



Revue
recherche
de
l'information

Revue
de
la
société
de
la
science
de
la
communication

BELGIAN JOURNAL OF OPERATIONS RESEARCH, STATISTICS AND COMPUTER SCIENCE (JORBEL)

PRINCIPAL EDITORS

M. DESPONTIN
Vrije Universiteit Brussel
Statistiek en Operationeel
Onderzoek
Pleinlaan 2
B - 1050 Brussel - Belgium

J. TEGHEM Jr.
Faculté Polytechnique de Mons
Département Information
et Gestion
rue de Houdain 9
B - 7000 Mons - Belgium

Ph. VINCKE
Université Libre de Bruxelles
Institut de Statistique
CP 210
Bd du Triomphe
B - 1050 Bruxelles - Belgium

BELGIAN ASSOCIATE EDITORS

J.P. BRANS, Vrije Universiteit Brussel,
Pleinlaan 2, 1050 Brussel
F. BROECKX, Rijksuniversitair Centrum,
Middelheimlaan 1, 2020 Antwerpen
Chr. DE BRUYN, Université de Liège, Bât. 31,
Sart-Tilman, 4000 Liège
F. DROESBEKE, Université Libre de Bruxelles, CP
210, bd du Triomphe, 1050 Bruxelles
J. FICHEFET, Facultés Universitaires Notre-Dame
de la Paix, rue Grandgagnage 21, 5000 Namur
L. GELDERS, Katholieke Universiteit Leuven,
Celestijnenlaan 300B, 3030 Leuven-Heverlee
F. JUCKLER, Université Catholique de Louvain, av.
de l'Espinette 16, 1348 Louvain-la-Neuve
L. KAUFMAN, Vrije Universiteit Brussel,
Laarbeeklaan 103, 1090 Brussel
J. LORIS-TEGHEM, Université de l'Etat, place
Warocqué 17, 7000 Mons
H. MULLER-MALEK, Rijksuniversiteit, Sint
Pietersnieuwstraat 41, 9000 Gent
H. PASTIJN, Ecole Royale Militaire, avenue de la
Renaissance 30, 1040 Bruxelles
R. SNEYERS, Institut Royal Météorologique,
avenue Circulaire 3, 1180 Bruxelles

INTERNATIONAL ASSOCIATE EDITORS

O.D. ANDERSON, University of Western Ontario,
London, Canada
J.F. BENDERS, Technical University, Eindhoven,
The Netherlands
R.E. BURKARD, University of Technology, Graz,
Austria
Chr. CARLSSON, Abo Academy, Finland
G.R. d'AVIGNON, Université Laval, Québec,
Canada
B. FICHET, Université d'Aix-Marseille II, France
T.J. HODGSON, North Carolina State University,
Raleigh, N.C., USA
J. KRARUP, University of Copenhagen, Denmark
M. NEUTS, University of Arizona, Tucson, USA
K. OHNO, Nagoya-Institute of Technology, Nagoya,
Japan

B. RAPP, Institute of Technology, Linköping,
Sweden
A.H.G. RINNOOY KAN, Erasmus Universiteit,
Rotterdam, The Netherlands
B. ROY, Université de Paris-Dauphine, Paris,
France
R. SLOWINSKI, Technical University, Poznan,
Poland
J. SPRONK, Erasmus Universiteit, Rotterdam, The
Netherlands
P. TOTH, University of Bologna, Italy

EDITORIAL POLICY

Operations Research, Statistics and Computer Science are of an ever increasing importance as Scientific Methods of Management. They often are linked both in Theory and Practice.

It is the main purpose of Jorbel, the Belgian Journal of Operations Research, Statistics and Computer Science, to promote Research, Applications and Cooperation in these fields.

Therefore the Editorial Policy is to encourage the publication of two kinds of papers: *scientific papers*, such as new theoretical contributions, original applications, computational studies, and on the other hand *didactic papers* such as surveys, tutorials.

Each issue normally includes a tutorial paper. All the papers are refereed. For the sake of effective communication it is recommended to the authors to have their papers written in English. Instructions to the authors are given at the end of this issue.

PUBLICATION

The Belgian Journal of Operations Research, Statistics and Computer Science is published by the Sogesci, rue de la Concorde 53, 1050 Bruxelles - B.V.W.B., Eendrachtstraat 53, 1050 Brussel. It is supported by the "Ministerie van de Vlaamse Gemeenschap" and the "Ministère de l'Education, de la Recherche et de la Formation, le service des études et de la recherche scientifique de la Communauté Française de Belgique".

**BELGIAN JOURNAL OF OPERATIONS RESEARCH,
STATISTICS AND COMPUTER SCIENCE**

**REVUE BELGE DE RECHERCHE OPERATIONNELLE,
DE STATISTIQUE ET D'INFORMATIQUE**

**BELGISCH TIJDSCHRIFT VOOR OPERATIONEEL
ONDERZOEK, STATISTIEK EN INFORMATICA**

Vol. 31 n° 1-2

March-June 1991

Vol. 31 nr 1-2

CONTENTS

- 3 L.F. GELDERS and D.G. CATTRYSSSE: Public Waste Collection: a Case Study.
- 17 W.G. HALLERBACH: The Stationarity of CAPM-Beta in a Changing Economic Environment.
- 34 M. KAYEMBE: Application of Neuts' Method to Vacation Models with Bulk Arrivals.
- 49 B. PELEGRIN: Th P-Center Problem in R^n with Weighted Tchebycheff Norms.
- 63 H. HASSLER and G.J. WOEGINGER: A Variation of Graham's LPT-Algorithm.
- 69 D.J.E. BAESTAENS: Modelling Corporate Vulnerability: Yet Another Empirical Attempt.
- 91 D. TRAPPERS: Another Counterexample to the Rank-Colouring Conjecture.

INSTRUCTIONS TO THE AUTHORS

1. JORBEL accepts papers in the fields of Operations Research, Statistics and Computer Science.
2. Theoretical, applied and didactical papers, as well as documented computer programs are considered.
3. As far as possible, the journal will include TUTORIAL PAPERS. Authors are invited to submit also such papers giving didactical surveys on specific themes. These papers can be theoretical or applied. They should be comprehensible for non specialists and should establish the link between research and practice.
4. For the sake of effective communication, the papers should be written in *English*.
5. All submitted papers will be refereed. Only the authors will be responsible for their text.
6. As the journal is printed by offset, the manuscript should be **camera ready**, including figures and tables. Pages should be numbered on the back.
The *first page* is composed by the printer. It should **only** include:
 - The title of the paper
 - The names and affiliations of the authors (address included)
 - An abstract of no more than 8 lines in English.

It is recommended to have the text of the manuscript typed double spaced with a prestige Elite type-head and the paragraph headings with a Script.
The formulas should also be typed and numbered on the right-hand side.
At each paragraph an 8 character jump should be observed.
7. The first author of published papers will receive 25 copies.
8. Papers should be submitted in **three** copies to one of the principal editors.
The authors should mention whether it is a tutorial paper.

ANNUAL SUBSCRIPTION RATE (4 issues)

BELGIUM: 1.000 BF. OTHER COUNTRIES: 1.200 BF

Sogesci: 53, rue de la Concorde, B-1050 Bruxelles, Belgium

B.V.W.B.: Eendrachtstraat 53, B-1050 Brussel, Belgium

Bank Account: 000-0027041-75

SOGESCI - B.V.W.B. BULLETIN

A Bulletin giving information on Operations Research, Statistics and Computer Science events is published quarterly by the SOGESCI-B.V.W.B.

Information on such events are to be mailed to Ph. Van Asbroeck, ULB, CP210, Boulevard du Triomphe, B-1050 Brussels, Belgium

**BELGIAN JOURNAL OF OPERATIONS RESEARCH,
STATISTICS AND COMPUTER SCIENCE**

**REVUE BELGE DE RECHERCHE OPERATIONNELLE,
DE STATISTIQUE ET D'INFORMATIQUE**

**BELGISCH TIJDSCHRIFT VOOR OPERATIONEEL
ONDERZOEK, STATISTIEK EN INFORMATICA**

Vol. 31 n° 1-2

March-June 1991

Vol. 31 nr 1-2

CONTENTS

- 3 L.F. GELDERS and D.G. CATTRYSSSE: Public Waste Collection: a Case Study.
- 17 W.G. HALLERBACH: The Stationarity of CAPM-Beta in a Changing Economic Environment.
- 34 M. KAYEMBE: Application of Neuts' Method to Vacation Models with Bulk Arrivals.
- 49 B. PELEGRIN: Th P-Center Problem in R^n with Weighted Tchebycheff Norms.
- 63 H. HASSLER and G.J. WOEGINGER: A Variation of Graham's LPT-Algorithm.
- 69 D.J.E. BAESTAENS: Modelling Corporate Vulnerability: Yet Another Empirical Attempt.
- 91 D. TRAPPERS: Another Counterexample to the Rank-Colouring Conjecture.



Printed in Belgium by
Robert Louis, 1050 Brussels
Tel (02) 640 10 40

PUBLIC WASTE COLLECTION: A CASE STUDY

**Ludo F. GELDERS
and
Dirk G. CATTRYSE**

**Katholieke Universiteit Leuven
Afdeling Industrieel Beleid
Celestijnenlaan 300 A, B-3001 Leuven-Heverlee, Belgium**

ABSTRACT

This paper deals with the waste collection in the N.E. area of Brussels. It summarizes a study carried out as a master's thesis project. Until now the collection scheme was based upon experience. Management felt that a more normative and systematic approach was needed. This paper discusses the modeling of the real-life problem based on the capacitated arc routing problem. The solution approach is based on the path scanning algorithm. The problem solver was coded in Pascal and linked with dBase-files which contain all information on the collection area. A reduction of approximately 15% in distance travelled was achieved.

1. PROBLEM DESCRIPTION.

The current problem deals with the waste collection of a set of 5 municipalities in the N.E. area of Brussels. Between 4 and 11 waste collection trucks have to serve these communities daily. On a yearly basis, these lorries total 360.000 km. These 360.000 km are divided into 60.000 km for the actual waste collection and 300.000 km for trips to the waste disposal area, which is located south of Brussels. The distance from the last collection point to the disposal area varies between 25 and 50 km.

Until now, the collection scheme was based upon experience. Each truck was assigned a set of streets to be served. In the current system, a truck crew is free to go home whenever the job is finished. This work organization relies upon the hypothesis that the work force will try to optimize its routing. Under current operations, the company employs 41 people, 34 being crew members of the trucks.

As always in garbage collection problems, the experience based work schemes take into account a very broad and complicated set of constraints, e.g. :

- one-way streets
- fluctuating quantities to be collected
- given collection frequency (twice per week)
- capacity limitation of trucks
- truck maintenance
- special collection services due to special events
- etc...etc...

Management felt however that a more normative and systematic approach might give some insight to do the job more efficient, through a better use of resources (decreasing the total travel distance, the number of trucks used, etc...). Moreover a need was felt to have an interactive decision tool which might provide quick answers to certain urgent problems (e.g. unavailability of a crew or a truck).

2. MODELING THE PROBLEM.

In literature one can find discussions on similar problems, e.g. : sanitation vehicle routing [2], vehicle routing for municipal waste collection [1] and routing electric meter readers [9].

The above waste collection problem can be modelled as the Capacitated Arc Routing Problem (CARP) : given an undirected network $G(N,E,C)$ with arc demands $q_{ij} \geq 0$ for each arc (i,j) which must be satisfied by one of a fleet of vehicles of capacity W , find a number of cycles each of which passes through the depot (node 1) which satisfy demands at minimal total cost. The CARP can be formulated as follows (see Golden and Wong [7]) :

$$\min \sum_{i=1}^N \sum_{j=1}^N \sum_{k=1}^K c_{ij} x_{ijk} \quad (1)$$

$$\text{s.t.} \quad \sum_{j=1}^N x_{jik} - \sum_{j=1}^N x_{ijk} = 0 \quad \begin{array}{l} i = 1, \dots, N \\ k = 1, \dots, K \end{array} \quad (2)$$

$$\sum_{k=1}^K (l_{ijk} + l_{jik}) = \left\lceil q_{ij}/W \right\rceil \quad (i,j) \in E \quad (3)$$

$$x_{ijk} \geq l_{ijk} \quad (i,j) \in E, k = 1, \dots, K \quad (4)$$

$$\sum_{i=1}^N \sum_{j=1}^N l_{ijk} q_{ij} \leq W \quad k = 1, \dots, K \quad (5)$$

$$\sum_{i \in Q} \sum_{j \in Q} x_{ijk} - N^2 y_{1qk} \leq |Q| - 1 \quad (6)$$

$$\sum_{i \in Q} \sum_{j \notin Q} x_{ijk} + y_{2qk} \geq 1$$

$$y_{1qk} + y_{2qk} \leq 1$$

$$y_{1qk}, y_{2qk}, x_{ijk}, l_{ijk} \in \{0,1\}$$

$k = 1, \dots, K$
 $q = 1, \dots, 2^{N-1} - 1$
 and every nonempty
 subset Q of $\{2, 3, \dots, N\}$

(7)

where N = the number of nodes,
 K = the number of available vehicles,
 q_{ij} = the demand on arc (i,j) ,
 W = the vehicle capacity ($W \geq \max q_{ij}$),
 c_{ij} = the length of arc (i,j) ,
 $x_{ijk} = 1$, if arc (i,j) is traversed by vehicle k , 0 otherwise,
 $l_{ijk} = 1$, if vehicle k services arc (i,j) , 0 otherwise,
 $\lceil z \rceil$ = the smallest integer greater than or equal to z ,
 Q = subset of N with cardinality between 2 and $|N|$.

The objective function (1) seeks to minimize total distance travelled. Equations (2) ensure route continuity. Equations (3) state that each arc with positive demand is serviced exactly once. Equations (4) guarantee that arc (i,j) can be serviced by vehicle k only if it covers arc (i,j) . Vehicle capacity is not violated on account of equations (5). Equations (6) prohibit the formation of (infeasible) subtours. Every index q corresponds to a set Q . Integrality restrictions are given in (7).

3. SELECTION OF SOLUTION METHODOLOGY.

Given the computational complexity of the CARP (which is an NP-hard problem (Golden and Wong [7])) it becomes necessary to apply approximate solution techniques or heuristics. Golden, De Armon and Baker [6] compare some heuristics developed to solve the CCPP (Capacitated Chinese Postman Problem). The CCPP is a special form of the CARP : $q_{ij} > 0$ for all i and j (instead of $q_{ij} \geq 0$). In general, 3 different approaches are available when dealing with multi-vehicle problems , i.e.,

- Cluster First-Route Second : the area is divided in subareas in such a way that every subarea can be serviced by one vehicle. For every subarea a Chinese Postman Problem (CPP) is solved.
- Route First-Cluster Second : first the problem is solved as a CPP. The complete cycle is then split in several routes which do not violate the capacity constraints of the vehicles and minimize the distance travelled.
- Route and Cluster Together : routing and clustering is done at the same time.

All three approaches have advantages and disadvantages. The first method creates smaller problems which are easy to solve and requires less data handling but yields suboptimal solutions and does not take the stochastic character of the demand into account. The second method has the advantage that the first phase remains the same even if demand varies. However a larger network has to be optimized and the relation between the tours and the depot (start node) and the different tours is gone (remember that the distance to the depot contributes the most to the total travelling distance). This method works rather well when every tour consists of only a few arcs [9]. This is not the case in this study. The third approach is well adapted to problems where a vast physical area contains a very large number of arcs. It allows for a good overall and integrated view of the problem. This third approach is chosen and three algorithms in that class are briefly discussed :

- Path Scanning [6] : the basic idea is to construct a tour at a time by adding arcs sequentially till the capacity is exhausted. Then the shortest return path to the depot is followed. The criteria for choosing an arc (i,j) are : the distance, c_{ij} , per unit remaining demand is minimized (a) or maximized (b), the distance from node j back to the depot is minimized (c) or maximized (d), if the vehicle is less than half-full maximize the distance (e); (a) looks for a large and quick payoff while (b) seeks to incur the larger expenses early; (c) tends to obtain shorter cycles; whereas (d) in general, yields longer cycles; (e) represents a hybrid approach. The set of cycles with smallest total distance is selected as outcome for this simple and rather fast algorithm.

- Construct and Strike [4, 6] : this algorithm repeatedly constructs feasible cycles and then strikes or removes them. The following steps are required :
 - step 1 : construct a cycle fulfilling the capacity constraint of the vehicle.
 - step 2 : set the demand of the arcs belonging to this cycle equal to 0.
 - step 3 : repeat steps 1 and 2 till there are no arcs, to be serviced, left in the network.

This procedure has one drawback : how to construct good cycles ?

- Augment - Merge [3, 7] : this procedure merges several smaller cycles in larger ones :
 - step 1 : start with as much cycles as arcs to be served.
 - step 2 : starting with the largest cycle available, see if a demand arc on a smaller cycle can be serviced on a larger cycle (= AUGMENT).
 - step 3 : evaluate the merging of any two cycles, subject to capacity constraints or additional restrictions. Merge the two cycles which yield the largest positive savings (= MERGE).
 - step 4 : repeat step 3 until finished.

Computational experience done by Golden, De Armon and Baker [6] indicates that the augment-merge strategy yields the best results on average, but is much more complicated and time consuming than the two other types. The path scanning algorithm has the advantages of simplicity and minor CPU-time requirements. This algorithm can easily be adapted to all the necessary needs and criteria as required by the problem stated above. The construct and strike algorithm has the same advantages but yields results which are worse.

Since management needs an interactive decision tool which provides quick answers to urgent problems, an approach based on the path scanning algorithm is chosen. Moreover, a number of different criteria for choosing the next arc in every step of the algorithm exist (they are mentioned earlier). By simulating some scenarios it is found that the following combination of two criteria leads to the best results.

Given a provisional path (start node = 1) arrives in node i, we add the arc (i,j) satisfying one of the following criteria :

- if the truck load is less than 50 %, maximize the shortest distance from j to the end-node, distance of arc (i,j) included.
- if the truck load is 50 % or more, minimize the shortest distance from j to the end-node, distance of arc (i,j) included.

This combination leads to cycles with a reasonable length without augmenting markedly the distance from the last collection point back to the start or end (waste disposal) node.

4. IMPLEMENTATION ISSUES.

4.1 CLUSTERING.

Although the path scanning algorithm is a route and cluster together method, we will start dividing the total collection area in a number of sectors. These sectors are much larger than the clusters in the cluster first - route second approach. For a particular day, all lorries will be sent to the same sector. Our major concern is to minimize the number of trips from and to the waste disposal node. In this way a lorry can take over a part of the tour of another one without driving a long distance (interaction between tours). Demand is stochastic and a vehicle can reach its maximum allowable load before finishing its cycle. Adding up the collection amounts to obtain the total daily collection quantity per sector, levels out a large number of these stochastic elements. Nevertheless some time-dependent parameters remain. There is a seasonal dependency because people produce more garbage in summer time than in winter. Also on the long term, the total amount of garbage to collect yearly increases. To start, we will define the sectors by aggregating the present sectors.

4.2 INPUT DATA.

It was an enormous task to collect all input data necessary for solving this problem. And as the accuracy of these data will influence to a large extent the accuracy of the solution, it is of major concern to keep these inputs up-to-date. So a well developed database structure had to be set up. The input data can be divided into three classes:

- data related to the network :
 - an adjacency list for every node
 - cost c_{ij} for every arc, expressed as:
 - 1. a collection cost (time)
 - 2. a driving cost (time to drive through the arc without collection)
 - quantity garbage q_{ij} to be collected for arc (i,j)
 - one-way streets or not
 - both sides of the streets can or can not be served simultaneously
 - collection arcs and arcs connecting one node to another without collecting garbage
 - start and end node for every sector.

- general data :
 - number of lorries available for the different types of garbage
 - maximum capacity (load) for every lorry and an average lorry capacity W as used in the program
 - as the density of garbage is not always the same, not weight but volume can be the binding constraint. We do not only need an average W but also a margin on W , to express the variability
 - seasonal coefficient on quantity of garbage produced.
- a number of parameters tested in the different scenarios to evaluate their sensitivity on the result.

The most important data in this model are the costs c_{ij} and the quantities q_{ij} related to every arc. c_{ij} can be expressed as the time necessary to serve an arc, with or without collection. The time can be derived by measuring the distance (l_{ij}) and multiplying it by the average velocity of a lorry. Distances of the arc were measured on detailed maps of the area and corrected with information from the drivers who noted down distance travelled. There are 2 velocities:

$$v_d = \text{velocity with collection} \quad \rightarrow \quad t_{dij} = l_{ij} / v_d$$

$$v_b = \text{velocity without collection} \quad \rightarrow \quad t_{bij} = l_{ij} / v_b$$

In the article of Bodin et al. [2] the velocity v_b is even divided into 4 classes depending on the types of streets.

The quantity q_{ij} is the most difficult element in gathering the data, because q_{ij} has a stochastic nature and a direct measurement (estimation) of q_{ij} is impossible for practical reasons.

q_{ij} was estimated in the following way. From historical data, total weekly collection volumes V per sector S will be determined.

$$q_{ij} = f_{ij} * V \text{ with } 0 \leq f_{ij} \leq 1,$$

($f_{ij} = 0$ if arc (i,j) belongs to the sector S and has to be serviced by another sector S' or arc (i,j) does not belong to the sector S)

When garbage is collected in mini-containers, f_{ij} can be easily estimated as follows :

$$f_{ij} = \frac{\text{\# containers for arc } (i,j)}{\text{\# containers for sector } S \text{ (arc } (i,j) \in S)}$$

It is clear that the costs c_{ij} can be estimated more accurately than the collection quantities q_{ij} per arc. Nevertheless, with regard to the accuracy of the data, following equation is very important to keep in mind :

$$e_r(q) = e_r(q_{ij}) * n^{-1/2} \quad \text{with: } q = \text{collection quantity per trip}$$

$$n = \text{number of arcs in one trip}$$

$$e_r(X) = \text{relative error on X}$$

Analogously we have:

$$e_r(l) = e_r(l_{ij}) * n^{-1/2} \quad \text{with: } l = \text{driving distance per trip}$$

$$l_{ij} = \text{driving distance per arc}$$

Example:

$$l = 20.000 \text{ m}$$

$$l_{ij} = 200 \text{ m} \quad \text{-----> } e_r(l) = e_r(l_{ij}) * 10$$

$$n = 1/l_{ij} = 100$$

This means that when a relative error of 5 % on l with a confidence interval of 95 % is required, a relative error on l_{ij} up to 50 % is allowed. As a consequence the way we estimated both distances and collection quantities per arc leads to rather accurate results when aggregated over a collection trip.

5. COMPUTATIONAL ASPECTS.

The software program can be separated into two parts:

- a general database program (in dBase III), to store and update all input data.
- a route generation program written in Pascal.

5.1 GENERAL DATABASE PROGRAM.

Every arc is defined by two nodes, the start- and end-node. The following information is given per arc :

- two costs : a collection cost and an empty cost (cost for driving through an arc without collecting garbage)
- capacity q_{ij} : quantity of garbage to collect for arc (i,j)
- type : 1 = directed arc
- 2 = undirected arc

- ten fields S1, S2 S10, corresponding to ten sectors with each field containing one of the following codes :
 - 1 = a collection arc for the corresponding sector.
 - 2 = arc belongs to the sector, but no collection is necessary.
 - 3 = arc does not belong to the sector.

This is a general database, containing all the possible streets in the collection problem. From this data base we can easily select the arcs belonging to one or more sectors (1 ... 10), these arcs will be copied into a sector data base and used as input for the route generation program. Whenever there is a change in the street pattern or collection quantities, the general data base will be updated and new partial data bases per sector will be derived. Organizing the input structure in this way will never lead to data inconsistency.

5.2 THE ROUTE GENERATION PROGRAM.

The program is built in a modular way. The first module reads all the network data out of the sector data base and doubles the undirected arcs. The second module reads the general data related to the sector (start-node, end-node, collection quantity) and other data like number of lorries, their capacity, etc...

After the user has chosen an optimization criterion, one arc at a time will be selected in module 3 to generate the different trips. In every node the path scanning algorithm is used when one of the related arcs is not yet served. The Dijkstra [5] algorithm is used in all other cases, to find the shortest way to another unserved arc.

When lorry capacity is reached, the program also uses the Dijkstra algorithm to find the shortest way back to the end-point.

A new trip will start in the end- or start-node of the network and module 3 will again generate a new routing. New trips will be generated as long as unserved arcs in the network exist. After every trip, data with regard to collection quantity q and length of the trip (l) are stored.

5.3 FLEXIBILITY OF THE PROGRAM.

The program is built in such a way that it is very easy to generate a number of scenarios. Following options are possible :

- changing the area and shape of the sectors

- choosing a start- and end-point of a trip in one node or in two different nodes. In this way a lorry can start its first trip at the depot and end it at the waste collection area (from where the second trip starts)
- different criteria to execute the path scanning algorithm
- introduction of a seasonal coefficient to adapt collection quantities per arc over the year
- a lorry, close to the end-point of the network has two options. It can start a new cycle, or it can drive directly to the waste collection area. A preference for one of these options can be implemented by adding a parameter α . Instead of having a lorry capacity W , we will give the truck in the end-point a capacity αW . If $q < \alpha W$, the truck will start a new cycle. If $q > \alpha W$, the truck will go directly to the collection area.
By giving α a number between 0 and 1 we can influence thoroughly the length and the number of trips
- simultaneous garbage collection on both sides of the street
- blind alleys are automatically added to an adjacent arc. In this way the network is simplified and empty driving time is minimized.

Elements not included in the program are : time constraints, holidays, rough garbage, traffic jams, right turns.

6. CONCLUSIONS

A combination of six current sectors was selected as a testing area. These sectors give a good representation of the whole collection area. Total weekly collection in the testing area is approximately 5000 tons.

6.1 INFLUENCE OF LORRY CAPACITY (W).

Initially W was fixed on 9.330 Kg. The total driving distance decreases by increasing W . From fig.1 we see that the decrease in driving distance is not continuous. The explanation is very simple. Whenever a lorry A reaches its capacity constraint W , it drives directly to the waste disposal area. Another lorry B (operating in the neighbourhood) will continue and serve the remaining collection quantity of the arc (i,j) . Increasing W will only decrease total driving distance when arc (i,j) can be collected completely by lorry A.

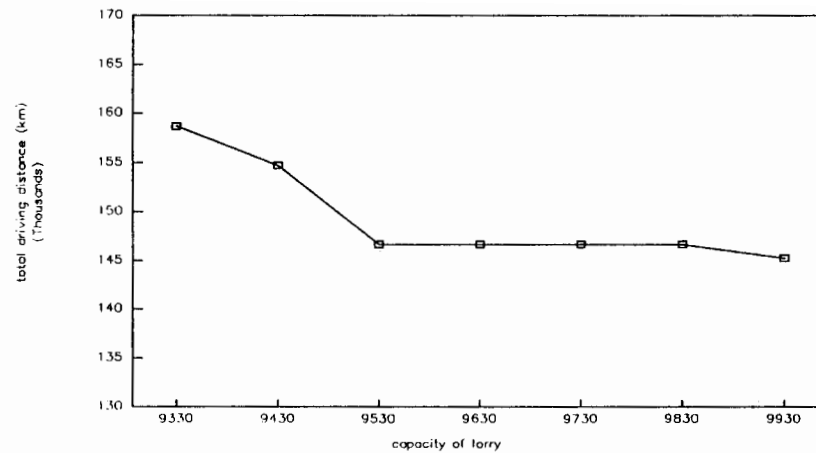


Fig.1 : Sensitivity of the lorry capacity on the total driving distance.

6.2 OPTIMAL CRITERION FOR OPEN AND CLOSED TRIPS.

A closed trip starts and ends in the same point of the network whereas, an open trip has a different start- and end-point. Intuitively it is clear that open trips give better results than closed trips for all possible criteria. However for combination tours, open as well as closed trips, using one and the same criterion for both types of trips gives better results than using two different criteria. This was a rather remarkable result we found from the simulations of the testing area.

6.3 INFLUENCE OF PARAMETER α .

As mentioned before, parameter α is of major influence on the number of trips and on the collection distance. The smaller α , the larger the number of trips (trips are shorter). It is more important to reduce the number of trips than to reduce the total collection distance, as the waste collection point is situated a long way from the collection area. Therefore parameter α close to 1 yields the best results in this case.

6.4 COMPARISON WITH THE CURRENT SYSTEM.

Using the aggregated collection area of the six current sectors, an overall reduction in driving distance of 15% was found. Due to a

reduction in the number of trips (10 trips/week against minimum 12 trips/week in the current system) total driving distance to and from the collection area was reduced by 14.5%. Another 4.5% was saved on collection distance. For the individual sectors there was no clear improvement. Whereas the Path Scanning Algorithm was developed for a CCPP problem, the scenario of individual sectors is more related to a CPP problem as one trip is mostly sufficient for collecting all waste in a sector. Aggregating sectors gives a better overall result as it reduces the number of partly loaded trucks which has an influence on the total number of trips.

References

- [1] Beltrami, E.J., Bodin, L.D., Networks and vehicle routing for municipal waste collection, Networks, 4 (1974), pp 65-94.
- [2] Bodin, L., Fagin, G., Welebny, R., Greenberg, J., The design of a computerized sanitation vehicle routing and scheduling system for the town of OYSTER BAY, New York, Computer & Operations Research, 16 (1989), pp 45-54.
- [3] Bodin, L., Golden, B., Assad, A., Ball, M., Routing and scheduling of vehicles and crews: the state-of-the-art, Computer & Operations Research, 10 (1983), pp 63-211.
- [4] Christofides, N., The optimal traversal of a graph, Omega, 1.6 (1973), pp 719-732.
- [5] Dijkstra, E., A note on two problems in connection with graphs, Numerische Mathematik, 1 (1959), pp 269-271.
- [6] Golden, B., De Armon, J., Baker, E., Computational experiment with algorithms for a class of routing problems, Computer & Operations Research, 10.1 (1983), pp 47-59.
- [7] Golden, B., Wong, R., Capacitated arc routing problems, Networks, 11 (1981), pp. 305-315.
- [8] Minieka, E., Optimisation algorithms for network and graphs, Marcel Dekker, New York, 1978.
- [9] Stern, H., Dror, M., Routing electric meter readers, Computer & Operations Research, 6 (1979), pp 209-223.
- [10] de Favereau de Jeneret, H., Oreins, R., Optimalisatie van de huisvuilophaling, Eindwerk Industrieel Beleid, K.U.Leuven, 1989.

THE STATIONARITY OF CAPM-BETA IN A CHANGING ECONOMIC ENVIRONMENT

Winfried G. HALLERBACH

Department of Finance
Erasmus University Rotterdam
P.O. Box 1738, NL-3000 DR Rotterdam, The Netherlands

ABSTRACT

In this paper, we specify a theoretical multi-factor return generating process of securities by means of which fundamental shifts in the structure of the economy can be modelled. We then analyze the conditions under which CAPM's risk measure 'beta' will be constant in this changing economy. These conditions are very restrictive and not likely to be met in practice.

1. INTRODUCTION

The concept of systematic risk is of central importance in Modern Portfolio Theory (MPT). Within the standard Capital Asset Pricing Model (Sharpe [1964], Lintner [1965], Mossin [1966]), the measure for this market risk is defined as:

$$\beta_i = \frac{\text{Cov}(r_i, r_m)}{\text{Var}(r_m)} \quad (1)$$

where r_i = the return on security i ;

r_m = the return on the market portfolio m .

Cov and Var denote the covariance and variance operator, respectively.

In applications of MPT, historical data are used to estimate the relevant *ex ante* β by means of the market model (cf. Fama [1976]). In its most simple form, this regression model can be written as:

$$r_{it} = \alpha_i + \beta_i r_{mt} + \varepsilon_{it}, \quad t \in T, i \in N \quad (2)$$

where α_i is a constant and ε_{it} is a zero-mean error term, uncorrelated with r_{mt} . T and N denote the sample of time periods and the set of N securities in the market portfolio, respectively. For the sake of brevity, the time index t will be suppressed.

The use of this market model involves a number of more or less arbitrary decisions, for example the choices with respect to:

- the particular form of the model: the simple form eq. (2), the excess return (or risk premium) form or the two factor model (cf. Blume & Friend [1973]);
- the index that serves as proxy for the market portfolio (for the effect of index choice, see for example Frankfurter [1976] and Roll [1977]);
- the total interval T chosen to perform the regression (cf. Baesel [1974], Levhari & Levy [1977] and Theobald [1981]);
- the appropriate length of the observation intervals within the sample of time periods T (cf. Hawawini [1983] and Handa, Kothari & Wasley

[1989]).

The sensitivity of the estimated β for the latter intervallling aspect determines its stability. Especially when daily data are used, nonsynchronous trading gives rise to an errors-in-variables problem and causes biased estimators (Scholes & Williams [1977], Smith [1978], Scott & Brown [1980], Dimson & Marsh [1983] and more recently Shanken [1987]).

In addition, as the *ex ante* value of β is relevant, it is important whether β is constant over time. This brings us to the issue of the intertemporal stationarity¹⁾ of β : is a β estimated from data in period T_1 equal to the β estimated in period T_2 ?

Blume [1971, 1975] found that β 's were non-stationary. More specific, he found that, over time, β 's tend to drift to the market average of one. Klemkosky & Martin [1975] examined various techniques by means of which β 's could be adjusted for this drift. Elgers et al. [1979] however, showed that whether β 's were calculated moving forward or backward in time, they similarly drifted towards the value one. This β -drift thus appears to be caused by statistical aberrations.

Fabozzi & Clark [1978], Sunder [1980], Chen [1981] and Simonds et al. [1986], among others, test a random coefficient model and also present evidence of non-stationarity. Fisher & Kamin [1985] suggest some improvements in estimating β 's by gradually decreasing the weight of each observation with the passage of time. To determine these weights, they use a Kalman filter.

Other studies only focus on part of the β -coefficient in eq. (1). Myers [1973] finds significant changes in the variability of the 'market factor' from period to period and concludes that, as a result, β 's are not stationary. Elton et al. [1978] focus on the numerator and investigate the stationarity of the correlation structure of security returns. They conclude that the best way to estimate securities' correlations is just to use the average correlation coefficient for the entire universe of securities ('overall mean'): additional efforts to refine the estimated correlations appeared to be fruitless.

1) Note that we explicitly make a difference between *stability* and *stationarity*.

The above mentioned studies are of an empirical nature and concentrate on the form and statistical properties of the non-stationarity. On the other hand, there are studies of a theoretical nature that concentrate on the identification of micro- or macroeconomic variables that influence β -stationarity. Officer [1973] investigated changes in the variability of the market factor (just as Myers [1973] did) and found that these changes appear to be related to changes in the Index of Industrial Production (a measure of business activity in the US) and changes in the M2 money supply. Levy [1971] related shifts in β to alternating bull and bear market conditions. Fabozzi & Francis [1977] reject the influence of bull-bear market forces on β , but conclude later that the intertemporal non-stationarity of β appears to result from changes which are associated with business cycle economics (Francis & Fabozzi [1979]). Using the concept of duration, Bildersee & Roberts [1981] relate changes in β to changing interest rates. More recently, DeJong & Collins [1985] start from the joint Option Pricing Model/ Capital Asset Pricing Model framework and find a significant relation between β -non-stationarity and the degree of leverage of the respective firm and changes in the risk-free interest rate.

McDonald [1985] considers major structural changes in the economy. These changes were accompanied by large changes in interest rates and yielded shifts in β 's. Further evidence concerning these changes in the structure of the economy and the variance-covariance structure of security returns is presented by Bollerslev et al. [1988], Harvey [1989] and Van Der Meulen [1987, 1989]. Analyzing the covariance structure of deflated returns of general classes of assets (including currencies), Van Der Meulen detects several large shifts in the covariances between these assets over time. As these changes affect the composition of the assets' covariance matrix, the β 's of these asset classes (and of the individual assets therein) presumably will not be constant.

In this paper, we accept the strong empirical evidence supporting the existence of β -non-stationarity, and consider the theoretical effect of fundamental shifts in the structure of the economy on β . We hypothesize that the returns of assets or securities are generated by a

factor model and that intertemporal changes in the variances of the factors provide the 'missing link' between the shifts in the covariance structure of these assets. In section two, we specify the theoretical multi-factor return generating process of securities by means of which the fundamental changes in the structure of the economy can be modeled. The conditions under which security- β 's will be constant in a changing economy form the central issue of this paper and will be dealt with in section three. Section four contains our conclusions and directions for future research.

2. MULTI-FACTOR MODELS AS RETURN GENERATING PROCESSES

To model the dynamic economic environment in which the returns of the securities are generated, we reduce the dimensionality of the interrelationships therein. We therefore introduce a mechanism that generates the return for each of the N securities in the market portfolio. More specific, we assume that there exists a function $g_i(\cdot)$ that links the individual returns r_i to a set of K common underlying variables or factors $(\delta_j)_{j \in K}$. In a general form, this return generating process (henceforth: RGP) can be expressed as:

$$r_{it} = g_i(\delta_{jt})_{j \in K} + \epsilon_{it}, \text{ with } E(\epsilon_{it}) = 0, i \in N \quad (3)$$

To incorporate randomness, the common factors are assumed to be stochastic and a random disturbance term ϵ_i is added to incorporate the influence of idiosyncratic factors on each individual return. We assume that the idiosyncratic factors ϵ_i and the common factors $(\delta_j)_{j \in K}$ are mutually independent. Furthermore, the idiosyncratic factors are assumed to be truly security specific and thus also mutually independent. With these assumptions, the dimensionality of the economic environment is reduced from $\frac{1}{2}N(N-1)$ relationships between the returns of securities to (i) the $\frac{1}{2}K(K-1)$ relationships between the common factors and (ii) the $N \times K$ relations between the securities and these factors. The latter relations can be explicitized by means of sensitivity coefficients.

Although the function $g_i(\cdot)$ is unknown, we can apply the multivariate version of Taylor's theorem and expand $g_i(\cdot)$, for example, around the value it takes when the common factors are set to their mean values. We then can rewrite eq. (3) as

$$r_i = g_i[E(\delta_j)] + \sum_j [\delta_j - E(\delta_j)] \frac{\partial g_i}{\partial \delta_j} + \sum_{h=2}^{\infty} \frac{1}{h!} \left[\sum_j [\delta_j - E(\delta_j)] \frac{\partial}{\partial \delta_j} \right]^h g_i + \epsilon_i \quad (4)$$

where the symbolic power between the large square brackets on the right is first to be expanded formally by the binomial theorem and then the powers $\partial/\partial \delta_j$ multiplied by g_i are to be replaced by the corresponding n^{th} derivatives $\partial^n g_i / \partial \delta_1^n$, $\partial^n g_i / (\partial \delta_1^{n-1} \partial \delta_2)$ &c., evaluated at the spanning point $[E(\delta_1), E(\delta_2), \dots]$.

Note that $g_i[E(\delta_j)]$ is not necessarily equal to $E(r_i)$ when g_i is a non-linear function of the δ_j 's (cf. Jensen's inequality for the univariate case). For simplicity, we can assume that terms of second and higher order are small and be neglected. Eq. (4) then reduces to:

$$r_i = g_i[E(\delta_j)] + \sum_j [\delta_j - E(\delta_j)] \frac{\partial g_i}{\partial \delta_j} + \epsilon_i \quad (5)$$

A stronger assumption involves the (local) linearity of g_i in the δ_j 's. The sensitivity coefficients of the returns for the common factors are then constants b_{ij} (for a specific range of δ_j):

$$\frac{\partial r_i}{\partial \delta_j} = \frac{\partial g_i}{\partial \delta_j} = b_{ij}, \quad \forall i \in N, j \in K \quad (6)$$

Under these assumptions, the RGP eq. (3) reduces to:

$$r_i = E(r_i) + \sum_j [\delta_j - E(\delta_j)] b_{ij} + \epsilon_i, \quad i \in N \quad (7)$$

If the identities of the factors are known, this linear relation can be estimated by means of OLS regression. As easily can be seen from eqs. (5) and (7), the OLS estimation does not provide a Taylor series approximation. The argument is that the choice of the point of expansion is arbitrary (see eq. (4)).

White [1980] and Van Praag [1981], however, note that the OLS estimators (b_{ij}) provide the best linear approximation of the dependent random variable (here: r_i) by the independent random variables (here: the δ_j 's): the b_{ij} 's result from minimizing the mean square error of the approximation. For the least squares approximation, White assumes that there exists a functional model, although that may be unknown. Van Praag considers this regression approximation as 'model free', so that we even can interpret the sensitivities b_{ij} without assuming any underlying functional model at all.

If there should exist a return generating model $g(\delta_j)_{j \in K}$ that is non-linear, it is important to note that problems can arise. As can be inferred from eqs. (5) and (7), the crucial point then is that the response coefficients b_{ij} are only correct for small deviations of the factors. If the variability of the factors changes, this may imply that the sensitivities for these factors also change. In other words, the functional form eq. (7) then changes and does not apply any longer for different time periods (i.e. different levels of the factors). In this paper, however, we'll treat the factor sensitivities as constant technical coefficients; we'll model structural shifts in the economy as changes in the variances of the factors.

3. THE STATIONARITY OF β WHEN RETURNS FOLLOW A LINEAR MULTI-FACTOR RGP

In this section, we investigate the implications of changing factor variances in a K-factor RGP for the β -coefficients of the securities. We start from the following assumptions:

(A1) For all securities, the RGP is linear in its parameters:

$$r_i = E(r_i) + \sum_j b_{ij} f_j + \epsilon_i, \quad \forall i \in N, \quad (8)$$

where f_j denotes $\delta_j - E(\delta_j)$.

(A2) The common factors are mutually uncorrelated and uncorrelated with the idiosyncratic factors:

$$\text{Cov}(f_j, f_{j'}) = \text{Cov}(f_j, \epsilon_i) = 0, \quad \forall j \neq j' \quad (9)$$

(A3) The idiosyncratic factors are truly security specific, so

$$\text{Cov}(\epsilon_i, \epsilon_{i'}) = 0, \quad \forall i \neq i'. \quad (10)$$

As the sensitivities b_{ij} link the (variability of the) security returns to the (variability of the) factors, we can say that these coefficients measure the *factor risk* of the securities. The variability of the error term ϵ_i induces the *idiosyncratic risk*.

The CAPM- β links the (variability of the) security returns to the (variability of the) return on the market portfolio, so β measures the *market risk* of the securities. As the market portfolio m is a convex combination of the N securities, its RGP can be expressed as:

$$r_m = \sum_i m_i r_i = E(r_m) + \sum_j b_{mj} f_j + \epsilon_m \quad (11)$$

where m_i is the proportion of the total market value of security i relative to the market value of the market portfolio ($\sum_i m_i = 1$).

Given the RGP with its assumptions, the covariances between the securities' returns and the return on the market portfolio can be expressed as:

$$\text{Cov}(r_i, r_m) = \sum_j b_{ij} b_{mj} \text{Var}(f_j) + m_i \text{Var}(\epsilon_i), \quad \forall i \in N \quad (12)$$

The variance of the return on the market portfolio equals:

$$\text{Var}(r_m) = \sum_j b_{mj}^2 \text{Var}(f_j) + \text{Var}(\epsilon_m) \quad (13)$$

Using eqs (12) and (13) in (1), we can rewrite β_i in terms of factor variances:

$$\beta_i = \frac{\sum_j b_{ij} b_{mj} \text{Var}(f_j) + m_i \text{Var}(\epsilon_i)}{\sum_j b_{mj}^2 \text{Var}(f_j) + \text{Var}(\epsilon_m)}, \quad \forall i \in N \quad (14)$$

The central issue is how β behaves if the economic environment (in which the returns of the securities are generated) changes. For this purpose, we additionally assume that:

(A4) fundamental shifts in the structure of the economic environment arise because the variance of one or more common return generating factors changes.

Using eqs. (12) and (13), the changes in the relevant variance-covariance structure of the security returns can then be modelled as:

$$d\text{Cov}(r_i, r_m) = \sum_j b_{ij} b_{mj} d\text{Var}(f_j), \quad \forall i \in N \quad (15)$$

and

$$d\text{Var}(r_m) = \sum_j b_{mj}^2 d\text{Var}(f_j) \quad (16)$$

Given eqs. (15) and (16), the question arises how β changes when the factor variances undergo small changes. If the changes in the numerator of β in a way offset the changes in the denominator, β will not change. But under what conditions will β indeed be stationary?

To answer these questions, we have to analyze the change in β as a result of a marginal change in a factor's variance. From eq. (14) it follows that:

$$\frac{\partial \beta_i}{\partial \text{Var}(f_j)} = \frac{b_{ij} b_{mj} \text{Var}(r_m) - \text{Cov}(r_i, r_m) b_{mj}^2}{[\text{Var}(r_m)]^2}, \quad \forall i \in N, j \in K \quad (17)$$

3.1. Single-factor RGP

We first consider the case in which the RGP only consists of one factor, factor k , and where the variance of this factor exhibits a marginal change.

THEOREM 1: *A security's β will only be stationary under a marginal structural shift in the one-factor economy if this β is equal to the security's factor sensitivity b_{ik} scaled to the of the market portfolio's factor sensitivity b_{mk} .*

Proof: It will be clear that β_i is invariant under a marginal change in $\text{Var}(f_k)$ if:

$$\frac{\partial \beta_i}{\partial \text{Var}(f_k)} = 0 \quad (18)$$

Combining eqs. (17) and (18) and using the definition of β_i yields the restriction

$$\beta_i = \frac{b_{ik}}{b_{mk}}, \quad (19)$$

provided that the market portfolio has a non-zero sensitivity for the factor k . In a one-factor RGP, this is of course very likely. ■

The next two corollaries give insight in the conditions under which eq. (19) is satisfied.

COROLLARY 1: *If the market portfolio is an equally weighted portfolio of a finite number of N securities, the β -stationarity condition (19) requires the equality of the ratio of the security's factor risk to the security's idiosyncratic risk and the ratio of average factor risk to average idiosyncratic risk.*

Proof: Substituting eq. (14) in (19), we get

$$\frac{b_{ik}b_{mk}\text{Var}(f_k) + m_i\text{Var}(\epsilon_i)}{b_{mk}^2\text{Var}(f_k) + \text{Var}(\epsilon_m)} = \frac{b_{ik}}{b_{mk}} \quad (20)$$

or

$$\frac{b_{ik}}{b_{mk}} = m_i \frac{\text{Var}(\epsilon_i)}{\text{Var}(\epsilon_m)} \quad (21)$$

When the market portfolio is equally weighted ($m_i = 1/N$, $\forall i$), we have

$$\begin{aligned} \text{Var}(\epsilon_m) &= \text{Var}\left(\sum_{s \in m} m_s \epsilon_s\right) = \text{Var}\left(N^{-1} \sum_s \epsilon_s\right) = N^{-2} \sum_s \text{Var}(\epsilon_s) \\ &= \frac{1}{N} \overline{\text{Var}(\epsilon)} \end{aligned} \quad (22)$$

where the bar over $\text{Var}(\cdot)$ denotes its average value.

Using this result, we can rewrite eq. (21) as

$$\frac{b_{ik}}{\text{Var}(\epsilon_i)} = \frac{b_{mk}}{\overline{\text{Var}(\epsilon)}} \quad (23) \quad \blacksquare$$

For a typical security, the idiosyncratic variance of its return equals the average of all securities' idiosyncratic variances. For its β to be stationary, its sensitivity to factor k must equal the market average sensitivity b_{mk} to this factor. Consequently (according to eq. (19)), its β must equal one. For a security with a larger (smaller) than typical idiosyncratic variance, b_{ik} must be larger (smaller) than b_{mk} , which implies that its β must be larger (smaller) than one.

COROLLARY 2: *If the market portfolio consists of a (countably) infinite number of securities, each with an infinitesimally small weight, then all β 's will be stationary in the one-factor economy.*

Proof: If the number of securities in the market portfolio increases more and more, it follows from eq. (22) that the market portfolio

becomes well-diversified with respect to the return generating factor: $\text{Var}(\epsilon_m) \rightarrow 0$ as $N \rightarrow \infty$. Without variability there can exist no covariability, so $\text{Cov}(\epsilon_i, \epsilon_m) = (1/N)\text{Var}(\epsilon_i) \rightarrow 0$ likewise. As a result of this naive diversification process, we can rewrite eq. (20) as

$$\frac{b_{ik}b_{mk}\text{Var}(f_k)}{b_{mk}^2\text{Var}(f_k)} = \frac{b_{ik}}{b_{mk}} \quad (24)$$

Hence, the stationarity condition (19) is satisfied by any value of b_{ik} and b_{mk} . ■

Note that the β -stationarity conditions, as stated in the corollaries above, are independent of $\text{Var}(f_k)$. Consequently, these conditions do not only apply for marginal changes in the factor's variance, but also for any arbitrary large changes.

3.2. Multi-factor RGP

Empirical evidence (as early as King [1966]) indicates that a single-factor RGP is an oversimplification of reality and that a multi-factor RGP is more appropriate. In this general case, the economic environment is allowed to change in more than one dimension, i.e. the variances of K factors are allowed to change ($K > 1$).

THEOREM 2: A security's β will only be stationary under any marginal structural shift in the multi-factor economy iff (i) the security's factor sensitivities $\{b_{ij}\}_{j \in K}$ are (pairwise) equal to the market portfolio's factor sensitivities $\{b_{mj}\}_{j \in K}$, and (ii) its β equals one.

Proof: The condition that a change in β as a result of marginal changes in the factor variances is zero can be expressed as:

$$d\beta_i = \sum_j \frac{\partial \beta_i}{\partial \text{Var}(f_j)} d\text{Var}(f_j) = 0 \quad (25)$$

Using eq. (17), this is equivalent to

$$\frac{1}{\text{Var}(r_m)} \sum_j [b_{ij}b_{mj} - \beta_i b_{mj}^2] d\text{Var}(f_j) = 0 \quad (26)$$

Solving for β_i , this yields

$$\beta_i = \frac{\sum_j b_{ij}b_{mj} d\text{Var}(f_j)}{\sum_j b_{mj}^2 d\text{Var}(f_j)} \quad (27)$$

It follows that β_i is only stationary under independent marginal changes in the variances of the return generating factors if

$$(i) \quad b_{ij} = b_{mj}, \quad \forall j \in K$$

According to eq. (27), we then also must have

$$(ii) \quad \beta_i = 1. \quad \blacksquare$$

The derivation of these conditions from eq. (27) does not depend on the magnitude of the factor variance changes. Consequently, these conditions apply for *any* arbitrary large changes in the factor variances.

Given eq. (14), the stationarity conditions (i) and (ii) imply

$$m_i \text{Var}(\epsilon_i) = \text{Var}(\epsilon_m) \quad (28)$$

If the market portfolio is an equally weighted portfolio of a finite number of securities, eq. (22) applies. Consequently, only *typical* securities can satisfy these stationarity conditions (cf. corollary 1). If, on the other hand, the market portfolio consists of a (countably) infinite number of securities, each with an infinitesimally small weight, then idiosyncratic risk is no longer relevant. Condition (ii) becomes redundant and the pairwise equality of the security's and market

portfolio's factor sensitivities then forms the necessary and sufficient stationarity condition.

Condition (i) places a severe restriction on the variance-covariance structure of the security returns. If this condition is satisfied, the covariances between the returns of all "intertemporal stationary β " securities would be identical. In a large, well-diversified capital market, these covariances would equal the variance of the return on the market portfolio. Also, the variances of the returns on these individual securities must then all be larger than the variance of the market portfolio's return (this follows from eq. (13): the market portfolio is well-diversified, but the security returns can contain an idiosyncratic component).

4. CONCLUSIONS

We hypothesize that fundamental shifts in the structure of the uncertain economic environment, in which the returns of the securities are generated, arise because the variance of one or more common underlying return generating factors changes. Using this return generating process, by means of which changes in the relevant covariance structure of the security returns are modelled, we analyzed the conditions under which a security's β -coefficient will be stationary.

In a single-factor context, all β 's will be stationary if the market portfolio is large, equally weighted and hence well-diversified with respect to the return generating factor. Should the market portfolio in contrast contain idiosyncratic risk, then for a security- β to be stationary there should exist a restrictive relationship between the security's factor risk and idiosyncratic risk. More specific, for that security the ratio factor risk to average (=market) factor risk must equal the ratio idiosyncratic risk to average idiosyncratic risk.

In a more realistic multi-factor context, security β 's will only be stationary if the large, equally weighted market portfolio is well-diversified with respect to the factors and the factor risks of the securities are pairwise equal to the factor risks of the market portfolio. This implies a severe restriction on the variance-covariance structure

of security returns, which is not likely to be satisfied in reality. In case the market portfolio should contain idiosyncratic risk, we have the additional stationarity condition that β 's must equal one. Hence, only securities with average factor risks and average idiosyncratic risk would exhibit stationary β 's.

The relevance of results is two-fold. From an *ex ante* point of view, the ability to predict changes in estimated β 's could be enhanced by empirically identifying the factors that generate the returns and thus contribute to β -nonstationarity. The ability to identify factors contributing to intertemporal changes in β also has potential value in an *ex post* context. Event studies rely on constant- β market models to estimate residual returns. If β changes over time, an error is made in estimating the residuals. It would be more correct to adjust the residuals for the changes in β . In the evaluation of investment performance, errors may result when relying on constant- β market models; adjustments of β according to changes in the economic environment would be adequate. Although empirical research has yielded some candidates for the factors (cf. Chen, Roll & Ross [1986]), it would be worthwhile to investigate the intertemporal behavior of their variability and link these results to the behavior of β .

ACKNOWLEDGEMENTS

The author wishes to thank Jaap Spronk, Jan Van Der Meulen and two anonymous referees of the journal for their critical comments on an earlier version of the paper. Of course, any remaining errors are the responsibility of the author.

REFERENCES:

- Baesel, J., 1974, On the Assessment of Risk: Some Further Considerations, *The Journal of Finance* December, pp. 1491-1494
- Bildersee, J.S. & G.S. Roberts, 1981, Beta Instability When Interest Rate Levels Change, *Journal of Financial and Quantitative Analysis* 16/3 September, pp. 375-380
- Blume, M.E. & I. Friend, 1973, A New Look at the Capital Asset Pricing Model, *The Journal of Finance* 28/1 March, pp. 19-33
- Chen, S.N., 1981, Beta Non-Stationarity, Portfolio Residual Risk and Diversification, *Journal of Financial and Quantitative Analysis* 16/1, pp. 95-111
- Chen, N.F., R. Roll & S.A. Ross, 1986, Economic Forces and the Stock Market, *The Journal of Business*, pp. 383-403
- DeJong, D.V. & D.W. Collins, 1985, Explanations for the Instability of Equity Beta: Risk-Free Rate Changes and Leverage Effects, *Journal of Financial and Quantitative Analysis* 20/1 March, pp. 73-94
- Dimson, E. & P.R. Marsh, 1983, The Stability of UK Risk Measures and The Problem of Thin Trading, *The Journal of Finance* 38/3 June, pp. 753-783
- Elgers, P.T., J.R. Haltiner & W.H. Hawthorne, 1979, Beta Regression Tendencies: Statistical and Real Causes, *The Journal of Finance* 34 March, pp. 261-263
- Elton, E.J., M.J. Gruber & T.J. Urich, 1978, Are Betas Best?, *The Journal of Finance* 33 December, pp. 1375-1384
- Fabozzi, F.J. & F. Clark, 1978, Beta as a Random Coefficient, *Journal of Financial and Quantitative Analysis* 13/1, pp. 101-116
- Fabozzi, F.J. & J.C. Francis, 1977, Stability Tests for Alphas and Betas Over Bull and Bear Market Conditions, *The Journal of Finance* September, pp. 1093-1099
- Fama, E.F., 1976, *Foundations of Finance*, Basic Books, New York
- Fisher, L. & J.H. Kamin, 1985, Forecasting Systematic Risk: Estimates of "Raw" Beta that Take Account of the Tendency of Beta to Change and the Heteroskedasticity of Residual Returns, *Journal of Financial and Quantitative Analysis* 20/2, pp. 127-149
- Francis, J.C. & F.J. Fabozzi, 1979, The Effects of Changing Macroeconomic Conditions on the Parameters of the Single Index Model, *Journal of Financial and Quantitative Analysis* 14/2 June, pp. 351-360
- Frankfurter, G.M., 1976, The Effect of Market Indexes on the Ex-Post Performance of the Sharpe Portfolio selection Model, *The Journal of Finance* 31 June, pp. 949-955
- Handa, P., S.P. Kothari & C. Wasley, 1989, The Relation between the Return Interval and Betas, *Journal of Financial Economics* 23, pp. 79-100
- Harvey, C.R., 1989, Time-Varying Conditional Covariances in Tests of Asset Pricing Models, *Journal of Financial Economics* 24, pp. 289-317
- Hawawini, G., 1983, Why Beta Shifts as the Return Interval Changes, *Financial Analysts' Journal* 39 May-June, pp. 73-77
- King, B.F., 1966, Market and Industry Factors in Stock price Behavior, *Journal of Business* 39/1, pp. 139-170
- Klemkosky, R.C. & J.D. Martin, 1975, The Adjustment of Beta Forecasts, *The Journal of Finance* 30 September, pp. 1123-1128
- Levy, R.A., 1971, On the Short-Term Stationarity of Beta Coefficients, *Financial Analysts' Journal* 27/5 November-December, pp. 55-62

- Lintner, J., 1965, The Valuation of Risk Assets and the Selection of Risky Investments in Stock Portfolios and Capital Budgets, *Review of Economics and Statistics* 47/1 February, pp. 13-37
- McDonald, B., 1985, Making Sense Out of Unstable Alphas and Betas, *The Journal of Portfolio Management* Winter, pp. 20-25
- Mossin, J., 1966, Equilibrium in a Capital Asset Market, *Econometrica* 34/4 October, pp. 768-783
- Myers, S.C., 1973, The Stationarity Problem in the Use of the Market Model of Security Price Behavior, *Accounting Review* 48 April, pp. 318-322
- Officer, R.R., 1973, The Variability of the Market Factor of the NYSE, *Journal of Business* 46/3 July, pp. 434-453
- Roll, R., 1977, A Critique of the Asset Pricing Theory's Tests, *Journal of Financial Economics* March, pp. 129-176
- Scholes, M. & J. Williams, 1977, Estimating Betas from Non-Synchronous Data, *Journal of Financial Economics* 5/3 December, pp. 309-327
- Scott, E. & S. Brown, 1980, Biased Estimators and Unstable Betas, *The Journal of Finance* 35/1 March, pp. 49-55
- Shanken, J., 1987, Nonsynchronous Data and the Covariance-Factor Structure of Returns, *The Journal of Finance* 42/2 June, pp. 221-231
- Sharpe, W.F., 1964, Capital Asset Prices: A Theory of Market Equilibrium under Conditions of Risk, *The Journal of Finance* 19/3 September, pp. 425-442
- Simmonds, R.R., L.R. LaMotte & A. McWorther, 1986, Testing for Non-stationarity of Market Risk: An Exact Test and Power Considerations, *Journal of Financial and Quantitative Analysis* 21/2 June, pp. 209-220
- Smith, K.V., 1978, The Effect of Intervalling on Estimating Parameters of the Capital Asset Pricing Model, *Journal of Financial and Quantitative Analysis* June, pp. 313-332
- Sunder, S., 1980, Stationarity of Market Risk: Random Coefficients Tests for Individual Stocks, *The Journal of Finance* 35/4 September, pp. 883-896
- Theobald, M., 1981, Beta Stationarity and Estimation Period: Some Analytical Results, *Journal of Financial and Quantitative Analysis* 16/5 December, pp. 747-757
- Van Der Meulen, J., 1987, The Missing Link, in Hallerbach, W.G. et al. (eds), *Finance, State of the Art 1987*, volume 10, Department of Finance, Erasmus University Rotterdam, pp. 73-90
- Van Der Meulen, J., 1989, *Portefeuilleperformance en Instabiliteit (Portfolio Performance and Instability)*, dissertation Erasmus University Rotterdam
- Van Praag, B.M.S., 1981, Model Free Regression, *Economic Letters* 7, pp. 139-144
- White, H., 1980, Using Least Squares to Approximate Unknown Regression Functions, *International Economic Review* 21/1 February, pp. 149-170

APPLICATION OF NEUTS' METHOD TO VACATION MODELS WITH BULK ARRIVALS

Matendo KAYEMBE

Université de Mons-Hainaut
Place Warocqué 17, B-7000 Mons, Belgium

ABSTRACT

This paper studies a single server infinite capacity queueing system with Poisson arrivals of customers groups of random size and a general service time distribution, the server of which applies a general exhaustive service vacation policy. A computational method is applied to obtain the steady state distributions of the queue of a post-departure or inactive phase termination epoch, at a post-departure epoch and at an arbitrary epoch. Relations between these distributions are given. As special cases, we consider two hybrid vacation models: the $(T(SV),N)$ -policy and the $(T(MV);N)$ -policy. In particular, when the service time distribution is of phase type, explicit results can be obtained for the N -policy. We show this by a simple example.

1. Introduction

This paper deals with the single server infinite capacity queueing system — considered in Loris-Teghem (1990a) — satisfying the following hypotheses:

H_1 : groups of customers arrive according to a Poisson process of parameter λ . The sizes of the groups are i.i.d. random variables, with probability distribution $\{d_k\}_{k \geq 1}$, generating function $D(z) \equiv \sum_{k \geq 1} d_k z^k$ ($|z| \leq 1$), finite expectation $\mathbb{E}[D]$ and second order moment $\mathbb{E}[D^2]$.

H_2 : The server alternates between active and inactive states. In the active state, the server provides service to customers, so that in what we call an “active phase”, the system is never empty. The epochs $0 \leq t_0 < t_1 < \dots < t_m < \dots$ at which the server “enters” the inactive state — and thus, becomes unavailable to the customers — are the times at which the system gets empty (exhaustive service). We denote by τ_m , $m \geq 1$, the epoch at which the inactive phase beginning at t_{m-1} terminates. Thus, if X_m , $m \geq 1$, is the number of groups of customers arriving in the time interval $]t_{m-1}, \tau_m]$, we have $X_m \geq 1$ a.s. We put $V_m = \tau_m - t_{m-1}$, $m \geq 1$, and we assume that the random variables V_m , $m \geq 1$, are i.i.d., with finite expectation and with distribution function $V(\cdot)$. We denote by (X, V) a random vector distributed as the (X_m, V_m) ($m \geq 1$).

H_3 : Customers are served individually in an order independent of their service times, which are i.i.d. random variables independent of the arrival process and the sequence $\{V_m\}_{m \geq 1}$, with distribution function $S(\cdot)$, Laplace-Stieltjes transform (L.S.T.) $\tilde{S}(\cdot)$, and with finite positive expectation $\mathbb{E}[S]$ and second order moment $\mathbb{E}[S^2]$.

H_4 : $\rho = \lambda \mathbb{E}[D] \mathbb{E}[S] < 1$.

For this model, Loris-Teghem (1990a, 1990b) gives, in terms of generating functions, a probabilistic proof of the stochastic decomposition property for the stationary queue length distribution, at a post-departure epoch and at an arbitrary epoch.

Our purpose here is to obtain computational expressions for the stationary queue length distributions.

Computable formulas for a single server queue with server vacations are given in Lucantoni, Meier-Hellstern and Neuts (1990). In the model considered in that paper, customers arrive (individually) according to a general arrival process, the Markovian Arrival Process — of which the Poisson Arrival Process is a very special case — and the vacation policy applied by the server is the “multiple vacation policy”, sometimes called the “ $T(MV)$ -policy”. The computable expressions obtained by the authors concern waiting

time distributions as well as queue length distributions, namely the stationary distribution of the queue length at a post-departure epoch, at an arrival epoch and at an arbitrary epoch.

In the present paper, we assume that groups of customers arrive according to a Poisson process. However, we consider a class of exhaustive service vacation policies which contains the $T(MV)$ -policy as a special case. We are interested in the queue length process, and we first consider the queue length at service completion or inactive phase termination epochs (section 2), for which we compute the steady-state distribution by Neuts' method. We then relate the steady-state distribution of the queue length at service completion epochs (section 3) and at an arbitrary epoch (section 4) to the former distribution. Section 5 is devoted to two particular policies: the $(T(SV); N)$ -policy and the $(T(MV); N)$ -policy. For $N = 1$, the latter reduces to the $T(MV)$ -policy considered in Lucantoni, Meier-Hellstern and Neuts. Our results for this policy agree with those obtained by particularizing the queue length results in Lucantoni, Meier-Hellstern, and Neuts to the Poisson Arrival Process.

We mention that a general batch arrival process, the Batch Markovian Arrival Process, has been considered recently by Lucantoni (1991). We plan to further investigate vacation models, by considering such an arrival process.

Notations

Let Y_m , $m \geq 1$, be the number of customers arriving in the time interval $[t_{m-1}, t_m]$; $\{Y_m\}_{m \geq 1}$ is a sequence of i.i.d. random variables with $Y_m \geq 1$ a.s. We denote by Y a random variable distributed as the Y_m , $m \geq 1$, with finite expectation $\mathbb{E}[Y]$ and second order moment $\mathbb{E}[Y^2]$.

For $t \geq 0$, $n \geq 0$, $\operatorname{Re}(s) \geq 0$ and $|z| \leq 1$, define:

$$\begin{aligned} P(n, t) &= \operatorname{Pr} \{n \text{ customers arrive in the time interval }]0, t]\} \\ &= \sum_{k=0}^n e^{-\lambda t} \frac{(\lambda t)^k}{k!} d_n(k) \end{aligned}$$

where $\{d_n(k)\}_{n \geq 0}$ is the k -fold convolution of the probability distribution $\{d_n\}_{n \geq 1}$.

$$A_n(t) = \int_0^t P(n, x) dS(x) \quad (1)$$

$$\check{A}_n(s) = \int_0^\infty e^{-sx} dA_n(x) = \int_0^\infty e^{-sx} P(n, x) dS(x)$$

$$\check{A}(t) = \sum_{n \geq 0} A_n(t) = S(t)$$

$$\dot{A}(z, s) = \sum_{n \geq 0} \dot{A}_n(s) z^n = \dot{S}[s + \lambda - \lambda D(z)]$$

$$A_n = A_n(+\infty) = \dot{A}_n(0) = \sum_{j=0}^n \psi_j d_n(j) \quad (2)$$

where $\psi_j = \int_0^\infty e^{-\lambda t} \frac{(\lambda t)^j}{j!} dS(t), j \geq 0$

$$A(z) = \sum_{n \geq 0} A_n z^n = \dot{A}(z, 0) = \dot{S}[\lambda - \lambda D(z)] \quad (3)$$

$$A_n^* = \int_0^\infty P(n, t)[1 - S(t)]dt = \sum_{j=0}^n \psi_j^* d_n(j) \quad (4)$$

where $\psi_j^* = \int_0^\infty e^{-\lambda t} \frac{(\lambda t)^j}{j!} [1 - S(t)]dt, j \geq 0$.

$$C_n(t) = Pr[V \leq t, Y = n] \quad (5)$$

$$\begin{aligned} (C_0(t) &\equiv 0) \\ \dot{C}_n(s) &= \int_0^\infty e^{-st} dC_n(t) \end{aligned}$$

$$\dot{C}(t) = \sum_{n \geq 1} C_n(t) = V(t)$$

$$\dot{C}(z, s) = \sum_{n \geq 1} \dot{C}_n(s) z^n$$

$$C_n = C_n(+\infty) = \dot{C}_n(0) = Pr[Y = n] \quad (6)$$

$$C(z) = \sum_{n \geq 1} C_n z^n = \dot{C}(z, 0)$$

Remark 1

(a) For $n \geq 1$, $C_n(t)$ depends on the specific type of vacation policy considered.

(b) Explicit expressions can be obtained for ψ_j and ψ_j^* , $j \geq 0$, when the service time distribution $S(\cdot)$ is of phase type with the (irreducible) representation $(\underline{\beta}, H)$ of order n . Recall that this means that $S(\cdot)$ is the distribution of the time until absorption in a Markov process with a finite number n of transient states $1, 2, \dots, n$ and a single absorbing state $n+1$. The infinitesimal generator of such a process is of the form

$$\begin{pmatrix} H & H^0 \\ 0 & 0 \end{pmatrix}$$

where H is a non singular $n \times n$ matrix which describes transitions between transient states, and $\underline{H}^0$ is a column vector of dimension n , which describes transitions from transient states to the absorbing state, with $\underline{H}^0 = -H\underline{e}$. Throughout this paper, $\underline{e}$ is a column vector of appropriate dimension with all its components equal to 1. $(\underline{\beta}, \beta_{n+1})$ is the initial probability vector and $S(x)$ is given by $S(x) = 1 - \underline{\beta} \exp(Hx)\underline{e}$, for $x \geq 0$. For simplicity, we consider only the case where $\beta_{n+1} = 0$.

Let

$$L \otimes M = \begin{pmatrix} L_{11}M & L_{12}M & \cdots & L_{1k_2}M \\ \vdots & \vdots & \ddots & \vdots \\ L_{k_1 1}M & L_{k_1 2}M & \cdots & L_{k_1 k_2}M \end{pmatrix}$$

be the Kronecker product of two matrices L and M of dimensions $k_1 \times k_2$ and $k'_1 \times k'_2$ respectively ($L \otimes M$ is a matrix of dimension $k_1 k'_1 \times k_2 k'_2$). Let $(1, (-\lambda))$ denote the simplest representation of the group arrival Poisson process and let I denote the identity matrix (of appropriate dimension). Then (see Neuts [1989, th. 5.1.5, pp. 244-246])

$$\psi_0 = -(I \otimes \underline{\beta})((- \lambda) \otimes I + I \otimes H)^{-1}(I \otimes \underline{H}^0), \quad (7)$$

$$\psi_k = \varphi \gamma^{k-1} \Omega, \quad k \geq 1,$$

and

$$\psi_0^* = -(I \otimes \underline{\beta})((- \lambda) \otimes I + I \otimes H)^{-1}(I \otimes \underline{e}), \quad (8)$$

$$\psi_k^* = \varphi \gamma^{k-1} \Omega_0, \quad k \geq 1,$$

where the matrices φ, γ, Ω and Ω_0 (of dimensions $1 \times n, n \times n$ and $n \times 1$ respectively) are given by

$$\begin{aligned} \varphi &= -(I \otimes \underline{\beta})((- \lambda) \otimes I + I \otimes H)^{-1}(\lambda \otimes I), \\ \gamma &= -(1 \otimes I)((- \lambda) \otimes I + I \otimes H)^{-1}(\lambda \otimes I), \\ \Omega &= -(1 \otimes I)((- \lambda) \otimes I + I \otimes H)^{-1}(I \otimes \underline{H}^0), \\ \Omega_0 &= -(1 \otimes I)((- \lambda) \otimes I + I \otimes H)^{-1}(I \otimes \underline{e}). \end{aligned} \quad (9)$$

From (7) and (8), it is clear that no numerical integrations are needed to compute ψ_k and ψ_k^* , $k \geq 0$. We refer to Latouche (1989) and Neuts (1989) for more information about

phase type distributions.

2. The queue length at service completion or inactive phase termination epochs

Let $N(t)$ be the number of customers in the system at time t . We are concerned with the discrete parameter process obtained by restricting the process $\{N(t), t \geq 0\}$ at epochs of service completion or inactive phase termination. Denote these instants by $\theta_n, n \geq 0$, and let $\zeta_n, n \geq 0$, be the number of customers arrived in the time interval $[\theta_n, \theta_{n+1}]$. Define

$$T_n = \theta_{n+1} - \theta_n \quad \text{and} \quad N_n = N(\theta_n+), n \geq 0.$$

Then, for $n \geq 0$,

$$N_{n+1} = (N_n - 1)^+ + \zeta_n$$

and

$$\begin{aligned} & Pr[N_{n+1} = j, T_n \leq x | N_0, T_0, N_1, T_1, \dots, T_{n-1}, N_n = i] \\ &= Pr[N_{n+1} = j, T_n \leq x | N_n = i] \quad (\equiv Q_{ij}(x)), x \geq 0 \end{aligned}$$

It follows that the sequence $\{(N_n, T_n)\}_{n \geq 0}$ is a M.R.P. on the state space $\{i \geq 0\} \times [0, +\infty[$. Its transition probability matrix $Q(\cdot)$ is given by

$$Q(x) = \begin{pmatrix} C_0(x) & C_1(x) & C_2(x) & \cdots \\ A_0(x) & A_1(x) & A_2(x) & \cdots \\ 0 & A_0(x) & A_1(x) & \cdots \\ 0 & 0 & A_0(x) & \cdots \\ \vdots & \vdots & \vdots & \ddots \end{pmatrix}, x \geq 0, \quad (10)$$

where $A_\nu(x)$ and $C_\nu(x), \nu \geq 0$ are the probability mass-functions defined in (1) and (5). Note that

$$\rho = \sum_{k \geq 1} k A_k < 1 \quad (H_4) \quad ,$$

$$\overset{\circ}{\alpha} \equiv \int_0^\infty x d\overset{\circ}{A}(x) = \int_0^\infty x dS(x) = \mathbb{E}[S] < +\infty \quad ,$$

$$\overset{\circ}{\beta} \equiv \int_0^\infty x d\overset{\circ}{C}(x) = \int_0^\infty x dV(x) = \mathbb{E}[V] < +\infty \quad ,$$

$$\sum_{k \geq 1} k C_k = C'(1) = \mathbb{E}[Y] < +\infty \quad .$$

Thus (see Neuts [1989, th. 1.3.1, pp. 12-13]), the M.R.P. $\{(N_n, T_n)\}_{n \geq 0}$ and the Markov chain $\{N_n\}_{n \geq 0}$ are positive recurrent. From (10), the transition probability matrix of the Markov chain $\{N_n\}_{n \geq 0}$ is given by

$$Q = Q(+\infty) = \begin{pmatrix} C_0 & C_1 & C_2 & \cdots \\ A_0 & A_1 & A_2 & \cdots \\ 0 & A_0 & A_1 & \cdots \\ 0 & 0 & A_0 & \cdots \\ \vdots & \vdots & \vdots & \ddots \end{pmatrix} \quad (11)$$

Clearly, this Markov chain is irreducible and aperiodic.

Let $\underline{x} = (x_i)_{i \geq 0}$ be the stationary probability vector, which satisfies

$$\begin{aligned} \underline{x} Q &= \underline{x} \\ \underline{x} \underline{e} &= 1 \end{aligned}$$

From (11), we have

$$x_i = x_0 C_i + \sum_{\nu=1}^{i+1} x_\nu A_{i+1-\nu}, \text{ for } i \geq 0, \quad (12)$$

2.1. Computation of x_0

From Neuts (1989, p. 15) we obtain

$$x_0 = \left[1 + \frac{1}{1-\rho} \sum_{\nu \geq 1} \nu C_\nu \right]^{-1} = \frac{1-\rho}{1-\rho + \mathbb{E}[Y]} \quad (13)$$

2.2. Computation of $x_i, i \geq 1$

From Neuts (1989, p. 17) we have the following recurrence formula:

$$x_i = A_0^{-1} \left[x_0 \hat{C}_{i-1} + \sum_{\nu=1}^{i-1} x_\nu \hat{A}_{i-\nu} \right], \quad i \geq 1, \quad (14)$$

where $\hat{A}_\nu = 1 - \sum_{r=0}^{\nu} A_r$, and $\hat{C}_\nu = 1 - \sum_{r=0}^{\nu} C_r$, for $\nu \geq 0$.

An analogous recurrence formula for the computation of $x_i, i \geq 1$, can be found in Ramaswami (1988).

2.3. Mean queue length at service completion or inactive phase termination epochs.

Put $X_v(z) = \sum_{i \geq 0} x_i z^i$ ($|z| \leq 1$). Equations (12) readily imply

$$X_v(z)[z - A(z)] = x_0[zC(z) - A(z)]$$

so that (see Neuts [1989, th. 1.4.1, pp. 17-18]) $\sum_{i \geq 1} ix_i \equiv X'_v(1)$ is given by

$$X'_v(1) = \frac{1}{2(1-\rho)} \{A''(1) + x_0[2C'(1) + C''(1) - A''(1)]\} \quad (15)$$

where x_0 is obtained in (13),

$$A''(1) = \sum_{\nu \geq 2} \nu(\nu-1)A_\nu = \lambda^2(\mathbb{E}[D])^2\mathbb{E}[S^2] + \lambda\mathbb{E}[S](\mathbb{E}[D^2] - \mathbb{E}[D]),$$

and

$$C''(1) = \sum_{\nu \geq 2} \nu(\nu-1)C_\nu = \mathbb{E}[Y^2] - \mathbb{E}[Y].$$

3. The queue length at post-departure epochs

By restricting the process $\{N(t), t \geq 0\}$ at service completion epochs, we obtain another M.R.P. on the state space $\{i \geq 0\} \times [0, +\infty[$, with transition probability matrix $Q_1(\cdot)$ given by

$$Q_1(x) = \begin{pmatrix} B_0(x) & B_1(x) & B_2(x) & \cdots \\ A_0(x) & A_1(x) & A_2(x) & \cdots \\ 0 & A_0(x) & A_1(x) & \cdots \\ 0 & 0 & A_0(x) & \cdots \\ \vdots & \vdots & \vdots & \ddots \end{pmatrix}, x \geq 0, \quad (16)$$

where $B_i(x) = \sum_{\nu=1}^{i+1} C_\nu(\cdot) \star A_{i+1-\nu}(x)$, for $i \geq 0$.

Thus $Q_1(x)$ differs from $Q(x)$ -given by (10)-only in the elements on the first line: $B_i(x)$ instead of $C_i(x)$, and the stationary probability vector $\underline{\pi} = (\pi_i)_{i \geq 0}$ of the corresponding Markov chain could be computed in a way similar to that described in section 2 for the computation of the stationary probability vector $\underline{x}$.

Note, however, that a simple probabilistic argument leads to the following relation

$$x_i = \frac{x_0}{\pi_0} \pi_i + x_0 C_i, \quad i \geq 1 \quad (17)$$

Indeed, as the system is not empty at an inactive phase termination, we have:

$$\begin{aligned}
x_0 &= Pr[\text{the system is empty after a transition}] \\
&= Pr[\text{the system is empty after a transition and this} \\
&\quad \text{transition is a service completion}] \\
&= Pr[\text{a transition is a service completion}] \times \\
&\quad \times Pr[\text{the system is empty after a transition, given} \\
&\quad \text{that this transition is a service completion}] \\
&= Pr[\text{a transition is a service completion}] \times \pi_0
\end{aligned}$$

Thus, $\frac{x_0}{\pi_0}$ is the probability that a transition is a service completion.
On the other hand,

$$\begin{aligned}
&Pr[\text{a transition is an inactive phase termination}] \\
&= Pr[\text{the system is empty after the previous transition}] \\
&= x_0
\end{aligned}$$

The result follows by application of the law of total probability, as C_i is the probability that i customers arrive in an inactive phase.

From (17) we have

$$\sum_{i \geq 1} x_i = \frac{x_0}{\pi_0} \sum_{i \geq 1} \pi_i + x_0 \sum_{i \geq 1} C_i$$

or

$$1 - x_0 = \frac{x_0}{\pi_0} (1 - \pi_0) + x_0$$

so that

$$\begin{aligned}
\pi_0 &= \frac{x_0}{1 - x_0} = \frac{1 - \rho}{\mathbb{E}[Y]} \\
\pi_i &= \frac{1}{1 - x_0} (x_i - x_0 C_i), i \geq 1
\end{aligned} \tag{18}$$

and, using (15),

$$\sum_{i \geq 1} i \pi_i = \rho + \frac{A''(1)}{2(1 - \rho)} + \frac{C''(1)}{2\mathbb{E}[Y]} \tag{19}$$

4. The stationary queue length at an arbitrary epoch

In this section, we are interested in computing the continuous time stationary queue length distribution, denoted by $\{p_i\}_{i \geq 0}$, by using one of the M.R.P. considered in the previous sections. As we noticed, the matrices $Q(\cdot)$ and $Q_1(\cdot)$ associated with these processes differ only by the elements on the first line. As the expression for these elements is simpler in $Q(\cdot)$ than in $Q_1(\cdot)$, we will relate the distribution $\{p_i\}_{i \geq 0}$ to the stationary distribution $\{x_i\}_{i \geq 0}$ of the queue length at a service completion or inactive phase termination epoch. (In Loris-Teghem (1990a, 1990b), a relation between the distributions $\{p_i\}_{i \geq 0}$ and $\{\pi_i\}_{i \geq 0}$ is obtained).

We first introduce the fundamental mean E of the M.R.P. $\{(N_n, T_n)\}_{n \geq 0}$ (i.e. the inner product of the stationary probability vector $\underline{x}$ and the vector of row sum means $\int_0^\infty x dQ(x)\underline{e}$ of the matrix $Q(\cdot)$). In the stationary version of the queue, E may be interpreted as the average time between two consecutive transitions (i.e. service completion or inactive phase termination). From Neuts (1989, p. 22) we have

$$\begin{aligned} E &= x_0 \overset{\circ}{\beta} + (1 - x_0) \overset{\circ}{\alpha} = \frac{\overset{\circ}{\beta}(1 - \rho) + \overset{\circ}{\alpha} \sum_{\nu \geq 1} \nu C_\nu}{1 - \rho + \sum_{\nu \geq 1} \nu C_\nu} \\ &= \frac{\lambda^{-1} \mathbb{E}[X] - \mathbb{E}[S](\mathbb{E}[D]\mathbb{E}[X] - \mathbb{E}[Y])}{1 - \rho + \mathbb{E}[Y]} \end{aligned} \quad (20)$$

We assume that time $t = 0$ corresponds to a transition epoch in the M.R.P. $\{(N_n, T_n)\}_{n \geq 0}$ and that $N_0 = i_0 \geq 0$.

For $t \geq 0$, let $M_{i_0 i}(t)$ denote the conditional expected number of visits to state i in the time interval $[0, t]$, given that $N_0 = i_0$, and let $M(t) = \{M_{i_0 i}(t), i_0 \geq 0, i \geq 0\}$ be the Markov renewal matrix of the M.R.P. $\{(N_n, T_n)\}_{n \geq 0}$.

The matrix $M(\cdot)$ is given by the convolution power series

$$M(t) = \sum_{n \geq 0} Q^{(n)}(t), \text{ for } t \geq 0.$$

Let $dM_{i_0 i}(t)$ denote the conditional probability that in the interval $]t, t + dt[$, the M.R.P. $\{(N_n, T_n)\}_{n \geq 0}$ enters state i , given that $N_0 = i_0$.

We now consider the continuous-parameter process $\{N(t), t \geq 0\}$ and define:

$$P_{i_0 i}(t) = Pr[N(t) = i | N_0 = i_0], \text{ for } i \geq 0,$$

and

$$K_{i_0 i}(t) = Pr[N(t) = i, \theta_1 > t | N_0 = i_0], \text{ for } i \geq i_0.$$

We have (see Çinlar (1969, 1975))

$$\begin{aligned}
 P_{i_0 i}(t) &= \sum_{j=0}^i \int_0^t dM_{i_0 j}(u) K_{ji}(t-u) \\
 &= \begin{cases} \int_0^t dM_{i_0 0}(u) e^{-\lambda(t-u)} & ; \quad \text{for } i = 0 \\ \int_0^t dM_{i_0 0}(u) K_{0i}(t-u) + \sum_{j=1}^i \int_0^t dM_{i_0 j}(u) [1 - S(t-u)] P(i-j, t-u), & \text{for } i \geq 1. \end{cases}
 \end{aligned}$$

The limits $p_i = \lim_{t \rightarrow \infty} P_{i_0 i}(t)$, $i \geq 0$, exist (and are independent of i_0) by virtue of the key renewal theorem and are given by

$$\begin{aligned}
 p_0 &= \frac{x_0}{E} \int_0^\infty K_{00}(t) dt = \frac{x_0}{\lambda E} \\
 &= \frac{\lambda^{-1}(1-\rho)}{\lambda^{-1} \mathbb{E}[X] - \mathbb{E}[S](\mathbb{E}[D] \mathbb{E}[X] - \mathbb{E}[Y])} \quad , \quad (21)
 \end{aligned}$$

$$p_i = \frac{1}{E} x_0 \int_0^\infty K_{0i}(t) dt + \frac{1}{E} \sum_{j=1}^i x_j \int_0^\infty P(i-j, t) [1 - S(t)] dt, \quad i \geq 1, \quad (22)$$

provided that each of the functions $K_{0i}(\cdot)$ — which depends on the specific type of vacation policy considered — be directly Riemann integrable.

Mean queue length at an arbitrary epoch

Relations (21) and (22) give, for the generating function $P_v(z) \equiv \sum_{i \geq 0} p_i z^i$ ($|z| \leq 1$):

$$P_v(z) = p_0 + \frac{1}{E} \left\{ x_0 \int_0^\infty K_0(t, z) dt + [X_v(z) - x_0] [1 - A(z)] [\lambda - \lambda D(z)]^{-1} \right\}, \quad (23)$$

where $K_0(t, z) = \sum_{i \geq 1} K_{0i}(t) z^i$.

Using this expression for $P_v(z)$, one could obtain the mean queue length $\sum_{i \geq 1} i p_i = P'_v(1)$.

Note, however, that this can be obtained by using the following decomposition property (Loris-Teghem (1990a, 1990b)):

$$P_v(z) = P_{nv}(z) \Delta_v(z) \quad , \quad |z| \leq 1 \quad ,$$

where $P_{nv}(z)$ is the generating function of the queue length at an arbitrary epoch in the bulk arrival model without vacations and $\Delta_v(z)$ is the following generating function:

$$\Delta_v(z) = \frac{1 - C(z)}{\mathbb{E}[X](1 - D(z))} \quad .$$

Thus

$$P'_v(1) = P'_{nv}(1) + \Delta'_v(1) \quad (24)$$

$$\text{with } P'_{nv}(1) = \rho + \frac{\lambda[\lambda(\mathbb{E}[D])^2\mathbb{E}[S^2] + \mathbb{E}[S](\mathbb{E}[D^2] - \mathbb{E}[D])]}{2(1 - \rho)}$$

$\Delta'_v(1)$ depends on the specific type of vacation policy considered.

5. Particular cases : the $(T(SV); N)$ -model and the $(T(MV); N)$ -model

In this section, we are interested in the functions $K_{0i}(t)$, $i \geq 1$, for two hybrid vacation models: the $(T(SV); N)$ -policy and the $(T(MV); N)$ -policy, introduced in Loris-Teghem (1985). In these models, upon becoming idle, the server leaves the system for a vacation of random length. In the first case, when he comes back to the system, the server becomes active immediately if he finds at least N customers present. Otherwise, he remains inactive, inspecting the queue, until N customers are present. In the second case, the vacations are repeated until the server finds at least N customers in the system upon return from a vacation. We refer the reader to Loris-Teghem (1990a) for more details about these policies. Note that in Loris-Teghem (1990b), the distribution $\{C_k\}_{k \geq 1}$ of Y is given for both the $(T(SV); N)$ -model and the $(T(MV); N)$ -model.

Let $U(\cdot)$ be the common distribution function of the vacation lengths, with finite expectation $\mathbb{E}[U]$.

5.1. The $(T(SV); N)$ -model

We have :

$$\begin{aligned} K_{0i}(t) &= P(i, t)[1 - U(t)], \text{ for } i \geq N \\ &= P(i, t) \quad \text{for } 1 \leq i \leq N - 1 \end{aligned} \quad (25)$$

Remark 2

For the N -policy (i.e. $U(\cdot) \equiv 1$), (25) reduces to

$$\begin{aligned} K_{0i}(t) &= 0, \text{ for } i \geq N \\ &= P(i, t), \text{ for } 1 \leq i \leq N - 1 \end{aligned}$$

Substituting into (22) yields

$$\begin{aligned} p_i &= p_0 \sum_{k=1}^i d_i(k) + \eta_i, \text{ for } 1 \leq i \leq N - 1 \\ &= \eta_i, \text{ for } i \geq N \end{aligned} \quad (26)$$

where η_i ($i \geq 1$) is given by

$$\eta_i = \frac{1}{E} \sum_{j=1}^i x_j A_{i-j}^* \quad (i \geq 1)$$

When the service time distribution is of phase type, explicit expressions can be obtained for p_i ($i \geq 1$). This follows from (4) and remark 1 (b). In the following example, we write the expressions for p_0 , p_1 , p_2 and p_3 .

Example

Consider the case where $n = N = 2$, $H = \begin{pmatrix} -\mu & \mu \\ 0 & -\mu \end{pmatrix}$ and $\underline{\beta} = (1, 0)$. We have

$$\begin{aligned} p_0 &= \frac{1 - \rho}{\mathbb{E}[X]} , \\ p_1 &= p_0 \left\{ d_1 + \frac{\lambda(\lambda + 2\mu)}{\mu^2} \right\} , \\ p_2 &= \lambda p_0 \left\{ \frac{(\lambda + \mu)^2(\lambda + 2\mu)}{\mu^4} - \frac{\lambda}{\mu^2} d_1 \right\} , \\ \text{and } p_3 &= \lambda p_0 \left\{ \frac{(\lambda + \mu)^4(\lambda + 2\mu)}{\mu^6} - \frac{(\lambda + \mu)(3\lambda^2 + 5\lambda\mu)}{\mu^4} d_1 \right. \\ &\quad \left. - \frac{(\lambda + 2\mu)}{\mu^2} (d_1^2 + d_2) - \frac{\lambda}{\mu^2} d_2 \right\} , \end{aligned}$$

where $\mathbb{E}[X] = 1 + d_1$.

Remark 3

The N -policy model with a phase type service time distribution was considered in Altioik (1987). In order to compute the p_i , $i \geq 1$, Altioik uses the following relations:

$$\begin{aligned} p_i &= \sum_{k=0}^n P_{i,k}^* , \text{ for } 1 \leq i \leq N - 1 \\ &= \sum_{k=1}^n P_{i,k}^* , \text{ for } i \geq N , \end{aligned} \tag{27}$$

where $P_{i,0}^*$ ($0 \leq i \leq N - 1$) and $P_{i,k}^*$ ($i \geq 1, k = 1, \dots, n$) are the steady-state probabilities that i customers are in the system and the server is idle, and to have i customers in the system and the server is in phase k respectively, for the computation of which Altioik develops algorithms.

5.2. The $(T(MV); N)$ -model

We have :

$$\begin{aligned} K_{0i}(t) &= \sum_{r \geq 0} \sum_{k=0}^{N-1} \int_0^t P(k, y) dU^{(r)}(y) P(i-k, t-y) [1 - U(t-y)] \quad , i \geq N \\ &= P(i, t) \quad , 1 \leq i \leq N-1 \end{aligned} \quad (28)$$

where $U^{(r)}(t)$ denotes the r -th power of convolution of $U(t)$.

Remark 4

Let $b_0 = \int_0^\infty e^{-\lambda t} dU(t)$ denote the probability that no arrival occur during a single vacation.

Putting $N = 1$ in

- (22) and (25) we get

$$p_i = \lambda p_0 \int_0^\infty P(i, t) [1 - U(t)] dt + \eta_i \quad (i \geq 1) \quad (29)$$

for the $T(SV)$ -model.

- (22) and (28) we get

$$p_i = \lambda p_0 (1 - b_0)^{-1} \int_0^\infty P(i, t) [1 - U(t)] dt + \eta_i \quad (i \geq 1) \quad (30)$$

for the $T(MV)$ -model.

When $S(\cdot)$ and $U(\cdot)$ are phase type distributions, explicit results can be obtained for p_i ($i \geq 1$) for both models. Note that p_0 appearing in (29) and (30) is given by

$$p_0 = (\mathbb{E}[X])^{-1} (1 - \rho)$$

where (see Loris-Teghem (1990a)) $\mathbb{E}[X] = \lambda \mathbb{E}[U] + b_0$ for the former model, while for the latter model $\mathbb{E}[X] = \lambda \mathbb{E}[U] (1 - b_0)^{-1}$.

References

- Altioek T. (1987). Queues with group arrivals and exhaustive service discipline, Queueing Syst., **2**, 307-320.
- Çinlar E. (1969). Markov renewal theory, Adv. App. Prob., **1**, 123-187.
- Çinlar E. (1975). Introduction to Stochastic Processes, Englewood Cliffs, NJ : Prentice-Hall.

- Latouche G. (1989). Distributions de type phase : tutorial, Cahiers du C.E.R.O., **31**, 3-11.
- Loris-Teghem J. (1985). Analysis of single server queueing systems with vacation periods, Belg. Journ. of Op. Res., Stat. and Comput. Sc., **25**, 47-54.
- Loris-Teghem J. (1990a). On vacation models with bulk arrivals, Belg. Journ. of Oper. Res., Stat. and Comput. Sci., **30(1)**, 53-66.
- Loris-Teghem J. (1990b). Remark on : On vacation models with bulk arrivals, Belg. Journ. of Oper. Res., Stat. and Comput. Sc., **30(4)**, 53-56.
- Lucantoni D.M. (1991). New results on the single server queue with a Batch Markovian Arrival Process, Stochastic Models, **7**, 1-46.
- Lucantoni D.M., Meier-Hellstern K.S. and Neuts M.F. (1990). A single-server queue with server vacations and a class of non-renewal arrival processes, Adv. Appl. Prob., **22**, 676-705.
- Neuts M.F. (1989). Structured Stochastic Matrices of M/G/1 Type and their Applications, Marcel Dekker, Inc., New York.
- Ramaswami (1988). A stable recursion for the steady state vector in Markov chains of M/G/1 type, Stochastic Models, **4**, 183-188.

THE P-CENTER PROBLEM IN R^n WITH WEIGHTED TCHEBYCHEFF NORMS

Bias PELEGRIN PELEGRIN

**Dpto. de Matemática Aplicada y Estadística
Universidad de Murcia, 30100-Espinardo. Murcia. Spain**

ABSTRACT

In this paper the p-center problem in R^n is studied when distances are measured by a weighted Tchebycheff norm. For the 1-center problem it is proved that the lower bound proposed by Dearing and Francis (1974) is attained and a one step algorithm is given to obtain an optimal solution. Then, an exact algorithm for $p > 1$ is proposed which generalizes the one given by Aneja et al. (1988) for the unweighted rectangular p-center problem on the plane. Finally, a new 2-approximation heuristic polynomial algorithm is given which is «best possible» since for $\delta < 2$ the existence of a δ -approximation polynomial algorithm would imply that $P = NP$.

1. INTRODUCTION

Let $M = \{ P_1, P_2, \dots, P_m \}$ be a given set of points in the Euclidean space $\mathbb{R}^n$. The p -center problem in $\mathbb{R}^n$ seeks p points $C_1, C_2, \dots, C_p$, $p < m$, so that the maximal weighted distance between each point P_i and its closest point C_j , $j=1, \dots, p$, is minimized. The problem is usually formulated as :

$$(P1) \quad \begin{array}{ll} \text{Minimize} & \max_{1 \leq i \leq m} \{ w_i \min_{1 \leq j \leq p} \{ N(P_i - X_j) \} \\ & X_1, \dots, X_p \in \mathbb{R}^n \end{array}$$

where N denotes a norm function and $w_i > 0$, $i=1, \dots, m$.

Alternatively, the problem can be seen as that of finding a partition $\alpha = \{ M_1, \dots, M_p \}$ of the set M into p disjoint subsets, so that the maximum among the maximal weighted distances between the best point (optimal solution to the 1-center problem with points P_i in M_j) and the points in each subset M_j is minimized. Then the problem is also formulated as :

$$(P2) \quad \begin{array}{ll} \text{Minimize} & r_\alpha = \max_{1 \leq j \leq p} \{ r(M_j) \} \\ & \alpha \in P(M, p) \end{array}$$

where $P(M, p)$ denotes the set of all partitions of M into p disjoint subsets and $r(M_j)$ denotes the optimal value of the 1-center problem associated with the points in M_j , $j=1, \dots, p$. Then the points $C_1, C_2, \dots, C_p$ are found as optimal solutions to the 1-center problems associated with the subsets of an optimal partition.

The problem arises typically for $n=2$ in the field of Location

Theory when one seeks locations for p facilities, so that the maximal distance or travel time between each demand point P_i and its closest facility C_j is minimized. Examples of this kind of problem are in the location of emergency services (like fire stations, ambulance and helicopter bases, police stations), location of radio and TV stations, messenger delivery services, etc. The problem can also be found in Cluster Analysis where the aim is to create p groups in a set of objects M , so that the maximal dissimilarity between each object P_i and the center C_j of its group is minimized. The objects are characterized by the values of n variables, and the dissimilarity is measured by a norm function. Examples are found in economics, social sciences, and almost all empirically based disciplines.

For $n=2$, in spite of being NP-hard (Megiddo and Supowit¹), the problem has recently been solved by exact algorithms for the euclidean norm (Drezner², Vijay³, Chen and Handler⁴), and for the rectangular norm (Watson-Gandy⁵, Drezner⁶, Aneja et al.⁷). For arbitrary n , due to its complexity, only heuristic algorithms have been used (Drezner², Spath⁸, Eiselt and Charlesworth⁹, Dyer and Frieze¹⁰), which can be applied for any norm N . Most of these heuristics require the 1-center problem to be solved many times, which is time consuming for $n > 2$. In this case, the centers $C_1, \dots, C_p$ are normally constrained to be points in M , as happens in other Cluster Analysis problems and in the classical p -median problem in Location Theory.

The aim of this paper is to study the problem when N is given by a weighted Tchebycheff norm, which is defined as:

$$N(X) = \max \{ \lambda_h |x_h|, h=1, \dots, n \}$$

where $X = (x_1, \dots, x_n)$ and $\lambda_h > 0, h=1, \dots, n$. For $n=2$, these norms are obtained by linear transformation from bidirectional polyhedral norms (Pelegrin¹¹), and they can then be used in Location Theory when the movement is restricted to two given directions. For instance, taking the transformation $x_1 = y_1 + y_2$, $x_2 = y_1 - y_2$ it follows that $|y_1| + |y_2| = \max\{|x_1|, |x_2|\}$, and the Tchebycheff norm is obtained from the rectangular norm. For any n , they are obtained by a scaling transformation from the Tchebycheff norm and can be used in Cluster Analysis as criterion distances.

To our knowledge, the problem is studied here for the first time for this type of norm. Firstly, we deal with the 1-center problem showing that the lower bound given in Dearing and Francis¹² is reached and solve the problem by a one step algorithm. Then, the algorithm given in Aneja et al.⁷, which solves optimally the problem in $\mathbb{R}^2$ with the rectangular norm and $w_i = 1, i=1, \dots, m$, is generalized to solve the problem in $\mathbb{R}^n$ with any weighted Tchebycheff norm and $w_i > 0, i=1, \dots, m$. Finally, some heuristic algorithms are considered and a 2-approximation polynomial algorithm is proposed which is "best possible" since, for any $\delta < 2$, the existence of a δ -approximation polynomial algorithm would imply that $P=NP$ as shown in Ref. 15 and 16.

2. THE 1-CENTER PROBLEM

The 1-center problem is formulated as follows :

$$(P3) \quad \underset{X \in \mathbb{R}^n}{\text{Minimize}} \quad R(X) = \underset{1 \leq i \leq m}{\text{Max}} \quad \{ w_i N(P_i - X) \}$$

where $X = (x_1, \dots, x_n)$, $P_i = (a_1^i, \dots, a_n^i)$, and $N(P_i - X) = \max_{h=1, \dots, n} \{ \lambda_h |a_h^i - x_h| \}$. Let X^* and r^* denote an optimal solution and the optimal value of (P3) respectively.

A lower bound of $R(X)$ for any norm N is given in Dearing and Francis¹² by :

$$B = \max_{i \neq k} \{ w_i w_k N(P_i - P_k) / (w_i + w_k) \}$$

The following property shows that the lower bound B is reached and gives the set S^* of optimal solutions to (P3).

Property 1

- i) $r^* = B$.
- ii) $S^* = \{ X \in \mathbb{R}^n : \max_{1 \leq i \leq m} \{ a_h^i - B/w_i \lambda_h \} \leq x_h \leq \min_{1 \leq i \leq m} \{ a_h^i + B/w_i \lambda_h \} \}$.

Proof:

Let $B(X, r)$ denote the ball centered at X and with radius r , i.e., $B(X, r) = \{ Y \in \mathbb{R}^n : \lambda_h |y_h - x_h| \leq r, h=1, \dots, n \}$, and let $B_h(X, r)$ denote the projection of $B(X, r)$ on the h -th axis, i.e., $B_h(X, r) = \{ y_h \in \mathbb{R} : \lambda_h |y_h - x_h| \leq r \}$. Then $B(X, r) = \prod_{h=1}^n B_h(X, r)$ because N is a weighted Tchebycheff norm (note that this holds for such norms only). Also $r \geq r^*$ iff $\cap \{ B(P_i, r/w_i) : P_i \in M \} \neq \emptyset$.

For $X_{ik} = (w_i P_i + w_k P_k) / (w_i + w_k)$, $i \neq k$, it follows that $B \geq w_i w_k N(P_i - P_k) / (w_i + w_k) = w_i w_k (N(P_i - X_{ik}) + N(P_k - X_{ik})) / (w_i + w_k) =$

$= w_i N(P_i - X_{ik}) = w_k N(P_k - X_{ik})$,
 therefore $X_{ik} = (x_1^{ik}, \dots, x_n^{ik}) \in B(P_i, B/w_i) \cap B(P_k, B/w_k)$ and
 $x_h^{ik} \in B_h(P_i, B/w_i) \cap B_h(P_k, B/w_k)$ for $h=1, \dots, n$. Hence on each h -axis
 $B_h(P_i, B/w_i) \cap B_h(P_k, B/w_k) \neq \emptyset$ for any $P_i, P_k \in M, i \neq k$. From the Helly
 property (Rockafellar¹³) , it follows that $\cap \{ B_h(P_i, B/w_i) : P_i \in M \} \neq \emptyset$.

For $h=1, \dots, n$, let x_h' be in $\cap \{ B_h(P_i, B/w_i) : P_i \in M \}$, then
 $X' = (x_1', \dots, x_n') \in \cap \{ B(P_i, B/w_i) : P_i \in M \}$. This implies that
 $R(X') = B$. Consequently, $r^* = B$ and $S^* = \{ X \in \mathbb{R}^n : R(X) = B \}$.

As $X \in S^*$ if and only if $w_i \max \{ \lambda_h | a_h^i - x_h | : h=1, \dots, n \} \leq B$
 for all $P_i \in M$, the proof is complete.

The lower bound B is also reached in $\mathbb{R}^2$ for bidirectional
 polyhedral norms, however this result is false for the L_1 norm
 in $\mathbb{R}^n, n > 2$, as it happens for $P_1 = (0,0,0)$, $P_2 = (1,1,0)$,
 $P_3 = (1,0,1)$, $P_4 = (0,1,1)$ and $w_i = 1, i=1, \dots, 4$.

As a consequence of property 1 , the following one step
 algorithm gives an optimal solution to (P3).

ALGORITHM 1

- Calculate :

$$\begin{aligned}
 B &= \max_{i \neq k} \{ w_i w_k / (w_i + w_k) \max_{1 \leq h \leq n} \{ \lambda_h | a_h^i - a_h^k | \} \} \\
 \alpha_h &= \max_{1 \leq i \leq m} \{ a_h^i - B/w_i \lambda_h \} \quad \text{for } h = 1, \dots, n \\
 \beta_h &= \min_{1 \leq i \leq m} \{ a_h^i + B/w_i \lambda_h \} \quad \text{for } h = 1, \dots, n
 \end{aligned}$$

- Choose $x_h^* \in [\alpha_h, \beta_h]$, e.g. $x_h^* = (\alpha_h + \beta_h)/2$ for $h = 1, \dots, n$.
- Set $X^* = (x_1^*, \dots, x_n^*)$ and $r^* = B$. Output X^* and r^* . Stop.

3. AN EXACT ALGORITHM FOR $P > 1$

In the following, the algorithm given in Aneja et al.⁷ is generalized to solve the problem (P2). Let r^* denote the optimal value of (P2) and $r_{ik} = w_i w_k N(P_i - P_k) / (w_i + w_k)$ for $i \neq k$. For any partition $\alpha = (M_1, \dots, M_p)$, from property 1 it follows that :

$$(1) \quad r(M_j) = \max_{i \neq k} \{ r_{ik} : P_i, P_k \in M_j \} .$$

Then, it is not necessary to solve any optimization problem to evaluate r_α , as happens with other norms, and the following property satisfies.

Property 2

$$r^* \in R = \{ r_{ik} : P_i, P_k \in M, i \neq k \} .$$

Let $r \in R$, and define the graph $G(r) = (V, E(r))$ as follows :

$$V = \{ P_1, \dots, P_m \}$$

$$E(r) = \{ (P_i, P_k) : r_{ik} \leq r \} .$$

Property 3

(P2) is equivalent to the problem of finding the minimum value of r in R , such that V is covered by a set of p cliques (maximal complete subgraphs) of $G(r)$.

Proof:

Let $\alpha = (M_1, \dots, M_p)$ be a partition such that $r_\alpha \leq r$, then $r_{ik} \leq r$

for $P_i, P_k \in M_j$, $j=1, \dots, p$. Therefore, $M_j, j=1, \dots, p$, are complete subgraphs of $G(r)$ from which a set of p cliques covering V can be generated adding (if necessary) points to each subset M_j . Conversely, if $V_1, \dots, V_p$ are cliques of $G(r)$ covering V , then a partition α can be generated, eliminating (if necessary) some points in $V_1, \dots, V_p$, which satisfy $r_\alpha \leq r$. Note that since this transformation of partitions into sets of covering cliques may be done in polynomial time, the problems are polynomially equivalent.

Let $V_1, \dots, V_q$ be all the distinct cliques of $G(r)$. There exist p cliques of $G(r)$ which cover V if and only if the optimal value of the following set covering problem, $SCP(r)$, is less than or equal to p :

$$\begin{aligned} SCP(r) \quad & \text{Minimize } z = \sum_j x_j \\ & \text{s.t.} \quad \sum_j a_{ij} x_j \geq 1, \quad i=1, \dots, m \\ & \quad \quad x_j \in \{0, 1\}, \quad j=1, \dots, q. \end{aligned}$$

where $x_j = 1$ if clique V_j is chosen, 0 otherwise; $a_{ij} = 1$ if $P_i \in V_j$, 0 otherwise. Then, the following algorithm solves (P2) optimally.

ALGORITHM 2

- Step 1 : Arrange all the distinct r_{ik} values, $i \neq k$, into an increasing sequence: $r_1 < r_2 < \dots < r_t$.
- Step 2 : By a binary search on the list $R = \{r_1, r_2, \dots, r_t\}$ find the smallest r in the list, i.e. r^* , for which the optimal value of $SCP(r)$ is p .

Step 3 : Output r^* and the cliques V_j^* corresponding to $x_j = 1$ in the optimal solution to $SCP(r^*)$.

An optimal partition $\alpha^* = (M_1^*, \dots, M_p^*)$ can easily be generated from the cliques V_j^* , $j=1, \dots, p$. Optimal centers $C_1, \dots, C_p$ are obtained by Algorithm 1 solving the 1-center problems associated with an optimal partition.

The above algorithm requires that all distinct cliques of $G(r)$ be found and then $SCP(r)$ solved, and this at most for $\log_2 t \approx 2\log_2 m - 1$ different values of r . It can be proved that the number of cliques (number of columns of $SCP(r)$) is bounded by m^n , so that the algorithm can only be used for very small values of n (in Aneja et al.⁷ problems with $n=2$ and $m \leq 100$ are optimally solved). For relatively large problems however, only heuristic algorithms can be used.

4. HEURISTIC ALGORITHMS

Most heuristic methods for other location-allocation problems can be modified to obtain approximate solutions to (P2) for any norm N . For instance, the "alternate location-allocation method" (Cooper¹³) and the "exchange method" (Spath¹⁴), both suggested for the p -median problem, can be used; the difference is that, due to the different objective function, it is necessary to use a 1-center algorithm instead of a 1-median algorithm as a subroutine. These types of method become more

efficient for a weighted Tchebycheff norm since it is enough to use the evaluation of $r(M_j)$ given in (1) as a subroutine during the iterative procedure instead of a 1-center algorithm, which is used only at the end to obtain the centers (Algorithm 1 can be used to obtain the centers from the generated partition).

Moreover, for (P2) some δ -approximation polynomial algorithms have been given , which generate a partition α such that $r_\alpha \leq \delta r^*$ for some given value of δ . For instance , Dyer and Frieze¹⁰ describe a simple heuristic for the p-center problem with an arbitrary metric which is a δ -approximation for $\delta = \min \{ 3 , w \}$, where w is the maximum ratio between the weights of the points in M , and Plesnik¹⁵ gives a 2-approximation polynomial algorithm ("best possible") for the p-center in graphs; both can be used to obtain approximate partitions of (P2) for any norm N from which the centers can be obtained by Algorithm 1.

We propose a new heuristic algorithm to generate an approximate partition ,based on properties 1 and 2 , as follows :

ALGORITHM 3

- Step 1 : - Generate any partition $\alpha_0 = (M_1^0, \dots, M_p^0)$.
- Calculate r_{α_0} .
 - Set $R = \{ r_{ik} : r_{ik} \leq r_{\alpha_0} , i \neq k \}$.

- Make a list L arranging all the distinct r_{ik} values of R into an increasing sequence.
- Set $r_0 = r_{\alpha_0}$ and $\alpha = \alpha_0$.

Step 2 : - If $|L| > 1$ take r as the median value of L . Make unlabelled all points in M . Go to step 3.

- Else $L = \{ r_0 \}$, STOP. Output r_0 and α .

Step 3 : - Choose an unlabelled point P_t of maximum weight and set $M_t' = \{ P_i : r_{it} \leq r \}$. Label all the points in M_t' .

- If all points in M are labelled go to step 4. Else go to step 3.

Step 4 : - Set $\alpha' = \{ M_t' : M_t' \text{ is generated in step 3} \}$.

- If $|\alpha'| \leq p$ set $L \leftarrow \{ r_{ik} \in L : r_{ik} \leq r \}$, $\alpha = \alpha'$ and $r_0 = r$. Else set $L \leftarrow \{ r_{ik} \in L : r < r_{ik} \}$.
- Go to step 2.

The complexity of step 1 is $O(m^2 \log m)$. Step 2 to step 4 is a binary search in the set L that finds the minimum value in L for which the partition α' generated in step 3 satisfies $|\alpha'| \leq p$. As step 3 is $O(m^2)$ and the binary search is $O(\log m)$ the complexity of the algorithm is $O(m^2 \log m)$.

Property 4

The above algorithm is a 2-approximation polynomial algorithm for (P2), and gives a lower bound r_0 of r^* .

Proof :

First, it will be shown that $r(M_t') \leq 2r$ for each set M_t' generated in step 3 and any $r \in L$. From the definition of M_t' and

(1), this is true if $|M_t'| = 1$ or $|M_t'| = 2$. Otherwise, let $P_i, P_k \in M_t'$, $i, k \neq t$, as $r_{it} \leq r$, $r_{kt} \leq r$, $w_i \leq w_t$ and $w_k \leq w_t$, then

$$\begin{aligned} r_{ik} &= w_i w_k / (w_i + w_k) N(P_i - P_k) \leq \\ &= w_i w_k / (w_i + w_k) (N(P_i - P_t) + N(P_k - P_t)) \leq \\ &= w_i w_k / (w_i + w_k) \left(\frac{(w_i + w_t)}{w_i w_t} + \frac{(w_k + w_t)}{w_k w_t} \right) r \leq \\ &= \left(1 + 2 \frac{w_i w_k}{(w_i + w_k) w_t} \right) r \leq \\ &= \left(1 + \max \{w_i, w_k\} / w_t \right) r \leq 2r, \end{aligned}$$

therefore from (1) it follows that $r(M_t') \leq 2r$.

Let $r \geq r^*$, that is there is a partition $\alpha = (M_1, \dots, M_p)$ with $r_\alpha \leq r$. Let the point P_t in an iteration of step 3 be in M_j ; as $r_{it} \leq r$ for any $P_i \in M_j$ it follows that $M_j \subset M_j'$. Then the partition α' generated in step 3 satisfies $|\alpha'| \leq p$. Consequently, if the partition α' generated in step 3 satisfies $|\alpha'| > p$ then $r < r^*$. Therefore $\min \{ r_{ik} : r_{ik} \in L \}$ is a lower bound of r^* in each iteration of step 4.

At the end $L = \{ r_0 \}$ and $|\alpha| \leq p$, as then $r_0 \leq r_\alpha \leq 2r_0 \leq 2r^*$ it follows that Algorithm 3 is a 2-approximation polynomial algorithm and r_0 is a lower bound of r^* .

The proposed algorithm is then "best possible" since for any $\delta < 2$ it is known that the existence of a δ -approximation polynomial algorithm would imply that $P=NP$ (Hsu and Nemhauser¹⁵, Hochbaum and Shmoys¹⁶). It generates a partition α , as happens with the other algorithms mentioned, so that the centers can be obtained from the generated partition by Algorithm 1. Also, an upper bound on the related error $(r_\alpha - r^*) / r^*$ can be obtained as $(r_\alpha - r_0) / r_0$.

REFERENCES

1. N. MEGIDDO and K.J. SUPOWIT (1984) On the Complexity of Some Common Geometric location Problems. SIAM J. on Comp.,13,182-196.
2. Z. DREZNER (1984) The p-center problem , Heuristics and Optimal algorithms.Jour. Oper.Res. Society, 35,741-748.
- 3.J. VIJAY (1985) An algorithm for the p-center problem in the plane.Transportation Science,19,235-245.
4. R. CHEN and G.Y. HANDLER (1987) Relaxation Method for the Solution of the Minimax Location-Allocation Problem in Euclidean Spaces. Naval Research Logistics,34,775-788.
5. C.D.T. WATSON-GANDY (1984) The Multi-facility Minimax Weber Problem. Eur. Journ. on Oper. Res.,18,44-50.
6. Z. DREZNER (1987) On the rectangular p-center problem. Naval Research Logistics,34,229-234.
7. Y.T. ANEJA,CHANDRASEKARAN and K.P.K. NAIR (1988) A note on the m-center problem with rectilinear distances. Europ. Journ. on Oper. Res.,35,118-123.
8. H. SPATH (1985) Cluster Dissection and Analysis. Ellis Horwood Limited, Chichester, England.
9. H. A. EISELT and G. CHARLESWORTH (1986) A note on p-center problems in the plane. Transportation Sci. 20,130-133.
10. M.E. DYER and A.M. FRIEZE (1985) A simple heuristic for the p-center problem. Oper. Res. Letters ,3,285-288.
11. B. PELEGRIN (1988) The p-center problem under bidirectional polyhedral norms. Proceedings III Meeting E.W.G. on Locational Analysis ,151-169, Sevilla (Spain).

12. P.M. DEARING and R.L. FRANCIS (1974) A minimax location problem on a network. *Transportation Science*, 8, 333-343.
13. L. COOPER (1964) Heuristic methods for location-allocation problems. *SIAM Review*, 6, 37-52.
14. H. SPATH (1977) Computational Experiences with the Exchange method. *Eur. Journ. on Oper. Res.*, 1, 23-31.
15. W.L. HSU and G.L. NEMHAUSER (1979) Easy and Hard bottleneck location problems. *Discrete Appl. Math.*, 1, 209-216.
16. D.S. HOCHBAUM and D.B. SHMOYS (1985) A best possible heuristic for the k-center problem. *Math. Oper. Res.*, 10, 180-184.

A VARIATION OF GRAHAM'S LPT-ALGORITHM

Hannes HASSLER

TU Graz

**Institut für Angewandte Informatik
und Kommunikationstechnologie
Klosterwiesgasse 32, A-8010 Graz, Austria**

and

Gerhard J. WOEGINGER

TU Graz

**Institut für Mathematik B
Kopernikusgasse 24, A-8010 Graz, Austria**

ABSTRACT

We consider the problem of scheduling a set of n jobs nonpreemptively on m identical machines. The goal is to minimize C^* , the maximum completion time over all jobs. This problem is known to be NP-complete. We present a class of new approximation algorithms that have low running time $O(n \log m)$ and guarantee the worst case performance of Graham's famous LPT-Algorithm.

1 Introduction

We consider the problem of scheduling a set of n jobs nonpreemptively on m identical machines. The goal is to minimize the maximum completion time over all jobs. We denote the maximum completion time in an optimal schedule by C^* . In 1969, Graham [2] suggested two simple heuristics for this problem. *List Scheduling*, the first heuristic, lists the jobs in any order and then assigns them in this order to the machines as they become free. The *Longest Processing Time* algorithm (or LPT, for short) first sorts the jobs by decreasing processing time and then applies list scheduling to them. List scheduling gives a schedule with maximum completion time at most $2 - \frac{1}{m}$ times C^* and LPT gives a schedule with maximum completion time at most $\frac{4}{3} - \frac{1}{3m}$ times C^* . Both bounds are tight. For the running time, it is easy to see that list scheduling takes $O(n \log m)$ time and LPT takes $O(n \log n)$ time. Thus, if we take the number of machines to be constant, the running time of list scheduling outperforms LPT by a $\log(n)$ -factor.

In this note, we present an algorithm that has the same worst case performance guarantee as the LPT-algorithm and has the low running time of list scheduling. More precisely, we construct a sequence of algorithms that lie between list scheduling and LPT and determine their exact worst case performance. The behaviour of the algorithms depends on the number m of machines and on some parameter k , $0 \leq k \leq 2m$. For $k \geq 2m$, our algorithm is at most the LPT-factor $\frac{4}{3} - \frac{1}{3m}$ off the optimum.

2 The Algorithm and the performance guarantee

Let J be a set of n jobs that must be scheduled to m machines and let k be some integer, $0 \leq k \leq n$. We may assume $m \leq n$ as otherwise the problem is trivial. Our approximation algorithm $A(m, k)$ proceeds as follows.

- (1) Determine the set J^k of the k jobs in J with largest processing times. Form a list that contains the jobs in J^k in arbitrary order and append the remaining jobs in $J - J^k$ to it.
- (2) Apply list scheduling to the constructed list.

It is well known that determining the k^{th} -largest processing time can be done in $O(n)$ time (see e.g. [1]). Thus, finding J^k takes $O(n)$ time and appending the remaining jobs takes $O(n)$ time, too. Step 2 is standard $O(n \log m)$ list scheduling, and so we get an overall time complexity of $O(n \log m)$ for algorithm $A(m, k)$.

Now let C_k^H denote the maximum completion time in a schedule constructed by the algorithm $A(m, k)$ and let C^* denote the maximum completion time in an optimal schedule. In Section 3, we will prove the following tight worst case bounds.

Range of k	Worst case Ratio of C_k^H/C^*	Ratio for $m = 2$
$0 \leq k \leq (m-1)/2$	$2 - 1/(m-k)$	$3/2$
$(m-1)/2 \leq k \leq m-1$	$2 - 2/(m+1)$	$4/3$
$m \leq k \leq (4m-1)/3$	$1.5 - 1/(4m-2k)$	$5/4$
$(4m-1)/3 \leq k \leq 2m-1$	$1.5 - 3/(4m+2)$	$6/5$
$2m \leq k$	$1.33 - 1/(3m)$	$7/6$

That means, the worst case ratio remains constant for all $k \geq 2m$. However, the running time deteriorates and so in practice we will not use large values for k . Numerical experiments demonstrated that the average performance of algorithm $A(m, k)$ improves with increasing k . Figure 1 graphically illustrates the worst case behaviour of $A(m, k)$ for $k \geq 0$.

3 The Proofs

Let m denote the number of machines and let k be the parameter of algorithm $A(m, k)$ as defined in the preceding section. As k is fixed we write C^H for C_k^H to simplify notation. By x we denote the job that is treated last by the algorithm.

First we consider the cases $0 \leq k \leq m-1$. The jobs in J^k all have length at least x . W.l.o.g. the jobs in J^k are scheduled to the last k machines. Consequently, at the time before job x is scheduled by the algorithm, the total load of the first

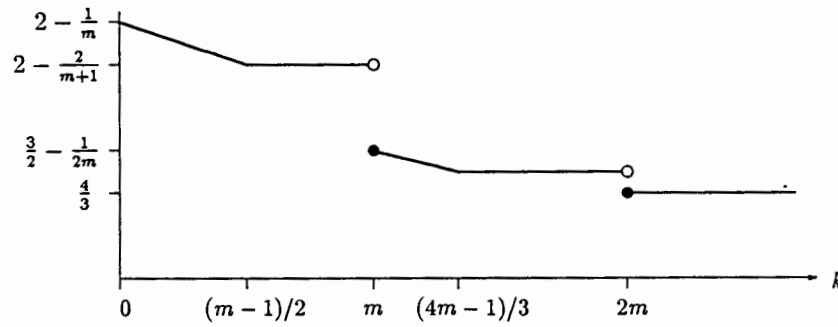


Figure 1: The worst case performance of algorithm $A(m, k)$

$m - k$ machines is at most $mC^* - (k + 1)x$. Job x is scheduled to the machine with smallest load at this time and the smallest load is less or equal the average load of the first $m - k$ machines. This gives

$$C^H \leq (mC^* - (k + 1)x)/(m - k) + x = (mC^* + (m - 2k - 1)x)/(m - k) \quad (1)$$

For $0 \leq k \leq (m - 1)/2$, the coefficient of x in inequality (1) is nonnegative. Using $x \leq C^*$, we get that $C^H/C^* \leq 2 - 1/(m - k)$ holds, the claimed result for $0 \leq k \leq (m - 1)/2$. For $(m - 1)/2 \leq k \leq m - 1$, the coefficient of x in inequality (1) is nonpositive. In this case for $x \geq mC^*/(m + 1)$, inequality (1) immediately implies $C^H/C^* \leq 2 - 2/(m + 1)$. For $x \leq mC^*/(m + 1)$, we use another averaging argument. At the time before job x is scheduled, the total load over all machines is less or equal $mC^* - x$. Hence

$$C^H \leq (mC^* - x)/m + x = C^* + (m - 1)x/m \leq (2 - 2/(m + 1))C^* \quad (2)$$

holds, and we have finished the proof for cases $0 \leq k \leq m - 1$.

Next, we treat the cases $m \leq k \leq 2m - 1$. We define $k' = k - m \geq 0$. First we observe that

$$x \leq C^*/2 \quad (3)$$

must hold. If x is greater than $C^*/2$, there are $k + 1 > m$ big jobs of length greater $C^*/2$. Then in the optimum schedule there exists a machine with two big jobs and load greater C^* , a contradiction. After the algorithm scheduled the $m + k'$ jobs in J^k , there are at least $m - k'$ machines that received exactly one job of J^k (Every machine gets at least one job. If there are only $m - k' - 1$ machines with one job, the remaining $k' + 1$ have at least two jobs each, and we get a total job number of at least $m + k' + 1 > k$ jobs in J^k , a contradiction). We rearrange the sequence of machines such that the $m - k'$ machines with one job of J^k range first. The last k' machines process the remaining $2k'$ jobs. Consequently, at the time before job x is scheduled by the algorithm, the total load of the first $m - k'$ machines is at most $mC^* - (2k' + 1)x$ and

$$C^H \leq (mC^* - (2k' + 1)x)/(m - k') + x = (mC^* + (m - 3k' - 1)x)/(m - k') \quad (4)$$

holds. For $0 \leq k' \leq (m - 1)/3$, the coefficient of x in the right hand side of inequality (4) is nonnegative. Using (3), we derive $C^H/C^* \leq 1.5 - 1/(2m - 2k')$ and we are finished. For $(m - 1)/3 \leq k' \leq m - 1$, the coefficient of x in the right hand side of inequality (4) is nonpositive. We distinguish between the two subcases $x \leq mC^*/(2m + 1)$ and $x \geq mC^*/(2m + 1)$. In the case $x \geq mC^*/(2m + 1)$, we can use inequality (4) again and we get $C^H/C^* \leq 1.5 - 3/(4m + 2)$. In the case

$x \leq mC^*/(2m+1)$, we consider the time before job x is scheduled. The total load over all machines is less or equal $mC^* - x$. We derive that

$$C^H \leq (mC^* - x)/m + x = C^* + (m-1)x/m \leq (1.5 - 3/(4m+2))C^* \quad (5)$$

holds, and the cases $m \leq k \leq 2m-1$ are finished, too.

Finally, we assume that $2m \leq k$ holds. It is easy to see that $x \leq C^*/3$ must hold. Otherwise, there are $k+1 \geq 2m+1$ jobs of length $> C^*/3$. In the optimum schedule, at least one machine must contain at least three of these jobs; this machine would have load $> C^*$. Therefore, we may reuse inequality (5) and plugging in $x \leq C^*/3$ instead of $x \leq mC^*/(2m+1)$, we derive $C^H/C^* \leq 4/3 - 1/3m$. This completes the proof of all worst case bounds claimed in Section 2. $\square$

To prove the tightness of our bounds, we exhibit the following five sets of examples. In each example sequence, the first k elements are the largest k elements in the sequence. Hence, algorithm $A(m, k)$ may skip step (1) and use the sequence exactly in the given ordering.

Example 1. ($0 \leq k \leq (m-1)/2$) First there are k jobs of length 1, then $(m-k)(m-k-1)$ small jobs of length $1/(m-k)$, and finally job x of length 1. By scheduling all small jobs on $m-k-1$ machines, we see that $C^* = 1$ holds. Our approximation algorithm equally schedules the small jobs on $m-k$ machines, and finally puts x on one of them. Hence, $C^H = 2 - 1/(m-k)$. $\square$

$\square$ ($(m-1)/2 \leq k \leq m-1$) Our list consists of $m-1$ big jobs of length $m/(m+1)$, m small jobs of length $1/(m+1)$ and the big job x of length $m/(m+1)$. The optimum solution simply matches every big job with a small job and gives $C^* = 1$. Our algorithm schedules all small jobs to one machine. Then job x makes $C^H = 2m/(m+1)$. $\square$

($m \leq k \leq (m-1)/3$) We use k big jobs of length $1/2$, then $(k-m)(k-m-1)$ small jobs of length $1/(2k-2m)$, and finally the big job x of length $1/2$. Similarly as in example 1, we see that $C^* = 1$ and that $C^H = 1.5 - 1/(4m-2k)$. $\square$

($(m-1)/3 \leq k \leq 2m-1$) The first $2m-1$ jobs are big jobs of length $m/(2m+1)$, then there are m small jobs of length $1/(2m+1)$ and job x of length $m/(2m+1)$. The optimum schedule assigns to each machine two big and one small job, the approximation schedule assigns all small jobs to one machine. Hence, $C^* = 1$ and $C^H = 1.5 - 3/(4m+2)$. $\square$

($2m \leq k$) Here we use Graham's lower bound example [2] for the

LPT-algorithm. We have jobs $p_1 \dots p_{2m}$ such that p_i has length $2m - \lceil i/2 \rceil$ and $p_{2m+1} = x$ has length m . It is easy to check $C^* = 3m$ and $C^H = 4m - 1$. $\square$

References

- [1] M.Blum, R.W.Floyd, V.R.Pratt, R.L.Rivest and R.E.Tarjan, Time bounds for selection, . Comp. yst. ciences 7, 1972, 448-461.
- [2] R.L.Graham, Bounds on multiprocessing timing anomalies, AM . Appl. Math. 17, 1969, 416-429.

MODELLING CORPORATE VULNERABILITY: YET ANOTHER EMPIRICAL ATTEMPT

D.J.E. BAESTAENS

**Erasmus University Rotterdam
Dept of Finance
Burg. Oudlaan 50, 3062 PA Rotterdam, The Netherlands**

ABSTRACT

This paper argues that potential weaknesses of MDA could be avoided by using the Mahalanobis Distance (d^2). The model's uniqueness allows for a more detailed analysis of corporate failure triggers (conversely shock absorbers) at the corporate level. Incidentally, the main input for the distance analysis, the sample's dispersion matrix, could be used to quantify the distance's sensitivity to changes in the (co)variances of so-called policy variables. Such analysis is likely to benefit central bank regulators (monetary & fiscal policy) as well as corporate management (interest rate sensitivity).

1. Introduction

Company failure is an emotionally event, unfortunately not really accounted for by a solid theoretical framework¹. As Barnes (1987) pointed out, researchers may well be testing the hypothesis that the event² may be explained by the statistical behaviour of ad hoc generated datasets.

The first attempts to quantify the bankruptcy process in the US were presented in the seminal Altman (1968,1977,1984) papers. Taffler (1977,1980,1982) was among the first to introduce and improve upon the Altman MDA-method³ in the UK corporate environment. These original contributions have since been the subject of elaborate testing procedures and of academic debate (see eg. Ezzamel and Mar-Molinero,1990; Karels and Prakash,1987) whereby the discussion focuses on the misspecification of the MDA model. Are researchers justified in applying a multivariate statistical technique while disregarding its restrictive assumptions ?

In theory, the use of MDA requires the prevalence of at least two easily identifiable and mutually exclusive groups of observations, multivariate normality of the cases under study, specification of the a priori probabilities of occurrence of each group and dispersion matrices' equality in case of linear discriminant analysis (LDA).

In this paper we assert that the first two requirements are very restrictive indeed. Because of these constraints the MDA method may not be the most appropriate method to describe and predict failure. Consequently, we will suggest a potential alternative to MDA. Finally, we will apply the alternative to a sample of UK firms from the Printing & Publishing Industry with a view to identifying entities deviating from the industry average.

2. Limitations of MDA

1. Group Composition

Although the first requirement does not usually constitute an obstacle, identifying distinct groups of failed and nonfailed firms may be more treacherous than it seems. How does one define failure when the failure process is conceptually nebulous ? A majority of studies appears to adopt some ex post empirical record as failure criterion, be it renegotiation with creditors, bond default, stock exchange delis-

¹ Donaldson's (1969) Financial Mobility Strategy attempted to integrate unanticipated events and cash flow patterns in a systematic manner.

² Or any other event, such as economic instability (Baestaens,1990)

³ Multiple Discriminant Analysis

ting or decision to file for reorganisation under Chapter XI of the Federal Bankruptcy Act. While analysis of the real world variable clearly indicates the category the firm belongs to, classification of the cases becomes a biased process as the failure category becomes unnecessarily stratified. To exemplify, we have not come across many MDA studies investigating the question why some firms are more successful than others in a given setting. Metaphorically, analysing failures may be like investigating the Black Death. What matters is not the detection of the pathogen by carrying out post-mortem analyses but the identification of the survivors' defence mechanisms. We fear that restraining the failure class sample to those entities that actually failed (whatever the failure definition) may lead to the repudiation of the evidence provided by those firms that were infected with the failure-disease but managed to get rid of the malady. We therefore prefer an a priori undefined approach to the ex post empirical classification routine. The advantages of the a priori approach are fourfold. First, we do not have to hunt for a specific, necessarily subjective (in the absence of a failure theory) failure criterion. Second, by pooling all firms we avoid the obligation of having to construct a non failed sample consisting of really healthy firms unlike Taffler (1982). Third, no tests are needed to verify the equality of the variance-covariance matrices across groups. Finally, our procedure possesses the implied benefit of not having to deal with the issue of determining a priori probabilities of the event occurrence.

2. Multivariate Normality

The normality condition sometimes appears to be interpreted as imperative (Eisenbeis, 1977; Karels and Prakash, 1987; Taffler, 1982) and sometimes as accessory (Barnes, 1982; Richardson and Davidson, 1983). This ambiguity may stem from the lack of a standard test for multivariate normality compounded by the fact that univariate normality does not guarantee multivariate normality. As non normality may be caused by the presence of outliers, the conventional statistical approach has always been to identify such outliers with a view to deleting⁴ or separating them from the rest of the data under study (Anscombe, 1960; Collett and Lewis, 1976; Hartwig and Dearing, 1979).

Our present approach on the contrary may be somewhat novel in the sense that outliers are treated as the major information source (Ezzamel and Mar-Molinero, 1990; Howell, 1989).

⁴ Sometimes euphemistically called Winsorization (changing an outlier's value to that of the closest non outlier) or Trimming (removal from the sample an equal number of the smallest and largest observations).

We assume that events in which either individual variables or the relationships between variables take unusual values or are distorted into unusual levels (as measured by statistical criteria) may reflect strain in or upon the economy, respectively the individual firm. This strain may or may not endure after the stress is removed³. Here we are interested in extreme states in their own right, though we accept that some of them will be due to spurious observations or chance events. Where there are many variables, as here, multivariate methods are well developed only for the normal distribution, as discussed by Bacon-Shone and Fung (1987), and this paper uses the Mahalanobis distance (d^2) and Hotellings T^2 as representative of such methods. Although d^2 assumes either a multivariate normal or elliptical distribution (Mitchell and Krzanowski, 1985), we are hoping that our data are not jointly normally distributed since we are actively seeking outliers from such a distribution. In this sense, we believe we are among the first to apply the Mahalanobis distance in its own right to the issue of identifying sick firms.

3. d^2 & T^2 : Alternatives to MDA ?

3.1. Presentation of d^2

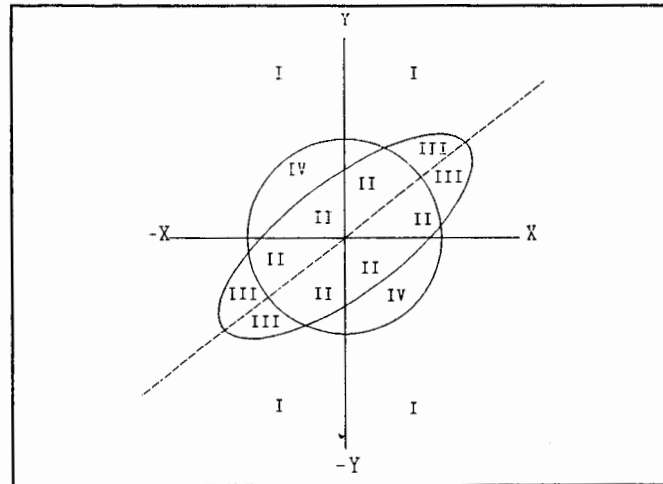
The joint distribution of normally distributed individual variables is often multivariate normal. Figure 1 shows a resulting ellipse of uniform probability density for two such variables X and Y, standardised to equal standard deviations.

A joint confidence region for X and Y is elliptical. This region is not the intersection of the two univariate confidence levels at the same significance level (circle in Figure 1). We call this circle an "uncorrelated" confidence region and points outside it "uncorrelated" outliers. Points in Regions I and II are respectively outliers and inliers for both the ellipse and the circle, whilst points in the Regions III are inliers to the ellipse but outliers to the square. Points in Regions IV are outliers to the ellipse but inliers to the square.

Conventional regression analysis does not yield confidence regions equivalent to the ellipse - projection into the X space makes predictions of Y conditional on X rather than absolute, so that the confidence region is hyperbolic and unbounded before X is observed, and "unusual" states of X are not recognised on observation. In general, the absence of a theoretical framework disallow researchers to give any particular set of X variables special status as independents.

³ In engineering disciplines, stress and strain denote distinct realities. Stress refers to the event that causes the distortion to happen while strain is the equivalent of the distortion.

Figure 1: Possible confidence regions using uncorrelated and correlated multivariate criteria.



To deal with this we note that points on the ellipse in Figure 1 share not only the same probability density, but also the same Mahalanobis Distance from the mean observation. The Mahalanobis Distance of a single multivariate observation from the mean observation of a sample of n observations can be estimated using:

$$d^2 = \sum_{i=1}^p \sum_{j=1}^p (x_i - \bar{x}_i) c^{ij} (x_j - \bar{x}_j) \quad (1)$$

where

$\bar{x}_i$ is the mean of the i th variable

and c^{ij} is the element in the i th row and j th column of the inverse of the variance-covariance matrix C^{-1} .

d^2 follows a Chi Squared distribution with p degrees of freedom (Manly, 1986). Figure 1 suggest that in correlated data sets this test will be the most useful and is likely to detect a set of outliers distinct from uncorrelated outliers. A chosen cut off value of the Mahalanobis Distance d^2 separates observations between Mahalanobis Outliers (Regions I plus IV in Figure 1) and Mahalanobis Inliers (Regions II plus III in Figure 1). We assume temporarily that a d^2 value has been chosen so as to give these regions convenient relative probabilities, and interpret them as follows. Points in Regions III are events where at least one variable is outside its individual confidence interval, but the joint value is not. Such events we call "structure preserving" in the sense that the expected correlation structure is pre-

served. In contrast, points in Regions IV we call "structure violating", because although neither variable is outside its individual confidence interval, the expected correlation of the two variables is violated. Events in Region I may or may not violate correlation structures.

3.2. Decomposition of the Mahalanobis distance

In a p dimensional observation each of the $p(p-1)/2$ pairs of variates may show either structure violating or structure preserving behaviour. It is desirable to have a means of inspecting this behaviour directly, and one which does not make use of an arbitrary d^2 to discriminate between Structure Violation and Structure Preservation.

We can express d^2 in equation (1) as the sum of all the elements of a $(p \times p)$ matrix R where $F_{ij} = x_i x_j V^{-1}_{ij}$ (Kendall and Stuart, 1983) so that

$$d^2 = \sum_{i,j} F_{ij} \quad (2)$$

Diagonal elements of F represent a weighting of a single squared deviation of the i th element (i.e. variable) from its mean, and off-diagonal elements represent a weighting of the product of the deviations of the i th and j th variables from their respective means. Since F is symmetrical we can conveniently combine off diagonal F_{ij} and F_{ji} in a single cell by defining T ($p \times p$) such that

$$\begin{aligned} \text{for } i = j, & \quad T_{ij} = F_{ij}, \\ \text{for } i < j, & \quad T_{ij} = 2F_{ij} \\ \text{for } i > j, & \quad T_{ij} = 0 \end{aligned}$$

Diagonal elements of T show the contribution to d^2 of the individual deviation of each variable in isolation, while subdiagonal elements show the specific contribution to d^2 from each variable's interaction with each single other variable. Diagonal elements in T and F are by definition positive, but off diagonal elements can take either sign.

A negative element T_{ij} for $i < j$ indicates that the joint deviations of x_i and x_j in this observation are less unlikely than the diagonal elements (their individual unlikelihoods) would suggest, and the fact that they are varying in the expected joint direction is "Structure Preserving". Such terms reflect the fraction of the joint variation of x_i and x_j that can be predicted from their correlation, and is akin to the "sum of squares explained" in ANOVA. They may correspond to events in Regions II and III of Figure 1 for the x_i and x_j concerned. Conversely a positive value for T_{ij} for $i < j$ indicates that an expected positive or negative correlation has been reversed, and this unexpected joint state of x_i and

x_j we can call "Structure Violating". It can correspond to an event in Regions IV or I (for variables x_i and x_j only).

A zero value of T_{ij} can occur for $i < j$ if x_i and x_j are uncorrelated in the sample as a whole ($T_{ij} = F_{ij} = 0$ for i not equal to j), or if standardised x_i or x_j or both are close to zero in this observation. Such observations are Structurally Neutral, or Uninformative.

3.3. Matrix simplification of F and T to F_s and T_s

So far, the triangular matrix T represents nothing more than a rearrangement of the full contribution matrix F . The value for d^2 was not at all affected. As matrices F and T in section above contain $p(p-1)/2$ different entries per observation, we attempted to simplify these matrices by setting to zero all elements in F (T) whose absolute value was smaller than a filter value s . The resulting matrix can be called F_s (T_s), and the approximated Mahalanobis Distance is d_s^2 , where

$$d^2 \approx d_s^2 = \sum_{i,j=1}^p F_{(s)ij} \quad (3)$$

and

$$d^2 \approx d_s^2 = \sum_{i,j=1}^p T_{(s)ij} \quad (4)$$

We used no theory to set s but investigated the sensitivity of d_s^2 to s , and selected the largest value of s for which d_s^2 seemed a close and stable approximation to d^2 . This simplified interpretation as well as providing a heuristic to avoid over interpreting the large noise content of a p dimensional observation.

3.4. Classification of entries in F_s

Given the simplest acceptable F_s (or T_s), we sorted the variables in descending order of the joint net contributions to d_s^2 (that is, the column, respectively row, totals of F_s). This partitioned F_s into three regions: Block A, where all totals were positive, Block B where the sums were zero and Block C where all aggregates were negative. We call all variables with nonzero diagonal entries Main Variables, since they matter in their own right.

Block A variables contribute positively to d^2 and are therefore acting in a Structure Violating way. These variables can be observed to reinforce d^2 in their own right (diagonal entry larger than zero resulting in the classification of this variable as Main Variable) and/or in interaction with other variables (off-diagonal entries larger than zero). The contribution of Block B variables remains neutral or uninformative as all diagonal and off-diagonal elements are set to zero. While the net contribution by Block C variables

can be classified as structure preserving (i.e. distance reducing), their individual contributions are never negative (diagonal elements are weighted squared deviations from respective means) implying that the variable interactions must more than offset these individual excursions. Again a nonzero positive deviation results in a classification as a Main Variable.

3.5. Hotelling's T^2 to search for longer lasting outlier episodes

Mahalanobis outliers for large enough d^2 are rare in the sample. They may at times be caused by chance events, or invalid measurements, or valid measurements of relationships that have no substantive economic importance. However they may also reflect important structural effects. For example the observed overall correlation may be due to forces that tend to suppress certain joint states of the variables. If the system is driven suddenly into a "non favoured" state by some shock or stress, and the corrective forces do not take full effect within a single sampling interval lagged effects may occur.

If such lagged effects are present, some Mahalanobis outliers will form part of extended episodes, in which extreme states in p space are only gradually approached and/or gradually departed from. While the squared Mahalanobis Distance was convenient for testing individual outlier observations, we used the Hotelling's T^2 to test whether the mean of a group of p dimensional observations differs from the mean of a second group (the remaining observations) assumed to have the same covariance matrix. For a null hypothesis of equal sample means the relevant test statistic reduces to an F distribution with p and $(n-p-1)$ degrees of freedom, in the notation of equation (5) below (Manley, 1986; Krzanowski, 1988).

$$T^2 = n_1 n_2 / (n_1 + n_2) \sum_{i=1}^p \sum_{j=1}^p (\bar{x}_{1i} - \bar{x}_{2i}) c^{ij} (\bar{x}_{1j} - \bar{x}_{2j}) \quad (5)$$

where

$\bar{x}_{1i}$ = the mean of the i th variable for group 1

p = number of variables

n = pooled sample size or $(n_1 + n_2)$, and

c^{ij} = the element in the i th row and j th column of the inverse of the pooled within group covariance matrix

In order to avoid bias by each outlier itself, we omitted it from each putative episode. We searched by trial and error for subsets of the total sample, contiguous with but not including the outlier(s), which differed significantly from the total sample.

Clearly there is some risk of "Data Mining" in such a search, and we have not attempted to derive an exact correction for it. Explicitly dynamic methods were not used because of the scarcity of degrees of freedom, and because we did not wish to assume uniform simple dynamics throughout the sample.

4. Summary of hypotheses

1. Extreme values or Outliers will be found in economic data sets which are wholly or mainly multivariate normal in their distribution.

2. Mahalanobis Distances of outliers (or inliers) can be decomposed to show the contributions made by each individual variable and by each pair of variables. The latter can in turn be divided into "Structure Preserving", "Structure Violating" and "Structurally Neutral" behaviour for this pair of variables. No specific predictions are made, but the structure of actual outliers are assumed to be of interest to economic modellers .

3. Mahalanobis Outliers will sometimes be a part of longer Dynamic episodes, in which either the approach to or the retreat from an outlier value in Mahalanobis space spreads over several sampling intervals.

4. Multivariate distributions of data on an economy will not be static, but will show, in addition to outliers, signs of sustained structural change (perhaps corresponding to intuitively identifiable periods of economic history). Such secular changes will increase the scatter of a fitted static distribution, and so reduce the power of tests for outliers, but Mahalanobis extreme values will still occur, and outliers strong enough to be detectable in these conditions may have substantive meaning.

5. Findings

5.1. Variable and Industry Selection

We selected a sample of 13 companies listed on the LSE and for which the FT provides daily financial coverage under the heading Newspapers & Publishers. Our sample constitutes about 34% of the total industry with strong emphasis (about 60%) on the newspapers & periodicals segment. The companies are given in table 1.

Table 1: Selected Companies and Code

1 Black (A. & C.)	BLAC
2 Maxwell Comm. Corp.	MAXC
3 Pearson	PSON
4 Trinity Int. Hld.	TRIN
5 Utd. Newspapers	UNWS
6 News Int. Spec. Div.	NEWS
7 Portsmouth & Sund.	PSUN
8 Reed International	REED
9 Haynes Pub.	HYNES
10 Bristol Eve. Post	BRTL
11 EMAP	EMAP
12 News Corp.	NEWSC
13 Daily Mail "A"	DMGT

Annual data from 1984 to 1989 (and where possible 1990) were compiled from Datastream. The Printing & Publishing industry was selected because of its alleged homogeneity in terms of cost structure and of the press coverage some industry members were/are receiving. The small number of companies (in absolute sense) excluded every attempt to restrict the sample to only those companies with the same year end. We agree with Gonedes (1973) that such restriction may be viewed as a sample stratification in the sense that the selected companies may share some characteristics that differ from those companies with different year ends.

Since we are interested in the health (or vulnerability) of individual firms over time relative to the industry considered, we pooled the cross-sectional data on individual firms resulting in 82 observations⁴.

Raw data were collected in seasonally adjusted form and where appropriate differenced to remove time trends. Table 2 lists the variables and codes. It can also be seen that most variables depart from a normal distribution on the Lilliefors test (1967), a variant on the Kolmogorov-Smirnov test when the population parameters are unknown, at the 5% significance level.

⁴ 13 firms times 6 years (84 to 89) + 4 firms times 1 year (since these firms report in April, accounting information for 1990 was available at the time of our data collection).

Table 2: Selected Variables

No	CATEGORY	VARIABLE	CODE	K-S TEST
1 COMPANY SPECIFIC				
1	PROFITABILITY	OPERATING MARGIN	PROFMAR	YES
2	CAPITAL STRUCTURE	BORROWING RATIO	BORRAT	NO
3	LIQUIDITY	CASH / CUR LIABILITIES	CASHCUR	NO
4	TURNOVER RATIO	SALES / NET CUR ASSETS	TURNCUAS	NO
5	PRODUCTIVITY	TAX / PRETAX PROFIT	TAXRAT	YES
2 PRINTING & PUBLISHING INDUSTRY RELATED				
6	EARNINGS INDEX	AVG EARNINGS EMPLOYEES P&P INDUSTRY TO AVG EARNINGS MANUFACTURING IND	EARNRAT	NO
7	TRADE TERMS	PPI OUTPUT OF PAPER PRODUCTS TO PPI MATERIALS PURCHASED BY IND	TRADET	NO
8	PRODUCTIVITY INDEX	OUTPUT INDUSTRY / OUTPUT MANUF IND VOLUME TERMS	RELOUTP	NO
3 ECONOMY WIDE				
9	MARKET CONFIDENCE	SP / FT ALL SHARE INDEX	SPFTA	NO
10	NATIONAL OUTPUT	GDP (SA, PERCENTAGE CHANGES)	GDP	YES

The bad results on the normality tests , while not very encouraging, must not be dramatised.

First, a non normal distribution represents a well known property of most financial ratios (Karels and Prakash, 1987; Lev and Sunder, 1979). Following Barnes (1982) we did therefore not attempt to apply standard Box and Cox transformations for right skewed data using the logarithm, square roots and cubic roots. Negative data values prevented the use of these transformations. Of course we could have added a constant where necessary or we could have applied the reciprocal transformation but we feared that too much data manipulation might have destroyed the ratios' informational content. The univariate non normality, of course, may have an adverse impact on the presence of multivariate outliers implying the need to consider the them with great caution, although univariate normality by itself does not guarantee multivariate normality.

Second, assuming the homogeneity of our sample, it may be argued that for a large sample such as ours a small violation of the normality assumption is sufficient for it to be rejected.

Unlike most MDA studies, we attempted to incorporate three possible determinants of corporate well-being and their interaction effects in one dataset. A firm's financial health can be viewed as a function of firm specific, industry related and/or economy wide elements. Obviously, each of these components could be emulated by numerous variables and ratios. In order to select a meaningful variable set for each component, factor analysis (PCA) may be called for to reduce the number of relevant variables. Here, we selected the variables on the basis of their perceived popularity in the literature.

5.2. Hypothesis 1

Table 3 shows the set of five outliers all at 1% significance. On the null hypothesis one might expect one or two such outliers in a sample of 82 observations.

Table 3: Identified Outlier Set

Case #	CODE	d ²	d
62	585	79.34	8.91
68	1085	57.33	7.57
67	1385	49.16	7.01
76	684	31.53	5.62
64	785	26.61	5.16
79	984	20.99	4.58
66	985	20.11	4.48
70	1285	17.24	4.15
69	1185	17.06	4.13
72	284	16.20	4.02

Cut-off value at the 1% level:
 >> 23.11 [χ^2 .01, 10]

Key to CODE:
 The last two digits refer to YEAR
 The first (two) digit(s) refers to COMPANY NUMBER in Table 1

All significant outliers (and for that matter the first ten outliers) refer to either 1984 or 1985. The nonuniform distribution of the outliers over time may indicate that the industry or parts thereof was under severe strain during that period. It would be useful to check the evidence of this assumed strain by examination of the relevant annual reports of the identified firms. Hypothesis 4 (see 18) appears to confirm our expectation of a nonstationary mean for the Mahalanobis model implying a structural break within the economy or industry under investigation. To the extent the break does not alter the estimated dispersion matrix, the model's validity remains safeguarded.

As aforementioned, the outliers may represent real phenomena or invalid measurements. The findings for hypothesis 2 on 16 may suggest measurement problems regarding the highest d^2 values.

Figure 2 and Figure 3 show all outliers over time ranked by firm. It can be seen that the outlier for firm 7 (PSUN) appears to be part of a longer dynamic episode in which both the approach to and the retreat from the outlier value spreads over several sampling intervals. In contrast, the values for firms 10 (BRTL) and 13 (DMGT) seem to belong to an episode in which the approach to, respectively the retreat from the outlier marks a rather abrupt process.

Figure 2: Distances Firms 1 to 7

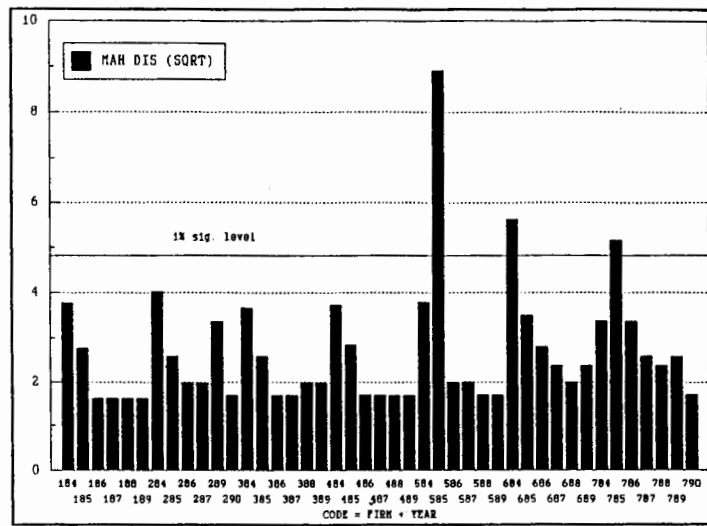
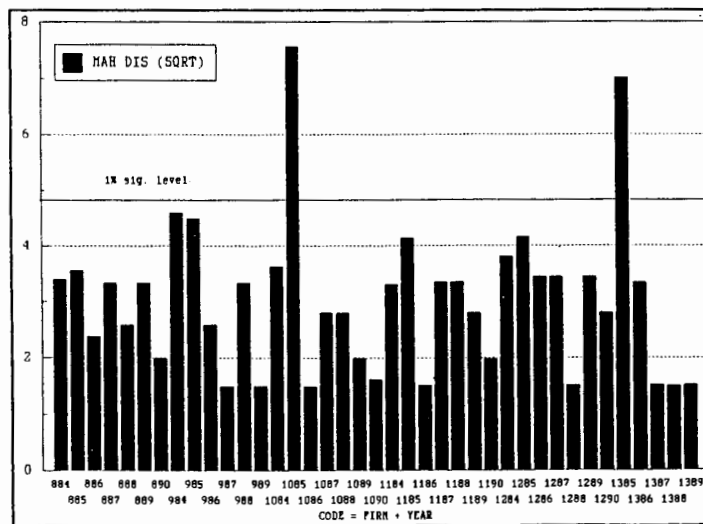


Figure 3: Distances Firms 8 to 13



5.3. Hypothesis 2: Internal Structure of Outliers

We now decompose the d^2 values by rearranging the variables in such a way that we gain insight in their individual and joint behaviour.

Since our variable set is not very large ($p=10$), the 45 pairs can be evaluated without necessarily simplifying the contribution matrix (F). Table 4 shows contribution matrices ranked by the net variable contribution to d^2 and d^2 . (ie. column, respectively row, labelled sum) for firm 7 (PSUN) for 1985.

The d^2 contribution matrix shows that all variables are active, that is display nonzero entries. This finding which was confirmed by the analysis of all F matrices supports our claim that it is dangerous to rely on univariate criteria when assessing corporate vulnerability. Moreover, the finding demonstrates the weakness of those techniques such as MDA that attempt to identify failure trigger points by the sole use of internal, company specific factors. External (industry and economy) factors are not to be omitted from the analysis.

It can be seen that both d^2 and d^2 appear to be driven (Block A) by the solitary contribution of the tax ratio variable and the joint contributions of the GDP variable. Structure preserving behaviour (Block C) is displayed by the trade terms, turnover to net current assets and liquidity ratios. The remaining variables do not contribute in any way to d^2 , and are therefore classified as Block B variables. All Block A variables are Main variables, i.e. important in their own right whereas one Block C variable (Trade terms) figures as a nonmain variable (i.e. unimportant in its own right).

The decomposed contribution matrix F could be used as input for the selection of appropriate "policy variables" by regulatory authorities or industry watchdogs. To understand the variables joint behaviour within the system, estimated variance and covariances could be altered to quantify the effects (that is, changes in d^2) of potential policy changes. The aim would then be to minimise d^2 or to bring it below some agreed reference value. To exemplify, we doubled the variance of the TAXRAT variable mentioned in Table 4 (authorities indirectly control this variable through the tax rate). Table 5 gives the result.

Table 5: Quantification of Intervention Impact PSUN 1985

	TAXRAT	RELOUTP	GDP	EARNRAT	SPFTA	PROFMAR	TURNCUAS	BORRAT	TRADET	CASHCUR	SUM
TAXRAT	7.84	-1.94	1.24	-0.25	-0.08	0.48	0.01	0.01	-0.65	-1.01	5.64
RELOUTP	-1.94	6.14	1.07	0.12	0.63	-0.91	0.01	0.01	-0.43	-0.09	4.60
GDP	1.24	1.07	2.05	-0.22	-0.04	0.34	0.01	0.01	-0.20	-0.75	3.49
EARNRAT	-0.25	0.12	-0.22	1.00	0.21	-0.39	-0.00	-0.01	0.28	0.01	0.74
SPFTA	-0.08	0.63	-0.04	0.21	0.66	-0.50	0.00	0.01	-0.21	-0.17	0.51
PROFMAR	0.48	-0.91	0.34	-0.39	-0.50	1.49	-0.00	0.00	0.12	-0.46	0.16
TURNCUAS	0.01	0.01	0.01	-0.00	0.00	-0.00	0.00	0.00	-0.00	-0.01	0.02
BORRAT	0.01	0.01	0.01	-0.01	0.01	0.00	0.00	-0.01	-0.01	-0.01	0.01
TRADET	-0.65	-0.43	-0.20	0.28	-0.21	0.12	-0.00	-0.01	0.89	0.17	-0.04
CASHCUR	-1.01	-0.09	-0.75	0.01	-0.17	-0.46	-0.01	-0.01	0.17	1.99	-0.33
SUM	5.64	4.60	3.49	0.74	0.51	0.16	0.02	0.01	-0.04	-0.33	14.81

Due to the intervention, this observation has become an inlier. Note that the TAXRAT variable is still the most important Main variable. Obviously, to fully assess the impact of government's intervention, it is necessary to compute all d^2 values again to ascertain whether or not outliers have become inliers or vice versa. We hope our simplified example has demonstrated the model's potential usefulness in identifying target, intermediary and policy variables without having recourse to monetarist or post-keynesian theories.

Potential measurement problems were signalled above. Analysis of the largest outlier (CODE 585:UNWS 1985, $d^2=79$) revealed that almost all of the distance value could be related to the individual contribution of one variable, the borrowing ratio. Table 6 gives the decomposition.

Table 6: Decomposition of the largest outlier, UNWS 1985

	BORRAT	PROFMAR	GDP	TURNCUAS	RELOUTP	TRADET	TAXRAT	EARNRAT	CASHCUR	SPFTA	SUM
BORRAT	88.96	0.11	0.22	0.12	-0.17	-1.90	-0.54	-3.56	-1.41	-1.73	80.10
PROFMAR	0.11	0.21	0.02	0.00	0.08	0.00	-0.05	-0.30	-0.15	0.24	0.17
GDP	0.22	0.02	0.03	0.01	-0.00	-0.05	-0.04	-0.07	-0.08	0.01	0.05
TURNCUAS	0.12	0.00	0.01	0.02	-0.01	-0.02	-0.01	-0.01	-0.03	-0.02	0.04
RELOUTP	-0.17	0.08	-0.00	-0.01	0.23	-0.01	-0.11	-0.10	-0.07	0.15	-0.00
TRADET	-1.90	0.00	-0.05	-0.02	-0.01	0.78	0.18	0.55	0.22	0.21	-0.04
TAXRAT	-0.54	-0.05	-0.04	-0.01	-0.11	0.18	0.20	0.15	0.20	-0.03	-0.07
EARNRAT	-3.56	-0.30	-0.07	-0.01	-0.10	0.55	0.15	3.56	0.12	-0.51	-0.17
CASHCUR	-1.41	-0.15	-0.08	-0.03	-0.07	0.22	0.20	0.12	0.92	0.11	-0.18
SPFTA	-1.73	0.24	0.01	-0.02	0.15	0.21	-0.03	-0.51	0.11	1.02	-0.56
SUM	80.10	0.17	0.05	0.04	-0.00	-0.04	-0.07	-0.17	-0.18	-0.56	79.34

Since the totals for all variables but the borrowing ratio hardly deviate from zero, we checked the accuracy of the BORRAT variable's value. As the raw value equalled -129.4, we had Datastream confirm this value and its sign⁷. The F matrix could be simplified such that almost all off-diagonal elements became zeros without impairing the original d^2 value. The importance of the diagonal entries (BORRAT and EARNRAT) indicates an outlier whereby deviations by individual variables were not compensated for by offsetting movements of other variables.

⁷ For the definition of the Borrowing ratio, see Datastream's Company Accounts Definitions Manual, 119).

Table 4: Contribution Matrices for PSUN 1985.

UNSIMPLIFIED SORTED CONTRIBUTION MATRIX:

TAXRAT	GDP	PROFMR	RELOUTP	SPFTA	EARNRAT	BORRAT	TURNCUAS	TRADET	CASHCUR	SUM
30.63	4.83	1.87	-7.57	-0.31	-0.99	0.04	0.02	-2.55	-3.94	22.04
4.83	2.62	0.56	0.18	-0.08	-0.34	0.01	0.01	-0.50	-1.21	6.08
1.87	0.56	1.57	-1.26	-0.52	-0.44	0.00	0.00	0.01	-0.64	1.15
-7.57	-1.26	-1.57	7.23	0.69	0.30	0.01	0.01	0.04	0.64	0.55
-0.31	-0.08	-0.52	0.23	0.69	0.30	0.01	0.01	-0.19	-0.34	0.34
-0.99	-0.34	-0.44	0.30	0.22	0.22	-0.01	-0.00	0.34	0.10	0.21
0.04	0.01	0.00	0.01	0.01	-0.01	0.00	0.00	-0.01	0.00	0.04
0.02	0.01	0.00	0.01	0.00	0.00	0.00	0.00	-0.00	-0.01	0.03
-2.55	-0.50	0.01	0.04	-0.19	0.34	-0.01	-0.00	0.00	-1.41	-1.41
-3.94	-1.21	-0.64	0.64	-0.14	0.10	-0.01	-0.01	0.41	2.37	-2.44
SUM	22.04	6.08	1.15	0.55	0.34	0.04	0.03	-1.41	-2.44	26.61

SIMPLIFIED SORTED CONTRIBUTION MATRIX:

FILTER s = 2 d¹ = 26.60765
 d² = 24.69636

TAXRAT	GDP	EARNRAT	BORRAT	TURNCUAS	SPFTA	PROFMR	RELOUTP	CASHCUR	TRADET	SUM	BLOCK A
30.63	4.83	0.00	0.00	0.00	0.00	0.00	-7.57	-3.94	-2.55	21.41	"
4.83	2.62	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	7.45	BLOCK B
0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	"
0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	"
0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	"
0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	"
0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	"
-7.57	0.00	0.00	0.00	0.00	0.00	0.00	7.53	0.00	0.00	-0.04	BLOCK C
-3.94	0.00	0.00	0.00	0.00	0.00	0.00	0.00	2.37	0.00	-1.57	"
-2.55	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	-2.55	"
SUM	21.41	7.45	0.00	0.00	0.00	0.00	-0.04	-1.57	-2.55	24.70	

BLOCK A

BLOCK B

BLOCK C

5.3. Dynamic episodes around each outlier

The search for simple dynamics could be examined by either a time series or a cross-sectional approach to our data.

The time series method focuses on the d^2 behaviour of individual firms. Episodes consisting of gradual approaches to an outlier constitute the necessarily input for failure prediction analyses. In other words, "sudden" outliers, those without any advance warning, are impossible to predict. Figure 2 on 14 shows a potentially predictable movement by firm 7 (PSUN) had observations before 1984 been included. For example, an episode consisting of years 1982 to 1984 could display a significantly different mean from that of a period running from, say, 1970 to 1981. Assuming increasing d^2 values over the 82-84 period, an even larger d^2 value (possibly an outlier) could reasonably have been anticipated.

The cross-sectional strategy centers on the distinction between years, not firms. The nonuniform distribution of the outliers over the sample period (all outliers dated from 1985 or 1984) implies a nonstationary mean. Table 7 shows the result of searching around each year for a continuous episode in which all the points were close together in Mahalanobis space and far away from the mean observation (at 1% significance).

Table 7: Search for episodes deviating from the mean state

HOTELLINGS TEST			
YEARS	VALUE	APPROX. F.	SIG. F.
84	544.3	3865	0 *
84 85	5.7	40.7	0 *
84 85 86	0.79	5.55	0 *
84 85 86 87	0.32	2.26	0.03
84 85 86 87 88	0.16	1.13	0.347

* DENOTES SIGNIFICANCE AT THE 1%

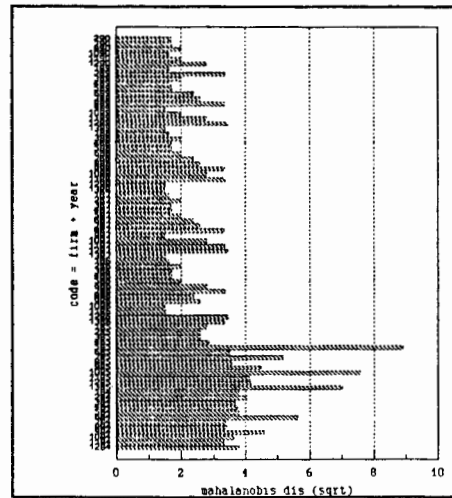
The decreasing T^2 statistic value suggests a strong but extremely short lived strain and long recovery period. For some reason, the mean of the observations for 1984 and for the period 84-85 appears to deviate significantly from that of the other years combined, marking a break in the pooled cross-sectional series pattern. These results reinforce our suspicion that something happened in the economy and/or in the industry during 84-85. The presence of such a break may reflect a sustained structural change in the economy. Hypothesis 4 deals with this problem.

5.4. Identification of Structural Changes

Figure 4 displays the outlier distribution over the sample period. All outliers occur during 1984 and 85. At first sight, the d^2 values appear to be autocorrelated over the

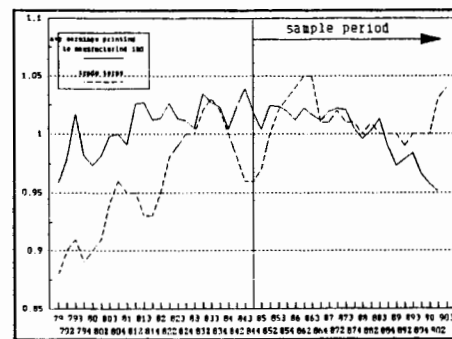
remaining sample period and may display significant dynamic behaviour. Given the extremely short observation period (86-89), we did not test for autocorrelation.

Figure 4: Outlier Distribution over Sample Period



Though formal procedures could be applied to test for the presence of a significant break (Chow test), we believe a qualitative analysis of the printing & publishing industry would reveal a trend reversal in 84-85. Figure 5 seems to support our feeling. The two indices stood more or less at their bottom value in 1984 and the industry's outlook improved from 1985.

Figure 5: Competitive position of Printing & Publishing Industry



6. Conclusions

We attempted to "explain" company failure by suggesting an alternative model to the classic MDA technique. Our model, the decomposition of the Mahalanobis Distance, identifies and analyzes the actual distortions rather than its origins. Indeed, it would be futile to try to isolate causal factors because a dynamic economic system behaves very much like a nonrecursive model in which the variables cannot be arranged in any hierarchical order. Consequently, no single variable can be picked as the sole or determining cause of the state of the system afterwards.

The Mahalanobis technique incorporates some advantages over the MDA. First, there is no need to worry about the conceptualisation of the categories the observations fall into and to resort to ad hoc criteria. Second, the Mahalanobis model stands for a unique design for each year and each firm whereas the discriminant function represents an aggregate system and therefore probably incorrect. Lastly, contrary to MDA analyses, we see no need to restrict the variables to those endogeneous, that is accounting or financial, components. External factors, such as industry and economy effects, need incorporated into the analysis.

The Mahalanobis model is, of course, not without its defects. We assumed a stable variance-covariance matrix. Our empirical test showed the presence of a nonstationary mean so the manifestation of an equally nonuniform variance cannot be dismissed. Furthermore, the predictive ability of the Mahalanobis model is restricted to those outliers that form part of longer lasting episodes. Unexpected or sudden outliers cannot be forecasted, but then how good is MDA at anticipating sudden corporate fragility ?

We believe that the comparative analysis of the contribution matrices (F) to the distance d^2 may yield valuable insights into the failure procedure and could contribute to a more systematic approach of the corporate collapse issue. Such analysis could be best performed by a pattern recognition technique such as an artificial neural system (ANS). Meanwhile, the Mahalanobis model is argued to assist industry regulators in the selection of "policy" variables and testing of various competing economic theories.

REFERENCES

- Altman, E.I. (1968), "Financial Ratios, Discriminant Analysis and the Prediction of Corporation Bankruptcy", Journal of Finance, 589-609
- Altman, E.I. (1984), "A Further Investigation of the Bankruptcy Cost Question", Journal of Finance, 1067-1089
- Altman, E.I., Avery, R.B., Eisenbeis, R.A. and Sinkey, J.F., (1981), Application of Classification Techniques in Business, Banking and Finance, Greenwich: JAI Press, Inc
- Altman, E.I., Haldeman, R.G. and Narayanan, P. (1977), "Zeta Analysis. A New Model to Identify Bankruptcy Risk of Corporations", Journal of Banking and Finance, 29-54
- Altman, E.I. and McGough, T.P. (1974), "Evaluation of a Company as a Going Concern", Journal of Accountancy, 50-57
- Altman, E. and Sametz, A. (eds) (1977), Financial Crises, Institutions & Markets in a Fragile Environment, New York: John Wiley & Sons
- Anscombe, F.J. (1960), "Rejection of Outliers", Technometrics, 2, 123-147
- Argenti, J. (1976), Corporate Collapse: The Causes and Symptoms, London: McGraw-Hill
- Bacon-Shone, J. and Fung, W.K. (1987) "A new graphical method for Detecting Single and Multiple Outliers in Univariate and Multivariate Data" Applied Statistics, 36:2, 153-162
- Baestaens, D.J.E. (1990), The Market Impact of Regulation on Financial Institutions, Ph.D. Thesis, Manchester Business School
- Barnes, P. (1982), "Methodological Implications of Non-Normally Distributed Financial ratios", Journal of Business Finance & Accounting, 9:1, 51-62
- Barnes, P. (1987), "The Analysis and Use of Financial Ratios: A Review Article", Journal of Business Finance & Accounting, 14:4, 449-61
- Bilderbeek, J. (1979), "An Empirical Study of the Predictive Ability of Financial Ratios in the Netherlands", Zeitschrift Fur Betriebswirtschaft, 5, 388-407
- Boot, A.W.A. and Wijn, M.F.C.M. (1991), "Insolventie, Vermogensstructuur en Vermogensmarkt", M.A.B., January-February, 22-32
- Collett, D. and Lewis, T. (1976), "The Subjective Nature of Outlier Rejection Procedures", Applied Statistics, 25, 228-237

- Donaldson, G. (1969), "Strategy for Financial Emergencies", H.B.R., Nov-Dec, 67-79
- El Hennawy, R.H.A. and Morris, R.C. (1983), "Market Anticipation of Corporate Failure in the UK", Journal of Business Finance & Accounting, 10:3, 359-372
- Eisenbeis, R.A., (1977), "Pitfalls in the Application of Discriminant Analysis in Business, Finance, and Economics", Journal of Finance, 23:3, 875-900
- Ezzamel, M. and Mar-Molinero, C. (1990), "The Distributional Properties of Financial Ratios in UK Manufacturing Companies", Journal of Business Finance & Accounting, 17:1, 1-29
- Gonedes, N.J. (1973), "Properties of Accounting Numbers: Models and Tests", Journal of Accounting Research, 212-237
- Hartwig F. and Dearing, B.E. (1979), Exploratory Data Analysis, London: Sage Publications
- Hotelling, H. (1931), "The generalisation of Student's ratio", Ann. Math. Stat., Vol 2, 360-78
- Howell S.D. (1989) "Quantitative Forecasting: Issues and Methods " in Ashton D. and Hopper T. (eds) Readings for Accounting practice, London: Philip Allen
- Karels, G.V. and Prakash, A.J. (1987), "Multivariate Normality and Forecasting of Business Bankruptcy", Journal of Business Finance & Accounting, 14:4, 573-593
- Lev, B. (1973), "Decomposition Measures for Financial Analysis", Financial Management, Spring, 56-63
- Lev, B. and Sunder, S. (1979), "Methodological Issues in the Use of Financial Ratios", Journal of Accounting and Economics, 1, 187-210
- Lilliefors, H.W. (1967), "On the Kolmogorov-Smirnov Test for Normality with Mean and Variance Unknown", Journal of the American Statistical Association, 62, 399-402
- Mahalanobis, P.C. (1930), "On Tests and Measures of Group Divergence", Journal of the Asiatic Society of Bengal, vol 26, 541-88
- Mahalanobis, P.C. (1936), "On the Generalised Distance in Statistics", Proceedings National Institute of Sciences of India, vol 2, 49-55
- Mahalanobis, P.C. (1948), "Historical Note on the D^2 -Statistic", Sankhya, vol 9, 237-40
- Manly, B.F.J. (1988), Multivariate Statistical Methods, A Primer, London: Chapman and Hall

Martikainen, T. and Ankelo, T. (1990), "On the Instability of Financial Patterns in Failed Firms and the Predictability of Corporate failure", Proceedings of the University of Vaasa, Discussion Papers 113, Vaasa: Finland

Mitchell A.F.S. and Krzanowski, W.J. (1985) "The Mahalanobis Distance and Elliptic Distributions", Biometrika, 72:2, 464-7

Martin, D. (1977), "Early Warning of Bank Failure", Journal of Banking and Finance, 1:3, 249-276

Richardson, F.M. and Davidson, L.F. (1983), "An Exploration into Bankruptcy Discriminant Model Sensitivity", Journal of Business Finance & Accounting, 10:2, 195-207

Taffler, R.J. (1977), "Finding Those Firms in Danger Using Discriminant Analysis and Financial Ratio Data: A Comparative UK Based Study", Working Paper No3, London: City University Business School

Taffler, R.J. (1982), "Forecasting Company Failure in the UK using Discriminant Analysis and Financial Ratio Data", Journal Royal Statistical Society Association, 145:3, 342-358

Taffler, R.J., (1984), "Empirical Models for the Monitoring of UK Corporations", Journal of Banking and Finance, 8, 199-227

Van Frederikslust, R.A.I. (1978), Predictability of Corporate Failure, Leiden-Boston: Martinus Nijhoff

ANOTHER COUNTEREXAMPLE TO THE RANK-COLOURING CONJECTURE

Dirk TRAPPERS

**Faculty of Applied Economics UFSIA
University of Antwerp, Prinsstraat 13, B-2000 Antwerpen, Belgium**

ABSTRACT

In this paper we extend Alan and Seymour's procedure to find a counterexample of the Rank Colouring Conjecture, which was proposed by Van Nuffelen. This conjecture states that the rank of the incidence matrix of a graph is an upper bound for the chromatic number of that graph.

We assume the reader is familiar with the paper of Alon and Seymour ([1]), in which they constructed a graph with rank 29 and chromatic number 32 as a counterexample to the Rank Colouring Conjecture. By generalizing this method one can construct at least two more graphs G with $\text{rk}(G) < \chi(G)$, one of these is however isomorphic to the counterexample of Alon and Seymour.

The conjecture was originally put forward by Van Nuffelen in [2]. He proved the conjecture for several classes of graphs but was not able to provide a general proof for arbitrary graphs. The hope for such a proof was finally abandoned after the publication of [1]. When one looks at the results of Van Nuffelen, it remains however very plausible that the counterexamples to the conjecture should belong to a limited class of graphs. Consequently, the main motivation for this paper is to provide a narrower description of the class of counterexamples.

Alon and Seymour found their counterexample by considering Cayley graphs, these graphs are constructed by taking a finite abelian group G as its vertex set and by connecting two vertices when their sum belongs to a certain predescribed subset K of G .

When we specialize this to $G = (\mathbf{Z}/2\mathbf{Z})^k$ (the k -dimensional vector space over the finite field with 2 elements), we can identify subsets of G with subsets of $\mathcal{P}(\{1, \dots, k\})$ (the powerset of $\{1, \dots, k\}$) by considering them as characteristic functions. The counterexample of Alon and Seymour is found by putting $k = 6$ and by taking K as the complement in $\mathcal{P}(\{1, \dots, k\})$ of the set

$$\{\emptyset, \{1\}, \{2\}, \{3\}, \{4\}, \{5\}, \{6\}, \{1, 2, 3, 4, 5, 6\}\}.$$

We will denote I for the set $\{1, \dots, k\}$ and $\mathcal{P}(I)$ for the powerset of I .

Definition : We will call a set $A \subset \mathcal{P}(I)$ symmetric in the numbers $\alpha_1, \dots, \alpha_n \in I$ if

- $\emptyset \in A$,
- every subset of $\mathcal{P}(I)$ with cardinality equal to one of the α_i is in A .

The set $E(H)$ in the counterexample of Alon and Seymour has $k = 6$, $n = 2$ and $\alpha_1 = 1$, $\alpha_2 = 6$.

We will denote the symmetric difference of two sets X and Y by $X * Y$, i.e.

$$X * Y = (X \cup Y) \setminus (X \cap Y).$$

A subset X of I uniquely corresponds to a vector $\varphi_X \in (\mathbf{Z}/2\mathbf{Z})^k$, where i -th component of φ_X equals 1 if $i \in X$ and equals 0 otherwise.

It is now easily seen that:

$$\varphi_{X*Y} = \varphi_X + \varphi_Y,$$

with addition in the vector space $(\mathbf{Z}/2\mathbf{Z})^k$.

Definition : A symmetric subset A of $\mathcal{P}(I)$ is called *colourful* if there are no $X, Y, Z \in A$ with $X * Y * Z = \emptyset$ or equivalently $X * Y = Z$.

Proposition 1: A colourful set A cannot be symmetric in an even number $n \leq \frac{2}{3}k$.

Proof:

Take A symmetric in n .

Put $X = \{1, \dots, n\}$ and $Y = \{n/2 + 1, \dots, 3n/2\}$, by the assumption on n it follows that $X, Y \in A$.

Furthermore

$$X * Y = \{1, \dots, n/2, n+1, \dots, 3n/2\} \in A$$

as the cardinality of $X * Y$ equals n . Consequently, A is not colourful.

□

Proposition 2: Take $x, y, z \in I$ with $x + y = z$ and take A colourful, we then have

- if A is symmetric in x and y , then A is not symmetric in z .
- if A is symmetric in x and z , then A is not symmetric in y ,

Proof:

The proof is analogous to proposition 1.

For the first statement, take $X = \{1, \dots, x\}$ and $Y = \{x + 1, \dots, z\}$, it follows that $X * Y = \{1, \dots, z\}$.

And for the second statement, put $X = \{1, \dots, x\}$ and $Z = \{1, \dots, z\}$, then as $z > x$ $X * Z = \{x+1, \dots, z\}$.

□

Take now the hypergraph H with $V(H) = I$ and $E(H) = A \subset \mathcal{P}(I)$, with A symmetric.

Define the graph G with vertex set equal to the vector space $V(G) = (\mathbf{Z}/2\mathbf{Z})^k$ and where u and $v \in V(G)$ are adjacent if $u * v \notin E(H)$.

By demanding that $E(H)$ is colourful, we get that every stable set of G has cardinality less or equal to 2 and so $\chi(G) \geq \frac{1}{2}|V(G)| = 2^{k-1}$. As mentioned in [1], if $E(H) \neq \{\emptyset\}$ then in fact $\chi(G) = 2^{k-1}$, as its complement contains a perfect matching.

To compute the rank of G , we have to count the number t of $X \subset V(H) = I$ such that $|X \cap E|$ is odd for precisely $\frac{1}{2}|E(H)|$ members E of $E(H)$.

By claim 2 in the paper of Alon and Seymour, it then follows that $\text{rk}(G) = 2^k - t$. To get a counterexample to the Rank-Colouring conjecture we will have to find graphs for which $t > 2^{k-1}$.

To this purpose, define a function $F : I \cup \{0\} \times I \cup \{0\} \rightarrow \mathbf{Z}$ by

$$F(a, b) = \sum_i \binom{b}{2i+1} \binom{k-b}{a-2i-1},$$

where the summation runs over all i such that $\max\{a+b-k, 1\} \leq 2i+1 \leq \min\{a, b\}$.

It is clear, by construction of F , that $F(a, b)$ equals the number of subsets of I with cardinality a that have an odd intersection with $\{1, \dots, b\}$.

Assume further that $E(H)$ is symmetric in $\alpha_1, \dots, \alpha_n$ then

$$|E(H)| = 1 + \sum_{i=1}^n \binom{k}{\alpha_i},$$

and so we have to find all $b \in I \cup \{0\}$ that satisfy

$$\sum_{i=1}^n F(\alpha_i, b) = \frac{1}{2} + \frac{1}{2} \sum_{i=1}^n \binom{k}{\alpha_i}. \quad (*)$$

If the $b_1, \dots, b_l$ are the solutions of $(*)$, then

$$t = 2^k - \text{rk}(G) = \sum_{j=1}^l \binom{k}{b_j}. \quad (**)$$

It is now an easy exercise in combinatorics to verify:

Proposition 3: For $a, b \in I \cup \{0\}$ we have:

- for a even : $F(a, b) = F(a, k - b)$,
- for a odd : $F(a, b) + F(a, k - b) = \binom{k}{a}$.

We then have:

Proposition 4: In order for $(*)$ to be satisfied in a way such that $\chi(G) > \text{rk}(G)$, at least one of the α_i should be even.

Proof:

Assume all α_i odd and take b a solution of $(*)$ such that $k - b$ also satisfies $(*)$.

Now $(*)$ and proposition 3 imply:

$$\sum_{i=1}^n \binom{k}{\alpha_i} - \sum_{i=1}^n F(\alpha_i, b) = \frac{1}{2} + \frac{1}{2} \sum_{i=1}^n \binom{k}{\alpha_i},$$

or

$$\sum_{i=1}^n F(\alpha_i, b) = -\frac{1}{2} + \frac{1}{2} \sum_{i=1}^n \binom{k}{\alpha_i}.$$

Which clearly contradicts $(*)$.

This allows us to conclude that the solution set $\{b_1, \dots, b_l\}$ of $(*)$, can be replaced by a subset $\{c_1, \dots, c_l\}$ of $\{0, \dots, \lfloor \frac{k}{2} \rfloor\}$ (in $(**)$ we may replace b by $k - b$).

But

$$\sum_{j=0}^l \binom{k}{c_j} \leq \sum_{i=0}^{\lfloor \frac{k}{2} \rfloor} \binom{k}{i} \leq 2^{k-1},$$

and consequently $\text{rk}(G) \geq 2^{k-1} = \chi(G)$.

□

Based upon these four proposition, it is possible to write a computer program, that generates sets $E(H)$ symmetric in $\alpha_1, \dots, \alpha_n$ that satisfy:

- at least one of the α_i is even,
- all the even α_i satisfy, $\alpha_i > 3k/2$,
- there are no three $\alpha_i, \alpha_j, \alpha_k$, such that $\alpha_i + \alpha_j = \alpha_k$.

The program itself can be found in Appendix 1. Appendix 2 contains the output of the program for $k = 4, \dots, 14$. On the upper part of the page, you find $F(a, b)$. Then each of the possible sets $\{0, \alpha_1, \dots, \alpha_n\}$ are listed together with the number $g = 1/2 + 1/2 \sum \binom{k}{\alpha_i}$ and $t = \sum \binom{k}{b_j}$.

I also verified the cases $k = 15, \dots, 18$, but these are not listed in the appendix.

The three results found are:

- $k = 6, n = 2$ and $\alpha_1 = 1, \alpha_2 = 6$.
- $k = 6, n = 2$ and $\alpha_1 = 5, \alpha_2 = 6$.
- $k = 7, n = 1$ and $\alpha_1 = 6$.

The corresponding sets $E(H)$ are then given by:

$$E(H_1) = \{\emptyset, \{1\}, \{2\}, \{3\}, \{4\}, \{5\}, \{6\}, \{1, 2, 3, 4, 5, 6\}\},$$

$$E(H_2) = \{\emptyset, \{1, 2, 3, 4, 5\}, \{1, 2, 3, 4, 6\}, \{1, 2, 3, 5, 6\}, \{1, 2, 4, 5, 6\}, \\ \{1, 3, 4, 5, 6\}, \{2, 3, 4, 5, 6\}, \{1, 2, 3, 4, 5, 6\}\},$$

$$E(H_3) = \{\emptyset, \{1, 2, 3, 4, 5, 6\}, \{1, 2, 3, 4, 5, 7\}, \{1, 2, 3, 4, 6, 7\}, \{1, 2, 3, 5, 6, 7\}, \\ \{1, 2, 4, 5, 6, 7\}, \{1, 3, 4, 5, 6, 7\}, \{2, 3, 4, 5, 6, 7\}\}.$$

It doesn't seem very likely anymore counterexamples will be found using this method, but presently I am not able to prove this.

The first graph G_1 is the one found by Alon and Seymour, it has 64 vertices. The second one G_2 also has 64 vertices and take the map

$$\Psi : G_1 \rightarrow G_2$$

defined by

$$\begin{aligned} & \Psi(x_1, x_2, x_3, x_4, x_5, x_6) \\ &= \begin{cases} (x_1, x_2, x_3, x_4, x_5, x_6) & \text{if } \sum x_i \text{ is even,} \\ (1 - x_1, 1 - x_2, 1 - x_3, 1 - x_4, 1 - x_5, 1 - x_6) & \text{otherwise.} \end{cases} \end{aligned}$$

It is then easily seen that Ψ is an isomorphism of graphs.

The third graph G_3 is essentially different as it contains 128 vertices. We have $\text{rk}(G_3) = 128 - 70 = 58$ and $\chi(G_3) = 64$. Although these numbers equal the ones found for the graph we get by taking G_1 twice and joining every vertex to every vertex of the other copy (see [1]), they are not isomorphic as the complement of G_3 is a connected graph.

This proves:

Theorem : *The method used by Alon and Seymour to construct a graph that doesn't satisfy the rank-colouring conjecture, only gives one other counterexample with less than 500000 vertices.*

References :

- [1] N. Alon, P.D. Seymour, A Counterexample to the Rank-Coloring Conjecture, Journal of Graph Theory, Vol.13, No.4, 523-525 (1989).
- [2] C. Van Nuffelen, A Bound for the Chromatic Number of a Graph, Am. Math. Month. 83 (1976) 265-266.

filleff(rcc.c)

```

# include <stdio.h>

# define      K      1L
# define      MAX      (long)(K + 1L)
# define      POWER      (long)(1L << (K - 1L))
# define      TTK      (long)(2L * K / 3L)

long  fac[13] =
{      1L,      1L,      2L,
      6L,      24L,      120L,
      720L,      5040L,      40320L,
      362880L,      3628800L,      39916800L,
      479001600L      };
long  F[MAX][MAX];

# define      max(a,b)      ((a) < (b) ? (b) : (a))

long  combination(x,y)
long  x,y;
{
    register      long  p, i;
    long  q;

    if(x < 13L)
        return (fac[x] / fac[y]) / fac[x - y];
    q = max(y,x - y);
    p = 1L;
    for(i = x;i > q;i--)
        p *= i;
    for(i = 1;i <= x - q;i++)
        p /= i;
    return p;
}

filleff()
{
    register      long  j, b, a;
    long  sum;

    for(a = 0;a < MAX;a++)
        for(b = 0;b < MAX;b++)
        {
            sum = 0;
            for(j = 1;j < MAX;j += 2)
                if(j <= b && j <= a && a - j <= K - b)
                    sum += combination(b,j)
                        * combination(K - b,a - j);
            F[a][b] = sum;
        }
}

```

```

long    A[MAX], B[MAX];

main()                                     main
{
    register    long    i, j;
    long        sum, number;
    char        name[20];
    FILE        *fp;
    sprintf(name,"graphs.%ld",K);
    if((fp = fopen(name,"w")) == NULL)
        exit(0);

    filfeff();

    fprintf(fp,"k = %2d\n-----\n\n",K);
    for(i = 0;i < MAX;i++)
    {
        for(j = 0;j < MAX;j++)
            fprintf(fp,"%4ld ",F[i][j]);
        fprintf(fp,"\n");
    }
    fprintf(fp,"\n");
    for(i = 0;i < MAX;i++)
        A[i] = 0;
next:
    for(i = 1;i <= MAX;i++)
    {
        if(i == MAX)
        {
            fclose(fp);
            exit(0);
        }
        if(A[i] == 0)
        {
            A[i] = 1;
            break;
        }
        A[i] = 0;
    }
    for(i = 1;i < MAX;i++)
        if(A[i] && i % 2 == 0)
        {
            if(i <= TTK)
                goto next;
            goto even;
        }
    /* A should contain at least 1 even number,
       but if 2*x <= TTK then 2*x is excluded */
    goto next;
even:
    for(i = 1;i < MAX;i++)
        for(j = 1;j < MAX;j++)

```

60

70

80

90

100

main(rcc.c)

```

{
    if(i + j >= MAX)
        continue;
    if(A[i] && A[j] && A[i + j])
        goto next;
}

number = 0;
for(i = 1; i < MAX; i++)
    if(A[i])
        number += combination(K,i);
if(number % 2 == 0)
    goto next;
number++;
number /= 2;

for(i = 0; i < MAX; i++)
    B[i] = 0;
for(i = 1; i < MAX; i++)
    if(A[i])
        for(j = 0; j < MAX; j++)
            B[j] += F[i][j];

sum = 0;
for(i = 0; i < MAX; i++)
    if(B[i] == number)
        sum += combination(K,i);

fprintf(fp,"%2d ",0);
for(i = 1; i < MAX; i++)
    if(A[i])
        fprintf(fp,"%21d ",i);
    else
        fprintf(fp," ");
fprintf(fp," --> (g = %61d) %61d",number,sum);
if(sum > POWER)
    fprintf(fp," !!!");
fprintf(fp,"\n");

goto next;

```

```

k = 4
-----
0 0 0 0 0
0 1 2 3 4
0 3 4 3 0
0 3 2 1 4
0 1 0 1 0

0 4 --> (g = 1) 8
0 1 4 --> (g = 3) 0
0 3 4 --> (g = 3) 0

k = 5
-----
0 0 0 0 0 0
0 1 2 3 4 5
0 4 6 6 4 0
0 6 6 4 4 10
0 4 2 2 4 0
0 1 0 1 0 1

0 4 --> (g = 3) 0
0 3 4 --> (g = 8) 15

k = 6
-----
0 0 0 0 0 0 0
0 1 2 3 4 5 6
0 5 8 9 8 5 0
0 10 12 10 8 10 20
0 10 8 6 8 10 0
0 5 2 3 4 1 6
0 1 0 1 0 1 0

0 6 --> (g = 1) 32
0 1 6 --> (g = 4) 35 !!!
0 5 6 --> (g = 4) 35 !!!

k = 7
-----
0 0 0 0 0 0 0 0
0 1 2 3 4 5 6 7
0 6 10 12 12 10 6 0
0 15 20 19 16 15 20 35
0 20 20 16 16 20 20 0
0 15 10 9 12 11 6 21
0 6 2 4 4 2 6 0
0 1 0 1 0 1 0 1

0 6 --> (g = 4) 70 !!!
0 5 6 7 --> (g = 15) 0

k = 8
-----
0 0 0 0 0 0 0 0 0
0 1 2 3 4 5 6 7 8
0 7 12 15 16 15 12 7 0
0 21 30 31 28 25 26 35 56
0 35 40 35 32 35 40 35 0
0 35 30 25 28 31 26 21 56
0 21 12 13 16 13 12 21 0
0 7 2 5 4 3 6 1 8
0 1 0 1 0 1 0 1 0

0 8 --> (g = 1) 128

```

```

0 1      8 --> (g = 5) 0
0      3 8 --> (g = 29) 0
0 1      3 8 --> (g = 33) 0
0      5 8 --> (g = 29) 0
0 1      5 8 --> (g = 33) 0
0      6 8 --> (g = 15) 0
0 1      6 8 --> (g = 19) 56
0      5 6 8 --> (g = 43) 8
0      7 8 --> (g = 5) 0
0      3 7 8 --> (g = 33) 0
0      5 7 8 --> (g = 33) 0
0      6 7 8 --> (g = 19) 56
0      5 6 7 8 --> (g = 47) 0
k = 9
-----

0 0 0 0 0 0 0 0 0
0 1 2 3 4 5 6 7 8 9
0 8 14 18 20 20 18 14 8 0
0 28 42 46 44 40 38 42 56 84
0 56 70 66 60 60 66 70 56 0
0 70 70 60 60 66 66 56 56 126
0 56 42 38 44 44 38 42 56 0
0 28 14 18 20 16 18 22 8 36
0 8 2 6 4 4 6 2 8 0
0 1 0 1 0 1 0 1 0 1

0      8 --> (g = 5) 0
0      3 8 --> (g = 47) 0
0      5 8 --> (g = 68) 0
0      7 8 --> (g = 23) 0
0      3 7 8 --> (g = 65) 0
0      5 7 8 --> (g = 86) 162

```


k = 10

0	0	0	0	0	0	0	0	0	0	0
0	1	2	3	4	5	6	7	8	9	10
0	9	16	21	24	25	24	21	16	9	0
0	36	56	64	64	60	56	56	64	84	120
0	84	112	112	104	100	104	112	112	84	0
0	126	140	126	120	126	132	126	112	126	252
0	126	112	98	104	110	104	98	112	126	0
0	84	56	56	64	60	56	64	64	36	120
0	36	16	24	24	20	24	24	16	36	0
0	9	2	7	4	5	6	3	8	1	10
0	1	0	1	0	1	0	1	0	1	0

0					8	-->	(g =	23)	0	
0	1				8	-->	(g =	28)	210	
0		3			8	-->	(g =	83)	0	
0	1	3			8	-->	(g =	88)	45	
0			5		8	-->	(g =	149)	0	
0	1		5		8	-->	(g =	154)	0	
0				7	8	-->	(g =	83)	0	
0		3		7	8	-->	(g =	143)	0	
0			5	7	8	-->	(g =	209)	0	
0				8	9	-->	(g =	28)	210	
0		3		8	9	-->	(g =	88)	45	
0			5	8	9	-->	(g =	154)	0	
0				7	8	9	-->	(g =	88)	45
0		3		7	8	9	-->	(g =	148)	0
0			5	7	8	9	-->	(g =	214)	45
0					10	-->	(g =	1)	512	
0	1				10	-->	(g =	6)	462	
0		3			10	-->	(g =	61)	252	
0	1	3			10	-->	(g =	66)	252	
0				7	10	-->	(g =	61)	252	
0	1			7	10	-->	(g =	66)	252	
0					9	10	-->	(g =	6)	462
0		3			9	10	-->	(g =	66)	252
0				7	9	10	-->	(g =	66)	252

k = 11

0	0	0	0	0	0	0	0	0	0	0	0
0	1	2	3	4	5	6	7	8	9	10	11
0	10	18	24	28	30	30	28	24	18	10	0
0	45	72	85	88	85	80	77	80	93	120	165
0	120	168	176	168	160	160	168	176	168	120	0
0	210	252	238	224	226	236	238	224	210	252	462
0	252	252	224	224	236	236	224	224	252	252	0
0	210	168	154	168	170	160	162	176	162	120	330
0	120	72	80	88	80	80	88	80	72	120	0
0	45	18	31	28	25	30	27	24	37	10	55
0	10	2	8	4	6	6	4	8	2	10	0
0	1	0	1	0	1	0	1	0	1	0	1

0					8	-->	(g =	83)	0	
0	1	3			8	-->	(g =	171)	0	
0			5		8	-->	(g =	314)	0	
0				7	8	-->	(g =	248)	0	
0			5	7	8	-->	(g =	479)	0	
0		3		8	9	-->	(g =	193)	0	
0		3		7	8	9	-->	(g =	358)	0
0					10	-->	(g =	6)	924	
0	1	3			10	-->	(g =	94)	0	
0				7	10	-->	(g =	171)	0	
0	1			8	10	-->	(g =	94)	0	

```

0      3      8      10    --> (g = 171) 462
0      3      9      10    --> (g = 116) 924
0      8      9      10    --> (g = 116) 462
0      7      8      9      10    --> (g = 281) 792
0      1      8      11    --> (g = 89) 0
0      1      5      8      11    --> (g = 320) 0
0      8      9      11    --> (g = 111) 0
0      5      8      9      11    --> (g = 342) 55
0      7      8      9      11    --> (g = 276) 462
0      5      7      8      9      11    --> (g = 507) 0
0      3      10      11    --> (g = 89) 0
0      8      10      11    --> (g = 89) 165
0      7      8      10      11    --> (g = 254) 0
0      9      10      11    --> (g = 34) 0
0      7      9      10      11    --> (g = 199) 0

```

k = 12

```

0 0 0 0 0 0 0 0 0 0 0 0 0
0 1 2 3 4 5 6 7 8 9 10 11 12
0 11 20 27 32 35 36 35 32 27 20 11 0
0 55 90 109 116 115 110 105 104 111 130 165 220
0 165 240 261 256 245 240 245 256 261 240 165 0
0 330 420 414 392 386 396 406 400 378 372 462 792
0 462 504 462 448 462 472 462 448 462 504 462 0
0 462 420 378 392 406 396 386 400 414 372 330 792
0 330 240 234 256 250 240 250 256 234 240 330 0
0 165 90 111 116 105 110 115 104 109 130 55 220
0 55 20 39 32 31 36 31 32 39 20 55 0
0 11 2 9 4 7 6 5 8 3 10 1 12
0 1 0 1 0 1 0 1 0 1 0 1 0

```

```

0      12    --> (g = 1) 2048
0      1      12    --> (g = 7) 0
0      3      12    --> (g = 111) 0
0      1      3      12    --> (g = 117) 0
0      5      12    --> (g = 397) 0
0      1      5      12    --> (g = 403) 0
0      3      5      12    --> (g = 507) 0
0      1      3      5      12    --> (g = 513) 0
0      7      12    --> (g = 397) 0
0      1      7      12    --> (g = 403) 0
0      3      7      12    --> (g = 507) 0
0      1      3      7      12    --> (g = 513) 0
0      9      12    --> (g = 111) 0
0      1      9      12    --> (g = 117) 0
0      5      9      12    --> (g = 507) 0
0      1      5      9      12    --> (g = 513) 0
0      7      9      12    --> (g = 507) 0
0      1      7      9      12    --> (g = 513) 0
0      10      12    --> (g = 34) 0
0      1      10      12    --> (g = 40) 495
0      3      10      12    --> (g = 144) 0
0      1      3      10      12    --> (g = 150) 0
0      7      10      12    --> (g = 430) 0
0      1      7      10      12    --> (g = 436) 0
0      9      10      12    --> (g = 144) 0
0      7      9      10      12    --> (g = 540) 495
0      11      12    --> (g = 7) 0
0      3      11      12    --> (g = 117) 0
0      5      11      12    --> (g = 403) 0
0      3      5      11      12    --> (g = 513) 0
0      7      11      12    --> (g = 403) 0
0      3      7      11      12    --> (g = 513) 0
0      9      11      12    --> (g = 117) 0
0      5      9      11      12    --> (g = 513) 0

```

```

0      7      9      11 12 --> (g = 513)      0
0      10 11 12 --> (g = 40)      495
0      3      10 11 12 --> (g = 150)      0
0      7      10 11 12 --> (g = 436)      0
0      9 10 11 12 --> (g = 150)      0
0      7      9 10 11 12 --> (g = 546)      0
k = 13
-----
0      0      0      0      0      0      0      0      0      0      0      0      0      0
0      1      2      3      4      5      6      7      8      9      10     11     12     13
0      12     22     30     36     40     42     42     40     36     30     22     12     0
0      66     110    136    148    150    146    140    136    138    150    176    220    286
0      220    330    370    372    360    350    350    360    372    370    330    220     0
0      495    660    675    648    631    636    651    656    639    612    627    792    1287
0      792    924    876    840    848    868    868    848    840    876    924    792     0
0      924    924    840    840    868    868    848    848    876    876    792    792    1716
0      792    660    612    648    656    636    636    656    648    612    660    792     0
0      495    330    345    372    355    350    365    360    343    370    385    220    715
0      220    110    150    148    136    146    146    136    148    150    110    220     0
0      66     22     48     36     38     42     36     40     42     30     56     12     78
0      12     2      10     4      8      6      6      8      4      10     2      12     0
0      1      0      1      0      1      0      1      0      1      0      1      0      1

0      1      10      --> (g = 150)      0
0      1      3      10      --> (g = 293)    1716
0      1      7      10      --> (g = 1008)    0
0      9 10      --> (g = 501)      0
0      3      9 10      --> (g = 644)      0
0      7      9 10      --> (g = 1359)    3003
0      9 10 11    --> (g = 540)      0
0      3      9 10 11    --> (g = 683)      0
0      7      9 10 11    --> (g = 1398)    0
0      12      --> (g = 7)          0
0      3      12      --> (g = 150)      0
0      1      5      12      --> (g = 657)      0
0      1      3      5      12      --> (g = 800)      0
0      7      12      --> (g = 865)      0
0      3      7      12      --> (g = 1008)    0
0      1      9      12      --> (g = 371)      0
0      5      9      12      --> (g = 1008)    0
0      1      7      9      12      --> (g = 1229)    0
0      10     12      --> (g = 150)      0
0      3      10     12      --> (g = 293)      0
0      7      10     12      --> (g = 1008)    0
0      11 12     --> (g = 46)        2002
0      3      11 12     --> (g = 189)      0
0      7      11 12     --> (g = 904)      0
0      3      7      11 12     --> (g = 1047)    0
0      5      9      11 12     --> (g = 1047)    0
0      3      10 11 12    --> (g = 189)      0
0      7      10 11 12    --> (g = 332)    2002
0      10     13      --> (g = 144)      0
0      7      10     13      --> (g = 1002)    0
0      10 11     13      --> (g = 183)      1716
0      7      10 11     13      --> (g = 1041)    0
0      12 13     --> (g = 651)      0
0      3      5      12 13     --> (g = 794)      0
0      7      9      12 13     --> (g = 365)      0
0      9      12 13     --> (g = 1223)    0
0      9 10     12 13     --> (g = 508)      0
0      7      9 10     12 13     --> (g = 1366)    1794
0      11 12 13    --> (g = 690)      0
0      3      5      11 12 13    --> (g = 833)      0
0      9      11 12 13    --> (g = 404)      286

```

```

0      7      9      11 12 13 --> (g = 1262)      0
0      9 10 11 12 13 --> (g = 547)      0
0      9 10 11 12 13 --> (g = 1405)      0
k = 14
-----

0      0      0      0      0      0      0      0      0      0      0      0      0      0      0
0      1      2      3      4      5      6      7      8      9      10     11     12     13     14
0     13     24     33     40     45     48     49     48     45     40     33     24     13     0
0     78     132    166    184    190    188    182    176    174    180    198    232    286    364
0    286    440    506    520    510    496    490    496    510    520    506    440    286     0
0    715    990   1045   1020   991    986   1001   1016   1011    982    957   1012   1287   2002
0   1287   1584   1551   1488   1479   1504   1519   1504   1479   1488   1551   1584   1287     0
0   1716   1848   1716   1680   1716   1736   1716   1696   1716   1752   1716   1584   1716   3432
0   1716   1584   1452   1488   1524   1504   1484   1504   1524   1488   1452   1584   1716     0
0   1287    990    957   1020   1011    986   1001   1016    991    982   1045   1012    715   2002
0    715    440    495    520    491    496    511    496    491    520    495    440    715     0
0    286    132    198    184    174    188    182    176    190    180    166    232    78    364
0     78     24     58     40     46     48     42     48     46     40     58     24     78     0
0     13      2     11      4      9      6      7      8      5     10      3     12      1     14
0      1      0      1      0      1      0      1      0      1      0      1      0      1      0

0      10      --> (g = 501)      0
0      10      --> (g = 508)      0
0      3       10      --> (g = 683)      0
0      1      3       10      --> (g = 690)      3003
0      7       10      --> (g = 2217)      0
0      1      7       10      --> (g = 2224)      0
0      9 10      --> (g = 1502)      3003
0      3       9 10      --> (g = 1684)      91
0      7       9 10      --> (g = 3218)      5005
0      10 11      --> (g = 683)      0
0      3       10 11      --> (g = 865)      0
0      7       10 11      --> (g = 2399)      0
0      9 10 11      --> (g = 1684)      91
0      3       9 10 11      --> (g = 1866)      2002
0      7       9 10 11      --> (g = 3400)      0
0      12      --> (g = 46)      4004
0      1       12      --> (g = 53)      0
0      3       12      --> (g = 228)      0
0      1      3       12      --> (g = 235)      0
0      5       12      --> (g = 1047)      0
0      1      5       12      --> (g = 1054)      0
0      3      5       12      --> (g = 1229)      0
0      1      3      5       12      --> (g = 1236)      0
0      7       12      --> (g = 1762)      4004
0      1      7       12      --> (g = 1769)      0
0      3      7       12      --> (g = 1944)      0
0      1      3      7       12      --> (g = 1951)      0
0      9       12      --> (g = 1047)      0
0      1      9       12      --> (g = 1054)      0
0      5      9       12      --> (g = 2048)      4095
0      1      5      9       12      --> (g = 2055)      0
0      7      9       12      --> (g = 2763)      0
0      1      7      9       12      --> (g = 2770)      0
0      11 12      --> (g = 228)      0
0      3      11 12      --> (g = 410)      4004
0      5      11 12      --> (g = 1229)      0
0      3      5      11 12      --> (g = 1411)      0
0      7      11 12      --> (g = 1944)      0
0      3      7      11 12      --> (g = 2126)      4004
0      9      11 12      --> (g = 1229)      0
0      5      9      11 12      --> (g = 2230)      0
0      7      9      11 12      --> (g = 2945)      0
0      10      13      --> (g = 508)      0
0      1      10      13      --> (g = 515)      0

```

0			7	10	13	-->	(g = 2224)	0			
0	1		7	10	13	-->	(g = 2231)	0			
0			9	10	13	-->	(g = 1509)	0			
0			7	9	10	13	-->	(g = 3225)	0		
0				10	11	13	-->	(g = 690)	3003		
0			7	10	11	13	-->	(g = 2406)	0		
0			9	10	11	13	-->	(g = 1691)	0		
0			7	9	10	11	13	-->	(g = 3407)	0	
0					12	13	-->	(g = 53)	0		
0	3				12	13	-->	(g = 235)	364		
0		5			12	13	-->	(g = 1054)	0		
0	3	5			12	13	-->	(g = 1236)	4004		
0			7		12	13	-->	(g = 1769)	0		
0	3		7		12	13	-->	(g = 1951)	364		
0			9		12	13	-->	(g = 1054)	0		
0		5	9		12	13	-->	(g = 2055)	0		
0			7	9	12	13	-->	(g = 2770)	0		
0				11	12	13	-->	(g = 235)	0		
0	3			11	12	13	-->	(g = 417)	0		
0		5		11	12	13	-->	(g = 1236)	0		
0	3	5		11	12	13	-->	(g = 1418)	0		
0			7		11	12	13	-->	(g = 1951)	0	
0	3		7		11	12	13	-->	(g = 2133)	0	
0			9		11	12	13	-->	(g = 1236)	0	
0		5	9		11	12	13	-->	(g = 2237)	0	
0			7	9	11	12	13	-->	(g = 2952)	0	
0					14	-->	(g = 1)	8192			
0	1				14	-->	(g = 8)	6435			
0		3			14	-->	(g = 183)	3432			
0	1	3			14	-->	(g = 190)	4433			
0			5		14	-->	(g = 1002)	3432			
0	1		5		14	-->	(g = 1009)	3432			
0		3	5		14	-->	(g = 1184)	3432			
0	1	3	5		14	-->	(g = 1191)	3432			
0			9		14	-->	(g = 1002)	3432			
0	1		9		14	-->	(g = 1009)	3432			
0		3	9		14	-->	(g = 1184)	3432			
0	1	3	9		14	-->	(g = 1191)	3432			
0				11	14	-->	(g = 183)	3432			
0	1			11	14	-->	(g = 190)	4433			
0		5		11	14	-->	(g = 1184)	3432			
0	1	5		11	14	-->	(g = 1191)	3432			
0			9	11	14	-->	(g = 1184)	3432			
0	1		9	11	14	-->	(g = 1191)	7436			
0			10	12	14	-->	(g = 547)	0			
0	1		10	12	14	-->	(g = 554)	0			
0		3	10	12	14	-->	(g = 729)	0			
0	1	3	10	12	14	-->	(g = 736)	0			
0			9	10	12	14	-->	(g = 1548)	0		
0			10	11	12	14	-->	(g = 729)	0		
0			9	10	11	12	14	-->	(g = 1730)	0	
0					13	14	-->	(g = 8)	6435		
0	3				13	14	-->	(g = 190)	4433		
0		5			13	14	-->	(g = 1009)	3432		
0	3	5			13	14	-->	(g = 1191)	7436		
0			9		13	14	-->	(g = 1009)	3432		
0	3		9		13	14	-->	(g = 1191)	3432		
0				11	13	14	-->	(g = 190)	4433		
0		5		11	13	14	-->	(g = 1191)	3432		
0			9	11	13	14	-->	(g = 1191)	3432		
0			10	12	13	14	-->	(g = 554)	0		
0			9	10	12	13	14	-->	(g = 1555)	0	
0			10	11	12	13	14	-->	(g = 736)	0	
0			9	10	11	12	13	14	-->	(g = 1737)	0

This is a picture of the complementary graph of a graph G that looks very much like the complementary graph of G_1 (and of G_2). To get a real picture of the complement of G_1 , you should also connect the vertices that are diametrically placed with respect to the central point of the drawing.

In a drawing of the complement of G_1 , one should label the vertices as follows:

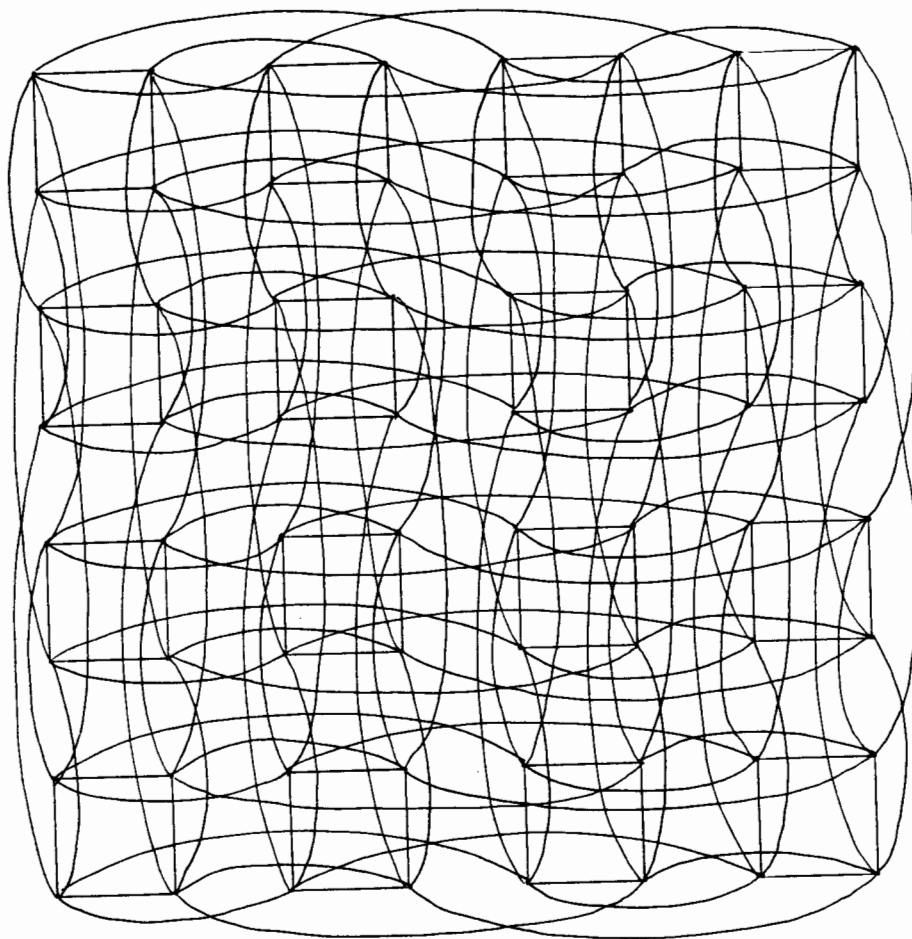
00	01	02	03	04	05	06	07
10	11	12	13	14	15	16	17
20	21	22	23	24	25	26	27
30	31	32	33	34	35	36	37
40	41	42	43	44	45	46	47
50	51	52	53	54	55	56	57
60	61	62	63	64	65	66	67
70	71	72	73	74	75	76	77

The labels are the octal numbers that correspond to the binary representation in $(\mathbb{Z}/2\mathbb{Z})^6$.

E.g.: The octal number 53 corresponds with the binary number 101011 and with the vector $(1, 0, 1, 0, 1, 1)$ of $(\mathbb{Z}/2\mathbb{Z})^6$.

To get a picture of the complement of G_2 , one should start from the same graph but choose the labels:

00	76	75	03	73	05	06	70
67	11	12	64	14	62	61	17
57	21	22	54	24	52	51	27
30	46	45	33	43	35	36	40
37	41	42	34	44	32	31	47
50	26	25	53	23	55	56	20
60	16	15	63	13	65	66	10
07	71	72	04	74	02	01	77



Review of the book : "Histoire de la Statistique"
by J.J. DROESBEKE & Ph. TASSI, P.U.F., Paris,
Collection "Que Sais-je ?" n° 2527, 1990

It is a challenge to write the history of a scientific discipline in the 128 pages of a "Que Sais-je ?". Though, that is what J.J. DROESBEKE and Ph. TASSI have done for Statistics. This little book is pleasant to read and it will give the reader (with some previous knowledge in the field) a general view of the roots and main streams of development of Statistical Theory. This book is all the more welcome that French literature is not so rich in histories of Statistics.

The first chapter deals with descriptive statistics. It is clear that Statistics is born of the need and desire of summarizing and interpreting large sets of collected data. Hence, graphical representations (first known bar chart in 1786), typical values (means, median, variance since Gauss and the mean squares method, ...), curve fitting, correlation and regression. Special attention is paid to the index numbers born three centuries ago (in England) for synthesizing price evolution.

Although many descriptive indices can be defined and used without probabilistic assumptions, probability is present in Statistics roughly since the mid XVIIIth century with the theory of errors of observation. As the main pieces of statistical theory rely heavily on Probability, the authors give a brief account of its history in the second chapter : a 16 pages trip from Pascal, Fermat and the Bernoulli Brothers to Kolmogorov. One learns for instance that the normal law, often attributed to Laplace and Gauss, can be traced back to the Ars Conjectandi of Jacques Bernoulli (1654-1705).

The main body of statistical knowledge is then split under the five following headings : Survey Sampling (Les sondages), The rise of statistical inference, Non parametric Statistics and robustness, Time series analysis, Data analysis. One may regret (but there is always something to regret in such a work) the absence of a chapter devoted to subjects like the general linear model and the analysis of variance or to the design of experiments. Anyway.

The fact that Survey Sampling appears first in the main chapters is probably due to its link with census (recensements) which are very ancient (perhaps 5000 B.C.). From the seventieth century on, the governments have shown growing interest in collecting data about the people and country they rule and, of course, at the cheapest possible cost. This led to survey sampling. The authors mainly depict the difficult genesis of the notion of representative sample.

The chapter on estimation and tests is centered on the competing theories of Fisher on one hand and Neyman-Pearson on the other hand. My regret here is that the respective positions of the bayesian and non-bayesian statisticians are not more extensively outlined. The authors would probably argue that the debate is not sufficiently mature by now to be able to give an objective presentation.

The next chapter is concerned with two types of more recently developed techniques : the theory of robust statistics and non-parametric statistics.

The fact that the observations depend on a very special and familiar parameter called time singularizes Time series analysis among stochastic processes and justifies special attention. The methods reviewed include graphical approaches, analysis in the frequency domain (based on the periodogram), methods of decomposition (like the famous CENSUS), analysis in the time domain (based on the correlograms), stochastic models including Box-Jenkins methodology.

The last but one chapter - the last one being devoted to a short biography of eight of the most prominent statisticians, those who made Statistics what it is - is about data analysis in the French meaning of the word. This chapter is "a gift of B. Fichet" to the authors. The history of the numerous methods is outlined : "metric methods" as principal component analysis, correspondance analysis and discriminant analysis (with or without distributional, i.e. normal, assumptions); non-metric methods, i.e. essentially, clustering.

In conclusion, the reading of this small book is recommended, as the authors say, to teachers of, searchers in and users of statistical methods. Several reading levels are possible : the general lines will be perceived by everyone while the details, especially in recent theories, will be understood by people with higher degree of education in the field.

Marc PIRLOT,
Faculté Polytechnique de Mons.