



Using the growth curve model in classification of repeated measurements

Dietrich von Rosen^{1,2} · Martin Singull²

Received: 8 March 2023 / Revised: 5 September 2023 / Accepted: 25 January 2024 /

Published online: 29 March 2024

© The Institute of Statistical Mathematics, Tokyo 2024

Abstract

In this paper, discrimination between two populations following the growth curve model is considered. A likelihood-based classification procedure is established, in the sense that we compare the two likelihoods given that the new observation belongs to respective population. The possibility to classify the new observation as belonging to an unknown population is discussed, which is shown to be natural when considering growth curves. Several examples and simulations are given to emphasize this possibility.

Keywords Discriminant analysis · Growth curve model · Likelihood-based classification

1 Introduction

Pearson (1915) discussed the problem of sexing bones, where the analysis was based on the assumption that the population consisted of two overlapping populations. The Coefficient of Racial Likeness (CRL) was presented by Pearson and colleagues [e.g., see Pearson (1926)]. Based on measurements of bones, scalps, etc., the idea was to discriminate between races and an indicator was created, i.e., the CRL. Some years after the CRL first appeared, Mahalanobis (1936) introduced the idea of a distance measure between two populations. In an interesting contribution to Mahalanobis' work Dasgupta (1993) noted that Mahalanobis had been inspired by Seal (1911), who discussed group distances and variables as coordinates on a plane. At this point, Hotelling (1931) must also be mentioned. Hotelling developed a test for group differences based on correlated variables, and was inspired by Pearson and the CRL.

✉ Martin Singull
martin.singull@liu.se

¹ Department of Energy and Technology, Swedish University of Agricultural Science, Uppsala, Sweden

² Department of Mathematics, Linköping University, 581 83 Linköping, Sweden

The father of discriminant analysis Sir Ronald Fisher (1936), formulated and solved the mathematical problem of separating groups via a linear function of correlated variables. He also showed how to assign an observation into one of the groups. Fisher (1938) repeated many of the statements from the earlier article but also connected his work to Mahalanobis' and Hotelling's contributions. Moreover, Fisher (1938) discussed the direction of the observations from each group. In fact Barnard (1935), who communicated with Fisher, wanted to discriminate between groups when a linear regression appeared over time for the observed variables, which partly can be interpreted as a choice between directions. Barnard's approach can be said to be a univariate approach. Our research considers similar problems but it will take place within a likelihood ratio approach based on dependent variables. Wald (1944) had a clear derivation of a discrimination rule and discussed in detail distributional problems related to classification. In a very interesting paper, Rao (1948) clearly indicated that one purpose of discriminant analyses is to assign individuals to different populations with the possibility to assign the individual into an unknown population. Our article will show how this phenomenon in a specific case can be handled. Rao also connected discrimination to the Neyman–Pearson theory. An informative survey of discriminant analysis up to 1950 is presented by Hodges (1950), and 25 years later by Huberty (1975). The above preamble of discriminant analysis gives the background for this article. In particular it can be noted that the following had been discussed: (i) the separation of groups using linear regression (Barnard 1935; ii) the classification into an unknown population (Rao 1948). Rao (1973, Sect. 8e.2), there is a brief discussion of how to approach the problem with a third unknown population.

Multivariate analysis of variance (MANOVA) became included in text books in the 1950s and one-way MANOVA is of course linked to discrimination among groups. Moreover, the generalized multivariate analysis of variance (GMANOVA) started to be studied, which ended up in a general model by Potthoff and Roy (1964), i.e., the growth curve model was born. The model will constitute a basis in the proposed classification analysis. Estimation of the parameters of the growth curve model was considered by Khatri (1966), who derived the maximum likelihood estimators. However, there is a prehistory of the model where growth curves were estimated, e.g., see Rao (1958, 1959). Furthermore, in Rao (1958) the classification of growth curves was briefly mentioned. For references connected to the growth curve model see von Rosen (2018).

It is interesting to note that Burnaby (1966) discussed the classification of growth curves based on geometrical arguments including directions. However, Burnaby looked upon growth as a nuisance parameter when discussing classification, i.e., Burnaby wanted to exclude any growth effect when classifying observations. Rao had an influence on Burnaby's article and Rao (1966) provided several alternatives for how one can construct discriminant functions involving linear growth curves. In both Burnaby's and Rao's work, the dispersion is supposed to be known, or it can be estimated from "other" data sources. To some extent our contribution in this article is a direct follow-up of Burnaby's results, but now the likelihood is used, which indeed can be used to estimate the dispersion jointly with the construction of classification rule.

Kshirsagar and Davis (1983a, 1983b) [see also Albert and Kshirsagar (1993)] considered several populations and the construction of discriminant function for the populations via canonical correlations. The works are a continuation of ideas which originated in the 1930s [see Hotelling (1936), Bartlett (1938), Bartlett (1939), Bartlett (1951) and Kshirsagar (1962)]. Note that Kshirsagar and Davis (1983a, 1983b) and others do not consider the classification of objects. The aim was to find a discriminant function via classical testing of hypothesis. In this work, no testing of hypothesis is carried out. Instead the likelihood is studied where the observation which is to be classified is part of the likelihood. The criteria used in this article are expressed in the original matrices, which gives a clear interpretation of the results. Mentz and Kshirsagar (2004) studied classification via the growth curve model but their approach is not likelihood-based and differs from the one presented in the subsequent sections. Recently, Ngailo et al. (2021) considered Fisher's linear classification function for repeated measurements using the mean structure connected to the growth curve model, and derived misclassification errors. There is also a Bayesian literature connected to the growth curve model and Lee (1982) used the posterior distribution to classify linear growth curves. More references on Bayesian classification of growth curves can be found in Lee (1982).

For completeness, the growth curve model (Potthoff and Roy 1964) is presented:

$$X = ABC + E, \quad E \sim N_{p,n}(\mathbf{0}, \Sigma, I), \quad (1)$$

where $X : p \times n$, $A : p \times q$, $q < p$, $B : q \times k$, $C : k \times n$, $k < n$ and $E \sim N_{p,n}(\mathbf{0}, \Sigma, I)$. The unknown parameters are B and Σ . Moreover, C is the between-individuals design matrix which is the same design matrix as in linear models or MANOVA models, and A is the within individuals design matrix. The design matrix C will in our applications always be of full rank and later it will be indicated that without loss of generality it can be assumed that the rank of A equals q , i.e., A is of full rank.

Over the years, several specialized text books on classification have been written, e.g., McLachlan (2004), where a huge number of references on the subject can be found. Moreover, most books on multivariate analysis include classification and discriminant analysis (e.g., see Srivastava and Khatri 1979; Muirhead 1982; Anderson 2003). However, none of these books have information about classification of linear growth curves (McLachlan 2004, has a short Sect. 3.7.6).

This article considers the discrimination between two populations based on linear growth curves analysis, i.e., we have $k = 2$. Let the populations be denoted by Π_1 and Π_2 and then growth curve models are given for each population by

$$\Pi_1 : X_1 = A\beta_1 \mathbf{1}'_{n_1} + E, \quad E \sim N_{p,n_1}(\mathbf{0}, \Sigma, I_{n_1}), \quad (2)$$

$$\Pi_2 : X_2 = A\beta_2 \mathbf{1}'_{n_2} + E, \quad E \sim N_{p,n_2}(\mathbf{0}, \Sigma, I_{n_2}), \quad (3)$$

where β_1 , β_2 and Σ are unknown parameters, $\mathbf{1}_{n_i}$ is a vector of n_i ones, $'$ stands for the transpose, $N_{p,n}(\mu, \Sigma, \Psi)$ denotes the matrix normal distribution, where in the present work always $\Psi = I$ meaning that there are independent columns in E , and $A : p \times q$, $q < p$, is the within individuals design matrix which describes the growth structure.

The dispersion matrix Σ is supposed to be positive definite. The population Π_1 consists of n_1 observations and Π_2 consists of n_2 observations. The purpose with the analysis is to decide if a new observation $\mathbf{x}$ can be assigned to Π_1 , Π_2 or to no one of these populations. A likelihood-based approach will be promoted which has not been taken place before when considering classification of linear growth curves.

2 Likelihood-based classification

Let X_i , $i \in \{1, 2\}$, be the random matrices corresponding to the observations from Π_i , $i \in \{1, 2\}$ and $\mathbf{x}$ is the random vector which corresponds to the observation which is to be classified. If $\mathbf{x} \in \Pi_1$, consider the following two models:

$$\begin{aligned}(X_1 : \mathbf{x}) &= A\beta_1 \mathbf{1}'_{n_1+1} + E, \quad E \sim N_{p, n_1+1}(\mathbf{0}, \Sigma, I_{n_1+1}), \\ X_2 &= A\beta_2 \mathbf{1}'_{n_2} + E, \quad E \sim N_{p, n_2}(\mathbf{0}, \Sigma, I_{n_2}),\end{aligned}$$

where X_1 , X_2 and $\mathbf{x}$ are jointly independently distributed. These two equations form the likelihood function when $\mathbf{x} \in \Pi_1$. However, if $\mathbf{x} \in \Pi_2$ the model is specified as

$$\begin{aligned}(X_2 : \mathbf{x}) &= A\beta_2 \mathbf{1}'_{n_2+1} + E, \quad E \sim N_{p, n_2+1}(\mathbf{0}, \Sigma, I_{n_2+1}), \\ X_1 &= A\beta_1 \mathbf{1}'_{n_1} + E, \quad E \sim N_{p, n_1}(\mathbf{0}, \Sigma, I_{n_1}),\end{aligned}$$

which then will be used to create the likelihood function. Let

$$\begin{aligned}X &= (X_1 : \mathbf{x} : X_2), \quad C_1 = \begin{pmatrix} \mathbf{1}'_{n_1+1} & \mathbf{0}'_{n_2} \\ \mathbf{0}'_{n_1+1} & \mathbf{1}'_{n_2} \end{pmatrix}, \quad C_2 = \begin{pmatrix} \mathbf{1}'_{n_1} & \mathbf{0}'_{n_2+1} \\ \mathbf{0}'_{n_1} & \mathbf{1}'_{n_2+1} \end{pmatrix}, \\ S_i &= X(I - C'_i(C_i C'_i)^{-1} C_i)X', \quad i \in \{1, 2\},\end{aligned} \quad (4)$$

where $\mathbf{0}_{n_i}$ denotes a vector of n_i zeroes. If $\mathbf{x} \in \Pi_i$, $i \in \{1, 2\}$, the corresponding likelihood can be written

$$L_i(\mathbf{B}, \Sigma) = (2\pi)^{-\frac{1}{2}(n_1+n_2+1)p} |\Sigma|^{-\frac{1}{2}(n_1+n_2+1)} e^{-\frac{1}{2}\text{tr}\{\Sigma^{-1}(X-ABC_i)(X-ABC_i)'\}},$$

where $\mathbf{B} = (\beta_1 : \beta_2)$, $q \times 2$. If $\mathbf{x} \in \Pi_i$, $i \in \{1, 2\}$, when taking into account all observations, the maximum likelihood estimators of ABC_i and Σ equal, without assuming any full rank conditions on A , (e.g., see von Rosen 2018, Theorem 3.1),

$$\begin{aligned}\widehat{AB}_i C_i &= A(A'S_i^{-1}A)^{-1}A'S_i^{-1}XC'_i(C_i C'_i)^{-1}C_i, \\ (n+1)\widehat{\Sigma}_i &= (X - \widehat{AB}_i C_i)(X - \widehat{AB}_i C_i)' \\ &= S_i + (I - A(A'S_i^{-1}A)^{-1}A'S_i^{-1})XC'_i(C_i C'_i)^{-1}C_iX' \\ &\quad \times (I - S_i^{-1}A(A'S_i^{-1}A)^{-1}A'),\end{aligned}$$

where “ $^{-}$ ” denotes an arbitrary generalized inverse and $n = n_1 + n_2$. The estimators $\widehat{\mathbf{B}}_i$ and $\widehat{\Sigma}_i$ denote estimators of $\mathbf{B}$ and Σ , respectively, when $\mathbf{x} \in \Pi_i$, $i \in \{1, 2\}$.

Given an inner product defined via S let $P_{A,S} = A(A'S^{-1}A)^{-1}A'S^{-1}$ be the projection on the column vector space of A . Note the well known fact that the projector $P_{A,S}$ as well as the above-given maximum likelihood estimators do not depend on the choice of generalized inverses. Moreover, $P_{A,S}$ will be the same for all matrices generating the column space of A . Thus, in principle, it can be assumed that A is of full rank. The likelihood ratio statistic for classifying x is equivalent to

$$\lambda = \frac{|\hat{\Sigma}_1|}{|\hat{\Sigma}_2|} = \frac{|S_1 + (I - P_{A,S_1})XC'_1(C_1C'_1)^{-1}C_1X'(I - P'_{A,S_1})|}{|S_2 + (I - P_{A,S_2})XC'_2(C_2C'_2)^{-1}C_2X'(I - P'_{A,S_2})|}. \quad (5)$$

Some manipulations of (5) yield

$$\begin{aligned} \lambda &= \frac{|S_1|}{|S_2|} \frac{|I + (S_1^{-1} - S_1^{-1}P_{A,S_1})XC'_1(C_1C'_1)^{-1}C_1X'(I - P'_{A,S_1})|}{|I + (S_2^{-1} - S_2^{-1}P_{A,S_2})XC'_2(C_2C'_2)^{-1}C_2X'(I - P'_{A,S_2})|} \\ &= \frac{|S_1|}{|S_2|} \frac{|I + (S_1^{-1} - S_1^{-1}P_{A,S_1})XC'_1(C_1C'_1)^{-1}C_1X'|}{|I + (S_2^{-1} - S_2^{-1}P_{A,S_2})XC'_2(C_2C'_2)^{-1}C_2X'|}, \end{aligned} \quad (6)$$

since $P'_{A,S_i}(S_i^{-1} - S_i^{-1}P_{A,S_i}) = 0$, $i \in \{1, 2\}$, and for any matrices U and V of proper sizes $|I + UV| = |I + VU|$.

Now $S_i^{-1} - S_i^{-1}P_{A,S_i} = A^o(A^oS_iA^o)^{-1}A^o$, $i \in \{1, 2\}$, where A^o is any matrix of full rank spanning the orthogonal complement to the column space of A . Moreover, $S_i + XC'_i(C_iC'_i)^{-1}C_iX' = XX'$, $i \in \{1, 2\}$. Thus the λ -ratio in (6) can be written

$$\lambda = \frac{|S_1| |(A^oS_1A^o)^{-1}| |A^oXX'A^o|}{|S_2| |(A^oS_2A^o)^{-1}| |A^oXX'A^o|} = \frac{|S_1| |(A^oS_1A^o)^{-1}|}{|S_2| |(A^oS_2A^o)^{-1}|}. \quad (7)$$

In order to classify an observation x into Π_1 , because $|\Sigma|^{-1/2}$ appears in the normal density, the ratio λ should be small, i.e., $\lambda < 1$.

The next step will be to study the ratio $|S_1|/|S_2|$. Put the pooled sum of squares matrix based on the joint samples from Π_1 and Π_2 , without taking into account that the new observation x belongs to one of these two populations, as

$$S = X_1X'_1 - n_1\bar{x}_1\bar{x}'_1 + X_2X'_2 - n_2\bar{x}_2\bar{x}'_2, \quad (8)$$

where the averages are given by $\bar{x}_i = n_i^{-1}X_i\mathbf{1}_{n_i}$ for $i \in \{1, 2\}$. The matrix S is Wishart distributed, with $n_1 + n_2 - 2$ degrees of freedom, which is denoted $S \sim W_p(\Sigma, n_1 + n_2 - 2)$.

The following lemma will be applied several times throughout the work.

Lemma 1 Let S_1 , S_2 and S be defined in (4) and (8). Then

$$S_i = S + \frac{n_i}{n_i + 1}(\bar{x}_i - x)(\bar{x}_i - x)', \quad i \in \{1, 2\}.$$

Calculations to establish the lemma are the same as when establishing a likelihood-based classification rule in MANOVA Anderson (2003, Sect. 6.5.5). The given relation in Lemma 1 implies that

$$\frac{|S_1|}{|S_2|} = \frac{|S + \frac{n_1}{n_1+1}(\bar{x}_1 - x)(\bar{x}_1 - x)'|}{|S + \frac{n_2}{n_2+1}(\bar{x}_2 - x)(\bar{x}_2 - x)'|} = \frac{1 + \frac{n_1}{n_1+1}(\bar{x}_1 - x)'S^{-1}(\bar{x}_1 - x)}{1 + \frac{n_2}{n_2+1}(\bar{x}_2 - x)'S^{-1}(\bar{x}_2 - x)}. \quad (9)$$

According to (7) the ratio $|A^{o'}S_2A^o|/|A^{o'}S_1A^o|$ should also be considered. Lemma 1 implies

$$A^{o'}S_iA^o = A^{o'}SA^o + \frac{n_i}{n_i+1}A^{o'}(\bar{x}_i - x)(\bar{x}_i - x)'A^o, \quad i \in \{1, 2\}$$

and $A^o(A^{o'}SA^o)^{-1}A^{o'} = S^{-1} - S^{-1}A(A'S^{-1}A)^{-1}A'S^{-1}$. Thus

$$\begin{aligned} |A^{o'}S_iA^o| &= |A^{o'}SA^o| \left(1 + \frac{n_i}{n_i+1}(\bar{x}_i - x)'S^{-1}(\bar{x}_i - x) \right. \\ &\quad \left. - \frac{n_i}{n_i+1}(\bar{x}_i - x)'S^{-1}A(A'S^{-1}A)^{-1}A'S^{-1}(\bar{x}_i - x) \right), \quad i \in \{1, 2\}. \end{aligned} \quad (10)$$

One should note that the factor $|A^{o'}SA^o|$ cancels when taking the ratio $|A^{o'}S_2A^o|/|A^{o'}S_1A^o|$. Returning to the likelihood ratio in (7), the next theorem follows from (7)–(10).

Theorem 1 *Let the “likelihood ratio” λ be given by (7). Then*

$$\lambda = \frac{H_2}{H_1},$$

where

$$H_i = 1 - \frac{\frac{n_i}{n_i+1}(\bar{x}_i - x)'S^{-1}A(A'S^{-1}A)^{-1}A'S^{-1}(\bar{x}_i - x)}{1 + \frac{n_i}{n_i+1}(\bar{x}_i - x)'S^{-1}(\bar{x}_i - x)}, \quad i \in \{1, 2\}.$$

The classification rule states that x will be classified into Π_1 if $\lambda < 1$ and thus it is of interest to look closer at the inequality $H_2 < H_1$, which is equivalent to

$$\frac{L_1}{L_2} < \frac{1 + \frac{n_1}{n_1+1}(\bar{x}_1 - x)'S^{-1}(\bar{x}_1 - x)}{1 + \frac{n_2}{n_2+1}(\bar{x}_2 - x)'S^{-1}(\bar{x}_2 - x)}, \quad (11)$$

where

$$L_i = \frac{n_i}{n_i+1}(\bar{x}_i - x)'S^{-1}A(A'S^{-1}A)^{-1}A'S^{-1}(\bar{x}_i - x), \quad i \in \{1, 2\}. \quad (12)$$

The relation in (11) can be shown by noting that $H_2 < H_1$ is equivalent to

$$\begin{aligned} & \frac{\frac{n_2}{n_2+1}(\bar{\mathbf{x}}_2 - \mathbf{x})' \mathbf{S}^{-1} \mathbf{A} (\mathbf{A}' \mathbf{S}^{-1} \mathbf{A})^{-1} \mathbf{A}' \mathbf{S}^{-1} (\bar{\mathbf{x}}_2 - \mathbf{x})}{1 + \frac{n_2}{n_2+1}(\bar{\mathbf{x}}_2 - \mathbf{x})' \mathbf{S}^{-1} (\bar{\mathbf{x}}_2 - \mathbf{x})} \\ & > \frac{\frac{n_1}{n_1+1}(\bar{\mathbf{x}}_1 - \mathbf{x})' \mathbf{S}^{-1} \mathbf{A} (\mathbf{A}' \mathbf{S}^{-1} \mathbf{A})^{-1} \mathbf{A}' \mathbf{S}^{-1} (\bar{\mathbf{x}}_2 - \mathbf{x})}{1 + \frac{n_1}{n_1+1}(\bar{\mathbf{x}}_1 - \mathbf{x})' \mathbf{S}^{-1} (\bar{\mathbf{x}}_1 - \mathbf{x})}. \end{aligned}$$

From von Rosen (2018, Theorem 3.1) it follows that the maximum likelihood estimators of β_i , $i \in \{1, 2\}$, based on the samples from Π_i , satisfy

$$\hat{\mathbf{A}}\hat{\beta}_i = \mathbf{A}(\mathbf{A}'\mathbf{S}^{-1}\mathbf{A})^{-1}\mathbf{A}'\mathbf{S}^{-1}\bar{\mathbf{x}}_i, \quad i \in \{1, 2\}.$$

If it is of interest to estimate $\hat{\beta}_i$ then $\mathbf{A}$ has to be of full rank. The ratio L_1/L_2 can be written

$$\frac{L_1}{L_2} = \frac{n_1(n_2+1)}{n_2(n_1+1)} \frac{(\hat{\mathbf{A}}\hat{\beta}_1 - \mathbf{P}_{\mathbf{A},\mathbf{S}}\mathbf{x})' \mathbf{S}^{-1} (\hat{\mathbf{A}}\hat{\beta}_1 - \mathbf{P}_{\mathbf{A},\mathbf{S}}\mathbf{x})}{(\hat{\mathbf{A}}\hat{\beta}_2 - \mathbf{P}_{\mathbf{A},\mathbf{S}}\mathbf{x})' \mathbf{S}^{-1} (\hat{\mathbf{A}}\hat{\beta}_2 - \mathbf{P}_{\mathbf{A},\mathbf{S}}\mathbf{x})}. \quad (13)$$

Now the relations in (11) and (13) will be utilized to classify an observation where the mean follows the growth curve model.

Proposition 1 *All matrices used below have been introduced in Sect. 2. The n_1 observations in Population 1 (Π_1) follow the model in (2), whereas the n_2 observations from Population 2 (Π_2) follow the model in (3). Let $\mathbf{x}$ be a new observation which according to the following criteria is to be classified into Π_1 or Π_2 or it is classified to belong to an unknown population.*

The unknown observation $\mathbf{x}$ is classified into Π_1 if the following two conditions are satisfied:

- (i) $\frac{n_1}{n_1+1}(\bar{\mathbf{x}}_1 - \mathbf{x})' \mathbf{S}^{-1} (\bar{\mathbf{x}}_1 - \mathbf{x}) < \frac{n_2}{n_2+1}(\bar{\mathbf{x}}_2 - \mathbf{x})' \mathbf{S}^{-1} (\bar{\mathbf{x}}_2 - \mathbf{x})$
- (ii) $(\hat{\mathbf{A}}\hat{\beta}_1 - \mathbf{P}_{\mathbf{A},\mathbf{S}}\mathbf{x})' \mathbf{S}^{-1} (\hat{\mathbf{A}}\hat{\beta}_1 - \mathbf{P}_{\mathbf{A},\mathbf{S}}\mathbf{x}) < c_0(\hat{\mathbf{A}}\hat{\beta}_2 - \mathbf{P}_{\mathbf{A},\mathbf{S}}\mathbf{x})' \mathbf{S}^{-1} (\hat{\mathbf{A}}\hat{\beta}_2 - \mathbf{P}_{\mathbf{A},\mathbf{S}}\mathbf{x}),$

where

$$c_0 = \frac{n_2(n_1+1)}{n_1(n_2+1)} \frac{1 + \frac{n_1}{n_1+1}(\bar{\mathbf{x}}_1 - \mathbf{x})' \mathbf{S}^{-1} (\bar{\mathbf{x}}_1 - \mathbf{x})}{1 + \frac{n_2}{n_2+1}(\bar{\mathbf{x}}_2 - \mathbf{x})' \mathbf{S}^{-1} (\bar{\mathbf{x}}_2 - \mathbf{x})}.$$

Remark 1 If only one inequality, either (i) or (ii), is fulfilled the new observation $\mathbf{x}$ is classified as belonging to some unknown population. Furthermore, if none of the inequalities are fulfilled the new observation $\mathbf{x}$ is classified into Π_2 .

A plain likelihood ratio criterion would only have been based on (ii), which is natural because there is no contribution to the likelihood from observations from an unknown population. However, through the proposed method and the likelihood, more information about the observations can be extracted. Unfortunately, deriving misclassification errors will be difficult, but one can achieve results via the derivation of misclassification errors in both (i) and (ii) and the Bonferroni inequality.

In a way, the proposal is natural because the growth curve model is a bilinear model and in the data there are two types of observed variations (deviation from the model). First, it is the variation of the observations around the sample means [treated via (i)] and second, it is the variation of the sample means around the estimated model [treated via (ii)]. This has led to the proposed method presented in Proposition 1. In Fig. 1, several cases are illustrated where the observation which is to be classified does not belong to any of the two populations [see (b)–(d)].

Already Rao (1948) noted the importance of classifying into an unknown population but this case has received little attention over the years. If in Proposition 1 both (i) and (ii) suggest classifying the observation in, say, Population 1, this may still not be correct. Therefore, Proposition 1 has to be complemented in order to accurately classify the cases in Fig. 1 (of course still with some kind of probabilistic misclassification error). The reason for this is that $\hat{\beta}_i$ is of size q and some components in $\hat{\beta}_i$ can suggest classifying x into population Π_1 , whereas other components suggest classifying x into population Π_2 . In Fig. 1b this phenomenon is illustrated where the

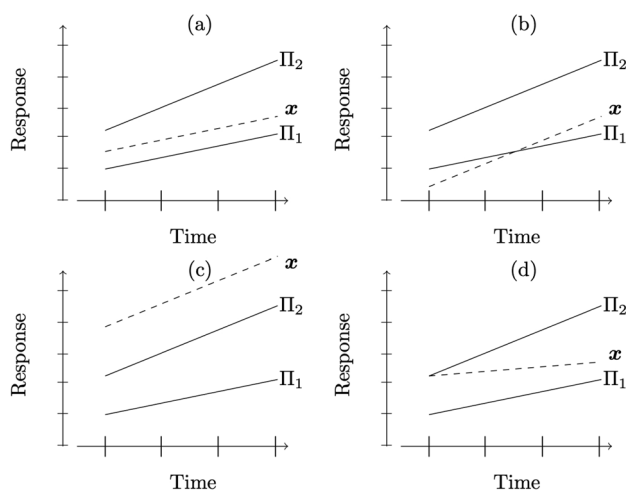


Fig. 1 In two populations, a response has been measured at four time points. In the present figures, Π_1 represents Population 1, Π_2 represents Population 2 and x (dashed line) stands for the new observation vector which is to be classified. In **a** it is shown that x belongs to Π_1 , whereas in **b** it is illustrated when x is close to the observations in Π_1 but the growth follows the one of Π_2 . Moreover, it is shown in **c** when x follows the growth of Π_2 but deviates from this population. In **d** it is indicated that x neither follows the growth of Π_1 nor the growth of Π_2 . Thus only in **a** the new observation can be assigned to a population

intercept suggests classifying $\mathbf{x}$ into one population but the slope suggests classifying into the other population. Note that when $\hat{\beta}_i$ is discussed it is assumed that $\mathbf{A}$ is of full rank, otherwise estimable functions have to be considered.

3 Comparison with other classification approaches

Over the years, many classification methods have been proposed. In this section four different methods, involving different classification functions, will be presented and compared with the approach in Proposition 1. These methods are all connected to the linear discriminant function. Following the derivation in Anderson (2003), with two normal populations and the same dispersion matrix, $N_p(\mu_1, \Sigma)$, $N_p(\mu_2, \Sigma)$, say Π_1 and Π_2 and an unknown observation $\mathbf{x}$, the following classification function appears:

$$d = \mathbf{x}' \Sigma^{-1} (\mu_1 - \mu_2) - \frac{1}{2} (\mu_1 + \mu_2)' \Sigma^{-1} (\mu_1 - \mu_2). \quad (14)$$

If $d > 0$, $\mathbf{x}$ is classified to belong to Π_1 and otherwise $\mathbf{x}$ is classified into Π_2 . Since μ_1 , μ_2 and Σ are unknown these parameters have to be estimated via samples from Π_1 and Π_2 . For example, if $X_1 \in \Pi_1$, $X_2 \in \Pi_2$,

$$\begin{aligned} X_1 &= \mu_1 \mathbf{1}'_{n_1} + E, & E &\sim N_{p, n_1}(\mathbf{0}, \Sigma, I_{n_1}), \\ X_2 &= \mu_2 \mathbf{1}'_{n_2} + E, & E &\sim N_{p, n_2}(\mathbf{0}, \Sigma, I_{n_2}), \end{aligned}$$

then maximum likelihood estimators are given by $\hat{\mu}_1 = \bar{x}_1$, $\hat{\mu}_2 = \bar{x}_2$ and $(n = n_1 + n_2)$

$$\begin{aligned} n \hat{\Sigma} &= S = (X_1 : X_2)(I_n - C'(CC')^{-1}C)(X_1 : X_2)', \\ C &= \begin{pmatrix} \mathbf{1}'_{n_1} & \mathbf{0}'_{n_2} \\ \mathbf{0}'_{n_1} & \mathbf{1}'_{n_2} \end{pmatrix}, \end{aligned}$$

where S is the same sums of squares matrix as in (8). Thus, $\mathbf{x}$ is classified into Π_1 if $D > 0$, where

$$n^{-1}D = \mathbf{x}' S^{-1} (\bar{x}_1 - \bar{x}_2) - \frac{1}{2} (\bar{x}_1 + \bar{x}_2)' S^{-1} (\bar{x}_1 - \bar{x}_2). \quad (15)$$

Burnaby (1966) considered the following samples:

$$\begin{aligned} X_1 &= A\beta_1 \mathbf{1}'_{n_1} + E, & E &\sim N_{p, n_1+1}(\mathbf{0}, \Sigma, I_{n_1}), \\ X_2 &= A\beta_2 \mathbf{1}'_{n_2} + E, & E &\sim N_{p, n_2}(\mathbf{0}, \Sigma, I_{n_2}), \end{aligned}$$

where A is the known within individuals design matrix modeling as before the responses' growth. Burnaby looked upon the growth as a nuisance factor, and to remove this factor, he introduced an orthogonal transformation (projection) perpendicular to the growth direction, i.e., $A'X_i$, which also implies that its dispersion equals $A' \Sigma A$. Although Burnaby's idea was to estimate Σ via some other data set, it is known today that $n^{-1}S$ (this is not the maximum likelihood estimator) can be used in the following classification function (Burnaby did not explicitly stated this result):

$$n^{-1}D_B = x'A^o(A'^oSA^o)^{-1}A'^o(\bar{x}_1 - \bar{x}_2) - \frac{1}{2}(\bar{x}_1 + \bar{x}_2)'A^o(A'^oSA^o)^{-1}A'^o(\bar{x}_1 - \bar{x}_2), \quad (16)$$

where x is classified into Π_1 if $D_B > 0$.

Mentz and Kshirsagar (2004) also considered (14) where the following estimators were inserted (A was supposed to be of full rank which is not necessary):

$$\hat{B} = (\hat{\beta}_1 : \hat{\beta}_2) = (A'A)^{-1}A(\bar{x}_1 : \bar{x}_2), \\ n\hat{\Sigma}_{MK} = ((X_1 : X_2) - A\hat{B}C)((X_1 : X_2) - A\hat{B}C)'$$

The classification in (14) with these estimators equals

$$D_{MK} = x'A(A'\hat{\Sigma}_{MK}A)^{-1}A'(\bar{x}_1 - \bar{x}_2) - \frac{1}{2}(\bar{x}_1 + \bar{x}_2)'A(A'\hat{\Sigma}_{MK}A)^{-1}A'(\bar{x}_1 - \bar{x}_2) \quad (17)$$

and if $D_{MK} > 0$ the observation x is classified into Π_1 . A modification of this result is to use maximum likelihood estimators

$$\hat{B}_{ML} = (\hat{\beta}_1 : \hat{\beta}_2) = (A'S^{-1}A)^{-1}A'S^{-1}(\bar{x}_1 : \bar{x}_2), \\ n\hat{\Sigma}_{ML} = ((X_1 : X_2) - A\hat{B}_{ML}C)((X_1 : X_2) - A\hat{B}_{ML}C)'$$

Since $A'\hat{\Sigma}_{ML}^{-1} = nA'S^{-1}$ the classification function equals

$$n^{-1}D_{ML} = x'S^{-1}A(A'S^{-1}A)^{-1}A'S^{-1}(\bar{x}_1 - \bar{x}_2) - \frac{1}{2}(\bar{x}_1 + \bar{x}_2)'S^{-1}A(A'S^{-1}A)^{-1}A'S^{-1}(\bar{x}_1 - \bar{x}_2). \quad (18)$$

If $D_{ML} > 0$ the observation x is classified into Π_1 . It can be noted that $D = D_B + D_{ML}$, which is in agreement with intuition because with D_{ML} growth curves are studied, whereas Burnaby wanted to remove the growth effects when classifying observations.

In the next example a number of scenarios are discussed.

Example 1 In this example, two populations are considered and the aim is to illustrate the proposed method in Proposition 1, as well as to compare the results with

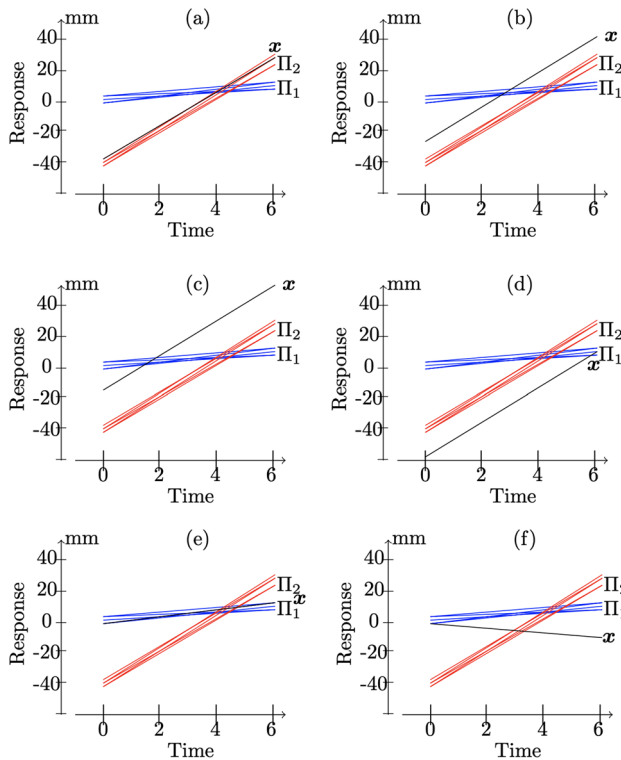


Fig. 2 In two populations a response deviating from baseline has been measured at four time points. In the present figures Π_1 (blue lines) represents Population 1, Π_2 (red lines) represents Population 2 and x (black line) stands for the new observation vector which is to be classified. In **a** it is shown that x belongs to Π_2 , whereas in **b** it is illustrated when the response for x is shifted a bit from the observations in Π_2 and has come closer to the observations in Π_1 , in comparison with **a**. In **c** the response for x has been shifted away from Π_2 more than in **b** whereas in **d** the response x has been shifted away from both the observations in Π_2 and Π_1 . In **e** x belongs to Π_1 and in **f** x does not belong to any of the two given populations (colour figure online)

the earlier proposed methods presented above. Simulated data are used and the ideas are illustrated in Fig. 2. The following matrices have been used to generate data:

$$A = (a_1, a_2) = \begin{pmatrix} 1 & 1 & 1 & 1 \\ 0 & 2 & 4 & 6 \end{pmatrix}', \quad B = (\beta_1, \beta_2) = \begin{pmatrix} 1 & -40 \\ 1 & 10 \end{pmatrix},$$

$$\Sigma = \begin{pmatrix} 10 & 5 & 2.50 & 1.25 \\ & 10 & 5 & 2.50 \\ & & 10 & 5 \\ & & & 10 \end{pmatrix}.$$

Using the matrices, 10 observations were generated from both populations which is in fact a small data set. From Population 1

$$X_1 = \begin{pmatrix} 1.180192 & 3.330091 & 2.284808 & 0.5316934 & -7.324091 & 4.287237 \\ 0.963614 & 3.172805 & 1.613572 & 4.5971772 & -2.381949 & 2.014789 \\ 4.265160 & -1.586984 & 2.916197 & 1.8198309 & 4.239860 & 4.136507 \\ 1.148597 & 2.255889 & 11.616286 & 4.9593502 & 2.364634 & 10.201689 \\ -5.4962367 & 2.2856576 & 3.613856 & 4.134668 & & \\ 0.4055098 & 0.2124452 & 7.358171 & 4.918653 & & \\ 6.2437675 & 1.6049673 & 7.070861 & 5.996032 & & \\ 10.8431665 & 6.2032100 & 7.062818 & 2.647788 & & \end{pmatrix}$$

and from Population 2

$$X_2 = \begin{pmatrix} -39.045120 & -40.929413 & -42.728002 & -31.026729 & -38.19440 & -40.21661 \\ -20.429371 & -20.086069 & -24.712534 & -18.288272 & -23.41076 & -21.42750 \\ -2.145189 & 0.276048 & -4.176433 & 6.028609 & -0.90917 & -1.54939 \\ 12.039531 & 20.310317 & 15.409145 & 25.366647 & 21.78737 & 22.63965 \\ -43.846127 & -41.193042 & -44.820454 & -40.031872 & & \\ -25.696265 & -19.383592 & -15.916768 & -19.625143 & & \\ -4.278519 & 5.790712 & 0.299045 & 6.308488 & & \\ 19.957922 & 24.849704 & 20.049941 & 20.884688 & & \end{pmatrix}$$

Using the above data, the following estimators were obtained (note that the number of observations is small in this study)

$$\hat{\beta}_1 = \begin{pmatrix} 0.69 \\ 0.83 \end{pmatrix}, \quad \hat{\beta}_2 = \begin{pmatrix} -40.67 \\ 10.13 \end{pmatrix}, \quad \hat{\Sigma} = \begin{pmatrix} 14.1 & 4.4 & 2.7 & 2.7 \\ & 7.7 & 4.0 & 1.5 \\ & & 10.9 & 5.9 \\ & & & 14.2 \end{pmatrix}.$$

Moreover it is noted that

$$S = \begin{pmatrix} 279.9 & 90.6 & 52.3 & 53.4 \\ & 151.3 & 83.6 & 30.9 \\ & & 210.6 & 117.2 \\ & & & 283.6 \end{pmatrix}.$$

In Fig. 2, the data in X_1 and X_2 are illustrated but for visibility reasons the real values are not presented (the variance indicated in the figures is too small). Corresponding to the six scenarios in Fig. 2, the observation which is to be classified equals

$$\text{Scenario 1, Fig. 2a : } \mathbf{x} = (-40, -20, 0, 20)'; \quad (19)$$

$$\text{Scenario 2, Fig. 2b : } \mathbf{x} = (-30, -10, 10, 30)'; \quad (20)$$

$$\text{Scenario 3, Fig. 2c : } \mathbf{x} = (-20, 0, 20, 40)'; \quad (21)$$

Table 1 The classification functions D , D_B , D_{MK} , D_{ML} and C_1/C_2 , respectively, defined in (15)–(18), (25) and (26) are evaluated for each of the six classification scenarios presented in (19)–(24)

Scenario	D	D_B	D_{MK}	D_{ML}	C_1	C_2
Scenario 1, Fig. 2a	−0.21	0.0002	−85.3	−0.21	−7.78	−8.55
Scenario 2, Fig. 2b	−0.15	0.0002	−59.2	−0.15	−5.29	−3.05
Scenario 3, Fig. 2c	0.0014	0.0002	−2.4	0.001	0.05	−0.001
Scenario 4, Fig. 2d	−0.28	0.0002	−111.4	−0.28	−10.27	−6.25
Scenario 5, Fig. 2e	0.22	0.0002	87.2	0.22	8.01	0.98
Scenario 6, Fig. 2f	0.26	−0.00005	106.1	0.26	9.56	0.91

Note that negative values for D , D_B , D_{MK} , D_{ML} support classification into Population 2 (Π_2) whereas a positive value for C_1 and C_2 also supports a classification into Population 2

$$\text{Scenario 4, Fig. 2d : } \mathbf{x} = (-50, -30, -10, 10)'; \quad (22)$$

$$\text{Scenario 5, Fig. 2e : } \mathbf{x} = (1, 3, 5, 7)'; \quad (23)$$

$$\text{Scenario 6, Fig. 2f : } \mathbf{x} = (1, -4, -7, -10)'. \quad (24)$$

The classification criteria, which are used when classifying $\mathbf{x}$ will now be listed. From Proposition 1 two criteria, C_1 and C_2 , are obtained as

$$C_1 = \frac{n_2}{n_2 + 1} (\bar{\mathbf{x}}_2 - \mathbf{x})' \mathbf{S}^{-1} (\bar{\mathbf{x}}_2 - \mathbf{x}) - \frac{n_1}{n_1 + 1} (\bar{\mathbf{x}}_1 - \mathbf{x})' \mathbf{S}^{-1} (\bar{\mathbf{x}}_1 - \mathbf{x}), \quad (25)$$

$$C_2 = c_0 (\mathbf{A} \hat{\boldsymbol{\beta}}_2 - \mathbf{P}_{A,S} \mathbf{x})' \mathbf{S}^{-1} (\mathbf{A} \hat{\boldsymbol{\beta}}_2 - \mathbf{P}_{A,S} \mathbf{x}) - (\mathbf{A} \hat{\boldsymbol{\beta}}_1 - \mathbf{P}_{A,S} \mathbf{x})' \mathbf{S}^{-1} (\mathbf{A} \hat{\boldsymbol{\beta}}_1 - \mathbf{P}_{A,S} \mathbf{x}), \quad (26)$$

where

$$c_0 = \frac{n_2(n_1 + 1)}{n_1(n_2 + 1)} \frac{1 + \frac{n_1}{n_1 + 1} (\bar{\mathbf{x}}_1 - \mathbf{x})' \mathbf{S}^{-1} (\bar{\mathbf{x}}_1 - \mathbf{x})}{1 + \frac{n_2}{n_2 + 1} (\bar{\mathbf{x}}_2 - \mathbf{x})' \mathbf{S}^{-1} (\bar{\mathbf{x}}_2 - \mathbf{x})}.$$

If both C_1 and C_2 are positive, $\mathbf{x}$ is classified to Π_1 and if both C_1 and C_2 are negative $\mathbf{x}$ is classified to Π_2 . Furthermore, if one is positive and the other negative, then the new observation $\mathbf{x}$ is classified to an unknown population. The six scenarios have been analyzed and in Table 1, and the results for D , D_B , D_{MK} and D_{ML} are reported together with C_1 and C_2 .

It can once again be remarked that D_B differs from the other criteria, which is natural because D_B treats the growth effect as a nuisance parameter. Moreover, D_B is constant for Scenario 1–5 because $\mathbf{A}'\mathbf{x} = \mathbf{0}$. Except for Scenario 3, the same

conclusions about assigning $\mathbf{x}$ appear for D , D_{MK} , D_{ML} and C_1/C_2 . In Scenario 3, C_1/C_2 classifies $\mathbf{x}$ into an unknown population, D and D_{ML} into Population 1, whereas D_{MK} assigns $\mathbf{x}$ into Population 2.

We continue with another example, which will be a small simulation study to compare and understand the performance for the different classifiers presented above.

Example 2 To compare the performance of the different classification methods, D , D_B , D_{MK} , D_{ML} and C_1/C_2 , we find the outcomes from each method, using the same simulated data set for all of them. Let $q = 2$, $p = 5$ with $n_1 = n_2 = 25$, i.e., linear growth, five repeated measurements, at time points $t_i = i$, $i \in \{1, \dots, 5\}$, for two groups, each consisting of 25 independent observations. The considered mean parameter matrix equals

$$\mathbf{B} = \begin{pmatrix} 0 & -2 \\ 0.75 & 1.5 \end{pmatrix},$$

and the covariance matrix $\Sigma = \mathbf{Q}'\mathbf{Q} + \sigma^2\mathbf{I}_p$, where $\mathbf{Q} = (q_{ij})_{i,j \in \{1, \dots, p\}}$ with q_{ij} uniformly distributed in $[-2, 2]$ and $\sigma^2 = 1$.

For 1000 replications we note: (i) the decision for a new observation from Population 1, and (ii) the decision for new observations from three different “unknown” populations, which are characterizing the mean models

Case 1: between the two populations, i.e., the mean of $\mathbf{A}\beta_1$ and $\mathbf{A}\beta_2$;

Case 2: “far away” from the populations, i.e., far from $\mathbf{A}\beta_1$ and $\mathbf{A}\beta_2$;

Case 3: not in the linear subspace given by the columns of $\mathbf{A}$, e.g., a quadratic mean model.

At first, we consider the case when we simulate new observations from Population 1. In Fig. 3 (left) we can see one example of such a simulation and in Table 2 the proportions of simulations, when the new observation is classified to Population 1, Population 2 or as unknown. By comparing the proportions, one can conclude that all the rules perform similar. However, C_1/C_2 , i.e., the proposed method in Proposition 1, has classified 6.0% of the simulations as unknown, which in this case is a

Table 2 The performance of the classification functions D , D_{MK} , D_{ML} and C_1/C_2 , respectively, defined in (15), (17), (18), (25) and (26)

Population	D (%)	D_{ML} (%)	D_{MK} (%)	C_1/C_2 (%)
1	74.8	76.0	71.6	72.5
2	25.2	24.0	28.4	21.5
Unknown	–	–	–	6.0

The proportions given in the table, are classification proportions for new observations from Population 1, i.e., if new observations are classified into Population 1, Population 2 or classified as belonging to an unknown population

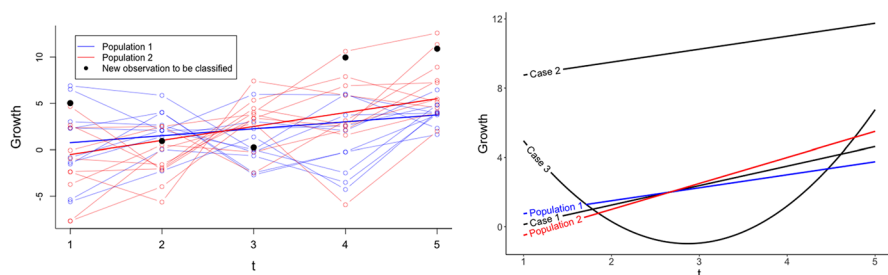


Fig. 3 *Left*: in Example 2 we simulated 25 profiles from each population. In this figure ten selected repeated measurement from each population is plotted together with the true growths (complete lines) and an observation which is to be classified. *Right*: the three different models, for which a new observation should be classified as unknown

wrong decision. The misclassification error for C_1/C_2 is the sum of the proportions for classifying into Population 2 and the proportion for classifying as unknown.

Secondly, we simulate new observations from the three different cases, discussed above and which can be seen in Fig. 3 (right). For those cases we would like our classifier to identify the new observations as belonging to an unknown population. In Table 3 we can see that this happens, but it seems not to be good enough and there is more information to gain from the data as we will see in Sect. 4. Furthermore, one can note that in Case 1, the classifications are roughly 50/50, which seems to be reasonable since the mean model for the new observation is just between the two populations. Similar result for Case 2, since the mean model is “far away” from the populations. In both cases the proposed rule with C_1/C_2 gives a proportion of

Table 3 The performance of the classification functions D , D_{MK} , D_{ML} and C_1/C_2 , respectively, defined in (15), (17), (18), (25) and (26)

Population	D (%)	D_{MK} (%)	D_{ML} (%)	C_1/C_2 (%)
<i>Case 1</i>				
1	51.4	51.1	50.8	45.4
2	48.6	48.9	49.2	42.8
Unknown	—	—	—	11.8
<i>Case 2</i>				
1	41.2	41.5	37.7	35.2
2	58.8	58.5	62.3	52.4
Unknown	—	—	—	12.4
<i>Case 3</i>				
1	93.5	94.3	70.1	91.6
2	6.5	5.7	29.9	3.9
Unknown	—	—	—	4.5

The proportions given in the table, are classification proportions for new observations from the three cases seen in Fig. 3 and described in Example 2, i.e., if the observations are classified into Population 1, Population 2 or as coming from an unknown population

simulations classified to an unknown population, roughly 12%, but lower (4.5%) when we consider the quadratic mean model.

4 More criteria for classifying into an unknown population

Already in the end of Sect. 2 it was noted, since the size of the mean parameters β_i , $i \in \{1, 2\}$, is q that many kind of differences can appear. In this section we propose some new ideas how to control for different effects in β_i , $i \in \{1, 2\}$.

It is informative to calculate the expectation of the quadratic forms in (13) (given S). The results presented in the next lemma will be used in what follows. The proof is based on knowledge about moments for quadratic forms in normally distributed variables and the first moment for the inverse Wishart distribution. Note, for $i \in \{1, 2\}$,

$$y_i = \sqrt{\frac{n_i}{n_i + 1}}(\bar{x}_i - \mathbf{x}) \sim N_p\left(\sqrt{\frac{n_i}{n_i + 1}}(\mathbf{A}\beta_i - E[\mathbf{x}]), \Sigma\right), \quad (27)$$

and

$$S \sim W_p(\Sigma, n - 2), \quad (28)$$

which are independent. Moreover, since $E[\mathbf{x}]$ equals either $\mathbf{A}\beta_i$ for $i \in \{1, 2\}$, $E[\mathbf{x}]$ belongs to the column space of $\mathbf{A}$.

Lemma 2 For y_i , $\mathbf{x}$ and S in (27) and (28), then

(i) for $\mathbf{Q}$: $p \times p$, and $\mathbf{A}$: $p \times q$, $q < p$,

$$E[y_i' \mathbf{Q} y_i] = \text{tr}\{\mathbf{Q}\Sigma\} + \frac{n_i}{n_i + 1}(\mathbf{A}\beta_i - E[\mathbf{x}])' \mathbf{Q}(\mathbf{A}\beta_i - E[\mathbf{x}]);$$

(ii) for $\mathbf{A}$: $p \times q$ of rank q and $n > p - 3$

$$E[S^{-1} \mathbf{A}(\mathbf{A}' S^{-1} \mathbf{A})^{-1} \mathbf{A}' S^{-1}] = \frac{1}{n - p + q - 3} \left(\frac{q}{n - p - 3} \Sigma^{-1} + \Sigma^{-1} \mathbf{A}(\mathbf{A}' \Sigma^{-1} \mathbf{A})^{-1} \mathbf{A}' \Sigma^{-1} \right).$$

Using a vector space decomposition, the next lemma can be proved (it is a special case of Theorem 1.2.25 in Kollo and von Rosen (2005)).

In order to handle the situation with many parameters, the following lemma will be utilized.

Lemma 3 Let $\mathbf{A} = (\mathbf{a}_1, \dots, \mathbf{a}_q)$ be of rank q and $\mathbf{A}_{(k)}$ equals $\mathbf{A}$ without the column $\mathbf{a}_k$. Let $\mathbf{T}_k = \mathbf{S}^{-1} \mathbf{A} \mathbf{H}_k^o (\mathbf{H}_k^{o'} \mathbf{A}' \mathbf{S}^{-1} \mathbf{A} \mathbf{H}_k^o)^{-1} \mathbf{H}_k^{o'} \mathbf{A}' \mathbf{S}^{-1}$, where $\mathbf{H}_k = \mathbf{A}' \mathbf{S}^{-1} \mathbf{A}_{(k)}$. Then, for $k \in \{1, \dots, q\}$,

$$T_k = S^{-1}A(A'S^{-1}A)^{-1}A'S^{-1} - S^{-1}A_{(k)}(A'_{(k)}S^{-1}A_{(k)})^{-1}A'_{(k)}S^{-1}.$$

Later it will be used that $A'_{(k)}T_k = \mathbf{0}$. The expectation $E[T_k]$ follows from Lemma 2 (ii).

Suppose the size of β_i , $i \in \{1, 2\}$, equals q , i.e., there are q different parameters which should be considered. The next proposed approach is not based on a factorization of the likelihood. Instead, if Proposition 1 suggests classifying into, say, Π_1 a strategy will be proposed which can support the outcome of the proposition.

In (12) the matrix $S^{-1}A(A'S^{-1}A)^{-1}A'S^{-1}$ appears and according to Lemma 3 it can be split into two terms of which one equals ($k \in \{1, \dots, q\}$)

$$T_k = S^{-1}AH_k^o(H_k^{o'}A'S^{-1}AH_k^o)^{-1}H_k^{o'}A'S^{-1},$$

where $H_k = A'S^{-1}A_{(k)}$ and $A_{(k)}$ is the matrix A without the k th element, as in Lemma 3. This can be done in q different ways, i.e., one split for each parameter. Note that T_k works like a filter and selects a_k when premultiplying by $A' = (a_1, \dots, a_q)'$. From (12) it follows that $\frac{n_i}{n_i+1}(\bar{x}_i - x)'T_k(\bar{x}_i - x)$ is of interest, i.e., it is natural to calculate

$$\frac{n_1}{n_1+1}(\bar{x}_1 - x)'T_k(\bar{x}_1 - x) - \frac{n_2}{n_2+1}(\bar{x}_2 - x)'T_k(\bar{x}_2 - x).$$

Therefore, the expectation of $\frac{n_i}{n_i+1}(\bar{x}_i - x)'T_k(\bar{x}_i - x)$, given S , is calculated. Given S the matrix T_k is constant and using Lemma 2(i), as well as independence between $\bar{x}_1 - x$ and S ,

$$\begin{aligned} E\left[\frac{n_i}{n_i+1}(\bar{x}_i - x)'T_k(\bar{x}_i - x)|S\right] &= \text{tr}\{T_k\Sigma\} \\ &+ \frac{n_i}{n_i+1}(A\beta_i - E[x])'T_k(A\beta_i - E[x]). \end{aligned} \quad (29)$$

Depending on which population $\bar{x}_i$ and x belong to the expression in (29) is of the following form:

- if $\bar{x}_i$ and x belong to the same population

$$E\left[\frac{n_i}{n_i+1}(\bar{x}_i - x)'T_k(\bar{x}_i - x)|S\right] = \text{tr}\{T_k\Sigma\}, \quad i \in \{1, 2\};$$

- if $\bar{x}_i \in \Pi_1$ and $x \in \Pi_2$

$$\begin{aligned} E\left[\frac{n_1}{n_1+1}(\bar{x}_1 - x)'T_k(\bar{x}_1 - x)|S\right] \\ = \text{tr}\{T_k\Sigma\} + \frac{n_1}{n_1+1}(\beta_1 - \beta_2)'A'T_kA(\beta_1 - \beta_2); \end{aligned}$$

- if $\bar{x}_i \in \Pi_2$ and $x \in \Pi_1$

$$E \left[\frac{n_2}{n_2 + 1} (\bar{\mathbf{x}}_2 - \mathbf{x})' \mathbf{T}_k (\bar{\mathbf{x}}_2 - \mathbf{x}) | \mathbf{S} \right] \\ = \text{tr} \{ \mathbf{T}_k \boldsymbol{\Sigma} \} + \frac{n_2}{n_2 + 1} (\boldsymbol{\beta}_1 - \boldsymbol{\beta}_2)' \mathbf{A}' \mathbf{T}_k \mathbf{A} (\boldsymbol{\beta}_1 - \boldsymbol{\beta}_2).$$

Based on Lemmas 2 and 3, it is possible to find the expectation of $\mathbf{T}_k$ in the above formulas but we gain nothing by doing so. Instead it is more fruitful to see that

$$(\boldsymbol{\beta}_1 - \boldsymbol{\beta}_2)' \mathbf{A}' \mathbf{T}_k \mathbf{A} (\boldsymbol{\beta}_1 - \boldsymbol{\beta}_2) = (\beta_{2k} - \beta_{1k})^2 \mathbf{a}'_k \mathbf{T}_k \mathbf{a}_k,$$

where β_{1k} and β_{2k} are the k th element in $\boldsymbol{\beta}_1$ and $\boldsymbol{\beta}_2$, respectively.

Based on the above moment calculations, it makes sense to set up, additionally to the conditions in Propositions 1, q new conditions to evaluate how an observation $\mathbf{x}$ should be classified.

Proposition 2 *Let the samples and the observation which is to be classified be as in Proposition 1.*

- (i) Suppose that the conditions in Proposition 1(i) and (ii) are satisfied. If for $k \in \{1, \dots, q\}$

$$\frac{n_1}{n_1 + 1} (\bar{\mathbf{x}}_1 - \mathbf{x})' \mathbf{T}_k (\bar{\mathbf{x}}_1 - \mathbf{x}) < \frac{n_2}{n_2 + 1} (\bar{\mathbf{x}}_2 - \mathbf{x})' \mathbf{T}_k (\bar{\mathbf{x}}_2 - \mathbf{x}),$$

then $\mathbf{x}$ is classified as belonging to Π_1 .

- (ii) Suppose that the conditions in Proposition 1(i) and (ii) are not satisfied. If for $k \in \{1, \dots, q\}$

$$\frac{n_1}{n_1 + 1} (\bar{\mathbf{x}}_1 - \mathbf{x})' \mathbf{T}_k (\bar{\mathbf{x}}_1 - \mathbf{x}) > \frac{n_2}{n_2 + 1} (\bar{\mathbf{x}}_2 - \mathbf{x})' \mathbf{T}_k (\bar{\mathbf{x}}_2 - \mathbf{x}),$$

then $\mathbf{x}$ is classified as belonging to Π_2 .

- (iii) If one of the above conditions in (i) or (ii) is not satisfied then $\mathbf{x}$ is classified as belonging to some unknown population.

Proposition 2 is based on $\mathbf{T}_k$ and from Lemma 3 it follows that also

$$\mathbf{S}^{-1} \mathbf{A}_{(k)} (\mathbf{A}'_{(k)} \mathbf{S}^{-1} \mathbf{A}_{(k)})^{-1} \mathbf{A}'_{(k)} \mathbf{S}^{-1} \quad (30)$$

can be of interest. This matrix can be split into a sum of rank one matrices which will lead to the possibility to investigate a number of new criteria.

In this article, however, only the case $q = 2$ is considered. If $q = 2$ and $k = 1$ the expression in (30) equals

$$S^{-1}a_2(a_2'S^{-1}a_2)^{-1}a_2'S^{-1}. \quad (31)$$

The expectation of the quadratic form, given S , equals

$$\frac{n_i}{n_i + 1}(\bar{x}_i - x)'S^{-1}a_2(a_2'S^{-1}a_2)^{-1}a_2'S^{-1}(\bar{x}_i - x), \quad (32)$$

which follows from (29) by replacing T_k with the expression in (31). Thus, if $\beta_i = (\beta_{1i}, \beta_{2i})'$, $i \in \{1, 2\}$, using the independence between $(\bar{x}_i - x)$ and S ,

$$\begin{aligned} E \left[\frac{n_i}{n_i + 1}(\bar{x}_i - x)'S^{-1}a_2(a_2'S^{-1}a_2)^{-1}a_2'S^{-1}(\bar{x}_i - x) | S \right] \\ = \text{tr}\{S^{-1}a_2(a_2'S^{-1}a_2)^{-1}a_2'S^{-1}\Sigma\} \\ + \frac{n_i}{n_i + 1}(A\beta_i - E[x])'S^{-1}a_2(a_2'S^{-1}a_2)^{-1}a_2'S^{-1}(A\beta_i - E[x]). \end{aligned} \quad (33)$$

Then, if $x \in \Pi_j$, for $j \in \{1, 2\}$, the last term in (33) equals

$$\begin{aligned} \frac{n_i}{n_i + 1}(\beta_i - \beta_j)'A'S^{-1}a_2(a_2'S^{-1}a_2)^{-1}a_2'S^{-1}A(\beta_i - \beta_j) \\ = \left(\beta_{1i} - \beta_{1j} + (\beta_{2i} - \beta_{2j}) \frac{a_2'S^{-1}a_2}{a_1'S^{-1}a_2} \right)^2 \frac{(a_1'S^{-1}a_2)^2}{a_2'S^{-1}a_2}. \end{aligned}$$

An unconditional expectation could have been obtained by applying Lemma 2 (ii) but this is not essential for the presentation.

As above, we use criteria C_1 and C_2 given in (25) and (26), respectively, to classify the new observation x . Moreover, if the criteria indicate that x belongs to either Π_1 or Π_2 it is suggested to continue with evaluating a number of other criteria, denoted C_3 to C_6 . From Proposition 2(i)

$$C_3 = \frac{n_2}{n_2 + 1}(\bar{x}_2 - x)'T_1(\bar{x}_2 - x) - \frac{n_1}{n_1 + 1}(\bar{x}_1 - x)'T_1(\bar{x}_1 - x), \quad (34)$$

where $T_1 = S^{-1}AH_1^o(H_1^{o'}A'S^{-1}AH_1^o)^{-1}H_1^{o'}A'S^{-1}$, and $H_1 = A'S^{-1}a_2$. If $C_3 > 0$, x is classified to come from Π_1 , otherwise Π_2 . Moreover, also from Proposition 2(i),

$$C_5 = \frac{n_2}{n_2 + 1}(\bar{x}_2 - x)'T_2(\bar{x}_2 - x) - \frac{n_1}{n_1 + 1}(\bar{x}_1 - x)'T_2(\bar{x}_1 - x), \quad (35)$$

where $T_2 = S^{-1}AH_2^o(H_2^{o'}A'S^{-1}AH_2^o)^{-1}H_2^{o'}A'S^{-1}$, and $H_2 = A'S^{-1}a_1$. If $C_5 > 0$, x is classified as coming from Π_1 , otherwise Π_2 . Finally there are the criteria which are “complementary” to C_3 and C_5 , see (32):

$$\begin{aligned} C_4 = \frac{n_2}{n_2 + 1}(\bar{x}_2 - x)'S^{-1}a_2(a_2'S^{-1}a_2)^{-1}a_2'S^{-1}(\bar{x}_2 - x) \\ - \frac{n_1}{n_1 + 1}(\bar{x}_1 - x)'S^{-1}a_2(a_2'S^{-1}a_2)^{-1}a_2'S^{-1}(\bar{x}_1 - x), \end{aligned} \quad (36)$$

$$C_6 = \frac{n_2}{n_2 + 1}(\bar{x}_2 - x)'S^{-1}a_1(a_1'S^{-1}a_1)^{-1}a_1'S^{-1}(\bar{x}_2 - x) \\ - \frac{n_1}{n_1 + 1}(\bar{x}_1 - x)'S^{-1}a_1(a_1'S^{-1}a_1)^{-1}a_1'S^{-1}(\bar{x}_1 - x). \quad (37)$$

If $C_4 > 0$ then x is classified as belonging to Π_1 and the same rule applies to C_6 . Thus, a proposed decision rule equals: if C_1, C_2 and C_3 to C_6 have the same sign, i.e., if all criteria are positive, classify x into Π_1 , if all are negative classify x into Π_2 , and otherwise classify x into an unknown population.

Example 3 Consider again the two populations and the six scenarios in Example 1. According to Table 4, it follows that according to the classification philosophy in this article, i.e., applying all six criteria ($C_1, \dots, C_6$), only x in Scenarios 1 and 5, will be classified into one of the two given populations, i.e., in Scenario 1 x will be assigned to Π_2 and in Scenario 5, to Π_1 . Moreover, it is only in Scenario 3 where C_1 and C_2 do not agree about the classification. In Scenarios 2, 4 and 6, it is C_3 to C_6 which hint that x should be assigned to an unknown population.

Example 4 With the same simulations as in Example 2 we can study the classifications using all six criteria ($C_1, \dots, C_6$). Hence, criteria $C_3, \dots, C_6$ will further help us to classify to an unknown population. Table 5 is the same as Table 2, but we have added the last column with the classification proportions using all six criteria ($C_1, \dots, C_6$). When using ($C_1, \dots, C_6$) one can note that a much smaller proportion of simulations are misclassified to Population 2. The proportion of simulations classified to an unknown population are larger than only using C_1 and C_2 .

Table 6 is the same as Table 3, but we have again added the last column with the classification rule using all six criteria ($C_1, \dots, C_6$). In all three cases the use of all six criteria are very promising, since the proportion of classifying the simulations as belonging to an unknown population has increased significantly. For Case 2, it is as high as 99.1.

Table 4 The classification functions C_1 to C_6 , respectively, defined in (25)–(37), are evaluated for each of the six classification scenarios presented in (19)–(24)

Scenario	C_1	C_2	C_3	C_4	C_5	C_6	Classify to
Scenario 1, Fig. 2a	−7.78	−8.55	−7.13	−0.66	−6.08	−1.71	Π_2
Scenario 2, Fig. 2b	−5.29	−3.05	−3.57	−1.74	−6.08	0.78	Unknown
Scenario 3, Fig. 2c	0.05	−0.001	3.20	−3.16	−4.74	4.78	Unknown
Scenario 4, Fig. 2d	−10.27	−6.25	−10.69	0.42	−6.08	−4.19	Unknown
Scenario 5, Fig. 2e	8.01	0.98	7.48	0.53	6.02	1.98	Π_1
Scenario 6, Fig. 2f	9.56	0.91	7.20	2.36	9.81	−0.25	Unknown

A negative value means that the criteria supports a classification in Population 2 (Π_2)

Table 5 The performance of the classification functions D , D_{MK} , D_{ML} and $(C_1, \dots, C_6)$, defined in (15), (17), (18) and (25)–(37), is presented for the data in Example 2

Population	D (%)	D_{MK} (%)	D_{ML} (%)	C_1/C_2 (%)	$(C_1, \dots, C_6)$ (%)
1	74.8	76.0	71.6	72.5	27.2
2	25.2	24.0	28.4	21.5	5.5
Unknown	–	–	–	6.0	67.3

The proportions given in the table, are classification proportions for new observations from Population 1, i.e., if it is classified into Population 1, Population 2 or to an unknown population

5 Comparisons of classification approaches using a data set

In this section we will consider a real data set comprising repeated measurements and compare the performance of the five different classification procedures which were discussed in the previous section. In the supplementary material, one can find studies of five more real data sets, with similar results and conclusions.

We will do it by a leave-one-out cross-validation procedure, i.e., removing one vector of repeated measurements at time, which then should be classified. The data set consists of data from two populations with n_1 and n_2 independent observations. Hence, we will classify $n = n_1 + n_2$ observations and find the proportions of correctly, misclassified and unknown classified measurements, respectively. A quadratic growth will be assumed.

The data set consists of five repeated measurements, at start (baseline) and after 6, 12, 20 and 24 months, of the serum cholesterol (Srivastava 2002, pp. 379–381). The patients constitute of two groups, the first group ($n_1 = 31$) received an active drug and

Table 6 The performance of the classification functions D , D_{MK} , D_{ML} , and $(C_1, \dots, C_6)$, defined in (15), (17), (18) and (25)–(37), is presented for the data in Example 2

Population	D (%)	D_{MK} (%)	D_{ML} (%)	C_1/C_2 (%)	$(C_1, \dots, C_6)$ (%)
<i>Case 1</i>					
1	51.4	51.1	50.8	45.4	17.9
2	48.6	48.9	49.2	42.8	13.7
Unknown	–	–	–	11.8	68.4
<i>Case 2</i>					
1	41.2	41.5	37.7	35.2	0.9
2	58.8	58.5	62.3	52.4	0.0
Unknown	–	–	–	12.4	99.1
<i>Case 3</i>					
1	93.5	94.3	70.1	91.6	36.7
2	6.5	5.7	29.9	3.9	0.7
Unknown	–	–	–	4.5	62.6

The proportions given in the table, are classification proportions for new observations from the three cases seen in Fig. 3, i.e., if it is classified into Population 1, Population 2 or to an unknown population

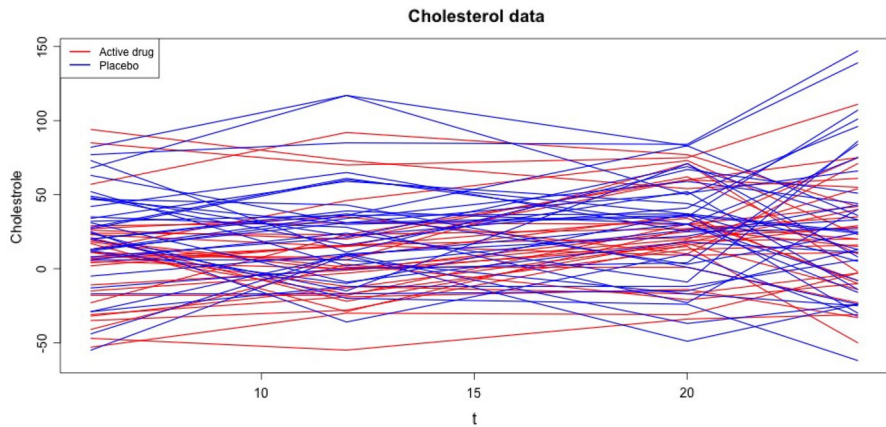


Fig. 4 Plot of cholesterol data. Four repeated measurements at 6, 12, 20 and 24 months for the two groups, an active drug intake ($n_1 = 31$) and placebo ($n_2 = 36$) are presented

Table 7 The performance of the classification functions D , D_{MK} , D_{ML} , $(C_1, \dots, C_6)$, respectively, defined in (15), (17), (18), and (25)–(37), i.e., the proportions of the patients based on the cholesterol measurements that are correctly, misclassified or classified to come from an unknown population, respectively

Population	D (%)	D_{MK} (%)	D_{ML} (%)	C_1/C_2 (%)	$(C_1, \dots, C_6)$ (%)
Correctly	64	63	60	52	19
Misclassified	36	37	40	30	7
Unknown	–	–	–	18	73

the second group ($n_2 = 36$) received placebo. To include the baseline effect, a new data set was created by individually subtracting baseline values from the original observations. This data set is exposed in Fig. 4, and in Table 7 we can find the proportions of correctly, misclassified and unknown unclassified patients, respectively.

As one can see in Table 7, the performance of the classification functions D , D_{MK} and D_{ML} are fairly similar and the number of correctly classified observations are as low as 60%. The data set is not designed for classification and the profiles for each group are fairly similar, which can be depicted in Fig. 4, hence the poor performance. Furthermore, when considering the two classifiers proposed in this paper, respectively, based on C_1/C_2 and $(C_1, \dots, C_6)$, several of the measurements are classified as unknown, which seems to be reasonable due to similar growths.

6 Concluding remarks

When classifying growth curves in multivariate linear models following a normal distribution, it is of interest and possible to include different aspects in the analysis. Hence, the mean structure for the specified populations can be on different levels, but they can

also follow different profiles, e.g., slopes. In this paper, we have developed a likelihood-based classification procedure for linear growth curves, i.e., a new observation of p repeated measurements is assigned to one of two specified populations. A key feature of the proposed classification procedure is that it can suggest that the new observation does not belong to any of the two populations, i.e., it should be classified to an unknown population. This can be the case, for example, when the new observation has the same level as one of the population, but different profiles or vice versa.

Given a new classification procedure, it is of interest to find out how well it will classify a new observation. However, in general, it can be difficult to find the exact expression of the misclassification probabilities. Even in the simplest case for Fisher's linear classification function with estimated parameters, the misclassification probabilities can only be found approximately. In this paper we have not discussed the misclassification probabilities for the proposed likelihood-based procedure, but the ideas can serve for future research on this problem.

We conclude that the proposed classifiers based on C_1/C_2 and $(C_1, \dots, C_6)$, have significant smaller number of misclassified observations compared to the other methods. Although, they deliver a smaller number of correctly classified objects. Relatively to the other classification procedures the approach of this paper is more "robust".

Supplementary Information The online version contains supplementary material available at <https://doi.org/10.1007/s10463-024-00900-1>.

Acknowledgements The authors would like to thank the anonymous reviewers for valuable and helpful suggestions which have significantly improved the paper.

References

- Albert, J. M., Kshirsagar, A. M. (1993). The reduced-rank growth curve model for discriminant analysis of longitudinal data. *Australian Journal of Statistics*, 35, 345–357.
- Anderson, T. W. (2003). *An introduction to multivariate statistical analysis*. Wiley series in probability and statistics (3rd ed.). Hoboken: Wiley.
- Barnard, M. M. (1935). The secular variations of skull characters in four series of Egyptian skulls. *Annals of Eugenics*, 6, 352–371.
- Bartlett, M. S. (1938). Further aspects of the theory of multiple regression. *Mathematical Proceedings of the Cambridge Philosophical Society*, 34, 33–40.
- Bartlett, M. S. (1939). A note on tests of significance in multivariate analysis. *Mathematical Proceedings of the Cambridge Philosophical Society*, 35, 180–185.
- Bartlett, M. S. (1951). The goodness of fit of a single hypothetical discriminant function in the case of several groups. *Annals of Eugenics*, 16, 199–214.
- Burnaby, T. P. (1966). Growth-invariant discriminant functions and generalized distances. *Biometrics*, 22, 96–110.
- Dasgupta, S. (1993). The evolution of the D^2 statistic of Mahalanobis. *Sankhyā*, 55, 442–459.
- Fisher, R. A. (1936). The use of multiple measurements in taxonomic problems. *Annals of Eugenics*, 7, 179–188.
- Fisher, R. A. (1938). The statistical utilization of multiple measurements. *Annals of Eugenics*, 8, 376–386. (Revised 1950, in Contributions to Mathematical Statistics, Wiley, New York).
- Hodges, J. L. (1950). Discriminatory analysis. I. Survey of discriminatory analysis (Report No. 1, Project No. 21-49-004). USAF School of Aviation Medicine, October 1950. (NTIS No. PB 102-690).
- Hotelling, H. (1931). The generalization of Student's ratio. *The Annals of Mathematical Statistics*, 2, 360–378.
- Hotelling, H. (1936). Relations between two sets of variables. *Biometrika*, 28, 321–377.

- Huberty, C. J. (1975). Discriminant analysis. *Review of Educational Research*, 45, 543–598.
- Khatri, C. G. (1966). A note on a MANOVA model applied to problems in growth curve. *Annals of the Institute of Statistical Mathematics*, 18, 75–86.
- Kollo, T., and von Rosen, D. (2005). *Advanced multivariate statistics with matrices*. Dordrecht: Springer.
- Kshirsagar, A. M. (1962). A note on the direction and collinearity factors in canonical analysis. *Biometrika*, 9, 255–259.
- Kshirsagar, A. M., Davis, T. M. (1983a). Direction and collinearity factors of Wilks' Λ associated with the Growth Curve model. *Australian Journal of Statistics*, 25, 467–481.
- Kshirsagar, A. M., Davis, T. M. (1983b). Discriminant functions from the dummy variable space for growth curve models. *Communications in Statistics-Theory and Methods*, 12, 427–440.
- Lee, J. C. (1982). Classification of Growth Curves. In *Classification, pattern recognition and reduction of dimensionality, Handbook of statist.* (Vol. 2, pp. 121–137). Amsterdam: North-Holland.
- Mahalanobis, P. C. (1936). On the generalized distance in statistics. *Proceedings of the National Institute of Science of India*, 2, 49–55.
- McLachlan, G. J. (2004). *Discriminant analysis and statistical pattern recognition*. Hoboken: Wiley.
- Mentz, G. B., Kshirsagar, A. M. (2004). Classification using growth curves. *Communications in Statistics-Theory and Methods*, 33, 2487–2502.
- Muirhead, R. J. (1982). *Aspects of multivariate statistical theory*. Wiley series in probability and mathematical statistics. New York: Wiley.
- Ngailo, E. K., von Rosen, D., Singull, M. (2021). Asymptotic approximation of misclassification probabilities in linear discriminant analysis with repeated measurements. *Acta et Commentationes Universitatis Tartuensis de Mathematica*, 25, 67–85.
- Pearson, K. (1915). On the problem of sexing osteometric material. *Biometrika*, 10, 479–487.
- Pearson, K. (1926). On the coefficient of racial likeness. *Biometrika*, 18, 105–117.
- Potthoff, R. F., Roy, S. N. (1964). A generalized multivariate analysis of variance model useful especially for growth curve problems. *Biometrika*, 51, 313–326.
- Rao, C. R. (1948). The utilization of multiple measurements in problems of biological classification. *Journal of the Royal Statistical Society, Series B*, 10, 159–203.
- Rao, C. R. (1958). Some statistical methods for comparison of growth curves. *Biometrics*, 14, 1–17.
- Rao, C. R. (1959). Some problems involving linear hypotheses in multivariate analysis. *Biometrika*, 46, 49–58.
- Rao, C. R. (1966). Discriminant function between composite hypotheses and related problems. *Biometrika*, 53, 339–345.
- Rao, C. R. (1973). *Linear statistical inference and its applications*. Wiley series in probability and mathematical statistics (2nd ed.). New York: Wiley.
- Seal, B. (1911). Meaning of race, tribe, nation. In G. Spiller (Ed.) *Inter-racial problems* (pp. 1–12). London: Universal Races Congress. <https://archive.org/details/papersoninterrac00univiala/page/n511/mode/2up?q=Seal>.
- Srivastava, M. S. (2002). *Method of multivariate statistics*. New York: Wiley-Interscience.
- Srivastava, M. S., Khatri, C. G. (1979). *An introduction to multivariate statistics*. New York-Oxford: North-Holland.
- von Rosen, D. (2018). *Bilinear regression analysis: An introduction*. *Lecture notes in statistics* (Vol. 220). New York: Springer.
- Wald, A. (1944). On a statistical problem arising in the classification of an individual into one of two groups. *The Annals of Mathematical Statistics*, 15, 145–162.

Publisher's Note Springer Nature remains neutral with regard to jurisdictional claims in published maps and institutional affiliations.

Springer Nature or its licensor (e.g. a society or other partner) holds exclusive rights to this article under a publishing agreement with the author(s) or other rightsholder(s); author self-archiving of the accepted manuscript version of this article is solely governed by the terms of such publishing agreement and applicable law.