-I)*X(I),VX)
          189 296 DO 297 I=1,19) (SUG(I)-I)*Y(I),VY)
          190 297 DO 298 I=1,19) (SUG(I)-I)*Z(I),VZ)
          191 298 DO 299 I=1,19) (SUG(I)-I)*X(I),VX)
          192 299 DO 300 I=1,19) (SUG(I)-I)*Y(I),VY)
          193 300 DO 301 I=1,19) (SUG(I)-I)*Z(I),VZ)
          194 301 DO 302 I=1,19) (SUG(I)-I)*X(I),VX)
          195 302 DO 303 I=1,19) (SUG(I)-I)*Y(I),VY)
          196 303 DO 304 I=1,19) (SUG(I)-I)*Z(I),VZ)
          197 304 DO 305 I=1,19) (SUG(I)-I)*X(I),VX)
          198 305 DO 306 I=1,19) (SUG(I)-I)*Y(I),VY)
          199 306 DO 307 I=1,19) (SUG(I)-I)*Z(I),VZ)
          200 307 DO 308 I=1,19) (SUG(I)-I)*X(I),VX)
          201 308 DO 309 I=1,19) (SUG(I)-I)*Y(I),VY)
          202 309 DO 310 I=1,19) (SUG(I)-I)*Z(I),VZ)
          203 310 DO 311 I=1,19) (SUG(I)-I)*X(I),VX)
          204 311 DO 312 I=1,19) (SUG(I)-I)*Y(I),VY)
          205 312 DO 313 I=1,19) (SUG(I)-I)*Z(I),VZ)
          206 313 DO 314 I=1,19) (SUG(I)-I)*X(I),VX)
          207 314 DO 315 I=1,19) (SUG(I)-I)*Y(I),VY)
          208 315 DO 316 I=1,19) (SUG(I)-I)*Z(I),VZ)
          209 316 DO 317 I=1,19) (SUG(I)-I)*X(I),VX)
          210 317 DO 318 I=1,19) (SUG(I)-I)*Y(I),VY)
          211 318 DO 319 I=1,19) (SUG(I)-I)*Z(I),VZ)
          212 319 DO 320 I=1,19) (SUG(I)-I)*X(I),VX)
          213 320 DO 321 I=1,19) (SUG(I)-I)*Y(I),VY)
          214 321 DO 322 I=1,19) (SUG(I)-I)*Z(I),VZ)
          215 322 DO 323 I=1,19) (SUG(I)-I)*X(I),VX)
          216 323 DO 324 I=1,19) (SUG(I)-I)*Y(I),VY)
          217 324 DO 325 I=1,19) (SUG(I)-I)*Z(I),VZ)
          218 325 DO 326 I=1,19) (SUG(I)-I)*X(I),VX)
          219 326 DO 327 I=1,19) (SUG(I)-I)*Y(I),VY)
          220 327 DO 328 I=1,19) (SUG(I)-I)*Z(I),VZ)
          221 328 DO 329 I=1,19) (SUG(I)-I)*X(I),VX)
          222 329 DO 330 I=1,19) (SUG(I)-I)*Y(I),VY)
          223 330 DO 331 I=1,19) (SUG(I)-I)*Z(I),VZ)
          224 331 DO 332 I=1,19) (SUG(I)-I)*X(I),VX)
          225 332 DO 333 I=1,19) (SUG(I)-I)*Y(I),VY)
          226 333 DO 334 I=1,19) (SUG(I)-I)*Z(I),VZ)
          227 334 DO 335 I=1,19) (SUG(I)-I)*X(I),VX)
          228 335 DO 336 I=1,19) (SUG(I)-I)*Y(I),VY)
          229 336 DO 337 I=1,19) (SUG(I)-I)*Z(I),VZ)
          230 337 DO 338 I=1,19) (SUG(I)-I)*X(I),VX)
          231 338 DO 339 I=1,19) (SUG(I)-I)*Y(I),VY)
          232 339 DO 340 I=1,19) (SUG(I)-I)*Z(I),VZ)
          233 340 DO 341 I=1,19) (SUG(I)-I)*X(I),VX)
          234 341 DO 342 I=1,19) (SUG(I)-I)*Y(I),VY)
          235 342 DO 343 I=1,19) (SUG(I)-I)*Z(I),VZ)
          236 343 DO 344 I=1,19) (SUG(I)-I)*X(I),VX)
          237 344 DO 345 I=1,19) (SUG(I)-I)*Y(I),VY)
          238 345 DO 346 I=1,19) (SUG(I)-I)*Z(I),VZ)
          239 346 DO 347 I=1,19) (SUG(I)-I)*X(I),VX)
          240 347 DO 348 I=1,19) (SUG(I)-I)*Y(I),VY)
          241 348 DO 349 I=1,19) (SUG(I)-I)*Z(I),VZ)
          242 349 DO 350 I=1,19) (SUG(I)-I)*X(I),VX)
          243 350 DO 351 I=1,19) (SUG(I)-I)*Y(I),VY)
          244 351 DO 352 I=1,19) (SUG(I)-I)*Z(I),VZ)
          245 352 DO 353 I=1,19) (SUG(I)-I)*X(I),VX)
          246 353 DO 354 I=1,19) (SUG(I)-I)*Y(I),VY)
          247 354 DO 355 I=1,19) (SUG(I)-I)*Z(I),VZ)
          248 355 DO 356 I=1,19) (SUG(I)-I)*X(I),VX)
          249 356 DO 357 I=1,19) (SUG(I)-I)*Y(I),VY)
          250 357 DO 358 I=1,19) (SUG(I)-I)*Z(I),VZ)
          251 358 DO 359 I=1,19) (SUG(I)-I)*X(I),VX)
          252 359 DO 360 I=1,19) (SUG(I)-I)*Y(I),VY)
          253 360 DO 361 I=1,19) (SUG(I)-I)*Z(I),VZ)
          254 361 DO 362 I=1,19) (SUG(I)-I)*X(I),VX)
          255 362 DO 363 I=1,19) (SUG(I)-I)*Y(I),VY)
          256 363 DO 364 I=1,19) (SUG(I)-I)*Z(I),VZ)
          257 364 DO 365 I=1,19) (SUG(I)-I)*X(I),VX)
          258 365 DO 366 I=1,19) (SUG(I)-I)*Y(I),VY)
          259 366 DO 367 I=1,19) (SUG(I)-I)*Z(I),VZ)
          260 367 DO 368 I=1,19) (SUG(I)-I)*X(I),VX)
          261 368 DO 369 I=1,19) (SUG(I)-I)*Y(I),VY)
          262 369 DO 370 I=1,19) (SUG(I)-I)*Z(I),VZ)
          263 370 DO 371 I=1,19) (SUG(I)-I)*X(I),VX)
          264 371 DO 372 I=1,19) (SUG(I)-I)*Y(I),VY)
          265 372 DO 373 I=1,19) (SUG(I)-I)*Z(I),VZ)
          266 373 DO 374 I=1,19) (SUG(I)-I)*X(I),VX)
          267 374 DO 375 I=1,19) (SUG(I)-I)*Y(I),VY)
          268 375 DO 376 I=1,19) (SUG(I)-I)*Z(I),VZ)
          269 376 DO 377 I=1,19) (SUG(I)-I)*X(I),VX)
          270 377 DO 378 I=1,19) (SUG(I)-I)*Y(I),VY)
          271 378 DO 379 I=1,19) (SUG(I)-I)*Z(I),VZ)
          272 379 DO 380 I=1,19) (SUG(I)-I)*X(I),VX)
          273 380 DO 381 I=1,19) (SUG(I)-I)*Y(I),VY)
          274 381 DO 382 I=1,19) (SUG(I)-I)*Z(I),VZ)
          275 382 DO 383 I=1,19) (SUG(I)-I)*X(I),VX)
          276 383 DO 384 I=1,19) (SUG(I)-I)*Y(I),VY)
          277 384 DO 385 I=1,19) (SUG(I)-I)*Z(I),VZ)
          278 385 DO 386 I=1,19) (SUG(I)-I)*X(I),VX)
          279 386 DO 387 I=1,19) (SUG(I)-I)*Y(I),VY)
          280 387 DO 388 I=1,19) (SUG(I)-I)*Z(I),VZ)
          281 388 DO 389 I=1,19) (SUG(I)-I)*X(I),VX)
          282 389 DO 390 I=1,19) (SUG(I)-I)*Y(I),VY)
          283 390 DO 391 I=1,19) (SUG(I)-I)*Z(I),VZ)
          284 391 DO 392 I=1,19) (SUG(I)-I)*X(I),VX)
          285 392 DO 393 I=1,19) (SUG(I)-I)*Y(I),VY)
          286 393 DO 394 I=1,19) (SUG(I)-I)*Z(I),VZ)
          287 394 DO 395 I=1,19) (SUG(I)-I)*X(I),VX)
          288 395 DO 396 I=1,19) (SUG(I)-I)*Y(I),VY)
          289 396 DO 397 I=1,19) (SUG(I)-I)*Z(I),VZ)
          290 397 DO 398 I=1,19) (SUG(I)-I)*X(I),VX)
          291 398 DO 399 I=1,19) (SUG(I)-I)*Y(I),VY)
          292 399 DO 400 I=1,19) (SUG(I)-I)*Z(I),VZ)
          293 400 DO 401 I=1,19) (SUG(I)-I)*X(I),VX)
          294 401 DO 402 I=1,19) (SUG(I)-I)*Y(I),VY)
          295 402 DO 403 I=1,19) (SUG(I)-I)*Z(I),VZ)
          296 403 DO 404 I=1,19) (SUG(I)-I)*X(I),VX)
          297 404 DO 405 I=1,19) (SUG(I)-I)*Y(I),VY)
          298 405 DO 406 I=1,19) (SUG(I)-I)*Z(I),VZ)
          299 406 DO 407 I=1,19) (SUG(I)-I)*X(I),VX)
          300 407 DO 408 I=1,19) (SUG(I)-I)*Y(I),VY)
          301 408 DO 409 I=1,19) (SUG(I)-I)*Z(I),VZ)
          302 409 DO 410 I=1,19) (SUG(I)-I)*X(I),VX)
          303 410 DO 411 I=1,19) (SUG(I)-I)*Y(I),VY)
          304 411 DO 412 I=1,19) (SUG(I)-I)*Z(I),VZ)
          305 412 DO 413 I=1,19) (SUG(I)-I)*X(I),VX)
          306 413 DO 414 I=1,19) (SUG(I)-I)*Y(I),VY)
          307 414 DO 415 I=1,19) (SUG(I)-I)*Z(I),VZ)
          308 415 DO 416 I=1,19) (SUG(I)-I)*X(I),VX)
          309 416 DO 417 I=1,19) (SUG(I)-I)*Y(I),VY)
          310 417 DO 418 I=1,19) (SUG(I)-I)*Z(I),VZ)
          311 418 DO 419 I=1,19) (SUG(I)-I)*X(I),VX)
          312 419 DO 420 I=1,19) (SUG(I)-I)*Y(I),VY)
          313 420 DO 421 I=1,19) (SUG(I)-I)*Z(I),VZ)
          314 421 DO 422 I=1,19) (SUG(I)-I)*X(I),VX)
          315 422 DO 423 I=1,19) (SUG(I)-I)*Y(I),VY)
          316 423 DO 424 I=1,19) (SUG(I)-I)*Z(I),VZ)
          317 424 DO 425 I=1,19) (SUG(I)-I)*X(I),VX)
          318 425 DO 426 I=1,19) (SUG(I)-I)*Y(I),VY)
          319 426 DO 427 I=1,19) (SUG(I)-I)*Z(I),VZ)
          320 427 DO 428 I=1,19) (SUG(I)-I)*X(I),VX)
          321 428 DO 429 I=1,19) (SUG(I)-I)*Y(I),VY)
          322 429 DO 430 I=1,19) (SUG(I)-I)*Z(I),VZ)
          323 430 DO 431 I=1,19) (SUG(I)-I)*X(I),VX)
          324 431 DO 432 I=1,19) (SUG(I)-I)*Y(I),VY)
          325 432 DO 433 I=1,19) (SUG(I)-I)*Z(I),VZ)
          326 433 DO 434 I=1,19) (SUG(I)-I)*X(I),VX)
          327 434 DO 435 I=1,19) (SUG(I)-I)*Y(I),VY)
          328 435 DO 436 I=1,19) (SUG(I)-I)*Z(I),VZ)
          329 436 DO 437 I=1,19) (SUG(I)-I)*X(I),VX)
          330 437 DO 438 I=1,19) (SUG(I)-I)*Y(I),VY)
          331 438 DO 439 I=1,19) (SUG(I)-I)*Z(I),VZ)
          332 439 DO 440 I=1,19) (SUG(I)-I)*X(I),VX)
          333 440 DO 441 I=1,19) (SUG(I)-I)*Y(I),VY)
          334 441 DO 442 I=1,19) (SUG(I)-I)*Z(I),VZ)
          335 442 DO 443 I=1,19) (SUG(I)-I)*X(I),VX)
          336 443 DO 444 I=1,19) (SUG(I)-I)*Y(I),VY)
          337 444 DO 445 I=1,19) (SUG(I)-I)*Z(I),VZ)
          338 445 DO 446 I=1,19) (SUG(I)-I)*X(I),VX)
          339 446 DO 447 I=1,19) (SUG(I)-I)*Y(I),VY)
          340 447 DO 448 I=1,19) (SUG(I)-I)*Z(I),VZ)
          341 448 DO 449 I=1,19) (SUG(I)-I)*X(I),VX)
          342 449 DO 450 I=1,19) (SUG(I)-I)*Y(I),VY)
          343 450 DO 451 I=1,19) (SUG(I)-I)*Z(I),VZ)
          344 451 DO 452 I=1,19) (SUG(I)-I)*X(I),VX)
          345 452 DO 453 I=1,19) (SUG(I)-I)*Y(I),VY)
          346 453 DO 454 I=1,19) (SUG(I)-I)*Z(I),VZ)
          347 454 DO 455 I=1,19) (SUG(I)-I)*X(I),VX)
```

FORTRAN IV G LEVEL 19

KEYIN

DATE = 71257

14/12/38

```

0039 NGG = 1
0040 CALL RCURS
0041 CALL RESET(GTEXT,GDSX,GDSY,GPOINT,GDSE,GDUSER,GINPUT)
0042 101 WRITE (4,1558)
0043 1658 FORMAT('PL PRESS ONE KEY CORRESPONDING TO THE AXIS.'/
1'P - VARIABLE YOU WISH TO SELECT, OR 3G IF.'/
2'P YOU WISH TO STOP,IF YOU WISH TO MAKE.'/
3'P CHANGES IN YOUR VECTOR PRESS KEY 28.')
CALL WRKFTS(GTEXT)
LP = PLCT(GTEXT)
60 NAXIS = DETAIN(NDET)
NAXIS = ADET(4)
IF (NAXIS.EQ.3G) RETURN
IF (NAXIS.EQ.28) GO TO 99
IF (NAXIS.LE.NV) GO TO 54
LP = UNPLOI(GTEXT)
CALL PESET(GTEXT,GDSX,GDSY,GPOINT,GDSE,GDUSER,GINPUT)
WRITE(4,52)
52 FORMAT('PL ERROR - PRESS PROPER KEY')
CALL WRKFTS(GTEXT)
LP = PLCT(GTEXT)
GO TO 60
54 LP = UNPLOI(GTEXT)
CALL RESET(GTEXT,GDSX,GDSY,GPOINT,GDSE,GDUSER,GINPUT)
RETURN
END

```

NOT REPRODUCIBLE

```

FORTRAN IV C LEVEL 19          RLISE          DATE = 71257          14/32/48

0001      SUBROUTINE ELLIPSE (INXIS)
0002      DIMENSION ELPSP(5,72),ELPSP(5,72),U(400),V(100)
0003      COMMON GDSX,GDSY,GJEXT,IP(2),GDSX,GDSY,X(50),Y(50),NV,NV,NGE,MSOL,
          INGL,CLMEANIS,EL,AL(36,11),XN(18),YIN(18),XLO,XLO,XHI,XHI,ADJET(5),
          ZR(16),XNANG(16),GJPUT,RI(16),SUS(10,8),INXSOL(50,2),IGLMAX(50
          3),ALG(5),GDSX(136),GDSY(136),D(5),D(5),D(5),D(5),D(5),D(5),D(5),D(5),
          EQUIVALENCE (IGMOR(1),ILP(1,1),V(1)),(ELPSP(1,1),U(1))
          DO 10 K=1,NV
0004      20 R(N) = C.O
0005      R(NXIS) = 1.0
0006      NV = NV+2
0007      KK = NV + 1
0008      DO 11 K=KK,NV
0009      11 R(N) = RI(16),NV
0010      AL(NV,NXIS) = C.O
0011      SS = 0.0
0012      DO 20 K=KK,NV
0013      20 20 K=KK,NV
0014      SS = SS + R(N)**2
0015      IF (SS.GT.C(1) GO TO 21)
0016      21 CONTINUE
0017      IF (SS.GT.C(1) GO TO 21)
0018      213 WRITE (14,1659)
0019      1659 FORMAT('P1.2A. INCONSISTENT PRESS KEY 20 AND NEW VECTOR')
0020      CALL NRPIS (GJEXT)
0021      IP = PLUGIGJEXT)
0022      RETURN
0023      DO 21 K=KK,NV
0024      R(N) = R(K)/SQRT(SS)
0025      21 CONTINUE
0026      NV = NV-(NV-1)/2
0027      DO 12 K=1,NV
0028      12 A(NANG(K)) = A(K,1)
0029      CALL XPRU(AMANG,AMANG,K,PC,NV,NV,1,C,2)
0030      CALL GIPR(K,N,NANG,AMANG,NV,2,2)
0031      S(1) = A(RNG(1))
0032      S(2) = A(RNG(2))
0033      S(3) = A(RNG(3))
0034      DO 13 K=1,NV
0035      13 U(1-K)=NP+K(1) = R(K,K)
0036      CALL GIPR(10,R,NV,NV,2)
0037      DO 14 K=1,NP
0038      U(K) = V(K)
0039      14 V(K) = V(NP+K)
0040      VMAX = V(1)
0041      VMIN = V(1)
0042      VMIN = V(1)
0043      DO 15 K=2,NP
0044      IF (V(K).GT.VMAX) VMAX = V(K)
0045      IF (V(K).LT.VMIN) VMIN = V(K)
0046      15 CONTINUE
0047      X(1) = 1-0.2*AMAX(XN(1),-5*MIN(XN(1),1))/1.6
0048      X(1) = (1-5*AMAX(XN(1),-5*MIN(XN(1),1))/1.6
0049      Y(1) = 1-0.2*AMAX(YN(1),-5*MIN(YN(1),1))/1.6
0050      Y(1) = (1-5*AMAX(YN(1),-5*MIN(YN(1),1))/1.6
0051      CALL UGRO (XLO,YLO,XHI,YHI)
0052      CALL PLAC(IGDSX,XMIN(NXIS),YMAX)
0053      CALL LINE(IGDSX,XMIN(NXIS),YMIN)

```

NOT REPRODUCIBLE

NOT REPRODUCIBLE

FORTRAN IV G LEVEL 19 ELLPSE DATE = 71257 14/12/48

```

0103 CALL POINT(GDSE,ELPSX(K,KK),ELPSY(K,KK))
0104 ICC CONTINUE
0105 LP = PLCT(GDSE)
0106 KK = NV + 1
0107 NVV = NV*2
0108 WRITE (4,1657)NAXIS, (R(K),K=KK,NVV)
0109 1657 FORMAT('PL',ICX,'VARIABLE',I2,'AGAINST VECTOR',/
0110 2:P',8F6.3)
0111 CALL XRFMIS(GTEXT)
0112 LP = PLUIGTEXT)
0113 RETURN
0114 END
0115
0116
0117
0118

```

NOT REPRODUCIBLE

NOT REPRODUCIBLE

```

FORTRAN IV G LEVEL 19      RELPSE      DATE = 71257      14/33/01

0051 RELPSY(K,KK)=RWRKNG(NG+X)+ELX*SQRT(1.-ANGLE**2)*SIG*ELV*ANGLE
0052 160 CONTINUE
0053 DO 170 K=1,NG
0054 DO 170 KK=1,72
0055 CALL POINT(GDUSER,RELPSX(K,KK),RELPSY(K,KK))
0056 170 CONTINUE
0057 LP = PLCT(GDUSER)
0058 RETURN
0059 211 NA = -15
0060 RETURN
0061 END

```