

AD-A123 758

MULTI-SAMPLE CLUSTER ANALYSIS USING AKAIKE'S
INFORMATION CRITERION(U) ILLINOIS UNIV AT CHICAGO
CIRCLE DEPT OF QUANTITATIVE METHODS H BOZDOGAN ET AL.
20 DEC 82 TR-82-6 N00014-80-C-0408

1/1

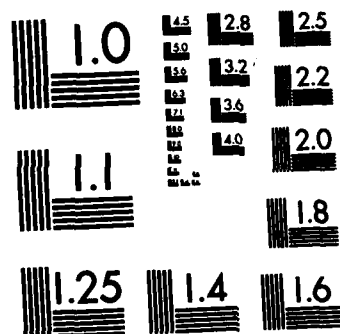
UNCLASSIFIED

F/G 12/1

NL

END

DTIC



MICROCOPY RESOLUTION TEST CHART
NATIONAL BUREAU OF STANDARDS-1963-A

(12)

MULTI-SAMPLE CLUSTER ANALYSIS
USING AKAIKE'S INFORMATION CRITERION*

by

HAMPARSUM BOZDOGAN
Department of Quantitative Methods
University of Illinois at Chicago

and

STANLEY L. SCLOVE
Department of Quantitative Methods
University of Illinois at Chicago

To appear in ANNALS OF INSTITUTE OF STATISTICAL MATHEMATICS, Paper 8

TECHNICAL REPORT NO. 82-6
December 20, 1982

PREPARED FOR THE
OFFICE OF NAVAL RESEARCH
UNDER
CONTRACT N00014-80-C-0408,
TASK NR042-443
with the University of Illinois at Chicago
Principal Investigator: Stanley L. Sclove

Reproduction in whole or in part is permitted
for any purpose of the United States Government.

Approved for public release; distribution unlimited

QUANTITATIVE METHODS DEPARTMENT
COLLEGE OF BUSINESS ADMINISTRATION
UNIVERSITY OF ILLINOIS AT CHICAGO
BOX 4348, CHICAGO, ILLINOIS 60680

DTIC
JAN 27 1983

A

*Presented by the first author as an Invited Paper, Special Session on
Cluster Analysis, 789th Meeting, American Mathematical Society, University of
Massachusetts, Amherst, MA, October 16-18, 1981.

83 01 27 002

ADA 191 233

DTIC FILE COPY

MULTI-SAMPLE CLUSTER ANALYSIS
USING AKAIKE'S INFORMATION CRITERION*

Hamparsum Bozdogan and Stanley L. Sclove
University of Illinois at Chicago

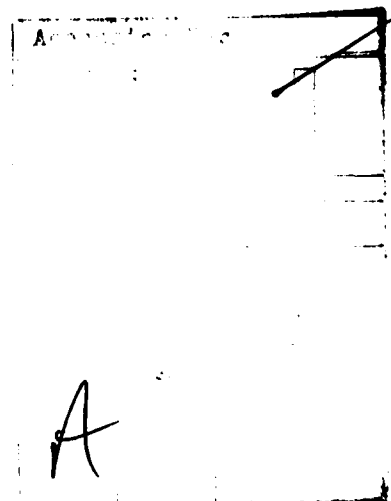
CONTENTS

Abstract; Key Words and Phrases

1. Introduction
2. The Multi-Sample Cluster Problem
3. Derivation of AIC for Two Multivariate Models
 - 3.1. AIC for the Multivariate Analysis of Variance (MANOVA)
Model: AIC (common Σ)
 - 3.2. AIC for the Multivariate Model with Varying Parameters:
AIC(varying μ and Σ)
4. Numerical Examples of Multi-Sample Cluster Analysis on Real Data Sets
5. Conclusions and Discussions

Acknowledgements

References



*Presented by the first author as an Invited Paper, Special Session on Cluster Analysis, 789th Meeting, American Mathematical Society, University of Massachusetts, Amherst, MA, October 16-18, 1981.

MULTI-SAMPLE CLUSTER ANALYSIS
USING AKAIKE'S INFORMATION CRITERION*

Hamparsum Bozdogan and Stanley L. Sclove
University of Illinois at Chicago

ABSTRACT

Multi-sample cluster analysis, the problem of grouping samples, is studied from an information-theoretic viewpoint via Akaike's Information Criterion (AIC). This criterion combines the maximum value of the likelihood with the number of parameters used in achieving that value. The multi-sample cluster problem is defined, and AIC is developed for this problem. The form of AIC is derived in both the multivariate analysis of variance (MANOVA) model and in the multivariate model with varying mean vectors and variance-covariance matrices. Numerical examples are presented for AIC and another criterion called w-square. The results demonstrate the utility of AIC in identifying the best clustering alternatives.

Key Words and Phrases: Multi-sample cluster analysis; w-square Criterion; Akaike's Information Criterion (AIC); MANOVA model; multivariate model with varying mean vectors and variance-covariance matrices; maximum likelihood.

*Presented by the first author as an Invited Paper, Special Session on Cluster Analysis, 789th Meeting, American Mathematical Society, University of Massachusetts, Amherst, MA, October 16-18, 1981.

MULTI-SAMPLE CLUSTER ANALYSIS
USING AKAIKE'S INFORMATION CRITERION*

Hamparsum Bozdogan and Stanley L. Sclove
University of Illinois at Chicago

1. Introduction

In this paper, we shall develop Akaike's Information Criterion (AIC) for multi-sample cluster analysis with common and also with varying variance-covariance matrices, since often in practice the assumption of equal variance-covariance matrices is a rather dubious requirement.

The problem of multi-sample cluster analysis arises when we are given a collection of samples (groups, treatments), to be clustered into homogeneous groups.

Many practical situations require the presentation of multivariate data from several structured samples for comparative purposes and the grouping of the heterogeneous samples into homogeneous sets of samples. Thus, it is reasonable to provide a practically useful statistical procedure that would use some sort of statistical model to aid in comparisons of various collections of samples, identify homogeneous groups of samples, telling us which samples should be clustered together and which should not.

Examples of multi-sample clustering situations are abundant in practice. We shall give two of these examples later and illustrate numerically.

The concept of multi-sample cluster analysis presented in this paper is relatively new. It has not been definitively studied before either using the conventional simultaneous test procedures (STP's) which are based on inference for the multivariate analysis of variance (MANOVA) model, or from an information-theoretic viewpoint, which we shall adopt in this paper via Akaike's

*Presented by the first author as an Invited Paper, Special Session on Cluster Analysis, 789th Meeting, American Mathematical Society, University of Massachusetts, Amherst, MA, October 16-18, 1981.

Information Criterion (AIC).

Multivariate analysis of variance (MANOVA) is a widely used model for comparing two or more multivariate samples with a common covariance matrix. In this model, the likelihood ratio principle leads to Wilks' [17] lambda, or in short Wilks' A Criterion as the test statistic. It plays the same role in multivariate analysis that F-ratio statistic plays in the univariate case.

Often, however, the formal analyses involved in MANOVA are not revealing or informative. Moreover, the test statistics used under this model are derived under the assumption of equal covariance matrices. If we have a reason to doubt equality of covariances, then we may first want to test the equality of covariances. In the multivariate case the equality of covariance matrices is certainly more hazardous. If the covariance matrices are unequal, a bias occurs in the test for equality of mean vectors. Therefore, for this reason we may want to first test the equality of covariance matrices instead of immediately leaping to the MANOVA hypothesis. This is an important option to use in clustering groups or samples when we are not willing to assume equal covariance matrices between the samples or groups in the multi-sample data.

Once the MANOVA hypothesis of equality of mean vectors is rejected at some prescribed significance level α , then it is necessary to study in detail the discrepancies between the null hypothesis and the data.

In the statistical literature, in the MANOVA case, there are a variety of conventional multiple comparison procedures for studying the discrepancies between the null hypothesis and the data. These test procedures are: Step-down Methods, Union Intersection Tests, and Simultaneous Confidence Intervals. For more details on these test procedures refer to Gabriel [7], Krishnaiah ([10], [11]), Srivastava [16], and others.

As noted in Consul [4], the exact distributions of these conventional test procedures are either unknown or are known for some particular cases only. Moreover, the problem of finding the percentage points of these statistics has become rather difficult. For these reasons, and for our purposes, these test procedures have little practical use. Furthermore, they create additional problems in terms of how to control the overall error rate α , since we can no longer use the same α to discover where the discrepancies between the null hypothesis and the data might occur.

In the case of testing the equality of covariance matrices, we find ourselves in the same situation as in the MANOVA model. For this problem also, there are in the statistical literature several test procedures. For example, one of the most commonly used tests is Box's M test despite the fact that it is very restrictive. For instance, Box's approximation seems to be only good if each sample size, n_g exceeds 20, and if the number of samples, K , and the number of variables, p , exceed 5. It is also very expensive to compute it on a high speed computer, even on an IBM 370.

Once the hypothesis of equality of covariance matrices is rejected at some prescribed significance level α , then again it is necessary to study in detail the discrepancies between the null hypothesis and the data.

Further reviewing the statistical literature, we see that there are no conventional simultaneous test procedures (STP's) in this case in studying the discrepancies between the null hypothesis and the data. One can perhaps construct a sequential likelihood ratio type test, but as is mentioned in Muirhead ([14], p. 296), the likelihood ratio test in testing the equality of covariance matrices has the defect that, when the sample sizes $n_1, n_2, \dots, n_K$ are not all equal, it is biased. Therefore, in the multi-sample cluster

problem with varying parameters, carrying out a sequence of likelihood ratio tests leaves much to be desired in identifying homogeneous groups of samples.

More recently, however, in the statistical literature, we see a likelihood based approach, called w-square criterion given in Mardia et al. [12] to aid in comparing various collections of samples, identifying homogeneous groups of samples, and telling which should be clustered together. For normal samples with equal covariance matrices, the w-square criterion is defined by

$$(1.1) \quad w_a^2 = \sum_{g=1}^K \sum_{\bar{x}_j \in C_g} n_j (\bar{x}_j - \bar{x}_g)' \hat{\Sigma}^{-1} (\bar{x}_j - \bar{x}_g)$$

where

C_g = the set of $\bar{x}_j$ assigned to the g th group, $g=1, 2, \dots, K$,

$\bar{x}_g$ = the weighted mean vector of the means in the g th group, or the cluster set C_g of groups,

$\hat{\Sigma}$ = $W/(n-K)$, the pooled estimate of Σ ,

$W = \sum_{g=1}^K A_g$ is the within-samples SSP matrix,

$n = \sum_{g=1}^K n_g$, and

K = the number of groups or samples to be clustered.

If the matrix of Mahalanobis distances D_{ij} given by

$$D_{ij}^2 = (\bar{x}_i - \bar{x}_j)' \hat{\Sigma}^{-1} (\bar{x}_i - \bar{x}_j)$$

is available, then for computational convenience, w_a^2 can be written as

$$(1.2) \quad w_a^2 = 1/2 \sum_{g=1}^K N_g^{-1} \sum_{C_g} n_i n_j u_{ij}^2$$

where

$$N_g = \sum_{x_j \in C_g} n_j.$$

Thus, when we are given multi-sample data and wish to cluster the samples, we compute w_a^2 in (1.1) or (1.2) for some or all of the alternative groupings of samples, and choose the minimum of w_a^2 to be the "best" alternative clustering of samples. This is appropriate, since maximizing the likelihood implies minimizing w_a^2 .

Even though the w_a^2 criterion is a step forward in identifying homogeneous groups of samples and evaluating multi-sample clusters, it has some disadvantages. For instance, it does not make any allowance for m , the number of parameters estimated within the model and the subsequent alternative submodels. It is always zero when the groups or samples are clustered as singletons, as we shall see later in Section 4. As it is given in (1.1), we can only work with w_a^2 criterion when we assume equal covariance matrices.

For the above stated reasons, and the problems encountered in the conventional test procedures which we discussed above, in this paper we shall propose Akaike's Information Criterion (AIC) as a new and unifying procedure for evaluating multi-sample clusters, and use it to identify the best clustering alternatives.

In 1971, Akaike first introduced an information criterion, referred to as

a Model Identification Criterion or Akaike's Information Criterion (AIC), for the identification and comparison of statistical models in a class of competing models with different numbers of parameters. It is defined by

$$(1.3) \quad \text{AIC} = (-2)\log_e(\text{maximized likelihood}) \\ + 2 (\text{number of free parameters within the model}).$$

It was obtained by Akaike ([2], [3]) based on the recognition that the classical method of maximum likelihood could be viewed as a method of identification of a statistical model realized by maximizing an estimate of the generalized entropy, or the expected log likelihood, of the model being fitted. It estimates minus twice the expected log likelihood of the model whose parameters are determined by the method of maximum likelihood. When several competing models are being compared or fitted, AIC is a simple procedure which measures the badness of fit or the discrepancy of the estimated model from the true model when a set of data is given. The first term in (1.3) stands for the penalty of badness of fit when the maximum likelihood estimators of the parameters of the model are used. The second term in the definition of AIC, on the other hand, stands for the penalty of increased unreliability or compensation for the bias in the first term as a consequence of increasing number of parameters. If more parameters are used to describe the data, it is natural to get a larger likelihood, possibly without improving the true goodness of fit by penalizing the use of additional parameters.

Thus, when there are several competing models, the parameters within the models are estimated by the method of maximum likelihood and the AIC-values are computed and compared to find a model with the minimum value of AIC. This

procedure is called the minimum AIC procedure. The model with the minimum AIC is called the minimum AIC estimate (MAICE) and is designated as the best model. Thus, in applying AIC the emphasis is on comparing the "goodness of fit" of various models with an allowance made for parsimony.

In Section 2, we shall define the general multi-sample cluster problem. In Section 3, we shall derive the AIC procedure both for the multivariate analysis of variance (MANOVA) model, and for the multivariate model with varying covariance matrices. We shall, in Section 4, give different numerical examples of multi-sample cluster analysis on different real data sets to demonstrate our results from applying minimum AIC procedures in different computer analyses. Finally, in Section 5, we shall present our conclusions and discussion.

2. The Multi-Sample Cluster Problem

Suppose each individual, object, or case, has been measured on p response or outcome measures (dependent variables) simultaneously in K independent groups or samples (factor levels). Let

$$(2.1) \quad \underline{X} \ (n \times p) = \begin{bmatrix} \underline{x}_1 \\ \underline{x}_2 \\ \vdots \\ \underline{x}_K \end{bmatrix}$$

be a single data matrix of K groups or samples, where $\underline{x}_g \ (n_g \times p)$ represents the observations from the g -th group or sample, $g=1,2,\dots,K$, and $n = \sum_{g=1}^K n_g$. The goal of cluster analysis is to put the K groups or samples into k homogeneous

groups, samples, or classes where k is unknown, but $k \leq K$.

Often individuals or objects have been sampled from $K > 1$ populations. The data matrix may be represented in partitioned form as above. Let n_g represent the number of individuals in the g -th (random) sample, $g=1,2,\dots,K$. The n_g are not restricted to being equal or proportional to other n_g 's. The total number of observations is $n = \sum_{g=1}^K n_g$. Let $\underline{x}_{gi}$ be the $p \times 1$ vector of observations in group $g=1,2,\dots,K$, and for individual $i=1,2,\dots,n_g$.

3. Derivation of AIC for Two Multivariate Models

3.1 AIC for the Multivariate Analysis of Variance (MANOVA) Model:

AIC (common $\underline{\Sigma}$)

We now turn our attention to consider situations with several multivariate normal samples.

For example, we may have multi-sample data with sample sizes $n_1, n_2, \dots, n_K$ which are assumed to come from K populations, the first with mean vector $\underline{\mu}_1$ and covariance matrix $\underline{\Sigma}$, the second with mean vector $\underline{\mu}_2$ and covariance matrix $\underline{\Sigma}$, ..., the K th with mean vector $\underline{\mu}_K$ and covariance matrix $\underline{\Sigma}$. Therefore, throughout this section we shall suppose that we may have independent data matrices $\underline{X}_1, \underline{X}_2, \dots, \underline{X}_K$, where the rows of $\underline{X}_g (n_g \times p)$ are independent and identically distributed (i.i.d.) according to a multivariate normal distribution, $N_p(\underline{\mu}_g, \underline{\Sigma})$, $g=1,2,\dots,K$. We may want to compare the K sample mean vectors given that all K distributions have a common covariance matrix $\underline{\Sigma}$. This is the well known multivariate analysis of variance (MANOVA) model. In terms of the parameters the MANOVA model is $\underline{\theta} = (\underline{\mu}_1, \underline{\mu}_2, \dots, \underline{\mu}_K, \underline{\Sigma})$ with $m = kp + p(p+1)/2$ parameters, where k is the number of groups, and p is the number of variables.

We shall derive the form of AIC for this model. Recall the definition of AIC from Section 1,

$$\begin{aligned} \text{AIC} &= -2 \log_e L(\hat{\theta}) + 2m \\ &= -2 \log_e (\text{maximized likelihood}) + 2m, \end{aligned}$$

where m denotes the number of free parameters within the model.

Consider K normal populations with different mean vectors μ_g , $g=1,2,\dots,k,\dots,K$. Let x_{gi} , $g=1,2,\dots,K$; $i=1,2,\dots,n_g$, be a random sample of observations from the g -th population $N_p(\mu_g, \Sigma)$. If the groups or samples can differ only in their mean vectors, we can write the multivariate one-way analysis variance (MANOVA) model as

$$(3.1.1) \quad x_{gi} = \mu_g + \varepsilon_{gi}, \quad g=1,2,\dots,K; \quad i=1,2,\dots,n_g,$$

where x_{gi} is the $(p \times 1)$ response or outcome vector in the g -th group for the i -th individual or object,
 μ_g are vector parameters, and
 ε_{gi} are independent $N_p(0, \Sigma)$ random vector errors.

Thus, the basic null hypothesis we usually are interested in testing is given by

$$(3.1.2) \quad H_0 : \mu_1 = \mu_2 = \dots = \mu_K.$$

The alternative hypothesis is given by

$$H_1 : \text{Not all } \mu_K \text{ are equal.}$$

Wilks' lambda is a general statistic for handling this problem. Although there are several other conventional statistics for this purpose, they all can be viewed as special cases of Wilks' Λ which we shall not discuss here.

For notational purposes, we shall denote by $\underline{I}$ the "total" sum of squares and products (SSP) matrix, by $\underline{W}$ the "within-group" or "within-sample" SSP matrix, and by $\underline{B}$ the "between-group" SSP matrix. Hence, it can be shown that

$$(3.1.3) \quad \underline{I} = \underline{W} + \underline{B} ,$$

where

$$(3.1.4) \quad \underline{I} = \sum_{g=1}^K \sum_{i=1}^{n_g} (\underline{x}_{gi} - \bar{\underline{x}})(\underline{x}_{gi} - \bar{\underline{x}})',$$

$$(3.1.5) \quad \underline{W} = \sum_{g=1}^K \sum_{i=1}^{n_g} (\underline{x}_{gi} - \bar{\underline{x}}_g)(\underline{x}_{gi} - \bar{\underline{x}}_g)',$$

and

$$(3.1.6) \quad \underline{B} = \sum_{g=1}^K n_g (\bar{\underline{x}}_g - \bar{\underline{x}})(\bar{\underline{x}}_g - \bar{\underline{x}})',$$

with

$$\bar{\underline{x}}_g = \frac{1}{n_g} \sum_{i=1}^{n_g} \underline{x}_{gi} , \quad g=1,2,\dots,K ,$$

$$\bar{\underline{x}} = \frac{1}{n} \sum_{g=1}^K \sum_{i=1}^{n_g} \underline{x}_{gi} , \quad n = \sum_{g=1}^K n_g .$$

Therefore, we can present multivariate one-way analysis of variance (MANOVA) table as follows.

TABLE 3.1. MANOVA TABLE

Source	d.f.	SSP matrix	Wilks' criterion
Between samples	K-1	$\underline{B}$	$\frac{ \underline{W} }{ \underline{I} }$
Within samples	n-K	$\underline{W}$	$-\Lambda(p ; n - K ; K - 1)$
Total	n-1	$\underline{I}$	

Now, we derive the form of Akaike's Information Criterion (AIC) for the MANOVA model given in (3.1.1), subject to the constraint given in (3.1.2).

The likelihood function of all the sample observations is given by

$$(3.1.7) \quad L(\underline{\mu}_g, \underline{\Sigma}_g; \underline{X}) = \prod_{g=1}^K L_g(\underline{\mu}_g, \underline{\Sigma}_g; \underline{X}_g),$$

or by

$$(3.1.8) \quad L = (2\pi)^{-np/2} \prod_{g=1}^K |\underline{\Sigma}_g|^{-n_g/2} \times$$

$$\exp \left\{ -1/2 \text{tr} \sum_{g=1}^K \underline{\Sigma}_g^{-1} \underline{A}_g - 1/2 \text{tr} \sum_{g=1}^K n_g \underline{\Sigma}_g^{-1} (\bar{\underline{x}}_g - \underline{\mu}_g)(\bar{\underline{x}}_g - \underline{\mu}_g)' \right\},$$

where $n = \sum_{g=1}^K n_g$ and $\underline{A}_g = \sum_{i=1}^{n_g} (\underline{x}_{gi} - \bar{\underline{x}}_g)(\underline{x}_{gi} - \bar{\underline{x}}_g)'$.

The log likelihood function is

$$(3.1.9) \quad l(\underline{\mu}_g, \underline{\Sigma}_g; \underline{X}) \equiv \log L(\underline{\mu}_g, \underline{\Sigma}_g; \underline{X})$$

$$\begin{aligned} &= - (np/2) \log(2\pi) - 1/2 \sum_{g=1}^K n_g \log |\underline{\Sigma}_g| - 1/2 \text{tr} \sum_{g=1}^K \underline{\Sigma}_g^{-1} \underline{A}_g \\ &\quad - 1/2 \text{tr} \sum_{g=1}^K n_g \underline{\Sigma}_g^{-1} (\bar{\underline{X}}_g - \underline{\mu}_g)(\bar{\underline{X}}_g - \underline{\mu}_g)' . \end{aligned}$$

Since the common covariance matrix is $\underline{\Sigma}$, the log likelihood function becomes

$$(3.1.10) \quad l(\{\underline{\mu}_g\}, \underline{\Sigma}; \underline{X}) \equiv \log L(\{\underline{\mu}_g\}, \underline{\Sigma}; \underline{X})$$

$$\begin{aligned} &= - (np/2) \log(2\pi) - (n/2) \log |\underline{\Sigma}| - 1/2 \text{tr} \underline{\Sigma}^{-1} \sum_{g=1}^K \underline{A}_g \\ &\quad - 1/2 \text{tr} \underline{\Sigma}^{-1} \sum_{g=1}^K n_g (\bar{\underline{X}}_g - \underline{\mu}_g)(\bar{\underline{X}}_g - \underline{\mu}_g)' , \end{aligned}$$

and the maximum-likelihood estimates (MLE's) of $\underline{\mu}_g$, and $\underline{\Sigma}$ are

$$(3.1.11) \quad \hat{\underline{\mu}}_g = \bar{\underline{X}}_g, \quad g=1, 2, \dots, K,$$

and

$$(3.1.12) \quad \hat{\underline{\Sigma}} = n^{-1} \underline{W},$$

where $\underline{W} = \sum_{g=1}^K \underline{A}_g$.

Substituting these back into (3.1.10) and simplifying, the maximized log likelihood becomes

$$\begin{aligned}
 (3.1.13) \quad l(\{\hat{\underline{u}}_g\}, \hat{\underline{\Sigma}}; \underline{X}) &\equiv \log L(\{\hat{\underline{u}}_g\}, \hat{\underline{\Sigma}}; \underline{X}) \\
 &= - (np/2) \log(2\pi) - (n/2) \log |n^{-1} \underline{W}| - (np/2),
 \end{aligned}$$

where $\underline{W}$ is the "within-group" SSP matrix.

Since

$$(3.1.14) \quad AIC = -2 \log_e L(\hat{\underline{\theta}}) + 2m,$$

where $m = kp + \frac{p(p+1)}{2}$ is the number of parameters, then AIC becomes

$$(3.1.15) \quad AIC \text{ (common } \underline{\Sigma}) = np \log_e(2\pi) + n \log_e |n^{-1} \underline{W}| + np + 2[kp + \frac{p(p+1)}{2}].$$

Since the constants do not affect the result of comparison of models, we could ignore them and reduce the form of AIC to a much simpler form

$$(3.1.16) \quad AIC^* \text{ (common } \underline{\Sigma}) = n \log_e |\underline{W}| + 2[kp + \frac{p(p+1)}{2}]$$

where $n = \sum_{g=1}^K n_g$ = the total sample size,

$|\underline{W}|$ = the determinant of "within-group" SSP matrix,

k = number of groups or samples compared,

p = number of variables.

However, for purposes of comparison we retain the constants and use

AIC (common $\underline{\Sigma}$).

3.2. AIC for the Multivariate Model with Varying Parameters:

AIC (varying $\underline{\mu}$ and $\underline{\Sigma}$)

As we mentioned in Section 1, the assumption of equality of covariance matrices in MANOVA can cause serious problems. For this reason we may want first test the equality of covariance matrices against the alternative that not all covariance matrices are equal, given no restriction on the population mean vectors. Therefore, throughout this section we shall suppose that we may have independent data matrices $\underline{X}_1, \underline{X}_2, \dots, \underline{X}_K$, where the rows of $\underline{X}_g$ ($n_g \times p$) are independent and identically distributed (i.i.d.) $N_p(\underline{\mu}_g, \underline{\Sigma}_g)$, $g=1, 2, \dots, K$. In terms of the parameters with varying mean vectors and covariance matrices, the multivariate model we shall consider is

$$\theta = (\underline{\mu}_1, \underline{\mu}_2, \dots, \underline{\mu}_K, \underline{\Sigma}_1, \underline{\Sigma}_2, \dots, \underline{\Sigma}_K)$$

with $m = kp + kp(p+1)/2$ parameters, where k is the number of groups, and p is the number of variables.

Thus, the basic null hypothesis we usually are interested in testing is given by

$$(3.2.1) \quad H_0: \underline{\Sigma}_1 = \underline{\Sigma}_2 = \dots = \underline{\Sigma}_K.$$

The alternative hypothesis is given by

$$H_1: \text{Not all } K \text{ covariance matrices are equal.}$$

In multivariate analysis this is known as the test of homogeneity of covariance matrices.

To derive Akaike's Information Criterion (AIC) in this case the log likelihood function is given by

$$(3.2.2) \quad l(\{\underline{\mu}_g, \underline{\Sigma}_g\}; \underline{X}) \equiv \log L(\{\underline{\mu}_g, \underline{\Sigma}_g\}; \underline{X})$$

$$\begin{aligned} &= - (np/2) \log(2\pi) - 1/2 \sum_{g=1}^K n_g \log |\underline{\Sigma}_g| - 1/2 \sum_{g=1}^K n_g \operatorname{tr} \underline{\Sigma}_g^{-1} \underline{A}_g \\ &\quad - 1/2 \sum_{g=1}^K n_g (\bar{\underline{x}}_g - \underline{\mu}_g)' (\bar{\underline{x}}_g - \underline{\mu}_g) . \end{aligned}$$

The MLE's of $\underline{\mu}_g$ and $\underline{\Sigma}_g$ are

$$(3.2.3) \quad \hat{\underline{\mu}}_g = \bar{\underline{x}}_g, \quad g=1,2,\dots,K,$$

and

$$(3.2.4) \quad \hat{\underline{\Sigma}}_g = \underline{A}_g/n_g.$$

Substituting these back into (3.2.2) and simplifying, the maximized log likelihood becomes

$$\begin{aligned} (3.2.5) \quad l(\{\hat{\underline{\mu}}_g, \hat{\underline{\Sigma}}_g\}; \underline{X}) &\equiv \log L(\{\hat{\underline{\mu}}_g, \hat{\underline{\Sigma}}_g\}; \underline{X}) \\ &= - (np/2) \log(2\pi) - 1/2 \sum_{g=1}^K n_g \log |n_g^{-1} \underline{A}_g| - (np/2). \end{aligned}$$

Since

$$(3.2.6) \quad \text{AIC} = -2 \log_e L(\hat{\underline{\theta}}) + 2m,$$

where $m = kp + kp(p+1)/2$ is the number of parameters, then AIC becomes

$$(3.2.7) \quad AIC(\text{varying } \underline{\mu} \text{ and } \underline{\Sigma}) = n \log_e (2\pi) + \sum_{g=1}^K n_g \log_e |\underline{A}_g|^{-1} + np + 2[kp + kp(p+1)/2].$$

Since the constants do not affect the result of comparison of models, we could ignore them and reduce the form of AIC to a much simpler form

$$(3.2.8) \quad AIC^*(\text{varying } \underline{\mu} \text{ and } \underline{\Sigma}) = \sum_{g=1}^K n_g \log_e |\underline{A}_g| + 2[kp + kp(p+1)/2],$$

where n_g = sample size of group or sample $g=1,2,\dots,K$,

$|\underline{A}_g|$ = the determinant of sum of squares and cross-products (SSCP) matrix for group or sample $g=1,2,\dots,K$,

k = number of groups or samples compared, and

p = number of variables.

However, for purposes of comparison we retain the constants and use AIC given by (3.2.7).

4. Numerical Examples of Multi-Sample Cluster Analysis on Real Data Sets

In this section we shall give two different numerical examples of multi-sample cluster analysis, cluster the samples, and choose the best clusterings by using Akaike's Information Criterion (AIC) as derived in Section 3.1 and 3.2. In example 4.2 we shall also present the numerical results of using the w-square criterion as an alternative approach. We shall briefly discuss the relative merits of AIC over w-square criterion. One should note that these criteria are qualitatively and quantitatively different.

Our computations were carried out for all the examples we shall present here on an IBM 4341, configured as a 370, by using a newly developed

statistical library software by the first author for multi-sample cluster analysis using AIC, called MSCA.AIC.

We shall illustrate our results first on the Fisher [5] iris data.

Example 4.1. Clustering of Irises by Groups: The iris data set is composed of 150 iris species belonging to three groups or species, namely Iris setosa (S), Iris versicolor (Ve), and Iris virginica (Vi) measured on sepal and petal length and width. Each group is represented by 50 plants.

This data set has been quite extensively studied in classification and cluster analysis since it was published by Fisher [5], and still today, is being used as a "testing ground" for classification and clustering methods proposed by many investigators such as Friedman and Rubin [6], Kendall [8], Solomon [15], Mezzich and Solomon [13], and many others, including the present authors.

For each of the 150 plants we already know the group structure of the iris species, namely $K=3$ groups or samples. Even though the two species, Iris setosa and Iris versicolor were found growing in the same colony, and Iris virginica was found growing in a different colony, Fisher reports in his linear discriminant analysis the separation of I. setosa completely from I. versicolor and I. virginica. Since then other investigators have shown similar results in their studies such as the ones we mentioned above.

With this in mind, we cluster $K=3$ samples (species) into $k=1,2$, and 3 groups on the basis of all the four variables. We obtain in total five possible clustering alternatives. (In general, the total number of possibilities is a Stirling Number of the Second Kind; see, e.g., Abramowitz and Stegun [1]). Denoting I. setosa by S, I. versicolor by Ve, and I. virginica by Vi, we have (S) (Ve) (Vi), (S, Ve) (Vi), (S, Vi) (Ve), (Ve, Vi)(S),

and (S, Ve, Vi) as the possible clustering alternatives. Using the MANOVA model and the multivariate model with varying parameters discussed in Sections 3.1 and 3.2 as our underlying models for clustering these three iris species, we obtained the following results.

TABLE 4.1. THE AIC'S FOR IRISES BY GROUPS ON ALL VARIABLES UNDER MANOVA MODEL

Alternative	Clustering	$n \log_e(2\pi)$	$n \log_e n^{-1}W $	np	k	2m	AIC (common Σ)
1	(S) (Ve) (Vi)	1,102.724	-1,504.2	600	3	44	242.524a
2	(S, Ve) (Vi)	1,102.724	-1,085.9	600	2	36	652.824
3	(S, Vi) (Ve)	1,102.724	- 988.39	600	2	36	750.334
4	(Ve, Vi) (S)	1,102.724	-1,299.6	600	2	36	439.124b
5	(S, Ve, Vi)	1,102.724	- 941.73	600	1	28	788.994

$n = 150$ plants, $p = 4$ variables

$m = kp + p(p+1)/2$ parameters

$AIC(\text{common } \Sigma) = n \log_e(2\pi) + n \log_e |n^{-1}W| + np + 2m$

^aFirst Minimum AIC

^bSecond Minimum AIC

TABLE 4.2. THE AIC'S FOR IRISES BY GROUPS ON ALL VARIABLES UNDER THE MODEL WITH VARYING PARAMETERS

Alternative	Clustering	$n \log_e(2\pi)$	$\sum_{g=1}^K n_g \log_e n_g^{-1}A_g $	np	k	2m	AIC (varying μ and Σ)
1	(S) (Ve) (Vi)	1,102.724	-1,653.895	600	3	84	132.829a
2	(S, Ve) (Vi)	1,102.724	-1,251.675	600	2	56	507.049
3	(S, Vi) (Ve)	1,102.724	-1,144.480	600	2	56	614.244
4	(Ve, Vi) (S)	1,102.724	-1,463.770	600	2	56	294.954b
5	(S, Ve, Vi)	1,102.724	- 941.580	600	1	28	789.144

$n = 150$ plants, $p = 4$ variables

$m = kp + kp(p+1)/2$ parameters

$AIC(\text{varying } \mu \text{ and } \Sigma) = n \log_e(2\pi) + \sum_{g=1}^K n_g \log_e |n_g^{-1}A_g| + np + 2m$

^aFirst Minimum AIC

^bSecond Minimum AIC

Looking at Tables 4.1 and 4.2, we see that, using all four variables simultaneously under both models, the MAICE clustering is (S) (Ve) (Vi). This indicates that indeed there are three types of species. Not surprisingly, the the second minimum AIC occurs at the alternative submodel 4 (Ve, Vi) (S), under both models, telling us that if we were to cluster any one of the two iris groups, we should cluster I. versicolor and I. virginica together as one homogeneous group, and we should cluster I. setosa completely separately. We note that the AIC values under submodel 2 and 3 are quite large indicating the inferiority of these submodels. We can see the effect of clustering I. setosa with I. versicolor in submodel 2, and also with I. virginica in submodel 3, by comparing the difference of AIC's in these submodels with that of submodel 4 in which I. versicolor and I. virginica were clustered together and I. setosa was clustered as a separate cluster on its own. According to AIC, we never cluster three iris species as one homegeneous group (submodel 5). Again by comparing the differences of AIC's of submodel 5 with that of submodels 4, 3, and 2, respectively, we can measure the amount of heterogeneity contributed by I. setosa, I. versicolor, and I. virginica, respectively, in each clustering alternative under the MANOVA model and the multivariate model with varying mean vectors and covariance matrices. The larger this difference, the greater the heterogeneity or separation of that group or sample from that of homegeneous groups or samples in each clustering alternative.

In comparing the AIC's in Tables 4.1 and 4.2, we further notice that AIC (varying μ and Σ) values are much less than the AIC (common Σ) values for each of the clustering alternatives except for the last clustering alternative (i.e., alternative 5) in clustering the iris groups or species. Since according to the definition of AIC, the model with the minimum AIC is chosen to be the best model,

then the above results suggest that when we are clustering iris data, and in general, we should use different covariance matrices rather than using equal covariance matrices.

We now present our results on the iris data by using the w-square criterion given by (1.1) in Section 1, when we assume equal covariance matrices between the iris groups or species. We should note here that in w-square criterion given by (1.1) and in Mardia et al. ([12], p. 367), the estimated pooled-within groups covariance matrix of $\underline{\Sigma}$ is computed only once across all the groups or samples to be clustered regardless of the number of clustering alternatives. In our version of w-square criterion we follow the same procedure, but we recompute the estimate of $\underline{\Sigma}$ in each clustering alternative when we vary the number of clusters of groups or samples, k, when we are given, K, the number of groups or samples to be clustered. We do this both under the assumption of equal and separate covariance matrices between the iris groups. Therefore, our numerical values on w-square criterion are quite different then the original w-square criterion given in Mardia et al. [12], despite the fact that we get the same results.

We give the computational results as follows.

TABLE 4.3. THE VALUES OF w_a^2 FOR IRISES BY GROUPS ON ALL VARIABLES

Alternative	Clustering	w_a^2 (common $\underline{\Sigma}$) ^a	w_a^2 (common $\underline{\Sigma}$) ^b	w_a^2 (varying $\underline{\Sigma}$) ^c
1	(S) (Ve) (Vi)	-----*	-----*	-----*
2	(S, Ve) (Vi)	2246.6046	137.9722	94.4149
3	(S, Vi) (Ve)	4484.6178	142.3345	96.0706
4	(Ve, Vi) (S)	430.0267**	109.5511**	76.8212**
5	(S, Ve, Vi)	4774.1661	175.2091	175.2091

$n = 150$ plants, $p = 4$ variables

^aOriginal w_a^2 given in (1.1)

^{b,c}Our version of w_a^2

* w_a^2 cannot be computed (always equal to zero)

**Minimum of w_a^2

Hence, we interpret the results in Table 4.3 in the same manner as we did for AIC's. We see that at the alternative submodel 1, w_a^2 cannot be computed and is always equal to zero when the iris groups are clustered as singletons. This is always the case in general. Certainly this is a definite disadvantage of w_a^2 as compared to AIC which has a value even if the iris groups are clustered as singletons, so that AIC can aid us in determining and understanding the amount of heterogeneity or separation of the groups on a unique scale. The minimum of w_a^2 occurs at the alternative submodel 4, telling us again that, if we were to cluster any one of the two iris groups, we should cluster I. versicolor and I. virginica together as one homogeneous group, and we should cluster I. setosa completely separate as one heterogeneous group.

In short, w-square criterion gives the same results as AIC does, but as we mentioned in Section 1, it does not make any allowance for m , the number of parameters estimated within the clustering alternatives. AIC makes such an allowance to achieve a parsimony when we compare "the goodness of fit" of various models as we do in comparing different clustering alternatives. W-square criterion is short of having this important feature. Also as we saw, when we have singleton clusters, it cannot be computed.

Therefore, in our next example, we shall only give our results on AIC, since our purpose is to introduce AIC in this paper as a new approach to be used in evaluating multi-sample clusters.

Example 4.2. Clustering Graduate Students by Their Classification Groups:

A data set for applicants to admission to a Graduate School of Business given in Johnson and Wichern ([9], p. 528) is composed of data for 85 applicants who were classified by the admissions officer as Admit (A), Not Admit (NA), and Borderline (B), based on undergraduate grade point average (GPA) and graduate management aptitude test (GMAT) scores. The group sizes are $n_1 = 31$, $n_2 = 28$, and $n_3 = 26$ applicants.

With this in mind, we cluster $K=3$ groups of applicants into $k=1,2$ and 3 homogeneous groups on the basis of the two variables. Using the MANOVA model and the multivariate model with varying parameters, our results are as follows.

TABLE 4.4. THE AIC'S FOR APPLICANTS BY THEIR CLASSIFICATION GROUPS UNDER MANOVA MODEL

Alternative	Clustering	$n \log_e(2\pi)$	$n \log_e n^{-1}W $	np	k	$2m$	AIC (common Σ)
1	(A) (NA) (B)	312.4391	406.1716	170	3	18	906.6107 ^a
2	(A,NA) (B)	312.4391	566.7477	170	2	14	1063.1868
3	(A,B) (NA)	312.4391	491.7043	170	2	14	988.1434 ^c
4	(B,NA) (A)	312.4391	474.0420	170	2	14	970.4811 ^b
5	(A, B, NA)	312.4391	581.9931	170	1	10	1074.4322

$n = 85$ applicants, $p = 2$ variables

$m = kp + p(p+1)/2$ parameters

$AIC(\text{common } \Sigma) = n \log_e(2\pi) + n \log_e |n^{-1}W| + np + 2m$

^a First Minimum AIC

^b Second Minimum AIC

^c Third Minimum AIC

TABLE 4.5. THE AIC'S FOR APPLICANTS BY THEIR CLASSIFICATION GROUPS UNDER THE MODEL WITH VARYING PARAMETERS

Alternative	Clustering	$n \log_e(2\pi)$	$\sum_{g=1}^K n_g \log_e n_g^{-1} \underline{A}_g $	np	k	2m	AIC (varying $\underline{\mu}$ and $\underline{\Sigma}$)
1	(A) (NA) (B)	312.4391	388.7472	170	3	30	901.1863 ^a
2	(A,NA) (B)	312.4391	509.4198	170	2	20	1011.8589
3	(A,B) (NA)	312.4391	480.2378	170	2	20	982.6769 ^c
4	(B,NA) (A)	312.4391	465.7116	170	2	20	968.1507 ^b
5	(A, B, NA)	312.4391	581.9931	170	1	10	1074.4322

n = 85 applicants, p = 2 variables

m = kp + kp(p+1)/2 parameters

$$AIC = (\text{varying } \underline{\mu} \text{ and } \underline{\Sigma}) = n \log_e(2\pi) + \sum_{g=1}^K n_g \log_e |n_g^{-1} \underline{A}_g| + np + 2m$$

^a First Minimum AIC

^b Second Minimum AIC

^c Third Minimum AIC

Hence, looking at Tables 4.4 and 4.5, we see that, under both models, the first minimum AIC occurs at the alternative submodel 1, that is, when (A) (NA) (B) are all clustered separately. This indicates that indeed there are three groups of applicants. Therefore, the MAICE is submodel 1. The second minimum AIC occurs at the alternative submodel 4 again under both models, telling us that if we were to cluster any one of the two groups, then we should cluster Borderline (B) and Not Admit (NA) groups together as one homogeneous group, and we should cluster Admit (A) group completely separate as one heterogeneous group. On the other hand, if we want to make a third choice, then the third minimum of AIC occurs at the alternative submodel 3, indicating to us the closeness of the Admit (A) group to the Borderline (B) group as one homogeneous cluster, and leaving Not Admit (NA) group on its own as a singleton cluster. Therefore, this way, we can check the significance of each of the the clustering

alternatives in the decision making process. In this example, we also never cluster the three groups as one homogeneous group (submodel 5).

In comparing the AIC's in Tables 4.4 and 4.5. for this example we also notice that, AIC (varying μ and Σ) values are less than the AIC (common Σ) values for each of the clustering alternatives except for the last clustering alternative (i.e., alternative 5) in clustering the applicant groups. These results suggest that we should use different covariance matrices. However, the values of AIC (varying μ and Σ) and AIC (common Σ) are significantly closer to one another that if we were to assume equal covariance matrices between the applicant groups a priori, it would not have been a dubious assumption for this particular data set.

Thus, it should be noted that via AIC we can now easily check the validity of our assumptions in terms of using equal covariance matrices as opposed to separate covariance matrices in a particular data set which is important in the multi-sample clustering situation, and in general.

5. Conclusions and Discussion

From our numerical results in Section 4, we see that AIC and consequently minimum AIC procedures can indeed successfully identify the best clustering alternatives when we cluster samples into homogeneous sets of samples both in the MANOVA model and the multivariate model with varying covariance matrices. We can measure the amount of homogeneity and heterogeneity in clustering samples. We can determine a priori whether we should use equal or varying covariance matrices in the analysis of a data set.

The fact that AIC does not require the table look-up, which is the case in conventional procedures, adds to the importance of the results obtained. This

is one of the important virtues that AIC breaks away from conventional procedures which try to test whether a parameter is "significant" or not using a significance level α which is essentially arbitrary. The other important virtue of AIC is that the penalty represented by the term $2 \times$ (number of free parameters) clearly demonstrates the necessity of choosing a class of models, at least one of which will be able to provide a good approximation to the distribution of data without adjusting too many parameters.

Thus, in concluding, we see that the use of AIC shows how to combine the information in the likelihood with an appropriate function of the number of parameters to obtain estimates of the information provided by competing alternative models. Therefore, the definition of MAICE gives a clear formulation of the principle of parsimony in statistical model building or comparison as we demonstrated by numerical examples. And MAICE provides a versatile procedure for statistical model identification which is free from the ambiguities inherent in the application of conventional statistical procedures.

Acknowledgements

This paper is a summary of some of the results in the Ph.D. thesis of the first author in the Department of Mathematics at the University of Illinois at Chicago, under the supervision of the second author. The first author is grateful to Professor Stanley L. Sclove for his valuable and useful comments. The suggestions of the Associate Editor, Dr. Noboru Ohsumi, and the referees were very helpful and are gratefully acknowledged. Appreciation is also extended to Mr. Ronald Ball for his programming assistance. This research was supported by Office of Naval Research Contract N00014-80-C-0408, Task NR042-443, and Army Research Office Contract DAAG29-82-K-0155, at the University of Illinois at Chicago.

REFERENCES

- [1] Ambramowitz, M., and Stegun, I.A. (1968), Handbook of Mathematical Functions with Formulas, Graphs and Mathematical Tables (Nat. Bur. of Stand. Appl. Math. Ser., No. 55), 7th printing. U.S. Govt. Printing Office, Washington, D.C.
- [2] Akaike, H. (1973), "Information Theory and an Extension of the Maximum Likelihood Principle," 2nd International Symposium on Information Theory, eds. B.N. Petrov and F. Csaki, Budapest: Akademiai Kiado, 267-281.
- [3] _____ (1974), "A New Look at the Statistical Model Identification," IEEE Transactions on Automatic Control, AC-19, 716-723.
- [4] Consul, P. C. (1969), "The Exact Distributions of Likelihood Criteria for Different Hypotheses," in P. R. Krishnaiah (Ed.), Multivariate Analysis-II, New York: Academic Press.
- [5] Fisher, R. A. (1936), "The Use of Multiple Measurements in Taxonomic Problems," Annals of Eugenics, 7, 179-188.
- [6] Friedman, H. P., and Rubin, J. (1967), "On Some Invariant Criteria for Grouping Data," Journal of the American Statistical Association, 62, 1159-1178.
- [7] Gabriel, K. R. (1969), "A Comparison of Some Methods of Simultaneous Inference in MANOVA," in P. R. Krishnaiah (Ed.), Multivariate Analysis-II, New York: Academic Press.
- [8] Kendall, M. G. (1966), "Discrimination and Classification," in P.R. Krishnaiah (Ed.), Multivariate Analysis, New York: Academic Press.
- [9] Johnson, R. A. and Wichern, D. W. (1982), Applied Multivariate Statistical-Analysis, Englewood Cliffs: Prentice-Hall, Inc.
- [10] Krishnaiah, P. R. (1969), "Simultaneous Test Procedures and General MANOVA Models," in P. R. Krishnaiah (Ed.), Multivariate Analysis-II, New York: Academic Press.
- [11] _____ (1979), "Some Developments on Simultaneous Test Procedures," in P. R. Krishnaiah (Ed.), Developments in Statistics, Vol. 2, New York: Academic Press.
- [12] Mardia, K. V., Kent, J. T., and Bibby, J. M. (1979). Multivariate Analysis, New York: Academic Press.
- [13] Mezzich, J.E., and Solomon, H. (1980), Taxonomy and Behavioral Science, New York: Academic Press.

- [14] Muirhead, R. J. (1982), Aspects of Multivariate Statistical Theory, New York: John-Wiley.
- [15] Solomon, H. (1971), Numerical Taxonomy, Mathematics in the Archaeological and Historical Sciences, Edinburgh: Edinburgh University Press, 62-81.
- [16] Srivastava, J. N. (1969), "Some Studies on Intersection Tests in Multivariate Analysis of Variance," in P. R. Krishnaiah (Ed.), Multivariate Analysis-II, New York: Academic Press.
- [17] Wilks, S.S. (1932), "Certain Generalizations in the Analysis of Variance," Biometrika, 24, 471-494.

Unclassified

SECURITY CLASSIFICATION OF THIS PAGE (When Data Entered)

REPORT DOCUMENTATION PAGE		READ INSTRUCTIONS BEFORE COMPLETING FORM
1. REPORT NUMBER Technical Report 82-6	2. GOVT ACCESSION NO. AD-A223758	3. RECIPIENT'S CATALOG NUMBER
4. TITLE (and Subtitle) Multi-Sample Cluster Analysis Using Akaike's Information Criterion		5. TYPE OF REPORT & PERIOD COVERED Technical Report
		6. PERFORMING ORG. REPORT NUMBER
7. AUTHOR(s) Hamparsum Bozdogan and Stanley I. Sclove		8. CONTRACT OR GRANT NUMBER(s) N00014-80-C-0408
9. PERFORMING ORGANIZATION NAME AND ADDRESS University of Illinois at Chicago Box 4348, Chicago, IL 60680		10. PROGRAM ELEMENT, PROJECT, TASK AREA & WORK UNIT NUMBERS
11. CONTROLLING OFFICE NAME AND ADDRESS —		12. REPORT DATE December 20, 1982
		13. NUMBER OF PAGES 27+ ii
14. MONITORING AGENCY NAME & ADDRESS (if different from Controlling Office) Office of Naval Research Statistics and Probability Arlington, VA 22217		15. SECURITY CLASS. (of this report) Unclassified
		15a. DECLASSIFICATION/DOWNGRADING SCHEDULE
16. DISTRIBUTION STATEMENT (of this Report) APPROVED FOR PUBLIC RELEASE: DISTRIBUTION UNLIMITED.		
17. DISTRIBUTION STATEMENT (of the abstract entered in Block 20, if different from Report) Unlimited distribution		
18. SUPPLEMENTARY NOTES		
19. KEY WORDS (Continue on reverse side if necessary and identify by block number) Multi-sample cluster analysis; W-square Criterion; Akaike's Information Criterion (AIC); MANOVA model; multivariate model with varying mean vectors and variance-covariance matrices; maximum likelihood.		
20. ABSTRACT (Continue on reverse side if necessary and identify by block number) Multi-sample cluster analysis, the problem of grouping samples, is studied from an information-theoretic viewpoint via Akaike's Information Criterion (AIC). This criterion combines the maximum value of the likelihood with the number of parameters used in achieving that value. The multi-sample cluster problem is defined, and AIC is developed for this problem. The form		

DD FORM 1 JAN 73 1473

EDITION OF 1 NOV 55 IS OBSOLETE
S/N 3102-UF-314-6601

Unclassified

SECURITY CLASSIFICATION OF THIS PAGE (When Data Entered)

Unclassified

SECURITY CLASSIFICATION OF THIS PAGE (When Data Entered)

(Abstract continued)

of AIC is derived in both the multivariate analysis of variance (MANOVA) model and in the multivariate model with varying mean vectors and variance-covariance matrices. Numerical examples are presented for AIC and another criterion called w-square. The results demonstrate the utility of AIC in identifying the best clustering alternatives.

Unclassified

SECURITY CLASSIFICATION OF THIS PAGE (When Data Entered)