

University of Leuven

A NONMETRIC ALGORITHM FOR ANALYZING PREFERENCE DATA ACCORDING TO THE UNFOLDING MODEL

GERRY EVERS-KIEBOOMS* & LUC DELBEKE

An algorithm for nonmetric internal unfolding analysis of a preference matrix is presented. It is based on the absolute value principle and intends to achieve a maximal mean rank-correlation between the rows of the data matrix and the corresponding rows of the distance matrix computed from the geometric representation. Using simulated data and a real data example, namely preferences for family compositions, the algorithm is compared with MINIRSA, an algorithm for unfolding based on the transformational principle.

INTRODUCTION: THE LINEAR QUADRATIC HIERARCHY OF MODELS

As a function of their complexity, most models for the multidimensional analysis of preference data can be structured in the linear-quadratic hierarchy of models (Carroll, 1972, 1980): the vector model, the unfolding model, the weighted unfolding model, and the general unfolding model. This classification is valid for the internal analysis of a preference matrix where a representation of the subjects and the stimuli in a joint multidimensional space is sought, as well as for the external analysis, where an a priori representation of the stimuli is given, so that only the locations of the subjects in the same space need to be calculated. A brief survey of the linear quadratic hierarchy of models, respectively for external and internal analysis, will permit us to situate our own algorithm among a number of available algorithms.

EXTERNAL ANALYSIS

For the different levels of the linear quadratic hierarchy of models, Carroll (1972) worked out techniques for metric external analysis. To obtain a nonmetric solution according to the unfolding model, Gleason (1969) used linear programming techniques to determine the coordinates

The first author now works as a research fellow of the Belgian Population and Family Study Center and is affiliated with the Center of Human Genetics at the University of Leuven. Both authors are indebted to Professor Srinivasan for comments on a previous version of this paper. Special gratitude is extended to Norman Verhelst, now at the University of Utrecht, and to Rob Stroobants for their valuable comments, programming and technical assistance. Editorial improvements in the text have been made thanks to the many valuable suggestions provided by this article's referee.

of the subjects. (His procedure for external analysis was inserted in a stepwise manner in an algorithm that also allows for internal analysis.) Srinivasan and Shocker (1973) developed a technique for external analysis according to the weighted unfolding model, using additional restrictions to ensure the non-negativeness of the weights. The generalization of this procedure to the general unfolding model was strongly inspired by the metric analysis of Carroll (1972).

INTERNAL ANALYSIS

Matrix solutions based on the vector model were proposed by Tucker (1960), Slater (1960), Delbeke (1968), Bechtel et al. (1971) and Carroll (1972). The one-dimensional unfolding model was presented by Coombs (1950) and representation algorithms were successively adapted by Goode (Coombs, 1964), Phillips (1971) and McClelland and Coombs (1975). The generalization to the multidimensional unfolding model was made by Bennett and Hays (1960). A metric version of the multidimensional unfolding model was described by Wang et al. (1975). Nonmetric procedures were developed by Lingoes (1966), Young and Torgerson (1967), Roskam (1968), Kruskal and Carroll (1969) and Gleason (1969). A probabilistic, multidimensional version of Coombs' unfolding model for paired-comparison data was proposed by Zinnes and Griggs (1974). Multidimensional representations of paired-comparison data can also be found in Carroll (1980) and Heiser and de Leeuw (1981). Srinivasan and Shocker (1973) described very briefly an iterative procedure for internal analysis according to the weighted unfolding model. Unfortunately, few details are given, and no applications to real nor to artificial data are described. The alternating least squares approach advocated by Takane, Young and de Leeuw (1977) can be used to perform multidimensional unfolding under several conditionality restrictions (matrix- or row-conditionality). Within their approach a weighted version of multidimensional unfolding is possible with replicated datasets (Young and Lewyckij, 1979). With some of the procedures mentioned above one can also do an external analysis by keeping the a priori stimulus configuration fixed. Finally, we refer to polynomial factor analysis (Carroll, 1972), which allows for an even more general model formulation than the one provided in the linear quadratic hierarchy of models.

In developing our own algorithm we have limited ourselves to the internal analysis of preference data leading to an unweighted multidimensional unfolding solution. Our main concern was to construct an algorithm for internal analysis based on the absolute value principle (to be introduced in the next section) and to compare it with an algorithm constructed on the basis of the transformational principle.

MAIN PURPOSE OF THE ALGORITHMS

When N persons, belonging to a set P , express their preferences for n stimuli (belonging to a set Q), then the $N \times n$ data matrix $\mathbf{H}$ with elements h_{ij} , representing the preference of person p_i for stimulus q_j is row-conditional (since only preference indices within the same row can be meaningfully compared with each other), and rectangular, i.e. off-diagonal (since the elements p_i and q_j belong to different sets). In what follows, we assume subjects rank order the stimuli, so that a higher preference is expressed by a lower ranking number. (If the data matrix $\mathbf{H}$ contains preference indices such that a higher preference is expressed by a higher preference index, then a decreasing monotone transformation on the elements of $\mathbf{H}$ is all that is needed before applying the algorithm presented here.)

All elements of P and Q are mapped in a m -dimensional Euclidean vectorspace (Re^m, d). The $N \times m$ matrix $\mathbf{X}$ contains the vectors with the coordinates of the subjects, the $n \times m$ matrix $\mathbf{Y}$ contains the vectors with the coordinates of the stimuli. The $N \times n$ matrix $\mathbf{D}$ contains the Euclidean distances d_{ij} between ideal point or subject point x_i and stimulus point or object point y_j .

Our purpose in doing an internal analysis of an off-diagonal, conditional preference matrix $\mathbf{H}$ is to construct a configuration of ideal points ($\mathbf{X}$) and stimulus points ($\mathbf{Y}$), so that for each subject point the rank order of distances to all stimulus points corresponds to the rank order of the data for this subject:

$$d_{ij} < d_{ik} \Leftrightarrow h_{ij} < h_{ik} \quad (1)$$

where $i = 1, \dots, N; j = 1, \dots, n$ and $j \neq k$.

From (1) it is clear that trivializations of strong order relations in the data are not permitted. The purpose of the algorithm can thus be formulated as follows: the rank-correlations between the corresponding rows of the data matrix $\mathbf{H}$ and the distance matrix $\mathbf{D}$ must be 1 to have a perfect solution. Since real data never satisfy the model perfectly and since we want a solution in a low-dimensional space, we cannot maintain N perfect rank-correlations as a realistic goal. We have to attenuate the purpose of "perfect" rank-correlations to "maximal" rank-correlations.

Since all nonmetric algorithms try to obtain optimal monotone relationships between corresponding rows in the data matrix and the distance matrix, they all can be evaluated by a criterion, analogous to a rank-correlation coefficient.

THE "TRANSFORMATIONAL PRINCIPLE" AND THE "ABSOLUTE VALUE PRINCIPLE"

Guttman (1968) distinguishes two different principles in the definition

of a function representing the lack of monotonicity: the "transformational principle" and the "absolute value principle". All the non-metric techniques mentioned in the introduction use the transformational principle. They simultaneously or successively solve a system of equations for two sets of unknowns: the distances in the m -dimensional space, the d_{ij} , which are a function of the coordinates of the ideal points and the stimulus points, and the target distances, the $\hat{d}_{ij}$, which are determined by an optimal monotone transformation of the elements within each row of the matrix $\mathbf{H}$, e.g. by regressing the distances monotonically on the data. The method of steepest descent is used to minimize a stress index, which is a measure of fit based on deviations between distances and target distances:

$$\text{Stress} = \left[\frac{1}{N} \sum_i \frac{\sum_j (d_{ij} - \hat{d}_{ij})^2}{\sum_j (d_{ij} - \hat{d}_i)^2} \right] \quad (2)$$

In nonmetric scaling based on the absolute value principle, a solution is sought for only one set of unknowns, namely the distances as a function of the point coordinates. The essential feature is the direct optimization of the monotone relationship between the data and the distances, thus, without the intermediate step of calculating target distances. Until now, algorithms of this type have been constructed only for the analysis of diagonal dissimilarity matrices (de Leeuw, 1970; Johnson, 1973; Croon, 1974; Möbus, 1976).

NONMETRIC SCALING OF PREFERENCE DATA BASED ON THE ABSOLUTE VALUE PRINCIPLE:

DEFINITION OF THE OBJECTIVE FUNCTION ϕ

We use the absolute value principle to construct an algorithm leading to a solution that tries to optimally satisfy condition (1). The algorithm amounts to the iterative solution of a homogeneous set of distance inequalities.

Starting from the rank ordering of n objects by subject p_i , $n(n-1)/2$ distance inequalities can be determined (if no ties in the data occur). To represent all the existing inequality restrictions for subject p_i we define the set of ordered triples L_{INEQ_i} :

$$L_{INEQ_i} = \{(p_i, q_j, q_k) | p_i \in P, q_j, q_k \in Q, j \neq k, \text{ and } h_{ij} < h_{ik}\} \quad (3)$$

The distance inequalities violated by a given subject and stimulus configuration are represented by a subset of L_{INEQ_i} , namely

$$L_i = \{(p_i, q_j, q_k) | p_i \in P, q_j, q_k \in Q, j \neq k, \text{ and } h_{ij} < h_{ik} \text{ but } d_{ij} \geq d_{ik}\} \quad (4)$$

Each ordered triple in L_i consists of the same subject p_i and two objects q_j and q_k for which the distance inequality is not satisfied. If $h_{ij} = h_{ik}$, the ordered triple (p_i, q_j, q_k) does not belong to L_{INEQ_i} and a fortiori not to L_i .

Consequently, no restriction is imposed upon the rank order of d_{ij} and d_{ik} when $h_{ij} = h_{ik}$. This corresponds to Kruskal's (1964, a, b) primary approach for handling ties.

Starting from the data matrix $\mathbf{H}$, $N \times n(n-1)/2$ distance inequalities can be determined (assuming there are no ties) and are represented in

$$L_{INEQ} = \bigcup_{i=1}^N L_{INEQ_i} \quad (5)$$

The violations against (1) are represented in

$$L = \bigcup_{i=1}^N L_i \quad (6)$$

Minimizing the number of violated distance inequalities amounts to minimizing the number of elements in the set L :

$$\phi = n(L) = \sum_{i=1}^N n(L_i) \quad (7)$$

Implied in the minimization of ϕ is the prevention of trivializations of strong order relations. Together with this ϕ -value, the mean Kendall tau can be calculated between the rows of the data matrix and the corresponding rows of the distance matrix derived from the final solution. If there are ties, a slightly adapted coefficient is adopted:

$$\bar{\tau} = \frac{n(L_{INEQ}) - 2n(L)}{n(L_{INEQ})} \quad (8)$$

By minimizing a stress coefficient as in (2), we do not minimize the number of inequality violations as such, because each inequality is weighted as a function of the discrepancy between d_{ij} and d'_{ij} . By minimizing the stress coefficient priority is given to eliminating large discrepancies. This may be an advantage of stress minimizing algorithms over ϕ -minimizing algorithms, but on the other hand, it will often lead to equating distances in the solution space in spite of strong order relations between the corresponding data elements. This eventually leads to a degenerate solution in which several points are collapsed together. Since trivializing a strong order relation in the data is penalized in the present algorithm, we expect that it will do better in recovering the configuration (and in avoiding degenerate solutions).

Our scaling algorithm OBJSUBJ is inspired to some extent by Srinivasan and Shocker's (1973) procedure. The name of the algorithm is due to the fact that it is based on the successive adaptations of object and subject point configurations.

THE OBJSUBJ ALGORITHM

Srinivasan and Shocker's [1973] algorithm for external analysis

according to the weighted unfolding model with their suggested extension to internal analysis, was the starting point for the construction of our algorithm. We tried out their proposed use of mixed integer linear programming (Hadley, 1964) to minimize the number of violated distance inequalities, but it did not work satisfactorily.

For this reason we developed an algorithm retaining the essential ideas of the suggested procedures, but instead of formulating mixed integer programming problems we use a much more simple and heuristic minimization procedure.

MINIMIZATION OF ϕ IN OBJSUBJ

Optimal subject point configuration. (a) Given a configuration for the subject and the object points from a previous iteration, we first try to satisfy the inequality restriction implied by subject p_i preferring stimulus q_j to stimulus q_k by shifting the subject point x_i along dimension a . This shift, called $e_{a(jk)}$, must result into the following inequality for the squared distances d^2 obtained after the shift:

$$\tilde{d}_{ik}^2 - \tilde{d}_{ij}^2 > 0 \quad (9)$$

Hence, $\tilde{d}_{ik}^2$ (and similarly $\tilde{d}_{ij}^2$) can be written as

$$\begin{aligned} \tilde{d}_{ik}^2 &= \sum_{\substack{p=1 \\ p \neq a}}^m (x_{ip} - y_{kp})^2 + [(x_{ia} + e_{a(jk)}) - y_{ka}]^2 \\ &= d_{ik}^2 + e_{a(jk)}^2 + 2x_{ia} e_{a(jk)} - 2e_{a(jk)} y_{ka} \end{aligned} \quad (10)$$

The inequality restriction in (9) then becomes:

$$d_{ik}^2 - d_{ij}^2 + 2e_{a(jk)} (y_{ja} - y_{ka}) > 0 \quad (11)$$

and it is satisfied if

$$\begin{aligned} e_{a(jk)} &> \frac{d_{ij}^2 - d_{ik}^2}{2(y_{ja} - y_{ka})} \text{ if } (y_{ja} - y_{ka}) \text{ is positive} \\ e_{a(jk)} &< \frac{d_{ij}^2 - d_{ik}^2}{2(y_{ja} - y_{ka})} \text{ if } (y_{ja} - y_{ka}) \text{ is negative.} \end{aligned}$$

The expressions on the right side of these inequalities denote the limiting values of the shift in x_{ia} which will eliminate the inequality violations. If $y_{ja} = y_{ka}$, a shift of x_i along dimension a cannot eliminate the violations of the corresponding distance inequality.

(b) We determine the limiting values of the other shifts of the same subject point x_i along the same dimension a which are necessary to satisfy all the other inequality restrictions in L_{INEQ_i} . Thus, for each ordered triple $(p_i, q_j, q_k) \in L_{INEQ_i}$ for which $y_{ja} \neq y_{ka}$ we determine

$$e_{a(jk)} = \frac{d_{ij}^2 - d_{ik}^2}{2(y_{ja} - y_{ka})} \quad (12)$$

and

$$s_{a(jk)} = -(y_{ja} - y_{ka}) \quad (13)$$

If $(y_{ja} - y_{ka}) > 0$, a shift of subject point x_i along dimension a , exceeding $e_{a(jk)}$ implies the satisfaction of an inequality restriction, which is violated as long as the shift remains smaller than $e_{a(jk)}$. Since the number of violations is reduced by 1 if the shift exceeds $e_{a(jk)}$, $s_{a(jk)} = -1$. If $(y_{ja} - y_{ka}) < 0$, a shift exceeding $e_{a(jk)}$ implies the violation of an inequality restriction, which is satisfied as long as the shift remains smaller than $e_{a(jk)}$. Since the number of violations is augmented by 1 if the shift exceeds $e_{a(jk)}$, $s_{a(jk)} = +1$.

The $e_{a(jk)}$ -values are put in an increasing order and the $s_{a(jk)}$ -values are arranged in the same order as the corresponding $e_{a(jk)}$. Then, the number of violations which remain after each possible shift of subject point x_i along dimension a is easily determined. If we shift x_i over a distance smaller than the smallest $e_{a(jk)}$ -value, the resulting number of violations is equal to the sum (S) of the number of ordered triples $\in L_{INEQ_i}$ for which $y_{ja} = y_{ka}$ and $d_{ij}^2 \geq d_{ik}^2$, and the number of ordered triples $\in L_{INEQ_i}$ for which $(y_{ja} - y_{ka})$ is positive. For a shift with a value between the smallest $e_{a(jk)}$ and the next $e_{a(lm)}$ in the series, the value $s_{a(jk)}$ is added to S . The same procedure is followed in stepping across all successive $e_{a(ts)}$ -values.

To know the resulting number of violations for an arbitrary shift e_a , we start with S and add all the $s_{a(jk)}$ -values corresponding to $e_{a(jk)}$ -values smaller than e_a . This makes it possible to determine the interval over which we can shift x_i along dimension a in order to obtain a minimal number of violations. Sometimes, there will be several intervals which will lead to the same minimal number of violations. For each such interval we determine a representative value e_a as follows:

- if the interval is bounded by a positive and a negative value:
 $e_a = 0$;
- if the interval contains only negative values:
 $e_a = \max \{\text{mean interval value}, u-0.1\}$, where u = upper bound of the interval;
- if the interval contains only positive values:
 $e_a = \min \{\text{mean interval value}, l+0.1\}$, where l = lower bound of the interval.

The representative value with the smallest absolute value is then taken as the optimal shift of the subject point x_i along dimension a .

(c) The optimal shift of x_i along all other dimensions is determined.

(d) Once the m optimal shifts and their corresponding minimal numbers of violations are known, we effectively move subject point x_i along the dimension for which the optimal shift results in the overall minimal number of violations.

(e) This procedure is reiterated for the same subject point x_i until the number of violations remains unchanged with respect to the previous iteration, with a maximum of 10 iterations.

(f) The same iterative procedure is used for all the other subject points.

Optimal object point configuration. (a) We first try to satisfy the inequality restriction implied by the person p_i preferring q_j to q_k by shifting the object point y_j along dimension a . This shift, $\zeta_{a(ijk)}$, must lead to the inequality:

$$\tilde{d}_{ik}^2 - \tilde{d}_{ij}^2 > 0. \quad (14)$$

Since only y_j is shifted along dimension a , $\tilde{d}_{ik}^2 = d_{ik}^2$. On the other hand:

$$\begin{aligned} \tilde{d}_{ij}^2 &= \sum_{\substack{p=1 \\ p \neq i}}^m (x_{ip} - y_{jp})^2 + [x_{ia} - (y_{ja} + \zeta_{a(ijk)})]^2 \\ &= d_{ij}^2 - (x_i - y_{ja})^2 + [x_{ia} - (y_{ja} + \zeta_{a(ijk)})]^2 \end{aligned} \quad (15)$$

Then the inequality restriction in (14) becomes:

$$-\zeta_{a(ijk)}^2 - 2(y_{ja} - x_{ia})\zeta_{a(ijk)} + d_{ik}^2 - d_{ij}^2 > 0 \quad (16)$$

To determine for which $\zeta_{a(ijk)}$ -value the inequality restriction is satisfied, we take the roots of the quadratic function in $\zeta_{a(ijk)}$. If the discriminant $[4(y_{ja} - x_{ia})^2 + 4(d_{ik}^2 - d_{ij}^2)] \leq 0$ a shift of y_j along dimension a cannot eliminate an already existing violation of the inequality restriction. If the discriminant is positive, the inequality restriction is satisfied if

$$\begin{aligned} &-(y_{ja} - x_{ia}) - [(y_{ja} - x_{ia})^2 + d_{ik}^2 - d_{ij}^2]^{1/2} < \zeta_{a(ijk)} \\ &< -(y_{ja} - x_{ia}) + [(y_{ja} - x_{ia})^2 + d_{ik}^2 - d_{ij}^2]^{1/2} \end{aligned}$$

If p_i prefers q_k to q_j , $\zeta_{a(ijk)}$, must be determined so that

$$\tilde{d}_{ij}^2 - \tilde{d}_{ik}^2 > 0 \quad (17)$$

$$d_{ij}^2 - d_{ik}^2 + 2(y_{ja} - x_{ia})\zeta_{a(ijk)} + \zeta_{a(ijk)}^2 > 0 \quad (18)$$

This inequality restriction is satisfied if

$$\begin{aligned} \zeta_{a(ijk)} &< (y_{ja} - x_{ia}) - [(y_{ja} - x_{ia})^2 + d_{ik}^2 - d_{ij}^2]^{1/2}, \text{ or} \\ \zeta_{a(ijk)} &> (y_{ja} - x_{ia}) + [(y_{ja} - x_{ia})^2 + d_{ik}^2 - d_{ij}^2]^{1/2} \end{aligned}$$

If the discriminant is negative, the inequality restriction is always satisfied.

(b) We determine the shifts of y_j along a which are needed to satisfy the other inequality restrictions in the set $O_j = JK_j \cup KJ_j$ where

$$JK_j = \{(p_i, q_j, q_k) | p_i \in P, q_j, q_k \in Q, i = 1, \dots, N, k = 1, \dots, n,$$

$$k \neq j \text{ and } h_{ij} < h_{ik}\},$$

$$KJ_j = \{(p_i, q_k, q_j) | p_i \in P, q_k, q_j \in Q, i = 1, \dots, N, k = 1, \dots, n,$$

$$k \neq j \text{ and } h_{ik} < h_{ij}\}$$

If there are no ties in the data, the number of inequality restrictions to be satisfied for each object point is equal to $N(n-1)$.

For each ordered triple $\in JK_j$, if the discriminant is positive, we determine the roots: $z_{1a(ijk)} < z_{2a(ijk)}$. We also determine the corresponding s -values: $s_{1a(ijk)} = -1$, since for a shift between $z_{1a(ijk)}$ and $z_{2a(ijk)}$ a violation disappears; $s_{2a(ijk)} = +1$, since a shift larger than $z_{2a(ijk)}$ introduces a violation. For each ordered triple $\in KJ_j$, if the discriminant is positive, we determine the roots $z_{1a(ijk)} < z_{2a(ijk)}$, or, if the discriminant is equal to 0, the unique root $z_{a(ijk)}$. We also determine the s -values. If the discriminant is positive: $s_{1a(ijk)} = +1$, since a shift between $z_{1a(ijk)}$ and $z_{2a(ijk)}$ introduces a violation; $s_{2a(ijk)} = -1$, since for a shift $> z_{2a(ijk)}$ the violation disappears again. If the discriminant is zero, then for a shift exactly identical to $z_{a(ijk)}$ we have to take into account that a violation appears, while for shifts which are smaller or larger, the number of violations remains unchanged, so that $s_{a(ijk)} = 0$.

After determining the $z_{a(ijk)}$ - and $s_{a(ijk)}$ -values, all the $z_{a(ijk)}$ -values are put in an increasing order, and the $s_{a(ijk)}$ -values are ordered in accordance with their respective $z_{a(ijk)}$ -values. Then, the number of violations resulting from an arbitrary shift z_a of object point y_j along dimension a is easily determined by adding all $s_{a(ijk)}$ -values corresponding to the $z_{a(ijk)} < z_a$ to the number of violations which corresponds to a shift smaller than the smallest $z_{a(ijk)}$ -value in the series. This last number of violations is equal to the number of triples in the set JK_j . The optimal shift of y_j along dimension a is determined by taking a representative value for the shift in the same way as we did for the subject point x_j in the last part of step (b) for the construction of an optimal subject configuration.

(c) By the same procedure we determine the optimal shift of y_j along the other dimensions.

(d) After the m optimal shifts and their corresponding minimal numbers of violations are known, we effectively move object point y_j along the dimension for which the optimal shift results in the overall minimal number of violations.

(e) After the coordinates of object point y_j are changed, we change the distance matrix accordingly and then we repeat the same procedure for all the other object points.

(f) The steps (a) to (e) are repeated iteratively until the number of violations does not decrease any more, with a maximum of 10 iterations.

Alternating adaptations of object and subject configurations. Starting with an initial configuration, we alternately adapt the object and subject configurations until the total number of violations remains unchanged for two consecutive adaptations.

THE INITIAL CONFIGURATION

Most authors emphasize the importance of a starting configuration

which will increase the speed of obtaining a final solution and which helps to avoid local minima. Using artificial data, we tested several starting configurations on their efficiency. From those we chose a type of initial configuration almost identical to the one proposed by Gleason (1969). Starting from the intercorrelations between subjects, the m principal axes of this correlation matrix are used to determine a subject configuration. If the preference indices in $\mathbf{H}$ are not observed directly as ranking numbers, then the elements in $\mathbf{H}$ are converted into ranking numbers r . Then we take the minimum of the N ranking numbers attributed to object q_j :

$$r_{\min j} = \min \{r_{1j}, \dots, r_{Nj}\} \quad (19)$$

The starting configuration of the stimulus points is then determined by taking as coordinates for each stimulus point y_j the centroid of the subject points x_i for which $r_{ij} = r_{\min j}$.

RESULTS OBTAINED WITH OBJSUBJ AND MINIRSA

ARTIFICIAL DATA

In order to compare the results obtained with OBJSUBJ with the results obtained with a stress minimizing algorithm, we applied both OBJSUBJ and MINIRSA (Roskam, 1968) to two kinds of artificial data. First, 20 data matrices were derived from two-dimensional artificial configurations with 10 subject and 5 object points. Each artificial configuration was obtained by constructing a 15×2 matrix with random numbers drawn from a rectangular distribution between 0 and 1. Then the origin was translated to the centroid and the configuration was normalized. The first 10 rows were considered as coordinates for the subject points, the last 5 rows as coordinates for the stimulus points. After the 10×5 distance matrix was calculated from the configuration, all the distances were multiplied by a constant and considered as preference indices. Another set of twenty 25×10 data matrices were generated in exactly the same way.

The second kind of artificial data consisted of twenty 10×5 data matrices in which all the data elements were generated directly by using random numbers from a rectangular distribution.

For each of the data matrices a solution was obtained with OBJSUBJ and MINIRSA. We always used the same type of initial configuration.

To compare the solutions obtained with the algorithms we shall use in the first place the total number of violations of distance inequalities, i.e. the final ϕ -value. This number thus includes the number of trivializations of strong order relations in the data. This ϕ -value is, of course, related to the mean adapted Kendall tau, which we have also computed. The other goodness of fit measure we are interested in is the stress index.

TAB. 1. SUMMARY OF RESULTS OF THE APPLICATION OF OBUSUBI AND MINIRSA TO ARTIFICIAL DATA

1. two-dimensional random configurations	mean number of violations	$\bar{r}$	stress ¹	congruence ¹	mean ² CPU time (msec)
1. a. 10 × 5 data matrices					
OBUSUBI-2D-solutions	.70(.50)	.986(.990)	.093	.90(.89/.92)	116
MINIRSA-2D-solutions	2.60(3.25)	.950(.935)	.004	.90(.91(.88/.89)	91
OBUSUBI-1D-solutions	7.95(8.35)	.840(.833)	.315	.9045	118
MINIRSA-1D-solutions	15.55(14.75)	.690(.705)	.200	.8765	293
1. b. 25 × 10 data matrices					
OBUSUBI-2D-solutions	16.10	.971	.094	.96/.96	10,372
MINIRSA-2D-solutions	29.05	.947	.005	.96/.96	1,449
OBUSUBI-1D-solutions	142.60	.745	.412	.9058	8,463
MINIRSA-1D-solutions	216.85	.612	.317	.8725	731
2. Random data matrices 10 × 5					
OBUSUBI-2D-solutions	6.65(7.05)	.866(.859)	.337		172
MINIRSA-2D-solutions	18.00(19.85)	.633(.603)	.062		732
OBUSUBI-1D-solutions	16.15(16.85)	.674(.663)	.473		134
MINIRSA-1D-solutions	24.00(25.65)	.567(.487)	.350		140

The numbers between brackets refer to results which were available from an earlier, but incomplete simulation study.¹ All reported coefficients are RMS-coefficients computed over the 20 matrices within each series.² All computations were performed in double precision on an IBM 3033 computer at the Computing Center of the University of Leuven.

As can be seen in Table 1 both coefficients, $\bar{r}$ and stress, were computed for each solution: the OBJSUBJ solution and the MINIRSA solution. For these data matrices which were derived from random configurations, we rotated the dimensions in the two-dimensional solution spaces towards maximal congruence with the synthetic configuration. The correlations between the point projections on the rotated dimensions and on their corresponding dimensions in the synthetic configuration were computed as indices for the recovery capabilities of both algorithms. For the one-dimensional solutions we computed the multiple correlations with the projections on the two dimensions in the synthetic configurations. Finally, we also registered the CPU time needed by each algorithm to reach the final solution for each data matrix. All these results have been summarized in Table 1 by giving mean values or RMS-values for each series of 20 matrices.

The difference in performance between OBJSUBJ and MINIRSA is in the expected direction: within each series of matrices, and within each dimensionality, OBJSUBJ does better in terms of $\bar{r}$, while MINIRSA is superior with respect to stress. All differences between the respective means are statistically significant on a level smaller than 0.1%, except for the difference in stress values for the two-dimensional solutions which were found for the 10×5 data matrices derived from the two-dimensional random configurations. For this last comparison a t -value was found of 2.46 ($P = 0.0118$). A typical characteristic of this series of solutions was that OBJSUBJ succeeded in 12 of the 20 cases in finding a solution within the region of the absolute minimum of the loss function, i.e. solutions with $\bar{r} = 0$, and hence also with a stress value equal to zero. For the same series of data matrices, MINIRSA found only four solutions which had a zero stress value. (In a previous simulation study, these numbers were respectively equal to 11 and 3.) The advantage of OBJSUBJ in finding two-dimensional solutions without any violations against the data which were generated from two-dimensional random configurations, disappears in the series with 25×10 matrices, where neither of both algorithms succeeded in finding such a solution.

The two-dimensional solutions obtained with the two algorithms do not show any superiority for one of the two algorithms with respect to the metric recovery of the synthetic configurations. The congruence coefficients computed for the one-dimensional solutions show a superiority for OBJSUBJ ($t = 1.18$ and $P = 0.13$ for the 10×5 matrices; $t = 4.18$ and $P = 0.0003$ for the 25×10 matrices).

The differences between the computed congruence coefficients however are small, so that we can say that, even with one-dimensional solutions, both algorithms lead to configurations which reflect to some degree some characteristics of the synthetic configurations (e.g. clusterings of subject and stimulus points). On the other hand, the computed coefficients are not very sensitive for the loss in recovery quality which may be due to the susceptibility for trivializations or near trivializations of strong order relations in the data. Therefore, in

Fig. 1, we show a few typical examples of the one-dimensional solutions as we find them in the two algorithms.

Each solution is represented by a separate graph. Within each graph, the asterisk disconnecting the scale line indicates the zero point of the scale. The identification numbers for subjects (ranging from 0 to 9) and for objects (ranging from 1 to 5) written above each other together with a brace sign have identical scale values. The numbers written above each other without a brace sign have nearly identical scale values, and in this latter case the more a number is separated from the scale line, the more extreme its scale value is with respect to the zero point on the scale.

Not too much importance should be attached to the equality or near-equality of the scale values for several subjects, because all equal scale values for subjects in the MINIRSA solutions which are shown in Fig. 1 are due to equal rank orders of the stimuli for these subjects in the data matrix. When we look at the stimulus configurations however, we see that there is a clear tendency in the MINIRSA solutions for clustering the stimuli in different groups of tied or nearly tied stimulus projections. So, for instance, in matrix 20, which contained exceptionally few different rank orders (only 4), the MINIRSA solution shows the grouping of 5 with 3 and 1 with 4. This solution however is not an optimal solution (stress = 0.0049; $\bar{r} = 0.90$) because OBJSUBJ leads to a perfectly fitting one-dimensional solution (the only one found in this series of matrices) which corresponds very well to the positions of the stimuli on an interconnecting curvilinear line which can be drawn in the original synthetic two-dimensional configuration space.

Table 1 shows very high CPU times for OBJSUBJ when it is applied to the 25×10 matrices. This is, of course, due to the heuristic character of the algorithm: the more points (and dimensions) in the configurations, the more steps have to be executed in the algorithm. These CPU times can be lowered somewhat by improving the program (e.g. by collapsing subjects with identical rank orders into one single subject point, which we have preferred not to do in the testing phase of the algorithm). These improvements will not change drastically this large difference between the two programs. When the number of points is very small, as with the 10×5 matrices, we get a shorter mean CPU time for MINIRSA, only when the solution is sought in a space having the true dimensionality. Therefore we conjecture that OBJSUBJ is a valuable alternative to at least one of the stress-minimizing algorithms, namely MINIRSA, when the number of points is small, and that it remains a valuable back-up algorithm, even when the number of points becomes relatively large, when other algorithms lead to solutions with a high number of trivializations or near trivializations. To illustrate this latter conjecture, we now give a real data example.

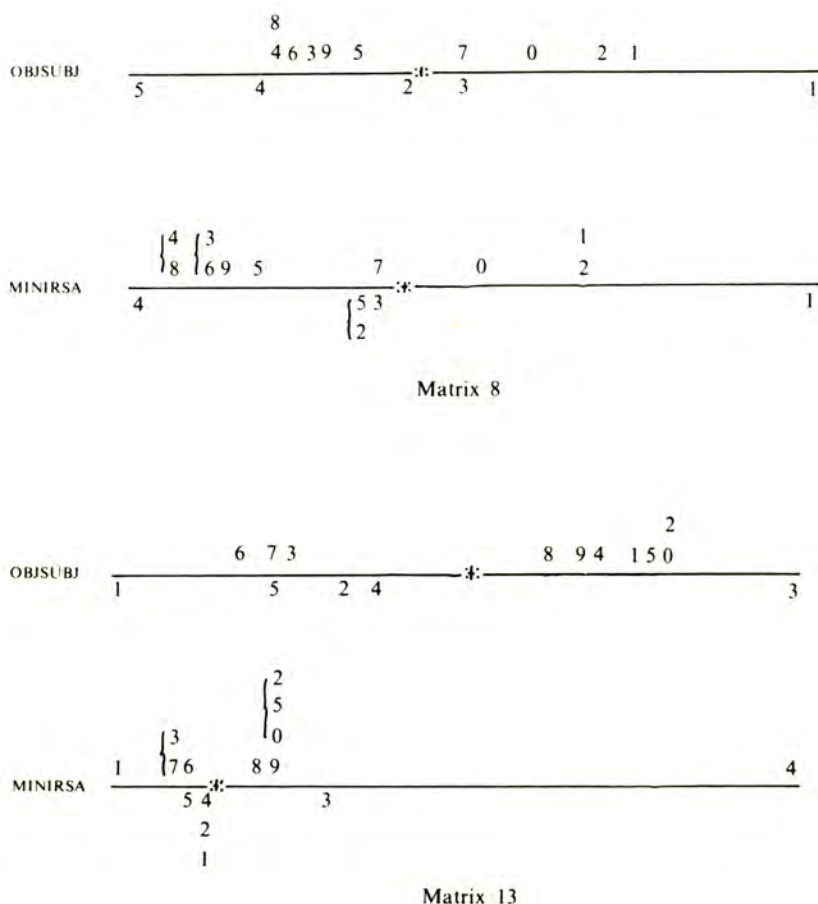
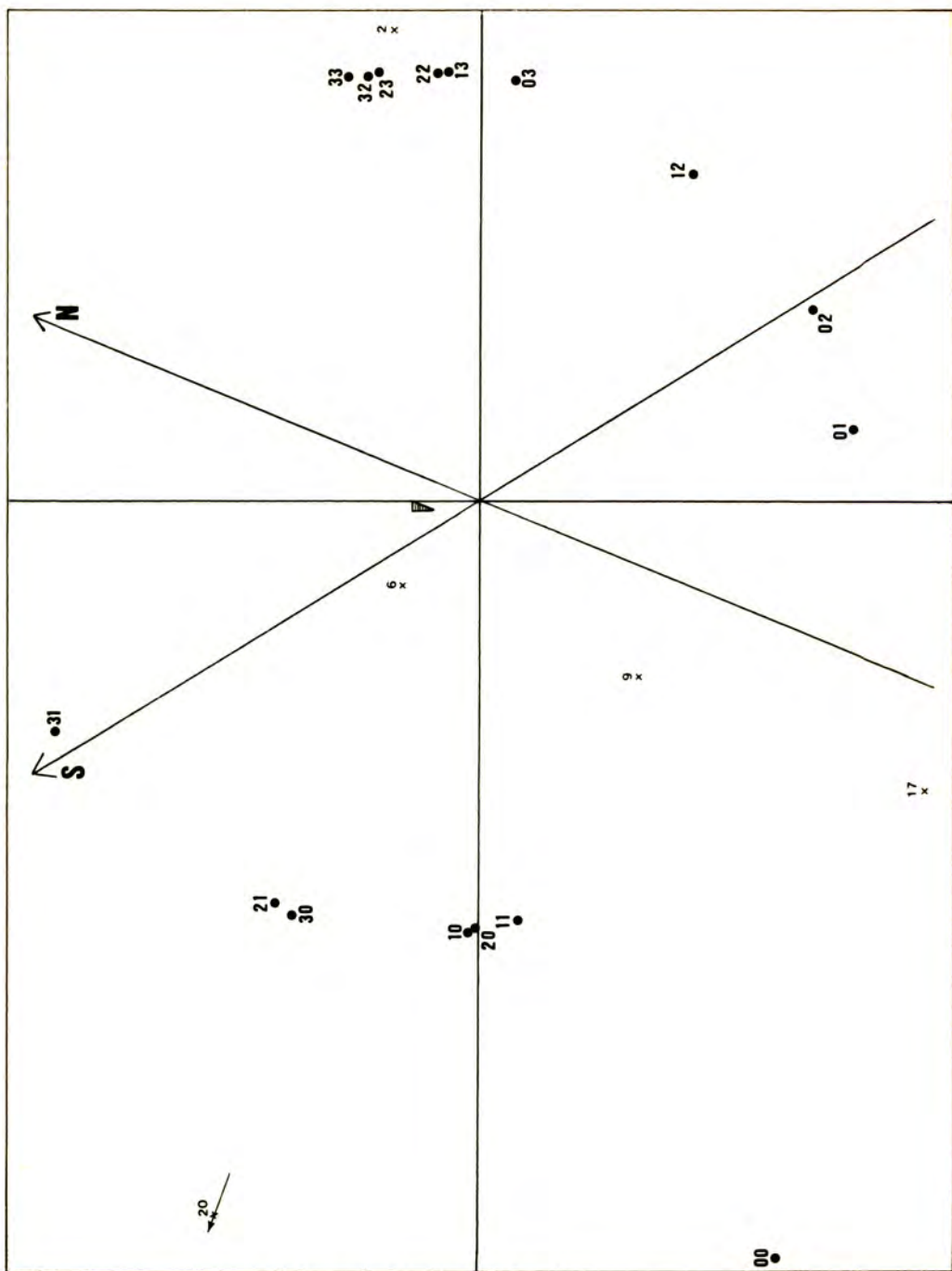


FIG. 1. GRAPHICAL REPRESENTATION OF FOUR ONE-DIMENSIONAL SOLUTIONS

REAL DATA : PREFERENCES FOR FAMILY COMPOSITIONS

The data consisted of a selected subgroup of 21 subjects whose individual preference rank orders for 16 family compositions could be fitted in terms of monotonicity between preference rank numbers and utility scale values for the stimuli, and where the scale values satisfied an additive conjoint measurement model. Each scale value could be decomposed in a utility component measured over the number of children attribute, and another utility component measured over the



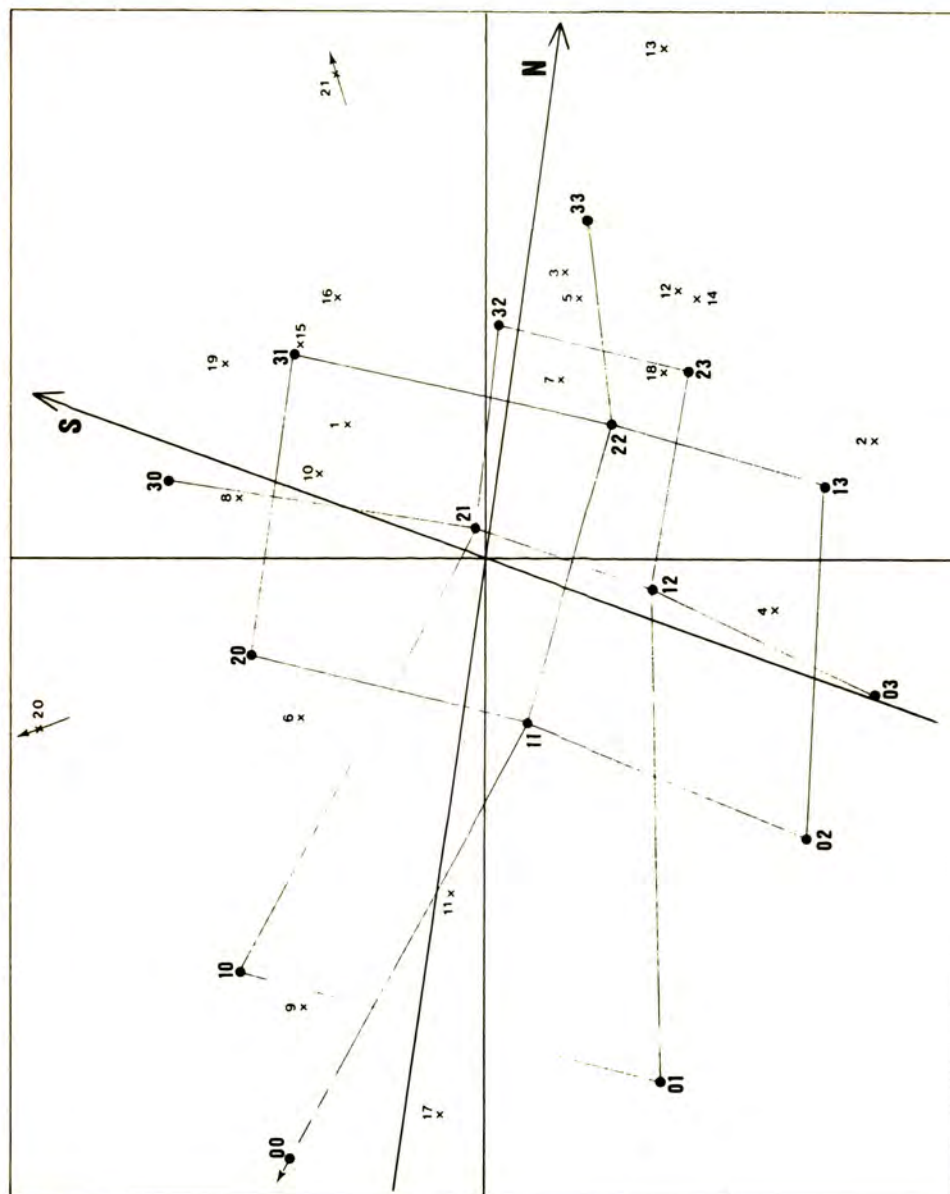


FIG. 2. TWO-DIMENSIONAL SOLUTIONS FOR FAMILY COMPOSITION PREFERENCES : a. MINIRSA, b. OBJSUBJ. The dots correspond to stimulus points, the crosses to subject points. The first bold faced integer indicates the number of boys in a family composition and the second indicates the number of girls. The small triangle in the MINIRSA solution contains several subregions with the localizations of the 16 subject points not shown in the graph.

TAB. 2. FITTING COEFFICIENTS FOR THE MINIRSA AND THE OBJSUBJ SOLUTIONS IN THE FAMILY COMPOSITION PREFERENCE ANALYSIS

	stress	$\bar{r}$	CPU time (IBM 158, in sec.)
MINIRSA	.0585	.74	52.38
OBJSUBJ	.1976	.89	271.15

Fig. 2 shows the results for the two algorithms. Neither of the solutions has a perfect fit, as can be seen in Table 2. The MINIRSA solution shows the stimuli on a nearly circular configuration, with the exception of the no-children family, which was clearly the most unpopular stimulus. The stimuli, however, are not evenly spread on the circle. They can be separated into two opposite groups: the families with the number of boys greater than the number of girls, together with the family with one boy and one girl, versus the remaining stimuli, including the family with three boys and two girls. Except for this last stimulus the opposition reflects the sex preference dimension. In Fig. 2(a) the *S* dimension has a maximum correlation between stimulus projections and sex bias levels running from +3 (the girls in a family are outnumbered by 3 by the boys) to -3 (the boys are outnumbered by 3 by the girls). This correlation coefficient is equal to 0.77. The *N* dimension on the other hand has a maximum correlation of 0.94 with the number of children in each family composition (from 6 to 0).

Although the solution in Fig. 2(a) reflects two important stimulus characteristics, it does not reveal the lattice structure we expected. Another undesirable aspect of the solution has to do with the subject points: 16 of the 21 subjects are localized within a very tiny subregion of the preference space. This region coincides with the center of the circular structure. Using the results of the separate one-dimensional unfolding analyses with respect to sex preference and number preference, each subject could be attributed a score for these preference scales according to the localization of their *I*-scales on the respective *J*-scales. These scores are labeled *IS* and *IN*, and can be used to evaluate the fit of the solutions in Fig. 2 with respect to the subject point configuration. In Table 3 we see that the correlation between the subject scores and the projections on the best fitting stimulus dimension is 0.60 for the number preference and only 0.12 for the sex preference. (These coefficients are equal to 0.53 and 0.45 when computed over the subgroup of the 16 subjects in the center of the configuration.)

The OBJSUBJ solution is shown in Fig. 2(b). Here, again, the best fitting *N* and *S* dimensions have been drawn in. Moreover, we have connected the stimuli with identical *N* and *S* levels, to give a clearer picture of the lattice configuration which we expected to find. The correlation between the objective stimulus characteristics and the

TAB. 3. CORRELATIONS BETWEEN SUBJECT SCORES (*IN* and *IS*) AND SUBJECT PROJECTIONS ON BEST FITTING DIMENSIONS (*N* and *S*)

		<i>S</i>	<i>IN</i>	<i>IS</i>
		.71	.60	-.18
	<i>N</i>	.99	.53	.32
		.08	.56	-.17
MINIRSA : <i>n</i> = 21			.17	.12
MINIRSA : <i>n</i> = 16	<i>S</i>		.50	.45
OBJSUBJ : <i>n</i> = 21			-.34	.76
				-.16
	<i>IN</i>			-.12
				-.16

projections on the best fitting dimensions are 0.93 and 0.96. It is also interesting to note that the stimulus projections on the *N* dimension show the same logarithmic type of scale as we found in the one-dimensional unfolding analysis for number preference: the higher the number of children, the smaller the difference in scale values. As we can see in Table 3 the correlations for the subjects now reach significant values for both dimensions: 0.56 and 0.76. These values are not much higher than those obtained with the subgroup of 16 subjects in the MINIRSA-solution, but the high intercorrelation between the *N* and *S* projections disappears in the OBJSUBJ solution. Therefore, we conclude that the latter solution is a much more adequate solution with respect to both the stimulus and the subject configuration. This is also reflected in the fact that not only the total number of rank order violations is much lower in the OBJSUBJ solution, namely 134 versus 330 in the MINIRSA solution, but that each individual rank order, without any exception, is fitted better by OBJSUBJ than by MINIRSA.

REFERENCES

- BECHTEL, G., TUCKER, L., & CHANG, W. A scalar product model for the multi-dimensional scaling of choice. *Psychometrika*, 1971, 36, 369-388.
- BENNETT, J., & HAYS, W. Multidimensional unfolding: determining the dimensionality of ranked preference data. *Psychometrika*, 1960, 25, 27-43.
- CARROLL, J. Individual differences and multidimensional scaling. In R. SHEPARD, A. ROMNEY & S. NERLOVE (Eds.), *Multidimensional scaling. Theory and applications in the behavioral sciences* (Vol. 1). New York: Seminar Press, 1972.
- CARROLL, J. Models and methods for multidimensional analysis of preferential choice (or other dominance) data. In E.D. LANTERMANN & H. FEGER (Eds.), *Proceedings of Aachen symposia on decision making and multidimensional scaling*. Berlin: Springer Verlag, 1980.
- COOMBS, C. Psychological scaling without a unit of measurement. *Psychological Review*, 1950, 57, 145-158.
- COOMBS, C. *A theory of data*. New York: Wiley, 1964.

- COOMBS, C. A note on the relation between the vector model and the unfolding model for preferences. *Psychometrika*, 1975, 40, 115-116.
- CROON, M. *Multidimensionele schaaltechnieken voor symmetrische dissimilariteitsmatrices*. Unpublished doctoral dissertation, University of Leuven, 1974.
- DELBEKE, L. *Construction of Preference Spaces. An investigation into the applicability of multidimensional scaling models*. Leuven: Leuven University Press, 1968.
- DELBEKE, L. Variables and methods of measurement. In J.A. MICHON, E.G. EIJKMAN & L.F. DE KLERK (Eds.), *Handbook of psychonomics* (Vol. 1). Amsterdam: North-Holland Publishing Company, 1979.
- DE LEEUW, J. *The positive orthant method for nonmetric multidimensional scaling*. (Research Note RN 001-70). University of Leyden, Department of Data Theory, 1970.
- EVERS-KIEBOOMS, G. *Multidimensionele analyse van voorkeurgegevens: ontwikkeling van drie algoritmes voor het opstellen van een spatiale representatie volgens het ontvouwingsmodel*. Unpublished doctoral dissertation, University of Leuven, 1977.
- GLEASON, T. *Multidimensional scaling of sociometric data*. Unpublished doctoral dissertation, University of Michigan, 1969.
- GUTTMAN, L. A general nonmetric technique for finding the smallest coordinate space for a configuration of points. *Psychometrika*, 1968, 33, 469-506.
- HADLEY, G. *Linear programming*. Reading, Mass.: Addison-Wesley, 1962.
- HEISER, W., & DE LEEUW, H. Multidimensional mapping of preference data. *Mathématiques et Sciences Humaines*, 1981, 19, 39-96.
- JOHNSON, R. Pairwise nonmetric multidimensional scaling. *Psychometrika*, 1973, 38, 11-18.
- KRUSKAL, J. Multidimensional scaling by optimizing goodness of fit to a nonmetric hypothesis. *Psychometrika*, 1964 (a), 29, 1-27.
- KRUSKAL, J. Nonmetric multidimensional scaling: A numerical method. *Psychometrika*, 1964 (b), 29, 115-129.
- KRUSKAL, J., & CARROLL, J. Geometrical models and Badness-of-Fit functions. In P. KRISHNAIAH (Ed.), *Multivariate analysis II*. New York: Academic Press, 1969.
- LINGOES, J. An IBM 7090 program for Guttman-Lingoes smallest space analysis — R II. *Behavioral Science*, 1966, 11, 322.
- MCCLELLAND, G., & COOMBS, C. ORDMET: a general algorithm for constructing all numerical solutions to ordered metric structures. *Psychometrika*, 1975, 40, 269-290.
- MÖBUS, C. Nonmetric multidimensional scaling without disparities and derivatives: A rankcorrelation-oriented approach through L_1 -approximation. *Archiv für Psychologie*, 1976, 128, 240-266.
- PHILLIPS, J. A note on representation of ordered metric scaling. *British Journal of Mathematical and Statistical Psychology*, 1971, 24, 239-250.
- ROSKAM, E. *Metric analysis of ordinal data in psychology*. Voorschoten: VAM, 1968.
- SLATER, P. The analysis of personal preferences. *British Journal of Statistical Psychology*, 1960, 13, 119-135.
- SRINIVASAN, R., & SHOCKER, A. Linear programming techniques for multidimensional analysis of preferences. *Psychometrika*, 1973, 38, 337-369.
- TAKANE, Y., YOUNG, F.W., & DE LEEUW, J. Nonmetric individual differences multidimensional scaling: An alternating least squares method with optimal scaling features. *Psychometrika*, 1977, 42, 7-67.
- TUCKER, L. Intra-individual and inter-individual multidimensionality. In H. GULLIKSEN & S. MESSICK (Eds.), *Psychological scaling. Theory and applications*. New York: Wiley, 1960, 155-167.
- WANG, M., SCHÖNEMANN, P., & RUSK, J. A conjugate gradient algorithm for the multidimensional analysis of preference data. *Multivariate Behavioral Research*, 1975, 10, 45-79.

- YOUNG, F.W., & LEWYCKYJ, R. *ALSCAL-4: User's guide*. University of North Carolina, Psychometric Laboratory, 1979.
- YOUNG, F., & TORGERSON, W. TORSAL: a Fortran IV program for Shepard-Kruskal Multidimensional scaling analysis. *Behavioral Science*, 1967, 12, 498.
- ZINNES, J., & GRIGGS, R. Probabilistic multidimensional unfolding analysis. *Psychometrika*, 1974, 39, 327-350.

Department of Psychology
University of Leuven
Tiensestraat 102
3000 Leuven

Received February 1982