



Qüestió

Quaderns d'Estadística
i Investigació Operativa

Any 1997, volum 21, núm. 1 i 2
Segona època

Entitats patrocinadores:

Universitat de Barcelona
Universitat Politècnica de Catalunya
Institut d'Estadística de Catalunya

Entitat col·laboradora:

International Biometric Society



Generalitat de Catalunya
**Institut d'Estadística
de Catalunya**

SUMARI

Editorial	3
<i>Estadística</i>	
Estimating the contamination level of data in the framework of linear regression analysis A. Rubio and J.Á. Věšek	9
Remarques sur le maximum de vraisemblance	37
C. Huber et M. Nikulin	
Diseño y análisis bayesianos de réplicas en la experimentación científica	59
M.J. Bayarri y M.A. Martínez	
Análisis factorial de tablas mixtas: nuevas equivalencias entre ACP normado y ACM M.I. Landaluce Calvo	99
Estimadores de razón: una revisión	109
J.A. Mayor Gallego	
<i>Investigació Operativa</i>	
Heurístico para los problemas de rutas con carga y descarga en sistemas LIFO	153
J.A. Pacheco	
Cálculo de la distribución del tiempo de vida de componentes mediante autopsia en sistemas binarios aditivos, serie-paralelo y paralelo-serie	177
F. Mallor Giménez, C. Azcárate Camio y A. Pérez Prados	
A tabu search algorithm to schedule university examinations	201
R. Alvarez-Valdés, E. Crespo and J.M. Tamarit	
<i>Estadística Oficial: Recerca estadística promoguda per Eurostat</i>	
Confidentiality in the European statistical system	219
Ph. Nanopoulos	
Protecting micro-data by micro-aggregation: the experience in Eurostat	221
D. Defays	
Statistical databases: the reference environment and three layers proposed by Eurostat	233
R. Dubois	
Redaction automatique de commentaires a partir de bases de données statistiques	241
J.L. Roos	
<i>Biometria</i>	
Some unresolved issues in non-linear population dynamics	269
J.N. Perry	
Statistical analysis of the generalized additive semiparametric survival model with ran- dom covariates	273
V. Bagdonavičius and M. Nikulin	
El problema de Behrens-Fisher en la investigación biomédica. Análisis crítico de un estudio clínico mediante simulación	293
E. Vegas Lozano	
<i>Secció docent i problemes</i>	
Problemas tipo en teoría de la decisión	321
R. Alonso Sanz	
<i>Cursos i congressos d'estadística</i>	



EDITORIAL

El volum 21 de *Qüestió*, corresponent a l'any 1997, presenta dues particularitats. D'una banda, el seu número 3 constituirà una edició especial en la qual s'hi inclourà una relació exhaustiva dels articles publicats a *Qüestió* en el període 1988-97 i dels avaluadors que han participat en aquests darrers deu anys de la revista, reprenent així la recopilació efectuada des del 1977, any fundacional de la primera època de *Qüestió*, fins el 1987. D'altra banda, els números 1 i 2 d'aquest volum s'apleguen en una edició única per tal d'actualitzar el període temporal de referència, tot fent coincidir l'ordre cronològic amb el moment efectiu de la seva publicació; d'aquesta manera es pretén subsanar el retardament d'un trimestre que actualment pateix la revista, encara que anualment s'hagi publicat l'equivalent als tres números de què consta cada volum.

En qualsevol cas, tant pel nombre d'articles com per l'extensió global del seu contingut, el present número 1-2 suposa la suma de dos números ordinaris, amb la presència d'articles originals a totes les seccions temàtiques de la revista. També cal advertir que, en aquesta edició, s'actualitzen i completen les normes per a l'acceptació d'articles originals i altres materials d'interès per part de *Qüestió*. En aquest sentit, es clarifiquen els processos de presentació i els requisits per a l'avaluació, es faciliten les eines i els formats adients per als autors i, al mateix temps, s'agiliten les tasques d'edició de la revista. Per últim, cal indicar que el preu de subscripció de *Qüestió* per enguany es manté igual que el de l'any 1996, en un esforç de contenció de la despesa a càrrec dels seus subscriptors.

Comentari de les seccions «Estadística», «Investigació Operativa» i «Biometria»

En aquest número doble, les seccions esmentades contenen un total d'onze articles. A la secció d'Estadística, l'article «Estimating the contamination level of data in the framework of linear regression analysis», d'A. Rubio i J.Á. Víšek, proposa un estimador del nivell de contaminació de les dades en el model lineal, amb resultats asintòtics i de simulació; a «Remarques sur le maximum de vraisemblance», de C. Huber i M. Nikulin, s'exposen tant els aspectes positius com els més conflictius de l'estimació per màxima versemblança, que és il·lustrada amb un seguit d'exemples i paradoxes; a «Diseño y análisis bayesianos de réplicas en la experimentación científica», de M.J. Bayarri i M.A. Martínez, s'analitzen, mitjançant tècniques baiesianes jeràrquiques, els resultats de diferents rèpliques d'experiències estadístiques amb importants conclusions sobre els resultats més probables a partir de les distribucions predictives obtingudes; l'article de M. Isabel Landaluze, «Análisis factorial de tablas mixtas: nuevas equivalencias entre ACP normado y ACM», és una aproximació a la representació de taules de dades mixtes via l'anàlisi de components principals

i de correspondències múltiples; per últim, a «Estimadores de razón: una revisión», de J.A. Mayor, es revisa l'estimació de paràmetres en poblacions finites mitjançant estimadors de raó, estudiant-se en especial el seu biaix.

Dels tres articles d'Investigació Operativa, el treball de J. Pacheco, «Heurístico para los problemas de rutas con carga y descarga en sistemas LIFO», conté un algorisme sobre problemes de càrrega i descàrrega de mercaderies en vehicles, aconseguint abordar problemes de mida gran però computacionalment viables; a «Cálculo de la distribución del tiempo de vida de componentes mediante autopsia en sistemas binarios aditivos, serie-paralelo y paralelo-serie», de F. Mallor, C. Azcárate i A. Pérez, s'ofereixen resultats coneguts, però sovint complicats, sobre la distribució del temps de vida d'un sistema i proposa un mètode d'obtenció més restringit però més senzill i pràctic; el tercer article, «A tabu search algorithm to schedule university examinations», de R. Álvarez, E. Crespo i J.M. Tamarit, presenta un algorisme per organitzar dates d'exàmens i assignar aules, evitant coincidències de dos o més exàmens en un mateix dia i espaiant el temps entre exàmens.

La secció de Biometria inclou també tres articles, el primer dels quals, «Some unresolved issues in non-linear population dynamics», de J.N. Perry, tracta dels problemes no resolts sobre la dependència del comportament dinàmic d'un sistema que, sovint, és utilitzat per la dinàmica de sèries temporals d'insectes; a «Analysis of the generalized additive semiparametric model with covariates», de V. Bagdonavičius i M. Nikulin, es proposen generalitzacions de funcions de risc additives i s'estudien estimacions de paràmetres quan les variables covariants són aleatòries; en el darrer treball, «El problema de Behrens-Fisher en la investigación biomédica. Análisis crítico de un estudio clínico mediante simulación», de E. Vegas, es fa un compendi del problema de Behrens-Fisher que s'il·lustra amb una aplicació a la recerca biomèdica mitjançant simulacions.

Carles Cuadras, Editor Executiu

Comentari de la secció «Estadística Oficial» i altres apartats

De nou, la secció d'Estadística Oficial agrupa els seus quatre originals a l'entorn d'una mateixa temàtica. En aquest cas s'ofereix una selecció de recerques recents realitzades o promogudes per l'Oficina Estadística de la Unió Europea (Eurostat) que varen avançar-se en el darrer COMPSTAT'96, a la sessió paral·lela sobre R+D en l'estadística oficial organitzada per l'Institut d'Estadística de Catalunya. Així, a la nota «Confidentiality in the European Statistical System», de P. Nanopoulos, es resumeix el marc d'actuació comunitària actual en termes jurídics i metodològics, especialment en la microagregació de dades estadístiques sensibles. La tasca directa d'Eurostat en aquest àmbit es detalla a «Protecting micro-data by micro-aggregation: the experience by Eurostat», de D. Defays, on s'apunten les propietats de l'agregació de dades mitjançant clusters de mida fixa i la seva aplicació pràctica a l'enquesta «Community Innovation Survey». Ambdós originals representen una continuïtat amb l'anterior número 3 del volum 20 (1996) i el volum 17 (1993), dedicats monogràficament a les tècniques de preservació del secret estadístic. Els altres dos originals se centren en l'optimització de les bases de dades estadístiques oficials; el primer, «Statistical databases: the reference environment and three layers», és de R. Dubois el qual emfasitza els criteris proposats per Eurostat per al processament adient dels resultats estadístics nacionals en l'àmbit dels seus sistemes d'informació. Per últim, «Redaction automatique de commentaires a partir de bases de données statistiques», de J.L. Roos, és una aportació de l'Institut d'Estadística de França (INSEE) al programa estadístic comunitari centrat en la generació automàtica de comentaris, d'interès en la producció regular de resultats estadístics residents a bases de dades i, especialment, per a l'anàlisi conjuntural d'informació socio-econòmica.

En darrer lloc, l'apartat dedicat a «Cursos i congressos d'estadística» fa referència a la propera celebració de la «Conference on Statistical Data Protection'98» (Lisboa, 25-27 de març de 1998), a càrrec d'Eurostat/Comissió Europea, convocant la presentació de treballs en els tres àmbits previstos: teoria, aplicacions i problemes computacionals en el control de la revelació estadística. Cal notar la voluntat dels organismes comunitaris en consolidar una nova disciplina de caràcter estadístico-computacional que, a més, serveixi per validar els resultats dels projectes ESPRIT aollits en el 4th Framework Program de la Unió Europea.

Enric Ripoll, Editor Executiu

Estadística

ESTIMATING THE CONTAMINATION LEVEL OF DATA IN THE FRAMEWORK OF LINEAR REGRESSION ANALYSIS^{*}

A. RUBIO^{*}

Universidad de Extremadura

J.Á. VÍŠEK^{**}

Academy of Sciences of the Czech Republic

An estimator of the contamination level of data is proposed in the framework of linear models and its asymptotic behavior is investigated. A numerical study illustrates its finite sample performance under an alternative.

Keywords: Huber's convex mixture model, contamination level estimation, linear model, asymptotic behavior, simulation study.

AMS Classification: Primary 62F35, secondary 62J05

* A. Rubio, Departamento de Matemáticas. Universidad de Extremadura. 10071 Cáceres.

** J.Á. Víšek, Department of Macroeconomics and Finance. Institute of Economic Studies. Faculty of Social Sciences. Charles University. Staré město. Smetanovo nábřeží 6, 110 01 Prague 1 & Department of Stochastic Informatics, Institute of Information Theory and Automation, Academy of Sciences of the Czech Republic. Pod vodárenskou věží 4. 182 08 Prague 8.

* This work was supported by Dirección General de Investigación Científica y Técnica (DGICYT), Ministry of Education and Sciences of Spanish Government

– Article rebut el maig de 1992.

– Acceptat l'octubre de 1996.

1. INTRODUCTION

It is well known that since the days of Sir Francis Galton (1886) the regression analysis has been used as an efficient tool of modelling the data. In the past it employed mostly the least squares (Legendre (1805), Gauss (1809)) although the method of the least absolute deviations has been proposed much earlier (Galilei (1632), Boscovisch (1757), Laplace (1793)). During the past quarter of this century a lot of other methods have appeared. Some of them produce in fact not a single estimator but the whole family of model estimators, in which every estimator corresponds to some value of the (tuning) parameter(s). Everything will be clear from the following example considering the family of the M -estimators of regression model corresponding to the family of Huber's ψ -functions, $\{\psi_k(z)\}_{k \in (0, \infty)}$ where

$$(1) \quad \psi_k(z) = \begin{cases} z & \text{for } |z| \leq k, \\ k \cdot \text{sign}(z) & \text{otherwise.} \end{cases}$$

The family of estimators $\{\hat{\beta}^{(n,k)}\}_{k \in (0, \infty)}$ is then given by

$$\hat{\beta}^{(n,k)} = \min_{\beta \in R^p} \left\{ \sum_{i=1}^n \rho_k(Y_i - X_i^T \beta) \right\}$$

where $\rho_k(z)$ is a criterial function with the derivative equal to $\psi_k(z)$ (for Y 's and X 's see (4) below). The value of the tuning constant k is in applications selected on the basis of experiences. We shall recall later (after introducing necessary notations) Huber's result which connects the optimal value of the tuning constant k with the mixture parameter ε in the Huber model of contamination. We shall also see that the «minimal» mixture parameter ε corresponds (in one-to-one way) to the contamination level $\varepsilon_{G,F}$ of data. It means that the optimal value k of the tuning constant depends on the contamination level $\varepsilon_{G,F}$. Now, if we select the value of k so that it «underestimates the contamination level» we may obtain wrong model of data. On the other hand «overestimating contamination level» leads to a needless loss of efficiency. Of course, a loss of efficiency is (typically) small and hence it is not important, see Věšek (1993).

There is however another problem. It is known that some robust methods focus themselves on a (too) restricted part of data considerably weighting down (or depressing completely) the influence of other part. Such behaviour may be observed especially for the methods with large breakdown point, i.e. for the methods assuming (extremely) high contamination level. So using methods designed for the high contamination level —as the least median of squares, the least trimmed squares or minimal biased estimators— we may obtain somewhat (or completely) misleading

estimate of model. Let us give (at least) one numerical example. First of all, let us recall definitions of the estimators which will be used. Let us put

$$(2) \quad r_i(\beta) = Y_i - \sum_{j=1}^p X_{ij} \beta_j, \quad i = 1, 2, \dots, n \quad h = \left\lfloor \frac{n}{2} \right\rfloor + \left\lfloor \frac{p+1}{2} \right\rfloor$$

and let $r_{(i:n)}^2(\beta)$ be the i -th order statistics among $r_i^2(\beta)$, $i = 1, 2, \dots, n$. Further, let us recall that (with h given by (2))

$$(3) \quad \begin{aligned} \hat{\beta}^{LS} &= \operatorname{argmin}_{\beta \in R^p} \sum_{i=1}^n r_i^2(\beta), & \hat{\beta}^{LMS} &= \operatorname{argmin}_{\beta \in R^p} r_{(h:n)}^2(\beta), \\ \hat{\beta}^{LTS} &= \operatorname{argmin}_{\beta \in R^p} \sum_{i=1}^h r_{(i:n)}^2(\beta), & \hat{\beta}^{(\rho_k)} &= \operatorname{argmin}_{\beta \in R^p} \sum_{i=1}^n \rho_k(r_i(\beta)) \end{aligned}$$

with ρ_1 – Huber’s function (with $\psi_1(t) = t$ for $|t| < c$, $\psi_1(t) = c \cdot \operatorname{sign} t$ otherwise) and ρ_2 – Hampel’s function (with $\psi_2(t) = \psi_1(t)$ for $|t| < 1.2c$, $\psi_2(t) = [c - \frac{5}{9}(t - 1.2c)] \cdot \operatorname{sign} t$ for $1.2c < |t| < 3c$ and zero otherwise; both with the tuning constant $c = 1.2$). Finally, let

$$\hat{\beta}^{L_1} = \operatorname{argmin}_{\beta \in R^p} \sum_{i=1}^n |r_i(\beta)| \quad \text{and} \quad \hat{\beta}^{TLS} = \operatorname{argmin}_{\beta \in R^p} \sum_{i \in I_\alpha} r_i^2(\beta)$$

where I_α is the index-set of points obtained by the symmetric trimming according to α -regression quantiles of Koenker and Bassett (1978) (value of α was 0.1).

Example. Demographic Data (49 cases, Chatterjee and Hadi (1988)). Dependence of the gross national product per capita on: the infant death per thousand live births, the number of inhabitants per physician, the population per km², the population per 10³ ha of agricultural land, the percentage of literate population over 15 years of age, the number of students enrolled in higher education per 10⁵ population. (The software which was used for evaluation was either prepared by the experienced statisticians as Jaromír Antoch (1991, 1992), Roger Koenker (1978) or Alvio Marazzi (1992), and we are grateful of possibility to utilize it, or by our colleagues and it was tested in an extensive numerical studies, see e. g. Vřšek (1996 a, 1997).)

(Value of α for *TLS* was 0.1.) It is not even difficult to find artificial (one-dimensional) data for which *LMS* and *LTS* estimates are orthogonal each to other and *S*-estimate divides the angel between them on two halves, see Vřšek (1994 a). (An objection may appear that perhaps in the example with Demographic data some regressors are insignificant. However, in the case of contaminated data, it is not simple task to say which regressors are significant and which not —for more arguments see

Víšek (1996 c).) Similarly, rather diverse results may be obtained when using different constants k when evaluating M -estimates.

Table 1. *Demographic Data*

Method	LS	LMS	LTS	TLS	L_1	Huber	Hampel
intercept	112.885	331.095	103.563	480.509	148.732	33.459	-146.586
Infant	-3.621	-2.774	-1.526	-6.764	-3.964	-3.025	-2.029
Inhabitants	0.009	-0.017	0.005	-0.013	0.032	0.015	0.032
Population	0.186	-1.024	0.009	-1.265	0.088	-0.052	0.242
Agriculture	0.003	0.072	-0.001	0.085	-0.005	0.000	-0.008
Literate	5.566	2.501	3.929	3.793	2.985	5.280	4.339
Higher	0.693	0.249	0.295	0.373	0.860	0.734	1.072

Table 2. *Demographic Data - Huber's estimator*

Tuning constant	0.6	0.8	1.0	1.2	1.5	2.0
intercept	110.136	69.369	42.526	33.459	67.058	113.059
Infant	-3.254	-3.191	-3.144	-3.025	-3.226	-3.545
Inhabitants	0.013	0.016	0.017	0.015	0.013	0.009
Population	0.018	0.013	-0.015	-0.052	-0.106	-0.165
Agriculture	-0.002	-0.001	-0.001	0.000	0.001	0.003
Literate	4.002	4.422	4.917	5.280	5.308	5.319
Higher	0.777	0.782	0.753	0.734	0.722	0.705

For other examples of such situations see Rubio *et al.* (1992) or Víšek (1994 a), (1994 b)).

So, one possible way how to begin any processing of data may be to estimate the contamination level of them and then to apply that method (from the family in question) which corresponds to the estimated contamination level. Of course, the problem with the selection of an adequate robust method is a complicated problem and we shall return to it briefly at the end of paper.

As we shall see later we will need for the estimation of the contamination level to estimate the density of data. It is simple to do it directly when we have at hand

a sample of i.i.d. r.v.'s, e.g. in the location problem, but naturally it cannot be applied directly on the response variable in the regression scheme. Nevertheless as we shall see below some preliminary considerations will open a straightforward way to a proposal of an estimator of contamination level in regression model, too.

One may learn from the text given below that the estimation of the contamination level is analogous to the estimation of the mixture parameter (see Rukhin(1994)). The difference is that in the case of the estimation of the contamination level we do not specify the «contaminating» distribution but only the «central» model. Of course, it implies that when we estimate the contamination level we may meet some additional (technical) difficulties in comparison with the estimation of the mixture parameter. So, taking into account that when looking for an estimator of the mixture parameter we are rarely able to prove more than the consistency, we cannot expect that for the estimator of the contamination level we will be able to give more than an asymptotic results about its behavior.

The results are accompanied by a simulation study. The reasons for the simulation study will be also discussed.

2. ESTIMATING CONTAMINATION LEVEL

Let N denote the set of all positive integers, R the real line and R^p the p -dimensional Euclidian space. We shall consider the linear model

$$(4) \quad Y = X \cdot \beta^0 + e$$

where for every $n \in N$ and some (fix) $p \in N$ we have $Y = (Y_1, Y_2, \dots, Y_n)^T$ (a response variable), $X = (x_{ij})_{i=1,2,\dots,n}^{j=1,2,\dots,p}$ (a design matrix), $\beta^0 = (\beta_1^0, \beta_2^0, \dots, \beta_p^0)^T$ (regression coefficients) and $e = (e_1, e_2, \dots, e_n)^T$ (random disturbances). We assume that the random variables in the sequence $\{e_i\}_{i=1}^\infty$ are i.i.d. and they are defined on a probability space $(\Omega, \mathcal{B}, P)$, and the carriers x_{ij} 's are fix and known while β^0 (the «true» value of the vector of regression coefficients) is unknown but also fix. (Let us remark here that in what follows all probabilistic assertions as «a. s.», «in probability» etc. will be understood with respect to P .) Assume moreover that there is a $K < \infty$ such that

$$(5) \quad \sup_{i \in N} \max_{j=1,2,\dots,p} |x_{ij}| < K$$

and

$$(6) \quad \lim_{n \rightarrow \infty} \frac{1}{n} X^T X = Q$$

where Q is a regular matrix. Sometimes we will use also an alternative notation $X_i = (x_{i1}, x_{i2}, \dots, x_{ip})^T$, so that the expression $X_i^T \beta$ will stay instead of the sum $\sum_{j=1}^p x_{ij} \beta_j$. Finally, let ν be a σ -finite measure defined on $(R, \mathcal{B}(R))$ (where $\mathcal{B}(R)$ is the Borel σ -algebra) such that distribution function F of e_1 is absolutely continuous with respect to ν , and then let us denote by $\mathcal{F}_\nu$ the set of all distribution functions which are absolutely continuous with respect to ν .

In Huber (1964) the convex combination model

$$(7) \quad G(z) = (1 - \varepsilon)F(z) + \varepsilon H(z)$$

($\varepsilon \in [0, 1], F, H \in \mathcal{F}_\nu$) was used to describe the contamination of data. So, the first idea how to define the contamination level may start with this model. One finds immediately that for given G the decomposition (7) is given uniquely only in the case when

$$\int_{\{z: g(z)=0\}} dF(z) > 0$$

where g is density of the distribution function G (with respect to ν). Then $\varepsilon = 1$ and $H(z) = G(z)$. If however (7) holds for some $\varepsilon \in [0, 1]$ then also for any $\varepsilon^* \in [\varepsilon, 1]$ we may write

$$G(z) = (1 - \varepsilon^*)F(z) + \varepsilon^* H^*(z)$$

where $H^*(z) = (\varepsilon^*)^{-1} \{(\varepsilon^* - \varepsilon)F(z) + \varepsilon H(z)\}$ and we evidently have $H^* \in \mathcal{F}_\nu$. It hints that the following definition of the contamination level, for the version using the densities, has to consider an essential minimum of possible values of ε 's, essential with respect to the measure ν (notice that for our case the result is same if we take «ess inf» with respect to F because f is involved in the characterization of $\varepsilon_{G,F}$).

Definition 1 For G and $F \in \mathcal{F}_\nu$ we shall define the contamination level as a value

$$\varepsilon_{G,F} = \inf \{ \varepsilon : G(z) = (1 - \varepsilon)F(z) + \varepsilon H(z), H \in \mathcal{F}_\nu \}$$

or equivalently

$$\varepsilon_{G,F} = \inf \{ \varepsilon : \nu \{ z : g(z) \neq (1 - \varepsilon)f(z) + \varepsilon h(z) \} = 0, H \in \mathcal{F}_\nu \}$$

where g, f and h are again the densities of G, F and of H with respect to ν , respectively.

Remark 1 The mixture

$$G(z) = (1 - \varepsilon)F(z) + \varepsilon H(z)$$

expresses the fact that the bulk of data, the «proper data», are distributed according to F and some portion of data (this portion being considered as contamination) is distributed according to H . The distribution H is not usually fully specified while the

distribution function F is selected (at least as a type of distribution) by the statistician when processing data. So in other words, having at hand data, we assume that they were generated by some unknown G , nevertheless, we would like to describe them (or explain them) by F . The reason for selecting the distribution F instead of the «true» distribution G may be, e.g., the fact that the distribution function G is expected to be too «wild» while the distribution function F is easy (or at least easier) to work with, and at the same time being a reasonable approximation of the distribution G . However we are aware that it need not be precisely the distribution which has generated the data, i.e. that $G \neq F$, and hence we admit that there is a distribution H and $\epsilon > 0$ so that $G = (1 - \epsilon)F + \epsilon H$. It means that we select F and hope that the «true» distribution of data lies in a neighborhood of F . $\square$

Remark 2 In the last twenty years a lot of others contamination models have appeared being based on different types of distances in the space of probability measures (e.g. on the Kolmogorov-Smirnov or Prokhorov metrics (see Huber (1981)), on the combination of convex combinations and total variation (see Rieder (1977)), on 2-alternating capacities (see Huber, Strassen (1973)), or on divergences (see Vajda (1989)). We believe that a straightforward generalization of the notion of contamination level is possible in any of these contamination models. We also believe that for such generalization the modifications of the theory which will be presented below, will be also straightforward (something has been already done for the location problem, see Věšek (1989)). $\square$

Lemma 1 (characterization of the contamination level $\epsilon_{G,F}$) Let $G(z)$ and $F(z) \in \mathcal{F}_v$. Then

$$\epsilon_{G,F} = \left\| \frac{f(z) - g(z)}{f(z)} I_{\{z: f(z) > g(z)\}} \right\|_{\infty}$$

$f(z)$ and $g(z)$ being again the densities of the distributions $F(z)$ and $G(z)$ with respect to v .

Proof: At first notice that

$$\left\| \frac{f(z) - g(z)}{f(z)} I_{\{z: f(z) > g(z)\}} \right\|_{\infty} \geq 0$$

and the lower bound is attained only when $v\{z : f(z) \neq g(z)\} = 0$. Let us denote by $\mathcal{E}_{G,F}$ the set of ϵ 's for which (7) holds, and let $\epsilon \in \mathcal{E}_{G,F}$. It implies that for some density $h(z)$ we have $g(z) = (1 - \epsilon)f(z) + \epsilon h(z)$, and hence $f(z) - g(z) \leq \epsilon f(z)$, i.e. $\epsilon \geq$

$\frac{f(z)-g(z)}{f(z)}$ for any $z \in R$ for which $f(z) > 0$. Then of course $\varepsilon \geq \left\| \frac{f(z)-g(z)}{f(z)} I_{\{z: f(z) > g(z)\}} \right\|_\infty$, and hence also $\varepsilon_{G,F} \geq \left\| \frac{f(z)-g(z)}{f(z)} I_{\{z: f(z) > g(z)\}} \right\|_\infty$.

On the other hand for any $\varepsilon \geq \left\| \frac{f(z)-g(z)}{f(z)} I_{\{z: f(z) > g(z)\}} \right\|_\infty$, $h_\varepsilon(z) = \frac{1}{\varepsilon} [g(z) - (1 - \varepsilon)f(z)]$ is a density, and

$$g(z) = (1 - \varepsilon)f(z) + \varepsilon h_\varepsilon(z).$$

It means that $\varepsilon \in \mathcal{E}_{G,F}$ and $\varepsilon_{G,F} \leq \left\| \frac{f(z)-g(z)}{f(z)} I_{\{z: f(z) > g(z)\}} \right\|_\infty$. ■

Remark 3 It follows from the proof of Lemma 1 that the infimum in Definition 1 is attained so that we may write

$$G(z) = (1 - \varepsilon_{G,F})F(z) + \varepsilon_{G,F}H_{\varepsilon_{G,F}}(z)$$

for some $H_{\varepsilon_{G,F}}(z)$. It may be of interest that for a convex mixture of two normal distribution, say

$$(8) \quad (1 - \varepsilon_H)(2\pi)^{-\frac{1}{2}} \exp\left(-\frac{z^2}{2}\right) + \varepsilon_H(2\pi)^{-\frac{1}{2}} \sigma^{-1} \exp\left(-\frac{z^2}{2\sigma^2}\right),$$

we obtain

$$(9) \quad \varepsilon_{G,F} = \varepsilon_H \quad \text{for } \sigma < 1$$

and

$$(10) \quad \varepsilon_{G,F} = (1 - \sigma^{-1})\varepsilon_H \quad \text{for } \sigma > 1.$$

Let us recall that Huber's result (proved in 1964, under condition that the logarithm of the density of central model is concave) says that the optimal selection of the ψ_k -function (see (1)) in the case that the data were generated by the mixture density (7) is given by $k = k(\varepsilon_H)$ where $k(\varepsilon_H)$ satisfies

$$(11) \quad (1 - \varepsilon_H)^{-1} = \int_{t_0(\varepsilon_H)}^{t_1(\varepsilon_H)} g(t) dt + \frac{g(t_0(\varepsilon_H)) + g(t_1(\varepsilon_H))}{k(\varepsilon_H)}$$

where $t_i(\varepsilon_H)$'s are given by

$$(12) \quad \frac{g'(t_i(\varepsilon_H))}{g(t_i(\varepsilon_H))} = (-1)^i \cdot k(\varepsilon_H), \quad i = 0, 1.$$

So an alternative notation for Huber's family of ψ -functions may be

$$(13) \quad \{\Psi_{k(\varepsilon)}\}_{\varepsilon \in (0,1)}.$$

We shall need it in the discussion later. □

Now we need to estimate $\|\frac{f(z)-g(z)}{f(z)}I_{\{z:f(z)>g(z)\}}\|_\infty$. Since the density $f(z)$ is known (see Remark 1) it is necessary to estimate the density $g(z)$. To be more precise, let $\{Z_k(\omega)\}_{k=1}^\infty$ be a sequence of i.i.d. random variables (defined on the probability space $(\Omega, \mathcal{B}, P)$). Assume that the corresponding distribution function $G_Z(z)$ belongs to $\mathcal{F}_V$ and denote by $g_Z(z)$ its density. We shall assume that $g_Z(z)$ vanishes outside an interval (c, d) , $-\infty \leq c < d \leq \infty$. Moreover, following Csörgö and Révész (1981) and Rosenblatt (1971), let us assume:

Conditions A *The bounded integrable kernel $w : R \rightarrow R$ vanishes outside an interval $(-a, a)$ with $-\infty \leq -a \leq c < d \leq a \leq \infty$ for some $a > 0$. Moreover, it is twice continuously differentiable with bounded first derivative $w'(z)$, say $\sup_{z \in R} |w'(z)| < L < \infty$ and with $(1+z^4)|w''(z)|$ also bounded. Finally, the kernel has a Fourier transform $\phi(t)$ with $(1+t^2) \cdot \phi(t)$ integrable and is symmetric, i. e. $w(z) = w(-z)$.*

Please notice that Conditions A imply that

$$(14) \quad \int_{-\infty}^{\infty} w'(z) dz = 0 \quad (\text{due to symmetry})$$

and

$$(15) \quad \left| \int_{-\infty}^{\infty} z \cdot w'(z) dz \right| < \infty.$$

Let us write $Z^T(\omega, n) = (Z_1(\omega), Z_2(\omega), \dots, Z_n(\omega))$. Then define for any $z \in R$, $\omega \in \Omega$ the kernel estimator of density by

$$(16) \quad \hat{g}_n(z, Z(\omega, n)) = \frac{1}{nc_n} \sum_{i=1}^n w(c_n^{-1}(z - Z_i(\omega)))$$

and in what follows we shall assume that the sequence of the bandwidths $\{c_n\}_{n=1}^\infty \searrow 0$.

For the simplicity we shall assume that in the rest of paper ν is the Lebesgue measure.

Lemma 2 (Csörgö, Révész). *Let us denote*

$$\pi_n(g) = \sup_{z \in R} |\mathbb{E} \hat{g}_n(z, Z(\omega, n)) - g_Z(z)|.$$

Then

$$\sup_{z \in (-a, a)} |n(\hat{g}_n(z, Z(\omega, n)) - g_Z(z)) - \Gamma_n(z)| = O(c_n^{-1} \log^2 n + n\pi_n(g)) \quad \text{a. s.}$$

where $\Gamma_n(z)$ is the following Gaussian process:

$$(17) \quad \Gamma_n(z) = \frac{1}{c_n} \int_0^1 \mathcal{K}(v, n) dw(c_n^{-1}(z - G^{-1}(v)))$$

with $\mathcal{K}(v, n)$ being a suitable Kiefer process.

The proof is a specification of the proof of the Theorem 6.1.1 of Csörgö, Révész (1981), page 223.

Lemma 3 Let $G(z)$ and $F(z)$ belong to $\mathcal{F}_v$. Further, let $\{S_n\}_{n=1}^\infty$ be a family of sets, $S_n \in \mathcal{B}(R)$, such that

$$(18) \quad S_n \nearrow \{z : f(z) > 0\}$$

and

$$(19) \quad [\inf_{z \in S_n} f(z)]^{-1} \max \left\{ \pi_n(g), n^{-1} c_n^{-1} \log^2 n, n^{-\frac{1}{2}} c_n^{-1} (\log \log n)^{\frac{1}{2}} \right\} = o(1).$$

Then

$$\sup_{z \in S_n} \frac{|\hat{g}_n(z, Z(\omega, n)) - g(z)|}{f(z)} = o(1) \quad \text{a. s.} \quad \text{as } n \rightarrow \infty.$$

Proof: Let us denote by $M_n = [\inf_{z \in S_n} f(z)]^{-1}$. Now from the Lemma 2 we have for some $\Delta < \infty$

$$\begin{aligned} -\Delta M_n (\pi_n(g) + c_n^{-1} n^{-1} \log^2 n) &\leq \frac{\hat{g}_n(z, Z(\omega, n)) - g(z)}{f(z)} - \frac{\Gamma_n(z)}{nf(z)} \\ &\leq \Delta M_n (\pi_n(g) + c_n^{-1} n^{-1} \log^2 n). \end{aligned}$$

Since the lower as well as the upper bound tends to zero it suffice to show that

$$(20) \quad \sup_{z \in S_n} \frac{|\Gamma_n(z)|}{nf(z)} = o(1).$$

However, due to (17) we have

$$c_n |\Gamma_n(z)| \leq \sup_{z \in R} |\mathcal{K}(z, n)| \int_0^1 dw(c_n^{-1}(z - G^{-1}(v))) \leq \Delta \cdot \sup_{z \in R} |\mathcal{K}(z, n)|$$

where $\mathcal{K}(z, n)$ is again an appropriate Kiefer process. Now, due to (19) for any $\varepsilon > 0$ we may find $n_\varepsilon \in N$ so that for any $n > n_\varepsilon$

$$\sup_{z \in S_n} \frac{(\log \log n)^{\frac{1}{2}}}{c_n n^{\frac{1}{2}} f(z)} \leq \varepsilon,$$

and we have

$$\sup_{z \in S_n} \frac{|\Gamma_n(z)|}{nf(z)} \leq \frac{\Delta}{c_n n} \sup_{z \in R} \frac{|\mathcal{K}(z, n)|}{f(z)} \leq n^{-\frac{1}{2}} \varepsilon \Delta \sup_{z \in R} |\mathcal{K}(z, n)| (\log \log n)^{-\frac{1}{2}}.$$

The proof of (20) then follows due to the law of iterated logarithm for the Kiefer process (see Corollary 1.15.1 of Csörgő, Révész (1981)).

■

As we have mentioned above we need to estimate $\|\frac{f(z)-g(z)}{f(z)} I_{\{z: f(z) > g(z)\}}\|_\infty$ and let us recall that

$$\|\frac{f(z)-g(z)}{f(z)} I_{\{z: f(z) > g(z)\}}\|_\infty = \inf \left\{ a : v\left(\left\{ \frac{z}{f(z)@} > a \right\}\right) = 0 \right\}.$$

A routine arguments yields that there are $\tilde{f}(z)$ and $\tilde{g}(z)$ such that $v(\{f(z) \neq \tilde{f}(z)\}) = 0$ and $v(\{g(z) \neq \tilde{g}(z)\}) = 0$, and

$$\|\frac{f(z)-g(z)}{f(z)} I_{\{z: f(z) > g(z)\}}\|_\infty = \sup_{z \in \{t: \tilde{f}(t) > 0\}} \frac{\tilde{f}(z) - \tilde{g}(z)}{\tilde{f}(z)}$$

and so we may assume that $f(z)$ and $g(z)$ are such versions of densities of $F(z)$ and of $G(z)$ that

$$\|\frac{f(z)-g(z)}{f(z)} I_{\{z: f(z) > g(z)\}}\|_\infty = \sup_{z \in \{t: f(t) > 0\}} \frac{f(z) - g(z)}{f(z)}.$$

Let us assume that we shall work in the rest of paper with such versions of densities. Then from the previous lemma it follows that we may estimate $\|\frac{f(z)-g(z)}{f(z)} I_{\{z: f(z) > g(z)\}}\|_\infty$ by

$$(21) \quad \sup_{z \in S_n} \frac{f(z) - \hat{g}_n(z, Z(\omega, n))}{f(z)}$$

due to the fact that

$$(22) \quad \begin{aligned} & \sup_{z \in \{t: f(t) > 0\}} \frac{f(z) - g(z)}{f(z)} - \sup_{z \in S_n} \frac{f(z) - \hat{g}_n(z, Z(\omega, n))}{f(z)} \\ & \leq \sup_{z \in S_n} \frac{\hat{g}_n(z, Z(\omega, n)) - g(z)}{f(z)} = o(1) \quad \text{a. s.} \quad \text{as } n \rightarrow \infty \end{aligned}$$

(realize that we have assumed $S_n \nearrow \{z : f(z) > 0\}$) and similarly

$$\sup_{z \in S_n} \frac{f(z) - \hat{g}_n(z, Z(\omega, n))}{f(z)} - \sup_{z \in \{t: f(t) > 0\}} \frac{f(z) - g(z)}{f(z)}$$

$$(23) \quad \leq \sup_{z \in S_n} \frac{g(z) - \hat{g}_n(z, Z(\omega, n))}{f(z)} = o(1) \quad \text{a. s.} \quad \text{as } n \rightarrow \infty.$$

But let us consider the case when $g(z) = f(z)$, i.e. when there is no contamination. Then, due to the random fluctuations of $\hat{g}_n(z, Z(\omega, n))$, we obtain positive value at (21) for every $\omega \in \Omega$. That is the reason why we shall not propose in the following theorem the estimator of $\hat{\epsilon}_{G,F}$ as $\sup_{z \in S_n} \frac{f(z) - \hat{g}_n(z, Z(\omega, n))}{f(z)}$ but in a little modified form.

Theorem 1 *Let G and F belong to $\mathcal{F}_V$. Further let $\{\hat{\beta}_n\}_{n=1}^\infty$ be a sequence of $\sqrt{n}$ -consistent estimators of β^0 and $\{V_n\}_{n=1}^\infty$ sequence of sets, $V_n \in \mathcal{B}(R)$, fulfilling the assumptions (18) and (19), and moreover let*

$$(24) \quad [\inf_{z \in V_n} f(z)]^{-1} c_n^{-1} \|\hat{\beta}_n - \beta^0\| = o_p(1).$$

Finally, let $\{f_n(z, \omega)\}_{n=1}^\infty$ be a sequence of random processes such that

$$(25) \quad \sup_{z \in V_n} |f_n(z, \omega) - f(z)| = o_p(\inf_{z \in V_n} f(z))$$

and put

$$\hat{\epsilon}_{G,F} = \max \left\{ 0, \min \left\{ \sup_{z \in V_n} \frac{f_n(z, \omega) - \hat{g}_n(z, r(\hat{\beta}_n))}{f(z)}, 1 \right\} \right\}$$

where we have denoted for any $\beta \in R^p$ by $r(\beta)$ the vector of residuals $(Y_1 - X_1^T \beta, Y_2 - X_2^T \beta, \dots, Y_n - X_n^T \beta)^T$. Then

$$(26) \quad \hat{\epsilon}_{G,F} - \epsilon_{G,F} = o_p(1).$$

Proof: If we had known the true value β^0 of the regression coefficients we may evaluate the «theoretical» residuals $Y_1 - X_1^T \beta^0, Y_2 - X_2^T \beta^0, \dots, Y_n - X_n^T \beta^0$ and to plug them into the kernel estimator of density and the (26) would have followed directly from Lemma 1 and 3, and (22), (23) and (25). Nevertheless, due to the fact that for any $z \in V_n$ and any $\omega \in \Omega$ (see (16) and (24))

$$(27) \quad \frac{|\hat{g}_n(z, r(\hat{\beta}_n)) - \hat{g}_n(z, r(\beta^0))|}{f(z)} \leq \frac{\sup_{t \in R} |w'(t)| \cdot p^{\frac{1}{2}} \cdot K \cdot \|\hat{\beta} - \beta^0\|}{c_n \cdot f(z)}$$

the (26) holds (for K see (5)).

■

In the next theorems we shall give an example of the «quantile» process $f_n(z)$, and some inequalities describing the asymptotic behavior of the estimator of the

contamination level. It will be clear from the proof of the inequalities that there is a hope that they are reasonably tight, of course, asymptotically. It means that the difference between the limit value and the lower bound is small. To be able to prove the assertions which were just mentioned, we shall need a result of Rosenblatt (1971). For the convenience of reader we are going to give it as a lemma.

Lemma 4 (Rosenblatt (1971), p. 1828). *Let Assumptions A be fulfilled and let a density $g(z)$ be continuously differentiable and bounded away from zero on $[0, 1]$. Then if $n^{-\frac{1}{2A}} = O(c_n)$ as $n \rightarrow \infty$, it follows that*

$$P \left\{ \max_{0 \leq z \leq 1} \left[\frac{c_n \cdot n}{\gamma g(z)} \right]^{\frac{1}{2}} \{ \hat{g}_n(z, Z(\omega, n)) - E \hat{g}_n(z, Z(\omega, n)) \} \right. \\ \left. \leq \left\{ 2 \log c_n^{-1} \right\}^{\frac{1}{2}} + \frac{A + \nu}{(2 \log c_n^{-1})^{\frac{1}{2}}} \right\} \rightarrow \exp\{-\exp\{-\nu\}\}$$

for $n \rightarrow \infty$ where G is the distribution function which corresponds to the density $g(z)$ and

$$\gamma = \int w^2(s) ds$$

and

$$A = \log \frac{B^{\frac{1}{2}}}{2\pi}, \quad B = -\frac{2}{\gamma} \frac{d^2}{ds^2} \left(\int w(u) w(u+s) du \right) \Big|_{s=0}.$$

Remark 4 The basic idea of the proof of Rosenblatt's lemma is as follows. At first, the process

$$\left[\frac{c_n \cdot n}{\gamma g(z)} \right]^{\frac{1}{2}} \{ \hat{g}_n(z, Z(\omega, n)) - E \hat{g}_n(z, Z(\omega, n)) \}$$

is approximated by an appropriate Wiener process. Then an assertion about the distribution of the supremum of Wiener process (Rosenblatt (1971) uses Cramér, Leadbetter (1967)) is applied. It implies that due to symmetry of Wiener process we have also

$$P \left\{ \max_{0 \leq z \leq 1} \left[\frac{c_n \cdot n}{\gamma g(z)} \right]^{\frac{1}{2}} \{ E \hat{g}_n(z, Z(\omega, n)) - \hat{g}_n(z, Z(\omega, n)) \} \right. \\ (28) \quad \left. \leq \left\{ 2 \log c_n^{-1} \right\}^{\frac{1}{2}} + \frac{A + \nu}{(2 \log c_n^{-1})^{\frac{1}{2}}} \right\} \rightarrow \exp\{-\exp\{-\nu\}\}$$

for $n \rightarrow \infty$. □

Remark 5 Rosenblat's result (28) is given in a somewhat unusual form. More usual form would be such which describes convergences of distribution functions, i. e.

$$(29) \quad P \left\{ (2 \log c_n^{-1})^{\frac{1}{2}} \max_{0 \leq z \leq 1} \left[\frac{c_n \cdot n}{\gamma g(z)} \right]^{\frac{1}{2}} \{ E \hat{g}_n(z, Z(\omega, n)) - \hat{g}_n(z, Z(\omega, n)) \} \right. \\ \left. - 2 \log c_n^{-1} - A \leq \nu \right\} \rightarrow \exp\{-\exp\{-\nu\}\}$$

for $n \rightarrow \infty$.

□

There are two things which we have to cope with to be able to use Rosenblat's result for our purposes. First of all, if we apply directly Rosenblat's result, there would be an inconvenient presence of $g^{\frac{1}{2}}(z)$ in the denominator of the above formula. Another difficulty is that the normalized difference contains $E \hat{g}_n(z, Z(\omega, n))$ and does not give so a «distance» from $g(z)$. Moreover, the difference between $E \hat{g}_n(z, Z(\omega, n))$ and $g(z)$ is proportional to a power of c_n (unfortunately not to power of n). A remedy for the all difficulties is a transformation of data (in our case the transformation of residuals). Let us assume that we have data $z_1, z_2, \dots, z_n$ which are realization of a sequence of i.i.d. random variables, distributed according to a distribution function (d.f.) $G(z)$ and that we have selected some another d.f. $F(z)$ to explain them. The first step will be to estimate the asymptotic distribution of $\hat{\epsilon}_{G,F}$ under the null hypothesis, i.e. under the hypothesis that $\epsilon_{G,F} = 0$ (realize that then $G(z) = F(z)$). Consider instead of $z_1, z_2, \dots, z_n$ the data $u_1, u_2, \dots, u_n$ such that $u_i = F(z_i)$ for $i = 1, 2, \dots, n$. Then the density $f^*(u)$ of the transformed data is equal to 1 over the interval $[0, 1]$ (and zero elsewhere). Moreover, in what follows let us assume that we have selected V_n (for V_n see Theorem 1) so that $F(V_n) \cap (h_n, 1 - h_n), h_n = c_n \cdot a$ (for a see Conditions A) where $F(V_n) = \{u : u = F(z), z \in V_n\}$ and let us compute the mean value of the kernel estimator for the transformed observations u_i 's. We obtain for $u \in F(V_n)$

$$(30) \quad E \hat{g}_n^*(u, U(\omega, n)) = \int \left\{ \frac{1}{n c_n} \sum_{i=1}^n w(c_n^{-1}(u - y)) f^*(y) \right\} dy \\ = \frac{1}{n} \sum_{i=1}^n \int_{-a}^a w(s) f^*(u - c_n s) ds = 1.$$

It means that the kernel estimator of the «transformed» density for any $u \in F(V_n)$ is unbiased. Let us assume that we have transformed the residuals using $F(z)$ and let us apply (28) on the transformed values. We obtain:

Theorem 2 Let Conditions A and the assumptions of Theorem 1 and of Lemma 4 be fulfilled. For any $p \in (0, 1)$ put $v^* = -\log(-\log p)$ and let

$$(31) \quad b_n = \left(\frac{\gamma}{c_n} \right)^{\frac{1}{2}} \left[\{2 \log c_n^{-1}\}^{\frac{1}{2}} + \frac{A + v^*}{(2 \log c_n^{-1})^{\frac{1}{2}}} \right].$$

Finally, let the density $f(z)$ be bounded and let $f_n(z) = f(z)(1 - b_n \cdot n^{-\frac{1}{2}})$. Then under the assumption that $\varepsilon_{G,F} = 0$ we have

$$(32) \quad \limsup_{n \rightarrow \infty} P\{\hat{\varepsilon}_{G,F} = 0\} \geq p \quad \text{for } n \rightarrow \infty$$

and

$$(33) \quad \limsup_{n \rightarrow \infty} P\left\{ (2 \log c_n^{-1})^{\frac{1}{2}} \left(\frac{c_n n}{\gamma} \right)^{\frac{1}{2}} \hat{\varepsilon}_{G,F} \leq v - v^* \right\} \geq \exp\{-\exp\{-v\}\}$$

for $(v - v^*)(2 \log c_n^{-1})^{-\frac{1}{2}} \left[\frac{\gamma}{c_n n} \right]^{\frac{1}{2}} \in (0, 1)$.

Proof: First of all, let us say that in the proof some constants, say $C_1, C_2, \dots$, will be used. Their definition will be assumed to hold only within the proof. We shall show that the proof follows immediately from (29). Let us consider the transformation $u = F(z)$ and let us denote the density of the transformed random variable by $g^*(u)$ and for any $\beta \in R^p$ put $r_u(\beta) = (r_{u1}(\beta), r_{u2}(\beta), \dots, r_{un}(\beta))^T$, with

$$(34) \quad r_{ui}(\beta) = F(Y_i - X_i^T \beta).$$

Further, notice that the level of contamination is invariant with respect to the transformation. It is clear either from the heuristic background or from the formal expression. Really, the contamination level represents the percentage of the observations (among the data) which are not distributed according to the central model. So that if we transform data, the «earlier» central model is transformed into some «new» central model (in our case uniform distribution over $[0, 1]$), and similarly, the «contaminating» distribution is transformed to some «new contaminating» distribution. So the percentage of «wrong» data is the same. On the other hand using the formal way, we see that the values of the fraction $\frac{f(z) - g(z)}{f(z)}$ are precisely the same as the values of $\frac{f(\text{inv}F(u)) - g(\text{inv}F(u))}{f(\text{inv}F(u))}$ at the corresponding point $u = F(z)$. Multiplying by the Jacobian of the transformation both the numerator and the denominator, we do not change the value of the ratio. But then the density of the central model will become the density of the uniform distribution over $[0, 1]$ and the density $g(z)$ is transformed on $g^*(u)$. So we obtain for the contamination level an equivalent expression

$$\sup_{u \in [0, 1]} \{1 - g^*(u)\}.$$

Taking into account the specification of $f_n(z)$ given in the theorem which reads for the transformed residuals as

$$(35) \quad f_{n,U}(u) = 1 - b_n \cdot n^{-\frac{1}{2}}$$

we have for the estimator

$$(36) \quad \hat{\epsilon}_{G,F} = \max \left\{ 0, \min \left\{ \sup_{u \in F(V_n)} \left\{ 1 - b_n n^{-\frac{1}{2}} - \hat{g}_n^*(u, r_u(\hat{\beta}_n)) \right\}, 1 \right\} \right\}$$

where we have denoted the kernel density estimator based on the transformed residuals by $\hat{g}_n^*(u, r_u(\hat{\beta}_n))$. Let us recall that for the general $\beta \in R^p$ we have

$$\hat{g}_n^*(u, r_u(\beta)) = \frac{1}{nc_n} \sum_{i=1}^n w(c_n^{-1}(u - r_{ui}(\beta))).$$

Now we may write

$$\begin{aligned} \limsup_{n \rightarrow \infty} P(\hat{\epsilon}_{G,F} = 0) &= \limsup_{n \rightarrow \infty} P\left(\sup_{u \in F(V_n)} \{1 - b_n n^{-\frac{1}{2}} - \hat{g}_n^*(u, r_u(\hat{\beta}_n))\} \leq 0 \right) \\ &= \limsup_{n \rightarrow \infty} P\left(\sup_{u \in F(V_n)} \{1 - \hat{g}_n^*(u, r_u(\hat{\beta}_n))\} \leq b_n n^{-\frac{1}{2}} \right) \\ &= \limsup_{n \rightarrow \infty} P\left(\left[\frac{c_n n}{\gamma} \right]^{\frac{1}{2}} \sup_{u \in F(V_n)} [1 - \hat{g}_n^*(u, r_u(\hat{\beta}_n))] \leq (2 \log c_n^{-1})^{\frac{1}{2}} + \frac{A + v^*}{(2 \log c_n^{-1})^{\frac{1}{2}}} \right) \\ &= \limsup_{n \rightarrow \infty} P\left((2 \log c_n^{-1})^{\frac{1}{2}} \left[\frac{c_n n}{\gamma} \right]^{\frac{1}{2}} \sup_{u \in F(V_n)} \{ \mathbb{E} \hat{g}_n^*(u, r_u(\beta^0)) - \hat{g}_n^*(u, r_u(\beta^0)) \} \right. \\ &\quad \left. + (2 \log c_n^{-1})^{\frac{1}{2}} \left[\frac{c_n n}{\gamma} \right]^{\frac{1}{2}} \sup_{u \in F(V_n)} [\hat{g}_n^*(u, r_u(\hat{\beta}_n)) - \hat{g}_n^*(u, r_u(\beta^0))] \leq (2 \log c_n^{-1}) + A + v^* \right) \\ &\geq \limsup_{n \rightarrow \infty} P\left((2 \log c_n^{-1})^{\frac{1}{2}} \left[\frac{c_n n}{\gamma} \right]^{\frac{1}{2}} \max_{u \in [0,1]} \{ \mathbb{E} \hat{g}_n^*(u, r_u(\beta^0)) - \hat{g}_n^*(u, r_u(\beta^0)) \} \right. \\ &\quad \left. + (2 \log c_n^{-1})^{\frac{1}{2}} \left[\frac{c_n n}{\gamma} \right]^{\frac{1}{2}} \max_{u \in [0,1]} [\hat{g}_n^*(u, r_u(\hat{\beta}_n)) - \hat{g}_n^*(u, r_u(\beta^0))] \leq 2 \log c_n^{-1} + A + v^* \right). \end{aligned}$$

So we shall need to estimate the difference

$$(37) \quad (2 \log c_n^{-1})^{\frac{1}{2}} \left[\frac{c_n \cdot n}{\gamma} \right]^{\frac{1}{2}} \cdot \max_{u \in [0,1]} \left| \hat{g}_n^*(u, r_u(\hat{\beta}_n)) - \hat{g}_n^*(u, r_u(\beta^0)) \right|.$$

Taking into account (34), we may write

$$\begin{aligned}
\hat{g}_n^*(u, r_u(\hat{\beta}_n)) - \hat{g}_n^*(u, r_u(\beta^0)) &= \frac{1}{nc_n} \sum_{i=1}^n \left[w(c_n^{-1}(u - r_{ui}(\hat{\beta}_n))) - w(c_n^{-1}(u - r_{ui}(\beta^0))) \right] \\
(38) \quad &= \frac{1}{nc_n^2} \sum_{i=1}^n w'(\xi_i) \left[F(Y_i - X_i^T \beta^0) - F(Y_i - X_i^T \hat{\beta}_n) \right] = \frac{1}{nc_n^2} \sum_{i=1}^n w'(\xi_i) f(\eta_i) X_i^T (\hat{\beta}_n - \beta^0)
\end{aligned}$$

where $\xi_i \in \left(c_n^{-1} \min \left\{ u - r_{ui}(\hat{\beta}_n), u - r_{ui}(\beta^0) \right\}, c_n^{-1} \max \left\{ u - r_{ui}(\hat{\beta}_n), u - r_{ui}(\beta^0) \right\} \right)$ and $\eta_i \in \left(\min \left\{ Y_i - X_i^T(\hat{\beta}_n), Y_i - X_i^T(\beta^0) \right\}, \max \left\{ Y_i - X_i^T(\hat{\beta}_n), Y_i - X_i^T(\beta^0) \right\} \right)$. Now, the expression (38) can be rewritten as

$$(39) \quad \frac{1}{nc_n^2} \sum_{i=1}^n \left[w'(\xi_i) - w'(c_n^{-1}(u - F(e_i))) \right] f(\eta_i) X_i^T (\hat{\beta}_n - \beta^0)$$

$$(40) \quad + \frac{1}{nc_n^2} \sum_{i=1}^n w'(c_n^{-1}(u - F(e_i))) [f(\eta_i) - f(e_i)] X_i^T (\hat{\beta}_n - \beta^0)$$

$$(41) \quad + \frac{1}{nc_n^2} \sum_{i=1}^n w'(c_n^{-1}(u - F(e_i))) f(e_i) X_i^T (\hat{\beta}_n - \beta^0).$$

Taking into account that

$$|w'(c_n^{-1}\xi_i) - w'(c_n^{-1}(u - F(e_i)))| \leq \sup_{z \in R} |w''(z)| c_n^{-1} K p^{\frac{1}{2}} \|\hat{\beta}_n - \beta^0\|,$$

and also the fact that we have assumed that $\sqrt{n} \|\hat{\beta}_n - \beta^0\| = O_p(1)$, we have

$$\begin{aligned}
\sup_{u \in R} \left\{ (2 \log c_n^{-1})^{\frac{1}{2}} \left[\frac{c_n n}{\gamma} \right]^{\frac{1}{2}} \left| \frac{1}{nc_n^2} \sum_{i=1}^n \left[w'(\xi_i) - w'(c_n^{-1}(u - F(e_i))) \right] f(\eta_i) X_i^T (\hat{\beta}_n - \beta^0) \right| \right\} \\
\leq C_1 \cdot c_n^{-\frac{5}{2}} n^{-\frac{1}{2}} (\log c_n^{-1})^{\frac{1}{2}}
\end{aligned}$$

where C_1 is a positive (and finite) constant and hence the expression (39) converges to zero as $n \rightarrow \infty$. Similarly

$$\sup_{u \in R} \left\{ (2 \log c_n^{-1})^{\frac{1}{2}} \left[\frac{c_n n}{\gamma} \right]^{\frac{1}{2}} \left| \frac{1}{nc_n^2} \sum_{i=1}^n w'(c_n^{-1}(u - F(e_i))) [f(\eta_i) - f(e_i)] X_i^T (\hat{\beta}_n - \beta^0) \right| \right\}$$

$$\leq C_2 \cdot c_n^{-\frac{5}{2}} n^{-\frac{1}{2}} (\log c_n^{-1})^{\frac{1}{2}}$$

where again C_2 is a positive (and finite) constant and hence the expression (40) converges also to zero as $n \rightarrow \infty$. It remains to cope with supremum of the expression (41) which may be bounded by

$$(42) \quad \frac{1}{n^{\frac{1}{2}} c_n^2} \sup_{u \in R} \left\{ n^{-\frac{1}{2}} \sum_{i=1}^n [w'(c_n^{-1}(z - F(e_i)))f(e_i) \right. \\ \left. - \mathbb{E} \{w'(c_n^{-1}(u - F(e_i)))f(e_i)\}] X_i^T (\hat{\beta}_n - \beta^0) \right\}$$

$$(43) \quad + \frac{1}{n^{\frac{1}{2}} c_n^2} \sup_{u \in R} \left\{ \mathbb{E} \{w'(c_n^{-1}(u - F(e_1)))f(e_1)\} \left\{ n^{-\frac{1}{2}} \sum_{i=1}^n X_i^T \right\} (\hat{\beta}_n - \beta^0) \right\}.$$

Taking into account once again that $\sqrt{n} \|\hat{\beta}_n - \beta^0\| = O_p(1)$ and following Csörgő and Révész (1981), theorem 6.1.1, we find that the expression (42) is of order $c_n^{-2} n^{-1}$ in probability, and hence after multiplication by the factor $(2 \log c_n^{-1})^{\frac{1}{2}} \left[\frac{c_n n}{\gamma} \right]^{\frac{1}{2}}$ we obtain order of this expression equal to $c_n^{-3} n^{-\frac{1}{2}} \log n \cdot (\log c_n^{-1})^{\frac{1}{2}}$. Now we may calculate

$$\mathbb{E} \{w'(c_n^{-1}(u - F(e_1)))f(e_1)\} = \int w'(c_n^{-1}(u - F(t)))f^2(t) dt \\ = \int w'(s)f^2(\text{inv}F(u - c_n s))c_n ds = \int w'(s)f^2(\text{inv}F(u))c_n ds - 2 \int w'(s)f(\zeta)f'(\zeta)c_n^2 s ds$$

where ζ is again an appropriately selected point. Taking into account (6), (14) and (15), we come to the conclusion that the term in (43) (after multiplication by $(\log c_n^{-1})^{\frac{1}{2}} \left[\frac{c_n n}{\gamma} \right]^{\frac{1}{2}}$) converges also to zero. So denoting

$$\kappa_n = (2 \log c_n^{-1})^{\frac{1}{2}} \left[\frac{c_n n}{\gamma} \right]^{\frac{1}{2}} \sup_{u \in F(V_n)} \left\{ \left[\hat{g}_n^*(u, r_u(\hat{\beta}_n)) - \hat{g}_n^*(u, r_u(\beta^0)) \right] \right\}$$

using (29) and taking into account that $v^* = -\log(-\log p)$ we have

$$\limsup_{n \rightarrow \infty} P(\hat{\varepsilon}_{G,F} = 0) = \\ = \limsup_{n \rightarrow \infty} P((2 \log c_n^{-1})^{\frac{1}{2}} \left[\frac{c_n n}{\gamma} \right]^{\frac{1}{2}} \sup_{u \in F(V_n)} \{ \mathbb{E} \hat{g}_n^*(u, r_u(\beta^0)) - \hat{g}_n^*(u, r_u(\beta^0)) \} + \kappa_n$$

$$\leq 2\log c_n^{-1} + A + v^* \longrightarrow \exp\{-\exp\{-(v^* + \kappa_n)\}\} \longrightarrow p.$$

It concludes the proof of (32). Repeating the same steps we arrive at

$$\begin{aligned} & \limsup_{n \rightarrow \infty} P \left\{ (2\log c_n^{-1})^{\frac{1}{2}} \left[\frac{c_n n}{\gamma} \right]^{\frac{1}{2}} \hat{\varepsilon}_{G,F} \leq v - v^* \right\} \\ & \geq \lim_{n \rightarrow \infty} P \left\{ (2\log c_n^{-1})^{\frac{1}{2}} \left[\frac{c_n n}{\gamma} \right]^{\frac{1}{2}} \max_{u \in [0,1]} \left\{ \mathbb{E} \hat{g}_n(u, r_u(\hat{\beta}_n)) - \hat{g}_n(u, r_u(\hat{\beta}_n)) \right\} \right. \\ & \quad \left. + \kappa_n \leq 2\log c_n^{-1} + A + v \right\} \longrightarrow \exp\{-\exp\{-v\}\} \end{aligned}$$

which concludes the proof. ■

Remark 6 It follows from the proof of the previous theorem that the lower bound in (33) is tight. □

Similarly for a sequence of local alternatives we may obtain:

Theorem 3 *Let the assumptions of Theorem 2 be fulfilled. Put for any $p \in (0, 1)$ again $v^* = -\log(-\log p)$ and let b_n be given by (31). Finally, let $f_n(z) = f(z)(1 - b_n \cdot n^{-\frac{1}{2}})$ and $H(x)$ any distribution such that $\varepsilon_{H,F} \neq 0$ and $\max_{u \in [0,1]} h(\text{inv}F(u)) < \infty$. Then for any $\varepsilon \in (0, 1)$ and the sequence of the local alternative $\{G_n(x)\}_{n=1}^\infty$, $G_n(x) = (1 - \varepsilon n^{-\frac{1}{2}})F(x) + \varepsilon n^{-\frac{1}{2}}H(x)$ we have*

$$\limsup_{n \rightarrow \infty} P \left\{ (2\log c_n^{-1})^{\frac{1}{2}} \left(\frac{c_n n}{\gamma} \right)^{\frac{1}{2}} (\hat{\varepsilon}_{G_n,F} - \varepsilon_{G_n,F}) \leq v - v^* \right\} \leq \exp\{-\exp\{-v\}\}$$

for $(v - v^*)(2\log c_n^{-1})^{-\frac{1}{2}} \left[\frac{c_n n}{\gamma} \right]^{\frac{1}{2}} \in (0, 1)$.

Proof: mimics the proof of Theorem 2. First of all, we shall make an idea about $\varepsilon_{G_n,F}$. Let again $f^*(u)$, $h^*(u)$ and $g^*(u)$ denote the transformed densities. Taking into account that $f^*(u) = 1$ and $h^*(u) \geq 0$, we have

$$\varepsilon_{G_n,F} = \sup_{u \in [0,1]} 1 - g^*(u) \leq 1 - 1 + \varepsilon n^{-\frac{1}{2}} = \varepsilon n^{-\frac{1}{2}}.$$

We will need also to estimate $\mathbf{E}\hat{g}_n^*(u, r_u(\hat{\beta}_n))$. Similarly as in (30) we utilize the fact that $f^*(u) = 1$ but now we need also to employ the assumption that $h^*(u)$ is bounded. Then we obtain

$$|\mathbf{E}\hat{g}_n^*(u, r_u(\hat{\beta}_n)) - 1| \leq \int \left| \hat{g}_n^*(u, r_u(\hat{\beta}_n)) - 1 \right| \left((1 - \varepsilon n^{-\frac{1}{2}}) + \varepsilon n^{-\frac{1}{2}} h^*(u) \right) du \leq C_1 n^{-\frac{1}{2}}$$

where C_1 is a finite constant. Now we may write

$$\begin{aligned} & P \left\{ (2 \log c_n^{-1})^{\frac{1}{2}} \left(\frac{c_n n}{\gamma} \right)^{\frac{1}{2}} (\hat{\varepsilon}_{G_n, F} - \varepsilon_{G_n, F}) \leq v - v^* \right\} \\ & \geq P \left\{ (2 \log c_n^{-1})^{\frac{1}{2}} \left(\frac{c_n n}{\gamma} \right)^{\frac{1}{2}} (\mathbf{E}\hat{\varepsilon}_{G_n, F} - \hat{\varepsilon}_{G_n, F}) - 2 \log c_n^{-1} - A \leq v - C_2 n^{-\frac{1}{2}} \right\} \end{aligned}$$

where again C_2 is a finite constant and the assertion of the theorem follows. ■

3. SIMULATION STUDY

All results derived in the previous section are of the asymptotic type. Moreover, the results were obtained in the asymptotic framework in which several parameters changed simultaneously. Except of the number of observations which increased to infinity also the width of window, quantile process and the set over which the supremum had been taken, converged to the corresponding limits. Anybody who had sometimes tried to use such results to approximate the corresponding probabilities for the finite samples, has find out that «an adjustment» of the parameters (the width of window, quantile process, etc.) needs some simulation studies. Sometimes we may even meet with the standpoint that such asymptotic results should be interpreted only as a guarantee of the consistency (or coherence, if you want) of our approach with the general «structure of mathematics». That is the reason why the behavior of the statistics which was proposed above should be studied for the finite samples by simulations. In this section we shall offer a very first experience in the case when the data are contaminated, i.e. for $\varepsilon_{G, F} \neq 0$.

For the numerical study we have simulated data in the following way: Considering the regression model

$$Y_i = 2 \cdot X_{i1} + 3 \cdot X_{i2} + 4 \cdot X_{i3} + e_i, \quad i = 1, 2, \dots, 30$$

we have generated 30 three-dimensional vectors uniformly distributed over $[0, 10]$ (used as the carriers $X_i = (X_{i1}, X_{i2}, X_{i3})^T$, $i = 1, 2, \dots, 30$) and 30 random numbers distributed according to the standard normal distribution (used as the errors $e_1, e_2, \dots, e_{30}$).

The normality was checked by chi-square and Kolmogorov-Smirnov tests accompanied by the test of skewness and of kurtosis (see Shapiro, Wilk (1965)). Then we have randomly selected from $\{e_i\}_{i=1}^{30}$ four numbers, say $e_{i_1}, e_{i_2}, e_{i_3}, e_{i_4}$, multiplied them by 4 (they have represented the contamination). According to the Remark 3 we have then

$$(44) \quad \epsilon_{G,F} = \frac{4}{30} \left(1 - \frac{1}{4}\right) = 0.1.$$

Then the least trimmed square (LTS) algorithm was used to estimate the coefficients of the regression model, i.e.

$$\hat{\beta}^{(LTS)} = \underset{\beta \in R^p}{\operatorname{argmin}} \sum_{i=1}^h r_{[i:n]}^2(\beta)$$

where $r_{[i:n]}^2$'s represent ordered squared residuals $r_i^2(\beta) = [Y_i - \sum_{j=1}^3 x_{ij}\beta_j]^2$, i.e. $r_{[1:n]}^2 \leq r_{[2:n]}^2 \leq \dots \leq r_{[n:n]}^2$, and $h = 17$ was selected to reach the maximal possible break-down point of the estimator (see Rousseeuw, Leroy (1987)). Finally the obtained residuals $r_i(\hat{\beta}^{(LTS)}) = Y_i - \sum_{j=1}^3 x_{ij}\hat{\beta}_j^{(LTS)}$, $i = 1, 2, \dots, 30$ were transformed (see Theorem 2). The whole procedure was 30 times repeated. In what follows let $(\hat{\epsilon}_{G,F})_k$ denote the estimate of the contamination level for the k -th sample.

In such a way set of 30 samples of transformed residuals with the contamination level equal to 0.1 (see (44)) was obtained (each sample contained 30 residuals). This collection of samples was used as a training set. As follows from (36) the value of $\hat{\epsilon}_{G,F}$ for any fix sample of data is a nonincreasing function of b_n , $\hat{\epsilon}_{G,F}(b_n) : [0, \infty) \rightarrow [0, 1 - \min_{u \in F(V_n)} \hat{g}_n(u, r_u(\hat{\beta}_n))]$. Since in our case we had $\min_{u \in F(V_n)} \hat{g}_n(u, r_u(\hat{\beta}_n)) < 0.9$, it was possible to find b_{30}^* so that

$$\frac{1}{30} \sum_{k=1}^{30} (\hat{\epsilon}_{G,F}(b_{30}^*))_k = 0.1$$

(see (44) once again). We have obtained $b_{30}^* = 1.80049$ (or in other words, we have learnt that the quantile process $f_{30,U}(u)$ (see (35)) for this type of data, this kernel etc. should be (approximately) equal to $1 - b_n^* n^{-\frac{1}{2}} = 1 - \frac{1.80047}{\sqrt{30}} = 0.67128$). The results of estimating regression coefficients and of the values of the estimates $\hat{\epsilon}_{G,F}$ (after assigning the value 1.80049 to b_{30}^*) have been collected in the Table 3.

$$\frac{1}{30} \sum_{k=1}^{30} (\hat{\epsilon}_{G,F})_k = 0.1 \quad \text{var}(\hat{\epsilon}_{G,F}) = 0.02961$$

Table 3. *Results of estimation of regression coefficients and level of contamination for the training set of samples*

case	$\hat{\beta}_1$	$\hat{\beta}_2$	$\hat{\beta}_3$	$\hat{\epsilon}_{G,F}$
1	1.878	3.116	4.027	0.0000
2	2.113	3.139	3.731	0.0000
3	1.898	3.028	4.100	0.1048
4	2.022	2.688	4.282	0.0073
5	2.172	2.837	4.046	0.1679
6	1.887	2.975	4.120	0.2185
7	2.086	3.154	3.796	0.0590
8	2.059	2.857	4.063	0.1817
9	1.924	3.090	4.009	0.0000
10	2.161	2.964	3.898	0.0000
11	2.003	3.129	3.896	0.0894
12	1.834	2.983	4.201	0.0000
13	1.936	2.893	4.150	0.2644
14	1.948	3.104	3.949	0.0000
15	2.297	3.016	3.766	0.1416
16	1.911	3.072	4.006	0.0967
17	1.836	3.087	4.110	0.0734
18	1.978	2.994	3.997	0.1159
19	2.007	2.871	4.174	0.1060
20	2.079	3.082	3.857	0.0000
21	2.075	2.856	4.039	0.3431
22	1.959	3.030	3.995	0.0555
23	1.945	3.082	4.017	0.1684
24	2.470	2.686	3.813	0.0835
25	1.969	3.091	3.973	0.0000
26	1.782	3.167	3.983	0.3292
27	1.910	3.134	3.980	0.0000
28	1.905	3.121	3.959	0.1230
29	1.871	2.941	4.183	0.0975
30	1.979	2.943	4.093	0.1882

Table 4. *Results of estimation of regression coefficients and level of contamination for the testing set of samples*

case	$\hat{\beta}_1$	$\hat{\beta}_2$	$\hat{\beta}_3$	$\hat{\epsilon}_{G,F}$
1	2.092	3.012	3.972	0.1610
2	2.003	2.992	4.029	0.2534
2 3	2.054	2.972	3.982	0.0000
4	1.969	3.134	3.914	0.0000
4 5	1.959	3.008	3.999	0.1306
6	2.014	3.037	3.917	0.0000
6 7	2.013	2.919	4.043	0.0138
8	2.198	2.943	3.803	0.1033
9	1.858	2.916	4.207	0.1440
10	1.862	3.035	4.071	0.1477
11	1.851	3.177	4.007	0.0000
12	2.086	2.889	4.024	0.2434
13	1.959	2.993	4.062	0.1582
14	1.979	3.024	4.032	0.2457
15	2.042	2.728	4.200	0.0000
16	1.802	2.986	4.192	0.0791
17	2.240	2.815	3.937	0.0412
18	2.063	3.040	3.888	0.3725
19	1.915	2.921	4.241	0.0977
20	1.928	2.981	4.073	0.1879
21	1.719	3.103	4.174	0.0000
22	2.237	2.913	3.844	0.0000
23	1.914	3.034	4.033	0.1315
24	1.975	2.969	4.081	0.2305
25	1.976	3.049	3.951	0.0000
26	1.845	2.963	4.172	0.1968
27	2.114	3.005	3.888	0.0000
28	1.987	2.882	4.173	0.0742
29	1.952	3.162	3.863	0.0000
30	1.979	3.120	3.918	0.0000

In the same way as the training set was prepared we have generated a testing set consisting again 30 samples, each of them containing 30 observations. Again LTS were applied and the corresponding results of the estimation of regression coefficients together with the results of estimation of contamination level of the residuals are given in the Table 4.

$$\frac{1}{30} \sum_{k=1}^{30} (\hat{\varepsilon}_{G,F})_k = 0.10042 \quad \text{var}(\hat{\varepsilon}_{G,F}) = 0.03047$$

4. CONCLUSIONS

The paper brings a (theoretical) background for the selection of some free parameters of the robust methods (of the linear regression analysis) via estimating the contamination level. As it was already discussed the asymptotic results, may be except of the consistency of the estimator, have mainly a theoretical importance of some coherence of our approach with the general principles of mathematics. For the practical applications we should rely (mainly) on the results of a simulation studies.

In more details, we may procede as follows.

At first we estimate contamination level.

Of course, to be able to do it we need to adjust some value to the quantile process, to the width of window, to select appropriately the set V_n etc. It may be done on the base of experiences with the data of the same or similar character, or using the results of «reasonably» organized simulation study. It is clear that the type of distribution of the random errors is relevant, and so we have to employ our ideas about the character of these disturbances.

Secondly, we select the «tuning» parameter(s) of the corresponding family of robust methods.

(As an example of such family may serve the family of Huber's ψ -functions $\{\psi_{k(\varepsilon)}\}_{\varepsilon \in (0,1)}$, see (13).) Such selection may be performed either according to a known formula, connecting the contamination level with the «tuning» parameter(s), or by means of some theoretical tool of the type of efficiency rate or the local deficiency. Let us give an example.

After evaluating the estimate $\hat{\varepsilon}_{G,F}$ of the contamination level $\varepsilon_{G,F}$ we may calculate the estimate $\hat{\varepsilon}_H$ of the appropriate Huber mixture parameter ε_H by means of (9) or of (10). Then using the relation (11) and (12) we assign the value of the tuning constant $k(\hat{\varepsilon}_H)$ for the Huber's ψ -function .

In the case of another type of robust procedure we may use e. g. the efficiency rate and/or the local deficiency to select the «tuning» constant corresponding to the estimated contamination level. An example of using the efficiency rate for the selection of proper α for α -estimators is given in Rubio, Vřšek (1992). Let us be again more explicit.

Let $\{T_\alpha\}_{\alpha \in (0,1)}$ be a family of α -estimators of a parameter θ . (For the simplicity let us assume that θ is scalar.) Let us recall that the α -estimators are defined as the minimal distance estimators minimizing the α -divergence between the empirical d. f. and a d. f. from a projection family of d. f.'s (see Vajda (1989)). Having evaluated the efficiency rate of T_α in the corresponding model of contamination, we find the optimal selection of α for the estimated contamination level $\hat{\epsilon}_{G,F}$, say $\alpha(\hat{\epsilon}_{G,F})$. (Let us recall that the efficiency rate was defined in Rubio, Vřšek (1992) as the derivative of the supremum of the (asymptotic) variances of the estimators; supremum is taken over the given model of contamination, usually over some neighborhood of a central model; derivative is evaluated with respect to the parameter of the family of estimators, in our example derivative with respect to α . Of course if an objective function would be other than supremum of the variances, we should evaluate derivative of this function in the model of contamination.) Then we use for the estimation of θ the estimator $T_{\alpha(\hat{\epsilon}_{G,F})}$ which then minimizes the supremum of the asymptotic variances for given contamination level.

We are aware that the selection of tuning constant on the basis of an estimate of contamination level is the selection within the limits of one type of estimators (e. g. Huber's ones). Selection among different types of estimators should be based on some general principles (e. g. homogeneity of residuals over factor space, see Rubio *et al.* (1993) or Rubio and Vřšek (1994), or subsample stability of the estimates, Vřšek (1996b)), and/or on some model oriented rules (explicitly) formulated by the expert who has collected data (see Vřšek (1995)).

The results of simulation study presented in Table 4 showed that the estimates of the contamination level are scattered quite near around the «true» value (see the estimate of the variance of the estimator). It supports a hope that the proposed estimator of the contamination level may work well. On the other hand, it is clear that the method belongs among computationally intensive ones.

5. ACKNOWLEDGEMENT

We would like to thank the anonymous referee for a very detailed suggestions of corrections of all omissions. His/her proposals led to the considerable improvement of the text.

6. REFERENCES

- [1] **Antoch, J.** and **Víšek, J.Á.** (1991). «Robust estimation in linear models and its computational aspects». *Contributions to Statistics: Computational Aspects of Model Choice*. Springer Verlag, (1992), ed. J. Antoch, 39–104.
- [2] **Antoch, J.** and **Vorlíčková, D.** (1992). *Selected methods of statistical analysis*. Academia, Praha.
- [3] **Boscovich, R.J.** (1757). «De litteraria expeditione per pontificiam ditionem, et synopsis amplioris operis, ac habentur plura eius ex exemplaria etiam sensorum impressa». *Bononiensi Scientiarum et Artium Instituto Atque Academia Commentarii*, **4**, 353–396.
- [4] **Chatterjee, S.** and **Hadi, A.S.** (1988). *Sensitivity Analysis in Linear Regression*. New York: J. Wiley & Sons.
- [5] **Cramér, H.** and **Leadbetter, M.R.** (1967). *Stationary and Related Processes*. Wiley, New York.
- [6] **Csörgö, M.** and **Révész, P.** (1981). *Strong Approximation in Probability and Statistics*. Akademia Kiadó, Budapest.
- [7] **Galilei, G.** (1632). *Dialogo dei masimi sistemi*.
- [8] **Galton F.** (1886). «Regression towards mediocrity in hereditary stature». *Journal of the Antropological Institute*, **15**, 246–263.
- [9] **Gauss F.C.** (1809). «Theoria molus corporum celestium». *Hamburg, Perthes et Besser*.
- [10] **Hewitt, E.** and **Stromberg, K.** (1965). *Real and Abstract Analysis*. Springer - Verlag, Berlin.
- [11] **Huber, P.J.** (1964). «Robust estimation of a location parameter». *Ann. Math. Statist.*, **36**, 1753–1758.
- [12] **Huber, P.J.** (1981). *Robust Statistics*. J. Wiley & Sons, New York.
- [13] **Huber, P.J.** and **Strassen, V.** (1973). «Minimax tests and Neyman-Pearson lemma for capacities». *Ann. Statist.*, **1**, 251–263.
- [14] **Koenker, R.** and **Bassett, G.** (1978). «Regression quantiles». *Econometrica*, **46**, 33–50.
- [15] **Laplace, P.S.** (1793). «Sur quelques points du systeme du mode». *Memoires de l'Academic Royale des Sciences de Paris*, 1–87.
- [16] **Legendre A.M.** (1805). «Nouvelles méthodes pour la détermination des orbites des comètes». *Paris, Courcier*.
- [17] **Marazzi, A.** (1992). *Algorithms, Routines and S Functions for Robust Statistics*. Wadsworth & Brooks/Cole Publishing Company, Belmont, California, 1992.

- [18] **Rosenblatt, M.** (1971). «Curve estimates». *Ann. Math. Statist.*, **42**, 1815–1842.
- [19] **Rousseeuw, P.J. and Leroy, A.M.** (1987). *Robust Regression and Outlier Detection*. J.Wiley & Sons, New York.
- [20] **Rubio, A.M., Aguilar, L. and Víšek, J.Á.** (1992). «Testing for difference between models». *Computational Statistics*, **8**, 57–70.
- [21] **Rubio, A. M., Quintana, F. and Víšek, J.Á.** (1993). «Diagnostics of non-linear regression». *Transactions of the Twelfth Prague Conference on Information Theory, Statistical Decision Functions and Random Processes, Prague, 1994*, 203–206.
- [22] **Rubio, A. M. and Víšek J.Á.** (1992). «Efficiency rate and local deficiency of Huber's estimator of location and of estimators». *Trabajos de Estadística*, **6**, 67–85.
- [23] **Rubio, A. M. and Víšek, J.Á.** (1994). «Diagnostics of regression model: Test of goodness of fit». *Proceedings of the Fifth Prague Symposium on Asymptotic Statistics*, eds. M. Hušková & P. Mandl, Physica Verlag, 423–432.
- [24] **Rukhin, A.L.** (1994). «Optimal estimation of the mixture parameter». *Transactions of the Twelfth Prague Conference on Information Theory, Statistical Decision Functions and Random Processes, Prague, 1994*, 214–216.
- [25] **Shapiro, S.S. and Wilk, M.B.** (1968). «A comparative study of various tests for normality». *J. Amer. Statist. Assoc.*, **63**, 1343–1372.
- [26] **Vajda, I.** (1989). *Theory of Statistical Inference and Information*. Kluwer Academic Publ., Dordrecht.
- [27] **Víšek, J.Á.** (1989). «Estimation of contamination level in the model of contamination with general neighbourhoods». *Kybernetika*, **25**, 278–297.
- [28] **Víšek, J.Á.** (1993). «On the role of contamination level and the least favorable behaviour of gross-error sensitivity». *Probability and Mathematical Statistics*, **14**, 2, 173–187.
- [29] **Víšek, J.Á.** (1994 a). «A cautionary note on the method of Least Median of Squares reconsidered». *Transactions of the Twelfth Prague Conference on Information Theory, Statistical Decision Functions and Random Processes, Prague, 1994*, 254–259.
- [30] **Víšek, J.Á.** (1994 b). «High robustness and an illusion of truth». *Transactions of ROBUST'94, J const CMF (Union of Czech mathematicians)*, eds. J. Antoch & G. Dohnal, ISBN 80-7015-492-6, 172–185.
- [31] **Víšek, J.Á.** (1995). «On the heuristics of statistical results». *Proceedings of 'PROBASTAT'94', Tatra Mountains Publ.*, **7**, 349–357, Bratislava
- [32] **Víšek, J.Á.** (1996 a). «n high breakdown point estimation». *Computational Statistics*, **11**, 137–146, Berlin.

- [33] **Víšek, J.Á.** (1996b). «Sensitivity analysis of M -estimates». *Annals of the Institute of Statistical Mathematics*, **48**, 469–495.
- [34] **Víšek, J.Á.** (1996c). «On the coefficient of determination: Simple but...». *Bulletin of the Czech Econometric Society*, **5**, 117–124.
- [35] **Víšek, J.Á.** (1997). «On the diversity of estimates». Submitted to *Computational Statistics and Data Analysis*.

REMARQUES SUR LE MAXIMUM DE VRAISEMBLANCE

CATHERINE HUBER*

Université de Paris V

MIKHAIL NIKULIN**

Université de Bordeaux II

Les bonnes et les mauvaises propriétés de la méthode du maximum de vraisemblance sont examinées dans cet article. En ce qui concerne les bonnes, elles sont emboîtées dans le cadre plus général de celles des M -estimateurs. En ce qui concerne les mauvaises, l'absence de robustesse dans les cas les plus usuels d'hypothèse gaussienne et la difficulté rencontrée dans la définition même de la vraisemblance sont les deux traits essentiels, avec celui de la qualité asymptotique de l'optimalité de la méthode, illustrée par un exemple paradoxal de jeu où la stratégie du maximum de vraisemblance n'est pas admissible.

Some remarks on the maximum likelihood

Keywords: Vraisemblance, M -estimateurs, paradoxe.

* C. Huber. Université de Paris V, 45 rue des Saints-Pères 75 270 Paris cedex 06

** M. Nikulin. Université de Bordeaux II, 146 rue Léo Saignat 33 800 Bordeaux

– Article rebut el juliol de 1995.

– Acceptat el març de 1996.

1. INTRODUCTION

La méthode du maximum de vraisemblance est à la fois l'une des plus utilisées et des plus controversées en statistique. Elle a en effet un attrait à la fois intuitif, parce que la vraisemblance semble bien contenir toute l'information fournie par les observations, et théorique, à cause des bonnes propriétés asymptotiques des estimateurs correspondants sous certaines conditions de régularité. Cependant, cette méthode a plusieurs inconvénients qui ont été mis en évidence en particulier par L. Le Cam [1975] et P. J. Huber [1980], mais aussi par d'autres auteurs tels que J. Berkson [1980], R. R. Bahadur [1958], D. Basu [1980], T. S. Ferguson [1982], M. Goldstein et J. V. Howard [1991], S. Dharmadhikari et K. Joag-Dev [1985], V. G. Voinov et M. S. Nikulin [1993] par exemple. Comme cette méthode est universellement applicable, on a tendance à l'utiliser parfois sans précautions, ce qui conduit à des résultats désastreux. Un des exemples les plus célèbres est celui de L. Le Cam qui exhibe un estimateur «du maximum de vraisemblance» qui converge vers deux fois la valeur du paramètre qu'il est censé estimer, simplement parce que le nombre des paramètres inconnus nuisibles croît à la même vitesse que la taille de l'échantillon [L. Le Cam, 1979]. La vraisemblance elle-même pose des problèmes lorsqu'il y a plusieurs versions de la densité de probabilité sur laquelle repose cette vraisemblance. De plus, chaque fois qu'il n'y a pas de propriété de convexité, il est difficile, voire impossible de trouver un maximum global. Même dans ce cas particulièrement favorable, le maximum de la vraisemblance peut être, et il l'est souvent, si «plat» que l'estimateur correspondant est numériquement très mal défini. B. Efron [1982] nous paraît avoir là-dessus un point de vue très sain: si l'on est suffisamment attentif aux conditions dans lesquels on applique la méthode, on n'aura pas de déboire majeur. Il n'en reste pas moins que deux défauts en quelque sorte intrinsèques subsistent:

1. L'absence de définition correcte d'un estimateur du maximum de vraisemblance lorsqu'il peut y avoir une ambiguïté sur la version de la densité à employer, ou bien lorsque le modèle n'est pas dominé. La tentative de F. W. Scholz [1980] pour donner une définition unifiée du maximum de vraisemblance est une réponse à cette question. Elle n'est malheureusement pas très connue et n'a pas eu l'impact auquel on aurait pu s'attendre.
2. Le manque de robustesse des estimateurs du maximum de vraisemblance les plus employés: ceux qui sont relatifs au modèle gaussien. Le plongement des estimateurs M-V dans un ensemble plus général, celui des M-estimateurs, par P. J. Huber [1980] traite de cette question car il permet de corriger l'absence de robustesse de certains estimateurs M-V.

Cet exposé se déroule en trois parties. Nous allons voir successivement:

- Les propriétés des estimateurs du maximum de vraisemblance considérés comme un cas particulier des M-estimateurs, sous des conditions de régularité

du modèle statistique supposé dominé. Trois propriétés essentielles sont considérées: la consistance, la normalité asymptotique et la sensibilité de ces estimateurs à des écarts par rapport au modèle. Ce point de vue sur le maximum de vraisemblance est celui qui est issu des études de robustesse [P.J. Huber, 1981] et il met en évidence l'instabilité des estimateurs usuels du maximum de vraisemblance.

- Un essai de théorie unifiée du maximum de vraisemblance, que le modèle soit ou non dominé. Cette partie de l'exposé est fondée sur un article assez peu connu de F.W. Scholz qui, bien conscient des ambiguïtés qui subsistent dans la définition des estimateurs du maximum de vraisemblance, défauts qui ont pour conséquence de mauvaises propriétés de ces estimateurs, a fait une très intéressante tentative de définition unifiée du maximum de vraisemblance.
- Quelques paradoxes du maximum de vraisemblance. On peut trouver de tels exemples en grand nombre chez L. Le Cam [1979], mais aussi dans les articles cités plus haut.

2. MAXIMUM DE VRAISEMBLANCE GÉNÉRALISÉ

Soit $X_1, X_2, \dots, X_n$ un échantillon d'une variable aléatoire X gouvernée par une probabilité $P_\theta \in H$, où H est un ensemble de mesures de probabilité absolument continues par rapport à une mesure μ :

$$H = \{P_t, t \in \Theta \subset \mathbb{R}^p\}$$

On note f la densité de P par rapport à μ , soit

$$f(., t) = \frac{dP_t}{d\mu}$$

et ℓ la dérivée par rapport à t du logarithme de f , soit

$$\ell(x, t) = \frac{\partial \text{Log} f(x, t)}{\partial t}.$$

Alors l'estimateur du maximum de vraisemblance de θ est $\hat{\theta}$ défini par l'équation

$$\sum_{i=1}^n \ell(x_i, \hat{\theta}) = 0$$

2.1. Définition d'un M-estimateur

Un estimateur du maximum de vraisemblance généralisé est solution d'une équation analogue, où la fonction ℓ est remplacée par une fonction ψ vérifiant de bonnes conditions de régularité. Ces conditions de régularité sont destinées à permettre à l'estimateur d'avoir toutes les propriétés des bons estimateurs du maximum de vraisemblance, c'est à dire consistance, et normalité asymptotique, plus celle d'être peu sensible à des écarts des observations par rapport au modèle. L'équation de la vraisemblance est alors remplacée par

$$\sum_{i=1}^n \psi(x_i, \hat{\theta}) = 0$$

Pour être plus rigoureux, on ne devrait pas parler d'un estimateur du maximum de vraisemblance, mais d'une suite $\{T_n\}_{n \in \mathbb{N}}$ d'estimateurs. Cette suite $\{T_n\}_{n \in \mathbb{N}}$ peut être considérée comme la restriction à

$$\mathcal{F}_n = \{\text{lois empiriques d'ordre } n\}$$

de T , fonctionnelle définie sur $\mathcal{F}$ par

$$\int \psi(x, T(F)) dF(x) = 0$$

Exemples:

1. Modèle de translation: l'équation s'écrit dans ce cas

$$\int \psi(x - T(F)) dF(x) = 0.$$

Si la fonction ψ est l'identité, l'estimateur obtenu est la moyenne et si c'est la valeur absolue, la médiane.

2. Modèle de régression linéaire: l'équation correspondante lorsque Y , variable aléatoire réelle est égale à une combinaison linéaire des composantes d'un vecteur x de dimension p plus une variable aléatoire S , l'erreur, de moyenne nulle, s'écrit

$$\int \psi(y - \langle \theta, x \rangle) x_j dF(x) = 0$$

$$j = 1, 2, \dots, p$$

2.2. Propriétés des M-estimateurs

La fonction $\psi(x, t)$ dépend des deux variables, éventuellement multidimensionnelles, t et x , désignant respectivement le paramètre et la variable. On définit:

$$m(\theta, t) = \int \psi(x, t) dP_\theta(x).$$

Theorem 1 (Consistance) *Si l'application $t \mapsto m(\theta, t)$ est nulle pour $t = \theta$ et décroissante dans un voisinage de θ , alors $\exists$ une suite $\{T_n\}_{n \in \mathbb{N}}$ de solutions de l'équation*

$$t \mapsto \int \psi(x, t) dF_n(x) = 0$$

qui converge presque sûrement vers θ .

Theorem 2 (Normalité asymptotique) *Sous les hypothèses du théorème 1 et si de plus, l'application $t \mapsto m(\theta, t)$ est continuellement dérivable, de dérivée $\frac{\partial \psi(x, t)}{\partial t}$ notée $\psi'(x, t)$ telle que:*

- $t \mapsto \psi'(x, t)$ soit continue, uniformément en x
- $0 < \int \psi^2(x, \theta) dP_\theta(x) = d^2(\theta) < \infty$
- $0 < \int \psi'(x, \theta) dP_\theta(x) = c(\theta) < \infty$

Alors, si $\{T_n\}_{n \in \mathbb{N}}$ est une suite de solutions de

$$\int \psi(x, T_n) dF_n(x) = 0$$

on a

$$\sqrt{n}(T_n - \theta) \xrightarrow{\mathcal{L}} N(0, \frac{d^2(\theta)}{c^2(\theta)})$$

Démonstration: Par définition de T_n , $0 = \psi_n(x, T_n) = \psi_n(x, \theta) + (T_n - \theta)\psi'(x, \theta_n)$ où θ_n tend presque sûrement vers θ . On peut donc écrire:

$$\begin{aligned} \psi'(x, \theta_n) &= [\psi'(x, \theta_n) - \psi'(x, \theta)] + [\psi'_n(x, \theta) - E_\theta \psi'(X, \theta)] + E_\theta \psi'(X, \theta) \\ &= A_n + B_n + c(\theta) \end{aligned}$$

Or A_n tend presque sûrement vers 0 à cause de la continuité en t de ψ' , qui est uniforme en x , et B_n tend aussi presque sûrement vers 0 à cause de la loi forte des grands nombres. Donc

$$\sqrt{n}(T_n - \theta) = \frac{-\sqrt{n}\psi_n(x, \theta)}{A_n + B_n + c(\theta)}$$

qui, d'après le théorème limite centrale appliqué à $n\psi_n = \sum_{i=1}^n \psi(X_i, \theta)$ est asymptotiquement normal de moyenne nulle et de variance $\frac{d^2(\theta)}{c^2(\theta)}$.

La sensibilité d'un estimateur à des écarts par rapport au modèle peut être mesurée par l'impact sur cet estimateur du retrait d'une petite masse ε de probabilité qui serait mise au point x , soit $\varepsilon\delta_x$. La courbe de sensibilité décrit, en fonction de x , la limite de cet impact lorsque ε tend vers 0. Cette courbe s'exprime simplement en fonction de ψ pour un M-estimateur comme le montre le théorème suivant (pour la démonstration et pour plus de détails, voir P. J. Huber [1980] ou C. Huber-Carol [1985]).

■

Theorem 3 (Sensibilité) *Soit F une fonction de répartition sur $(\mathbb{R}, \mathcal{B})$, et soit Y une variable aléatoire réelle définie sur $(\mathbb{R}, \mathcal{B})$. Si, pour tout t réel, l'équation $\int \psi(y, t) dF(y) = 0$ définit sans ambiguïté une fonctionnelle T sur*

$$\mathcal{G} = \{G = (1 - \varepsilon)F + \varepsilon\delta_x, x \in \mathbb{R}, 0 \leq \varepsilon \leq \varepsilon_0\}$$

pour un ε_0 de $[0; 1]$ et si, de plus, $t \mapsto \psi(y, t)$ est continue dérivable, de dérivée $\psi'(y, t)$ telle que $|\psi'(y, t)| < g(y)$ pour une fonction g de $L^2(F)$, alors la suite $\{T_n\}_{n \in \mathbb{N}}$ a pour courbe d'influence:

$$CI_{T, F}(x) = \frac{-\psi(x, T(F))}{\int \psi'(y, T(F)) dF(y)}$$

3. THÉORIE UNIFIÉE DU MAXIMUM DE VRAISEMBLANCE

3.1. Cas discret

C'est dans le cas discret que la définition d'un estimateur du maximum de vraisemblance ne pose en général pas de problème. En effet, si $X \sim P_\theta$, où P_θ est une loi discrète dont le paramètre θ appartient à un ensemble Θ , un estimateur du maximum de vraisemblance est défini comme

$$\hat{\theta} = \arg \max_{\theta \in \Theta} P(X = x | \theta)$$

3.2. Cas d'un modèle dominé

Prenons un modèle dominé par une mesure μ , par exemple la mesure de Lebesgue. Soit $P_\theta \ll \mu, \theta \in \Theta$. On considère une version $f_\theta = \frac{dP_\theta}{d\mu}$ de la densité de P_θ par rapport à μ et la vraisemblance de l'observation x devient $V(\theta) = f(x|\theta)$. Aussi la définition utilisée dans le cas discret ne peut elle plus être employée sans le risque d'avoir une ambiguïté: *l'estimateur du maximum de vraisemblance dépendra de la version choisie pour la densité*. Au lieu de prendre la valeur de la densité au point x observé, on va considérer un voisinage $N(x)$ de ce point et, en notant $V(N(x))$ le volume de ce petit voisinage de x , on aura

$$P(X \in V(x) | \theta) \simeq f(x|\theta)V(N(x))$$

3.3. Cas d'un modèle non dominé

Considérons la situation non paramétrique suivante: la famille de lois est celle des lois symétriques par rapport à 0. Le paramètre est un paramètre de translation. Dans ce cas, il n'y a pas de mesure dominante σ -finie. Il n'y a donc pas de densités à comparer. Dans ce cas, deux tentatives ont été faites pour généraliser la méthode du maximum de vraisemblance: l'une par Kiefer et Wolfowitz en 1956, l'autre par Kalbfleish et Prentice en 1980. F. W. Scholz en propose une troisième qui est plus satisfaisante.

1. Généralisation de Kiefer et Wolfowitz On prend les probabilités de la famille deux par deux. A ce moment-là, le couple $\{P, Q\}$ est toujours dominé par la mesure somme $P + Q$. Cependant, comme ci-dessus, la version des dérivées de Radon-Nicodém n'étant pas spécifiée, à chaque choix va correspondre une solution différente.

Exemple:

Soit $\varphi(x) = \frac{1}{\sqrt{2\pi}} \exp(-\frac{x^2}{2})$ la densité de la loi normale $N(0, 1)$ et

$$\varphi^*(x) = \begin{cases} \varphi(x), & \text{si } x \neq 1, \\ 10, & \text{si } x = 1. \end{cases}$$

Soit le modèle $\{P_\theta = f_\theta(x), \mu = \varphi(x - \theta), \mu; \theta \in \mathbb{R}\}$ où μ est la mesure de Lebesgue, c'est à dire que $P_\theta = N(\theta, 1)$. Supposons que l'on ait une seule observation x . Alors, l'estimateur de θ sera x avec la première version et $x - 1$ avec la deuxième, et cela pour tout x . Si maintenant on a n observations, $x_1, x_2, \dots, x_n$, ce sera $x_{(j)} - 1$ pour l'une des statistiques d'ordre $x_{(j)} - 1$. On voit donc que cette généralisation a les mêmes défauts que le cas continu: l'estimateur dépend de la spécification choisie pour la version de la densité, *à moins que la version de la densité ne soit, a priori, une donnée du problème*.

2. Généralisation de Kalbfleish et Prentice Cette généralisation a été proposée par Kalbfleish et Prentice dans le contexte des durées de survie et pour l'estimateur de Kaplan-Meier de la fonction de répartition. Elle consiste à transformer les données de la manière suivante:

- Discrétiser les données: elles sont de toute façon toujours discrètes, à cause de la limitation de la précision des mesures, de l'arrondi.
- Faire un maximum de vraisemblance sur la multinomiale correspondante.
- Espérer que le maximum de vraisemblance discret va converger vers une limite quand l'erreur d'arrondi tend vers 0. Ce dernier point n'est pas encore élucidé (cf Van der Laan, 1995): deux points restent obscurs, le premier est: y-a-t-il une limite ?, et le deuxième: cette limite, quand elle existe, dépend-elle du groupement ?

3. Théorie unifiée de F. W. Scholz Cette définition unifiée résulte du mariage des deux idées précédentes: d'une part, on prend les probabilités par couples, et, d'autre part, on considère un voisinage de l'observation x . On définit les quantités suivantes:

$\mathcal{X}$	un espace métrique de métrique d ,
$\mathcal{B}$	la sigma-algèbre des boréliens de $\mathcal{X}$,
$\mathcal{P}$	une famille de probabilités sur $(\mathcal{X}, \mathcal{B})$,

Alors, $\forall x \in \mathcal{X}$, on définit

Définition 1

$$\mathcal{N}_x = \{N_x : x \in \overset{o}{N}_x, N_x \in \mathcal{B}\}$$

et $D(N_x)$ est le diamètre de N_x .

Définition 2 (Dominance en un point x) Soit $P, Q, \in \mathcal{P}$. On dit que $P \overset{x}{\geq} Q$ si

$$\liminf_{\varepsilon \rightarrow 0} \left\{ \frac{P(N_x)}{Q(N_x)} : N_x \in \mathcal{N}_x, D(N_x) \leq \varepsilon \right\} \geq 1,$$

où, par convention $\frac{0}{0} = 1$.

Notation: On notera la quantité qui intervient à gauche dans la formule précédente $\frac{\lim_{N_x} P(N_x)}{Q(N_x)}$.

□

Définition 3 (Equivalence en un point x) On dit que P et Q sont équivalentes en x si $P \stackrel{x}{\geq} Q$ et $Q \stackrel{x}{\geq} P$ et on note $P \stackrel{x}{=} Q$.

Commentaires:

- La relation $\stackrel{x}{\geq}$ est réflexive $P \stackrel{x}{\geq} P \quad \forall P$.
- Elle est transitive: $P \stackrel{x}{\geq} Q$ et $Q \stackrel{x}{\geq} R$ impliquent que $P \stackrel{x}{\geq} R$.
- Si on définit $\{P\}_x$ comme la classe d'équivalence de P pour x , soit

$$\{P\}_x = \{Q \in \mathcal{P} : Q \stackrel{x}{=} P\}$$

et $\mathcal{P}_x$ comme l'ensemble de ces classes, pour P variant dans $\mathcal{P}$, alors $\{\mathcal{P}_x, \stackrel{x}{\geq}\}$ est un ensemble partiellement ordonné.

- $P \stackrel{x}{=} Q \Leftrightarrow \lim_{D(N_x) \rightarrow 0} \frac{P(N_x)}{Q(N_x)}$ existe et est égale à 1.

Définition 4 (Maximum de vraisemblance) La statistique $P_0 \in \mathcal{P}$ est un estimateur du maximum de vraisemblance par rapport à x et à $\mathcal{P}$ si $\forall Q \in \mathcal{P}$ tel que $Q \stackrel{x}{\leq} P_0$ alors $Q \stackrel{x}{=} P_0$ c'est à dire que P_0 est M-V si et seulement si l'une des deux conditions équivalentes suivantes est réalisée:

1. Il n'existe pas de loi Q dans $\mathcal{P} - \{P_0\}_x$ telle que $Q \stackrel{x}{\geq} P_0$
- 2.

$$\liminf \frac{Q(N_x)}{P_0(N_x)} < 1, \quad \forall Q \in \mathcal{P} - \mathcal{P}_0.$$

L'existence d'un estimateur M-V n'est pas garantie par cette définition

4. UN PARADOXE DE LA VRAISEMBLANCE

L'exemple suivant est considéré par Fraser(1984), Wolpert(1988) et Joshi (1989) comme mettant en question les principes de la vraisemblance. On considère, dans une urne, 6 boules numérotées respectivement k , k , $4k+1$, $4k+2$, $4k+3$ et $4k+4$. Il s'agit d'estimer le paramètre $\theta \in \mathbb{N}$, qui vaut k , après avoir tiré au hasard une boule numéroté $S = s$. La vraisemblance $V(s, \theta)$ vaut:

$$\begin{aligned}
V(s, \theta) &= 2/6 \text{ si } s = \theta \\
&= 1/6 \text{ si } s = 4\theta + 1 \\
&= 1/6 \text{ si } s = 4\theta + 2 \\
&= 1/6 \text{ si } s = 4\theta + 3 \\
&= 1/6 \text{ si } s = 4\theta + 4
\end{aligned}$$

Par conséquent, comme $\max_{\theta} V(s, \theta) = 1/3$ pour $\hat{\theta} = s$, l'estimateur M-V de θ est

$$u_L(s) = s.$$

Mais en fait, pour une valeur observée s de S , il n'y a que deux valeurs possibles de θ :

$$\begin{aligned}
\theta &= s && \text{si l'on a tiré l'une des deux boules numérotées } \theta \\
\theta &= \left\lfloor \frac{s-1}{4} \right\rfloor && \text{si l'on a tiré l'une des quatre autres boules} \\
&&& \text{numérotées } 4\theta + i \text{ pour } 1 \leq i \leq 4.
\end{aligned}$$

On peut écrire:

$$\begin{aligned}
V(\theta, s) &= 0 && \text{si } \theta \notin \{s; \left\lfloor \frac{s-1}{4} \right\rfloor\} \\
&= \frac{1}{3} && \text{si } \theta = s \\
&= \frac{2}{3} && \text{si } \theta = \left\lfloor \frac{s-1}{4} \right\rfloor
\end{aligned}$$

On peut donc envisager une deuxième stratégie

$$u_C(s) = \left\lfloor \frac{s-1}{4} \right\rfloor.$$

Le paradoxe provient alors du fait que la probabilité que la première stratégie u_L , qui est celle du maximum de vraisemblance, donne la bonne réponse vaut $1/3$ alors que la deuxième stratégie donne la bonne réponse avec la probabilité $2/3$, si θ est strictement positif. Dans le cas où θ est nul, cette dernière probabilité vaut même 1. Si l'on considère le tableau croisé de S et θ , on obtient le tableau suivant:

	0	1	2	3	4	5	6	7	8	9	10	11	12
0	2	1	1	1	1								
1		2				1	1	1	1				
2			2										
3				2						1	1	1	1
4					2								

On peut encore accroître cette disparité en généralisant l'exemple précédent de la manière suivante: on a dans une urne 99 boules marquées k , et 99^2 boules marquées

$99_k^2 + i$, i variant de 1 à 99^2 . Si les deux estimateurs de k sont

$$\begin{aligned} u_C(s) &= \left\lfloor \frac{s-1}{99^2} \right\rfloor \\ u_L(s) &= s \end{aligned}$$

la probabilité que u_C soit correct est supérieure ou égale à 0,99 alors que celle de u_L est de 0,01. Autrement dit, la vraisemblance de u_C est 99 fois plus élevée que celle de u_L .

Les principes en cause sont les suivants:

- LP (Likelihood Principle):
Ce principe est celui selon lequel, toute l'information sur le paramètre θ est contenue dans la fonction de vraisemblance de θ étant donnée l'observation s .

Ce principe n'est pas remis en cause par ce paradoxe.

- LPF (Likelihood Preference Principle):
Lorsqu'on fait de l'inférence statistique, on doit toujours préférer les valeurs de θ qui correspondent aux valeurs les plus grandes de la vraisemblance.

C'est ce principe qui est remis en cause par cet exemple.

- CPP (Confidence Preference Principle):
Si deux règles A et B donnent des intervalles de confiance pour θ de même taille, et si, pour tout θ , la probabilité que A soit correct est supérieure ou égale à celle de B, et, pour au moins un θ , strictement supérieure, on doit préférer A à B.
Ce principe donne la préférence à la stratégie u_C dans notre exemple. Cet exemple met donc en évidence une contradiction entre les deux principes LPF et CPP.

En fait, il semble que *aucun* de ces principes ne soit bon. On peut s'en rendre compte en cherchant une loi a priori sur $\Theta = \mathbb{N}$ qui mène, d'un point de vue bayésien, à la préférence de l'une ou de l'autre des deux stratégies u_C et u_L .

Soit la loi a priori sur $\mathbb{N}$

$$p_i = P(\theta = i), i \in \mathbb{N}.$$

Alors

$$\frac{P(\theta = u_C | s)}{P(\theta = u_L | s)} = \frac{q(s)}{1 - q(s)} = \frac{p_{u_C}}{2p_{u_L}}$$

pour $s \geq 0$.

Pour que cela conduise au choix de u_L il faut que

$$p_s \geq \frac{1}{2} p_{\lfloor \frac{s-1}{4} \rfloor}$$

Cependant, il y a beaucoup de lois a priori pour lesquelles u_C est préférable, par exemple:

$$p_i = \frac{1}{2^{i+1}}.$$

On peut se poser les questions suivantes:

- Existe-t-il une loi a priori qui conduirait à u_L quel que soit s ?
La condition à remplir est

$$p_s \geq (1/2) p_{\lfloor s-1/4 \rfloor}$$

pour $s = 1, 2, \dots$ qui entraînerait que

$$\sum_{s=1}^{\infty} p_s \geq 2 \sum_{s=0}^{\infty} p_s = 2$$

ce qui est impossible.

- Y a-t-il seulement un nombre fini de valeurs de s pour lesquelles u_C est préférable ?
Si c'était le cas, il existerait n au-delà duquel on aurait:

$$\sum_{s=n}^{\infty} p_s \geq 2 \sum_{s=n}^{\infty} p_s$$

ce qui est impossible, sauf si le support de θ n'était pas infini.

Il y a deux arguments contre le principe LPP:

- Aucune loi a priori sur $\Theta = \mathbb{N}$ ne conduit à la stratégie u_L pour tout s .
- Pour toute loi a priori sur $\Theta = \mathbb{N}$, il y a une infinité de valeurs de s pour lesquelles u_C est préférable.

Commentaires:

- Si le support de θ n'est plus $\Theta = \mathbb{N}$ mais $\Theta = \{0, 1, 2, \dots, N\}$, où N est fini fixé, tout change.

Exemple: $0 \leq \theta \leq 4$

Le tableau suivant donne les probabilités de ne pas se tromper pour $N=4$. C'est la stratégie u_C qui est choisie si le critère est maximin.

		θ					Minimum
		0	1	2	3	4	
Règle	C	1	2/3	2/3	2/3	2/3	2/3
Règle	L	1/3	1/3	1	1	1	1/3

- Une autre possibilité consiste à utiliser la règle mixte M_4 : choisir u_C avec la probabilité $2/3$ et u_L avec la probabilité $1/3$ lorsque l'observation s est comprise entre 1 et 4, bornes comprises. Dans les autres cas, c'est à dire $s = 0$ ou $s = 5$, les deux stratégies u_C et u_L sont de toute façon, identiques. On obtient alors le tableau suivant:

		θ					Minimum
		0	1	2	3	4	
Règle	C	1	2/3	2/3	2/3	2/3	2/3
Règle	L	1/3	1	1	1	1	1
Règle	M_4	7/9	7/9	7/9	7/9	7/9	7/9

5. UNE GÉNÉRALISATION PRATIQUE DU MAXIMUM DE VRAISEMBLANCE

Une autre approche, qui peut être considérée comme une généralisation de la méthode du maximum de vraisemblance, a été proposée par Weiss et Wolfowitz (1966), (1974), voir aussi Blyth (1983), Weiss (1986), Voinov et Nikulin (1993) etc. On fait ici quelques remarques concernant cette approche.

Soit $X = (X_1, \dots, X_n)^T$ un échantillon, $X_i \sim f(x; \theta)$, $\theta \in \Theta \subseteq \mathbb{R}^1$. Notons $L(X; \theta)$ la fonction de vraisemblance de X :

$$L(X; \theta) = \prod_{i=1}^n f(X_i, \theta),$$

et soit r un nombre positif, $r \in \mathbb{R}_+^1$. Dans ce cas la statistique

$$T_r = T_r(X) = \arg \max_t \int_{t-r}^{t+r} L(X; \theta) d\theta, \quad T_r \in \Theta. \quad (1)$$

s'appelle *l'estimateur du maximum de probabilité* pour θ . On remarque que en réalité la relation (1) détermine une famille $\{T_r, r > 0\}$ des estimateurs T_r pour θ , ce qui nous permet de choisir l'estimateur optimal dans un sens ou un autre. Il est clair que

$$T_0 = \lim_{r \rightarrow 0} T_r = \hat{\theta}_n, \quad (2)$$

où $\hat{\theta}_n$ est l'estimateur de maximum de vraisemblance, à condition que le maximum globale dans (1) soit un point de continuité de $L(X; \theta)$. Si $L(X; \theta)$ est discontinue dans ce cas la relation (2) peut être considérée comme une définition de l'estimateur du maximum de vraisemblance. Cette convention nous permet d'inclure T_0 dans la classe $\{T_r\}$ et par conséquent de considérer, suivant en cela C. Blyth, la méthode du maximum de probabilité comme une généralisation de la méthode du maximum de vraisemblance.

Exemple 1

Soit $X = (X_1, X_2, \dots, X_n)^T$ un échantillon,

$$X_i \sim f(x; \theta), \quad \theta \in \Theta = \mathbb{R}^1, \quad f(x; \theta) = \begin{cases} \exp-(x - \theta) & x \geq \theta, \\ 0, & x < \theta. \end{cases} \quad (3)$$

Il est bien connu que la statistique $X_{(1)} = \min(X_1, \dots, X_n)$ est l'estimateur du maximum de vraisemblance pour θ . Pour construire un estimateur du maximum de probabilité il nous faut trouver n'importe quelle statistique T_r pour laquelle

$$\int_{T_r - r}^{T_r - r} L(X; \theta) d\theta$$

atteint son maximum global. Blyth a montré (1983) que

$$T_r = X_{(1)} - r. \quad (4)$$

Par des calculs directs on peut montrer que le le risque quadratique

$$E(T_r - \theta)^2 = r^2 - \frac{2r}{n} + \frac{2}{n^2}$$

de l'estimateur T_r atteint son minimum quand $r = 1/n$,

$$E(T_{1/n} - \theta)^2 = \min_r E(T_r - \theta)^2 = \frac{1}{n^2}. \quad (5)$$

On remarque que

$$E(X_{(1)} - \theta)^2 = \frac{2}{n^2} \quad (6)$$

et donc de ce point de vue l'estimateur du maximum de probabilité est meilleur que l'estimateur du maximum de vraisemblance, qui est superefficace dans ce modèle et par conséquent $X_{(1)}$ est inadmissible par rapport à la fonction de perte quadratique.

On compare maintenant ces deux méthodes par rapport à la fonction de perte de Laplace. On trouve facilement que on minimise le risk

$$\min_r E |T_r - \theta| = \min_r \left\{ \frac{1}{n} (nr - 1 + 2e^{-nr}) \right\} = \frac{0.693}{n}$$

si on choisit $r = (\ln 2)/n$. Donc l'estimateur $T_{(\ln 2)/n}$ est le meilleur dans la classe $\{T_r\}$ par rapport à la fonction de perte de Laplace. On trouve en même temps que

$$E |X_{(1)} - \theta| = \frac{1}{n},$$

et on en tire de nouveau que l'estimateur du maximum de probabilité est meilleur que l'estimateur du maximum de vraisemblance et donc l'estimateur de maximum de vraisemblance est inadmissible par rapport à la fonction de perte de Laplace. $X_{(1)}$ est superefficace puisque sa variance est plus petite que la borne de Crámer-Rao. On peut trouver d'autres exemples intéressants on peut trouver chez Blyth (1983), Voinov et Nikulin (1993, 1996).

6. MÉTHODE DU MAXIMUM DE VRAISEMLANCE ET LA MÉTHODE DES MOMENTS

Soit $X = (X_1, X_2, \dots, X_n)^T$ un échantillon, et d'après l'hypothèse

$$H_0 : X_i \sim f(x; \theta), \quad \theta = (\theta_1, \dots, \theta_s)^T \in \Theta \subseteq \mathbb{R}^s,$$

$$f(x; \theta) = h(x) \exp \left\{ \sum_{k=1}^s \theta_k x^k + v(\theta) \right\}, \quad x \in \mathcal{X} \subseteq \mathbb{R}^1, \quad (1)$$

où $\mathcal{X}$ est un ensemble borelien dans $\mathbb{R}^1$,

$$\mathcal{X} = \{x : f(x; \theta) > 0\} \text{ pour tout } \theta \in \Theta.$$

La famille (1) est très riche, on y trouve, par exemple, la famille des lois normales $N(\mu, \sigma^2)$:

$$f(x; \theta) = \frac{1}{\sqrt{2\pi}\sigma} \exp \left[-\frac{(x-\mu)^2}{2\sigma^2} \right], \quad |x| < \infty, \quad \Theta = \{\theta = (\mu, \sigma^2)^T : |\mu| < \infty, \sigma > 0\},$$

et la famille des lois de Poisson $P(\theta)$:

$$f(x; \theta) = \frac{\theta^x}{x!} e^{-\theta}, \quad \lambda \in \{0, 1, \dots\}, \quad \theta \in \Theta =]0, \infty[,$$

appartiennent à (1) avec $s = 2$ et $s = 1$ respectivement.

Supposons que

1) Le support $\mathcal{X}$ ne dépend pas de θ ;

2) Le Hessien

$$H_v(\theta) = - \left\| \frac{\partial^2}{\partial \theta_i \partial \theta_j} v(\theta) \right\|_{(s \times s)} \quad (2)$$

de la fonction $v(\theta)$ soit défini positif sur Θ .

3) Le moment $a_s = \mathbf{E}X_1^s$ d'ordre s existe.

Dans ce modèle H_0

$$U_n = \left(\sum_{i=1}^n X_i, \sum_{i=1}^n X_i^2, \sum_{i=1}^n X_i^3, \dots, \sum_{i=1}^n X_i^s \right)^T \quad (3)$$

est la statistique exhaustive minimale. De (1) on trouve que

$$-grad v(\theta) = a(\theta) = (a_1(\theta), a_2(\theta), \dots, a_s(\theta))^T, \quad (4)$$

et donc la statistique

$$T_n = \frac{1}{n} U_n \quad (5)$$

est le meilleur estimateur sans biais (**MVUE**, minimum variance unbiased estimateur, voir par exemple Voinov et Nikulin, 1993) pour $a(\theta)$, car

$$E_{\theta} T_n \equiv a(\theta), \quad \theta \in \Theta, \quad (6)$$

et cette relation nous permet d'une façon unique d'obtenir l'estimateur θ_n^* pour θ par la méthode des moments de l'équation

$$T_n = a(\theta) \quad (7)$$

en fonctions de la statistique exhaustive U_n .

Par ailleurs, les conditions 1)-3) nous garantissent l'existence de l'estimateur de maximum de vraisemblance $\hat{\theta}_n$, qui est la racine unique (!) de la même équation

$$T_n = a(\theta).$$

On en déduit donc que pour la famille exponentielle (1) les méthodes de maximum de vraisemblance et des moments nous amènent au même estimateur, $\hat{\theta}_n = \theta_n^*$, de θ .

Exemple 1

On connaît bien que pour la loi normale $N(\mu, \sigma^2)$ l'estimateur de maximum de vraisemblance $\hat{\theta}_n$ pour $\theta = (\mu, \sigma^2)^T$ est $\hat{\theta}_n = (\bar{X}_n, s_n^2)^T$, où

$$\bar{X}_n = \frac{1}{n} \sum_{i=1}^n X_i, \quad s_n^2 = \frac{1}{n} \sum_{i=1}^n (X_i - \bar{X}_n)^2, \quad (8)$$

et donc $\hat{\theta}_n = \theta_n^*$.

Exemple 2

Soit $X = (X_1, \dots, X_n)^T$ un échantillon, et d'après l'hypothèse H_0 :

$$X_i \sim f(x; \theta) = \frac{\theta^x}{x!} e^{-\theta}, x \in \mathcal{X} = \{0, 1, 2, \dots\}.$$

On connaît bien que dans ce cas

$$\hat{\theta}_n = \theta_n^* = \bar{X}_n.$$

Enfin on note que cette remarque nous permet de considérer la méthode de vraisemblance comme un cas particulier de la méthode des moments, voir Huber et Nikulin (1993), Greenwood et Nikulin (1996).

RECONNAISSANCE

Les auteurs sont reconnus à un arbitre et à C.M. Cuadras par leur utiles commentaires.

7. REFERENCES

- [1] **Bahadur, R.R.** (1958). «Examples of inconsistency of maximum likelihood estimates». *Sankhyā*, **20**, pp 207–210.
- [2] **Barnett, V.D.** (1966). «Evaluation of the maximum-likelihood estimator where the likelihood equation has multiple roots». *Biometrika*, **53**, **1 and 2**, 151–165.

- [3] **Basu, D.** (1954). *An inconsistency of the method of maximum likelihood*.
- [4] **Berger J.O., Wolpert R.L.** (1988). *The likelihood principle*, Hayward : Institute of mathematical statistics.
- [5] **Berkson, J.** (1980). «Minimum chi-square, not maximum likelihood». *AS*, **8**, **3**, 457–487.
- [6] **Blyth C.R.** (1982). *Maximum probability estimation in small samples*, in A. Festschrift for Erik Lehmann, (Bickel and Doksum ed.), Wadsworth, pp 83–86.
- [7] **Dharmadhikari, S. and Joag-Dev, K.** (1985). «Examples of non unique maximum likelihood estimators». *The American Statistician*, **39**, **3**, 199–200.
- [8] **Efron, B.** (1980). «Discussion of a paper by Joseph Berkson on Minimum chi-square». *AS*, **8**, **3**, 457–487.
- [9] **Evans B., Fraser D.A.S. and Monette J.** (1986). «On principles and arguments to likelihood». *Canadian Journal of Statistics*, **14**, pp 181–199.
- [10] **Ferguson, T.S.** (1982). «An inconsistent maximum likelihood Estimate». *JASA*, **77**, **380**.
- [11] **Fraser D.A.S., Monette G. and Ng K.W.** (1984), *Marginalization, likelihood and structural models*, in Multivariate Analysis VI, (ed. P.R. Krishnaiah), pp 209–217, Amsterdam : North Holland.
- [12] **Goldstein, M. and Howard, J.V.** (1991). «A likelihood paradox». *JRSS*, **53**, **3**, 619–628.
- [13] **Huber, P.J.** (1980). *Robustness theory*.
- [14] **Huber-Carol, C.** (1985). *Théorie de l'inférence statistique robuste*. Springer-Verlag.
- [15] **Huber C., Nikulin M.** (1993). *Applications statistiques des transformations des variables aléatoires*, Preprint, U.F.R. M.I.2S, Université Bordeaux 2.
- [16] **Joshi V.M.** (1989). «A counter-example against the likelihood principle». *JRSS B*, **51**, 215–216.
- [17] **Kalbfleish J.D. and Prentice R.L.** (1980). *The statistical Analysis of failure time data*, Wiley, New York.
- [18] **Kempthorne, O.** (1989). «The fate worse than death and other curiosities and stupidities». *The American statistician*, **43**, **3**.
- [19] **Kiefer, J. and Wolfowitz, J.** (1956). «Consistency of the maximum likelihood estimator in the presence of infinitely many incidental parameters». *AMS*, **27**, 887–906.
- [20] **Le Cam, L.** (1979). «Maximum likelihood: an introduction». *Lecture notes*, **18**.
- [21] **Lehmann, E.L.** (1983). *Theory of point Estimation*. Wiley, N.Y.

- [22] **Norton, H.W.** (1956). «One likelihood adjustment may be inadequate». *Biometrics*, 79–81.
- [23] **Richards, F.S.** (1961). *A method of maximum likelihood estimation*.
- [24] **Russell, A.B.** and **Samaniego, F.J.** (1983). «Maximum Likelihood estimation for a discrete multivariate shock model». *JASA*, **78**, **382**, 445–448.
- [25] **Scholz, F.W.** (1980). «Towards a unified definition of maximum likelihood». *The Canadian Journal of Statistics*, **8**, **2**, 193–203.
- [26] **Van der Laan, M.** (1995). *Efficient and Inefficient Estimation in Semiparametric Models*. Mathematische Instituut van de Universiteit. Utrecht, 189 p.
- [27] **Voinov, V.G.** et **Nikulín, M.S.** (1993). *Unbiased estimators and their applications*. Vol 263, Kluwer Academic Publishers, Dordrecht, Boston, London.
- [28] **Weiss L., Wolfowitz J.** (1966). «Generalized maximum likelihood estimators». *Theory of Probability and its Applications*, **11**, **1**, 68–93.
- [29] **Weiss L., Wolfowitz J.** (1974). *Maximum Probability Estimators and Related Topics*, Lectures Notes in Mathematics, Springer-Verlag.
- [30] **Weiss L.** (1986). «A note on small-sample maximum probability estimation». *Statistics and Probability letters*, **4**, 109–111.

ENGLISH SUMMARY

SOME REMARKS ON THE MAXIMUM LIKELIHOOD

CATHERINE HUBER*

Université de Paris V

MIKHAIL NIKULIN**

Université de Bordeaux II

Some paradoxes on the maximum likelihood principle are presented and commented. We consider the properties of the maximum likelihood estimators as a particular case of the M -estimators. We propose and unified theory which includes non-dominated models. Several examples are given.

Keywords: M -estimators, likelihood paradoxes, unified maximum likelihood theory.

* C. Huber. Université de Paris V, 45 rue des Saints-Pères 75 270 Paris cedex 06

**M. Nikulin. Université de Bordeaux II, 146 rue Léo Saignat 33 800 Bordeaux

– Received July 1995.

– Accepted March 1996.

The maximum likelihood method is both one of the mostly used and one of the most controversial method in statistical estimation. It is intuitively appealing because the likelihood appears to capture the whole available and useful information contained in the data, and it has, too, a theoretical justification, via its excellent asymptotic properties under regularities conditions on the underlying model, which turn out very often to be even optimality properties.

Nevertheless, this method has several drawbacks which have been pointed out by L. Le Cam [1979] and P. J. Huber [1981], but also, among others, by J. Berkson [1980], R. R. Bahadur [1958], D. Basu [1980], T. S. Ferguson [1982], M. Goldstein and J. V. Howard [1991], S. Dharmadhikari and K. Joag-Dev [1985], V. G. Voinov and M. S. Nikulin [1993] among others. As this method is universally applicable, there is a tendency to use it sometimes without much care, which leads to disastrous results. One of the most famous examples is the one of Le Cam which exhibits a maximum likelihood estimates which is not consistent (as the number n of observations grows to infinity, it converges to twice the value of the parameter it is supposed to estimate, because the number of nuisance parameters involved in the model increases at speed n) [L. Le Cam, 1979]. The likelihood itself is a source of problems when there are several versions of the probability density. Moreover, convexity properties are needed in order to get a global maximum, otherwise it is difficult or even impossible to get an overall maximum. And even though the estimate is unique, in that very favourable case, it happens very often that the likelihood curve is very flat in a neighbourhood of the maximum so that a big change in the estimate results in a very tiny modification of the likelihood, which results in numerical problems and lack of reliability of the evaluation of the estimate.

It is thus necessary, as is recommended B. Efron [1980], to be rather careful when using this method, and, in that case, no major disaster is to be feared of. But, anyway two intrinsic drawbacks are to be taken into account:

1. First, the lack for a definition of a maximum likelihood estimate when there is an ambiguity upon the version of the density which has to be in use, or else when the model is not dominated. F. W. Scholz [1980] tentatively defined a unified version of the maximum likelihood. It seems that it is not well known and that this interesting idea did not spread as much as it should have done.
2. Second, the lack of robustness of the most usual maximum likelihood estimates, those related to the gaussian models. Imbedding those estimates in the more general M -estimates [P.J. Huber, 1981], is a way of correcting the lack of stability of certain M-L estimates.

This paper treats of three different topics:

- The properties of the maximum likelihood estimates considered as a special case of M estimates under regularity conditions of the statistical model: consistency, asymptotic normality and sensitivity.

An M estimate is a solution of equation

$$\sum_{i=1}^n \varphi(x_i, \hat{\theta}) = 0$$

where $\varphi(x, t)$ can be chosen to be equal to $\ell(x, t) = \frac{\partial \text{Log} f(x, t)}{\partial t}$, which is the special case of an M-L estimate regularity properties on φ ensures consistency, asymptotic normality and also robustness through the concept of sensitivity.

- A tentative unified theory of the maximum likelihood based on F.W. Scholz approach.
- A paradox arising from a measure of maximum likelihood is given on a finite probability space.

DISEÑO Y ANÁLISIS BAYESIANOS DE RÉPLICAS EN LA EXPERIMENTACIÓN CIENTÍFICA*

M.J. BAYARRI*

M.A. MARTÍNEZ*

Universitat de València

En la práctica estadística es frecuente la replicación de experiencias estadísticas, justificada por muy diversos intereses: ratificar conclusiones, estudiar la variación en las respuestas cuando se experimenta sobre una población distinta, validar un modelo, detectar sesgos, etc. Al disponer, antes de replicar, de los resultados de cierto estudio previo, es razonable pensar en buscar un buen diseño para la réplica utilizando la información que dicho primer estudio proporciona. Un diseño óptimo de la réplica consistirá en decidir, asumido un modelo, sobre el menor tamaño muestral que otorga al investigador suficientes garantías de concluir con éxito en su inferencia. Una modelización jerárquica bayesiana permite precisar fácilmente las relaciones entre datos y poblaciones (original y replicada) y derivar conclusiones sobre los resultados más probables (en función del tamaño de la réplica) a través de las distribuciones a posteriori y predictivas que se obtengan.

Bayesian design and analysis of replications in scientific experimentation

Keywords: Replicación, modelo jerárquico bayesiano, factor Bayes, simulación, aproximación Monte-Carlo.

*Este trabajo ha sido parcialmente subvencionado por la Consellería de Cultura, Educació i Ciència de la Generalitat Valenciana, bajo el proyecto de Investigación GV-1081/93.

*M.J. Bayarri y M.A. Martínez. Dpto. Estadística e I.O. Universitat de València. Dr. Moliner, 50. 46100 BURJASSOT. València.

—Article rebut el desembre de 1995.

—Acceptat el juny de 1996.

1. INTRODUCCIÓN

1.1 La replicación de experimentos. Contexto

A pesar de que la **replicación** es una práctica estadística habitual en la investigación experimental, no parece existir una modelización sistemática encaminada a sacar el máximo provecho de la información que puede proporcionar; en la bibliografía sólo hemos encontrado algunos apuntes acerca de su potencial en la investigación y diversas consideraciones sobre la interpretación «acertada» de los resultados de la réplica (ver Neuliep, 1990).

Entendemos por **réplica** o **replicación** la repetición de una experiencia estadística ϵ_o , consistente en una nueva recogida de datos y análisis de los mismos, con el objetivo último de inferir sobre cierto parámetro poblacional (referente a una media, una diferencia de medias, de efectos de tratamientos, un coeficiente de correlación, etc).

Para replicar una experiencia estadística previa, el investigador elegirá, en función del objetivo de su estudio, la misma población muestreada en ϵ_o u otra diferente; decidirá asimismo entre la reproducción de las condiciones de experimentación originales (replicación exacta) y su modificación (parcial o total), para inferir sobre la variabilidad de la característica que estudia.

Se hace uso de la replicación en todos aquellos campos de investigación en los que es común la repetición de análisis estadísticos para profundizar en el estudio de cierta característica poblacional; campos tan diversos como Medicina, Ecología, Economía, Sociología, Agricultura, Mejora Animal, etc.

1.2 Interpretación

A pesar de lo ampliamente debatido que ha sido el tema de la replicación (ver Neuliep, 1990), no existe un criterio claro sobre cómo valorar los resultados que las réplicas proporcionan. Para ilustrar esta problemática Tversky y Kahnemann (ver Utts, 1991) planteaban la replicación de una experiencia basada en la resolución de un contraste t sobre cierto parámetro de interés θ (contraste que nosotros adaptamos, por simplificar, a uno basado en una población *normal* con media θ y varianza conocida e igual a 1):

$$\begin{cases} H_0 & : \theta \leq 0 \\ H_1 & : \theta > 0. \end{cases}$$

A partir de una primera muestra de 15 observaciones, se había obtenido un valor para el estadístico de contraste de:

$$z_o = 2.19.$$

(La media muestral observada había sido $x_o = 0.5654$).

A continuación se formulaba la siguiente pregunta a un grupo de expertos: «Si replicáramos con el mismo tamaño muestral ($n_r = 15$), ¿cuál sería el máximo valor del estadístico z_r que podría ser considerado como *fracaso* al reproducir el primer estudio?».

La respuesta general fue un valor de:

$$z_r = 1.60.$$

Como hacían notar Tversky y Kahnemann (a los que nos referiremos por T&K a partir de ahora), las implicaciones de este estudio empírico resultaban un poco contradictorias. Por una parte, mientras que el resultado de la réplica ($z_r = 1.60$) se percibía a primera vista como un fracaso a la hora de intentar reproducir el experimento original, reduciendo así nuestra confianza en las conclusiones del primer estudio, si decidiéramos combinar los datos del primer y segundo estudios como pertenecientes a una misma muestra (llevando a cabo un análisis global sobre el efecto en cuestión), conseguiríamos sin embargo, ratificar con más fuerza las conclusiones de dicho primer estudio, al obtener para el estadístico combinado z un valor superior al de z_o :

$$z = 2.68,$$

en contra entonces, de lo que en principio hacía suponer la conclusión individual de la réplica.

La forma más adecuada de concluir sobre lo observado en la réplica y cómo definir entonces la consecución de «éxito» o fracaso al replicar, dependerá de cada problema específico. En ocasiones interesará combinar toda la información disponible (estudio original y réplica), mientras que en otras el interés residirá en la comparación de las conclusiones obtenidas individualmente en uno y otro estudios.

Para abordar la sistematización de las réplicas asumimos que el investigador ya ha planteado:

- (a) cuál es el problema concreto que pretende resolver con la réplica, esto es, cuáles son los *objetivos* inmediatos con que justifica esta «repetición» estadística;
- (b) cómo va a replicar: qué observará, qué modelo propondrá y al fin, qué metodología y procedimiento estadísticos seguirá para concluir en su estudio.

Los *objetivos* del experimentador para emprender la replicación de un experimento pueden ser muy diversos. Basados en las referencias mencionadas sobre replicación y en algunas indicaciones encontradas en réplicas publicadas, hemos intentado caracterizar los más frecuentes e importantes (ver Mayoral, 1995). Entre ellos podríamos resaltar:

1. la **confirmación** de las conclusiones obtenidas en el primer estudio, esto es, la validación (ratificación) de su inferencia, replicando sobre la misma población de ϵ_0 y bajo idénticas condiciones de experimentación;
2. la **generalización** de dichas conclusiones a otras poblaciones distintas a la muestreada en ϵ_0 ;
3. el estudio de la **variabilidad** de la característica de interés en función de la evolución de la población y/o las modificaciones experimentales introducidas respecto de ϵ_0 ; en definitiva, lograr concluir sobre el dominio de las conclusiones primeras y su sensibilidad a los cambios en la experimentación;
4. la **distinción** de poblaciones y/o **discriminación** de variables en la explicación de cierto efecto poblacional de interés, detectable únicamente cuando se repiten experimentos modificando de alguna forma las variables que los afectan, o simplemente alterando su valor (replicaciones no exactas);
5. la **detección de sesgo** en un primer estudio, ocasionado por la ausencia o deficiencias del diseño.

Aunque no es crucial al desarrollo de nuestro trabajo, supondremos por simplicidad que la réplica aborda un problema estadístico similar al del estudio original (fuera éste un contraste de hipótesis, un problema de estimación, etc), utilizando para su resolución la misma metodología (clásica, bayesiana, ...) y procedimiento (p-valores, estimadores, distribuciones a posteriori, probabilidades finales, factores Bayes, ...) empleados en el primer estudio.

La estadística bayesiana hace muy fácil la incorporación secuencial de información, hecho valorable en todo problema que involucra replicaciones sucesivas de un primer experimento de las que ir aprendiendo. Asimismo, permite una modelización sencilla para las relaciones entre los distintos experimentos por medio de las distribuciones de los parámetros (modelo jerárquico). La inferencia y la predicción se obtienen de una manera natural a partir de las distribuciones a posteriori y predictivas respectivamente. Constituye así la replicación, un marco idóneo para la aplicación de la metodología bayesiana. La consideración de los objetivos, inferencias y características particulares de cada investigación conducirán nuestra modelización y análisis bayesianos de las réplicas.

2. PLANTEAMIENTO DEL PROBLEMA

Llamaremos ε_r a la replicación de cierto análisis estadístico previo ε_o . X_o y X_r representarán las variables a observar en cada uno de ellos (en nuestro caso, estadísticos suficientes; en general identificarán a vectores de observaciones).

Asumimos el estudio sobre cierto parámetro poblacional, conceptualmente similar en ambos experimentos (una media para nosotros), representado por θ_o en el estudio original y θ_r en la réplica.

Trabajamos únicamente bajo tres supuestos básicos (una relación más exhaustiva junto con su correspondiente análisis estadístico puede encontrarse en Mayoral, 1995).

- S1. ε_o consiste en un contraste de hipótesis sobre el parámetro poblacional θ_o . Se propone el mismo contraste en la réplica, sobre una población *similar* a la de ε_o , con el objetivo de **reproducir** la conclusión original acerca del rechazo (o no) de la hipótesis nula.
- S2. En ε_o se infiere sobre un parámetro de interés θ_o . La réplica, llevada a cabo generalmente sobre una población distinta, se plantea para inferir sobre la **discrepancia** entre ambas poblaciones, esto es, sobre la diferencia de efectos $\delta = \theta_r - \theta_o$.
- S3. Se sospecha la **existencia de sesgo** (β) en un primer estudio en el que se pretendía inferir sobre cierto parámetro μ . Suponiendo poder controlar dicho sesgo en un estudio posterior sobre la misma población, se trata de inferir sobre él.

La flexibilidad del **modelo jerárquico bayesiano** permite una fácil modelización de las relaciones que surgen en replicación: la discrepancia entre los parámetros θ_o y θ_r ocasionada por el distanciamiento temporal, espacial, poblacional, etc, entre el primer estudio y la réplica; la presencia de ruidos no evitables en la implementación y que afectan a la percepción del efecto real a través de la muestra; la variación del modelo que las modificaciones experimentales en la réplica pueden ocasionar (inclusión/exclusión de variables de influencia, alteración del rango de valores posibles para las variables y parámetros involucrados,...), etc.

Utilizamos un modelo jerárquico sencillo con únicamente tres niveles: en el primer nivel modelizamos la distribución condicional de las variables observables (X_o y X_r); en el segundo nivel, las relaciones entre los parámetros que definen al modelo en el primer nivel (θ_o y θ_r); el tercer nivel es reservado para especificar las distribuciones a priori de los parámetros restantes (hiper-parámetros).

Trabajando con metodología bayesiana derivamos las correspondientes **distribuciones a posteriori** de interés, a través de las cuales podremos medir la similitud o

discrepancia entre los estudios original y réplica, y una *distribución predictiva* para el resultado de la réplica que nos permitirá, entre otras cosas, la planificación «óptima» de ésta.

En función del objetivo del investigador y de la inferencia a realizar se proponen distintas posibilidades de definir el *éxito al replicar*. Cada definición de éxito se traduce en cierta condición que habrá de satisfacer el resultado de la réplica, x_r . Puesto que éste aún no se ha observado (todo el análisis es previo a la replicación), la probabilidad de éxito se calculará utilizando la distribución de la variable aleatoria X_r . Dada la información proporcionada por el primer estudio, esta distribución es la predictiva $m(x_r|x_o)$, dependiente del tamaño muestral de la réplica, n_r . El análisis derivado del modelo propuesto proporcionará herramientas para diseñar una réplica (decidir un tamaño n_r) que otorgue suficientes «garantías de éxito», o simplemente sugerirá desestimar su práctica en caso de no rebasar éstas el nivel exigido. Fijado por el investigador dicho nivel de garantías M , el tamaño «óptimo» para replicar se obtendrá para el menor valor de n_r que proporcione una probabilidad de éxito tal que:

$$(1) \quad Prob(\text{éxito}) \geq M.$$

Estudiamos cada uno de los supuestos (S1, S2 y S3) modelizando y obteniendo las inferencias necesarias para, una vez planteado en cada caso lo que se entenderá por éxito, calcular la probabilidad de lograrlo con los datos del ejemplo propuesto por T&K sobre una «posible» replicación de un problema de contraste de hipótesis (planteado al inicio de la sección 1.2).

Introduciremos dos modelos que denotaremos por M12 y M3. El modelo M12 será la base para estudiar la reproducción de conclusiones en el contraste (S1) y para comparar las poblaciones (S2). El modelo M3 nos permitirá inferir sobre el sesgo «intuído» en el primer estudio (S3). Para ambos modelos se incorporará en la última parte de este trabajo, un grado más de incertidumbre (modelos M12A y M3A), con el objeto de robustecer los resultados ya obtenidos tras flexibilizar la modelización.

3. REPRODUCCIÓN Y COMPARACIÓN CON LA RÉPLICA

3.1. Modelo M12

Planteamos la repetición de cierta experiencia estadística previa, para lo cual asumimos la modelización siguiente:

I. DATOS

Suponemos que las muestras en el estudio original y en la réplica son (condicionalmente) independientes, y que en cada estudio el estadístico suficiente X_i tiene una distribución normal con varianza conocida:

$$(2) \quad \begin{aligned} f(x_o, x_r | \theta_o, \theta_r) &= f(x_o | \theta_o) \cdot f(x_r | \theta_r) \\ f(x_i | \theta_i) &= N(x_i | \theta_i, \gamma_i^2), \quad i = o, r, \end{aligned}$$

con $\gamma_i^2 = \sigma^2/n_i$, σ^2 conocida (=1 por simplificar) y n_i ($i = o, r$) el tamaño muestral utilizado en cada estudio.

II. PARÁMETROS

Suponemos (como una primera aproximación sencilla a la posible relación entre los dos estudios) que los parámetros poblacionales en el experimento original y en la réplica son intercambiables (la consideración de relaciones más complejas entre ambos experimentos —involucrando covariables, relaciones espaciales y/o temporales, etc—, se abordará en trabajos futuros):

$$(3) \quad \begin{aligned} p(\theta_o, \theta_r | \mu) &= p(\theta_o | \mu) \cdot p(\theta_r | \mu) \\ p(\theta_i | \mu) &= N(\theta_i | \mu, \tau^2), \quad i = o, r. \end{aligned}$$

La varianza τ^2 , que suponemos conocida, cuantifica la información a priori acerca de la semejanza o disparidad entre los dos estudios: cuanto menor es τ^2 , más homogéneas se consideran las poblaciones del experimento original y de la réplica. Estudiaremos en nuestro ejemplo de T&K la sensibilidad de las conclusiones a la magnitud de τ^2 utilizando un valor moderadamente alto, $\tau^2 = 1$, quince veces mayor que la varianza muestral del estudio original, y otro moderadamente pequeño, $\tau^2 = 1/5$, tan sólo tres veces mayor que $\gamma_o^2 (= 1/15)$.

III. HIPER-PARÁMETRO

Utilizamos una distribución mínimo-informativa a priori para el hiper-efecto poblacional μ (que puede interpretarse, si se desea, como la media de una hiperpoblación de la que «muestreamos» θ_o y θ_r):

$$(4) \quad p(\mu) \propto \text{cte}.$$

A la modelización clásica (primer nivel) se añade el elemento bayesiano $p(\theta_o, \theta_r)$ que modeliza la relación entre los dos experimentos. Tendremos así acceso tanto a las

conclusiones derivadas de análisis clásicos en ε_o y ε_r , como a las obtenidas de análisis bayesianos. Obviamente, resulta más natural llevar a cabo un análisis coherente y concluir sobre los parámetros de ambos estudios en base a sus correspondientes distribuciones a posteriori. Sin embargo, el análisis que proponemos es más amplio y podrá aplicarse también cuando las inferencias que se deseen realizar acerca de θ_r estén basadas en la metodología clásica. El hecho clave para inferir sobre la réplica es que su resultado X_r es una variable aleatoria a la que el modelo anterior permite dotar de una distribución predictiva, obtenida de las distribuciones finales para los parámetros:

$$(5) \quad \begin{aligned} m(x_r|x_o) &= \int f(x_r|\theta_r) \cdot p(\theta_r|x_o) d\theta_r \\ &= N(x_r|x_o, \gamma_o^2 + \gamma_r^2 + 2\tau^2), \end{aligned}$$

donde $p(\theta_r|x_o)$ contiene toda la información sobre el parámetro poblacional θ_r proporcionada por el primer estudio:

$$p(\theta_r|x_o) = \int p(\theta_r|\mu) \cdot p(\mu|x_o) d\mu,$$

con

$$p(\mu|x_o) \propto f(x_o|\mu) \cdot p(\mu) \propto \int f(x_o|\theta_o) \cdot p(\theta_o|\mu) d\theta_o.$$

3.2. Reproducción de conclusiones

Suponemos realizado un primer análisis estadístico ε_o que resolvía cierto contraste de hipótesis sobre una media poblacional. Se propone un segundo muestreo ε_r sobre una población que a priori se supone «próxima» a la inicial (podría tratarse de la misma) y, manteniendo fijas las condiciones de experimentación originales, se pretende reproducir las primeras conclusiones. Nótese que cuando la replicación se llevara a cabo sobre la misma población de ε_o , lograr este objetivo significaría «validar» el primer análisis. La desviación aleatoria entre las medias poblacionales θ_o y θ_r viene descrita por el modelo de intercambiabilidad. El contraste que en ε_r se plantea será análogo al del estudio original:

$$\begin{cases} H_0^i & : \theta_i \leq c_i \\ H_1^i & : \theta_i > c_i, \quad i = o, r. \end{cases}$$

Consideramos $c_r = c_o = c$, esto es, pretendemos concluir de la misma forma sobre la magnitud de θ_r que sobre la de θ_o . Sin pérdida de generalidad tomamos $c = 0$, puesto que cualquier contraste de este tipo lo podemos desplazar en torno al cero sin más que restar una constante.

En el ejemplo de T&K, la inferencia en el primer estudio se llevaba a cabo a través del p-valor observado (metodología clásica). Al observar en él un valor para

el estadístico de contraste de $z_o = 2.19$, se rechazaba la hipótesis nula a un nivel de significatividad $\alpha = 0.05$. Reproducir conclusiones significará entonces, **volver a rechazar** la hipótesis nula en la réplica:

$$\begin{aligned}\text{ÉXITO} &\equiv \text{Reproducir } \varepsilon_o \\ &\equiv \text{Rechazar } H_0^r.\end{aligned}$$

Una vez definido lo que consideraríamos **éxito** en la réplica, buscamos el «mejor» tamaño para lograrlo, a partir del análisis bayesiano derivado del modelo M12.

El rechazo de la hipótesis nula en la réplica admite sin embargo dos posibilidades (la disyuntiva que T&K cuestionaban):

(R1) Si se decide analizar la réplica de forma aislada, sin añadir la información que proporciona el primer estudio (primer supuesto en la situación paradójica de T&K), «**rechazar**» significará, siguiendo la metodología clásica para inferir en la réplica, obtener un p-valor p_r inferior a 0.05. En el escenario en que trabajamos, el p-valor (inferencia clásica) coincide con la probabilidad a posteriori de H_0^r (inferencia bayesiana). Así, el rechazo de H_0^r será la conclusión tanto de un clásico, como de un bayesiano que asumiera una función de pérdida que hiciera razonable rechazar dicha hipótesis cuando su probabilidad final resultara inferior a α :

$$p_r \leq \alpha \Leftrightarrow P(H_0^r | x_r) \leq \alpha.$$

Dada su equivalencia, es suficiente considerar uno de estos índices de evidencia para llevar a cabo nuestro análisis. Así, en la réplica tendremos **éxito R1** cuando:

$$(6) \quad p_r \leq 0.05.$$

Puesto que p_r depende del resultado de la réplica (también de su tamaño), $p_r = p_r(X_r, n_r)$, antes de replicar es una variable aleatoria cuya distribución viene determinada por la distribución de X_r . La distribución predictiva $m(x_r | x_o)$ será la herramienta natural para calcular la probabilidad de que la condición [6] se cumpla:

$$\begin{aligned}\text{Prob}(\text{éxito R1}) &= P[p_r \leq \alpha | x_o] = P[X_r \geq \frac{z_{1-\alpha}}{\sqrt{n_r}} | x_o] \\ &= \int_{z_{1-\alpha}/\sqrt{n_r}}^{+\infty} m(x_r | x_o) dx_r.\end{aligned}$$

(R2) Si se pretende combinar el resultado del estudio original con el de la réplica (segundo supuesto de T&K), es preciso especificar el tipo de análisis a realizar. Aunque obviamente es posible llevar a cabo un análisis clásico utilizando un p-valor combinado, nos decantamos por un análisis bayesiano que involucra toda la información disponible mediante las distribuciones a posteriori. Por mantener la analogía con R1, hablaremos de conseguir **éxito R2** rechazando H'_0 cuando la probabilidad final para la hipótesis nula una vez observada la réplica sea inferior a 0.05, esto es:

$$(7) \quad P(\theta_r \leq 0 | x_o, x_r) \leq 0.05.$$

De nuevo, fijado n_r obtendremos un rango de valores de X_r para los que se verifica la condición [7], y de ahí se podrá elegir el tamaño que otorgue suficientes «garantías de éxito» (una probabilidad de éxito razonablemente alta):

$$Prob(\text{éxito R2}) = \int_{\{x_r: P(\theta_r \leq 0 | x_o, x_r) \leq 0.05\}} m(x_r | x_o) dx_r \geq M,$$

con M el nivel de garantías exigido para replicar (impuesto por el investigador).

La elección de uno u otro tipo de rechazo en la inferencia a realizar con la réplica dependerá del investigador. Las probabilidades de éxito según R1 y R2 son bastante similares, como puede apreciarse en la Figura 1b, si bien el rechazo de H'_0 combinando información (R2) resulta algo más probable que el rechazo con sólo el resultado de la réplica (R1). Para una réplica de igual tamaño que ε_o ambas probabilidades de rechazo giran en torno a 0.5, sensiblemente inferiores a lo que intuitivamente cabía esperar.

Puesto que suponemos una **replicación fiel** del primer estudio, es razonable asumir a priori una **discrepancia pequeña** entre los parámetros poblacionales, es decir, un valor pequeño para $\tau^2 (= 1/5)$. Hemos hecho sin embargo, un estudio comparativo entre los resultados que obtendríamos partiendo de un valor pequeño ($\tau^2 = 1/5$) y uno grande ($\tau^2 = 1$), con el fin de estudiar la sensibilidad de las conclusiones respecto de la opinión inicial incluída en la modelización. Para los dos tipos de rechazo el comportamiento es el mismo (Figura 1a): en tamaños pequeños apenas divergen las probabilidades de éxito para los dos valores de τ^2 , si bien a partir de $n_r = 20$ resulta ya apreciable el efecto de la restricción inicial más fuerte sobre la proximidad entre θ_o y θ_r ($\tau^2 = 1/5$), que da a la réplica mayor probabilidad de reproducir las conclusiones de ε_o .

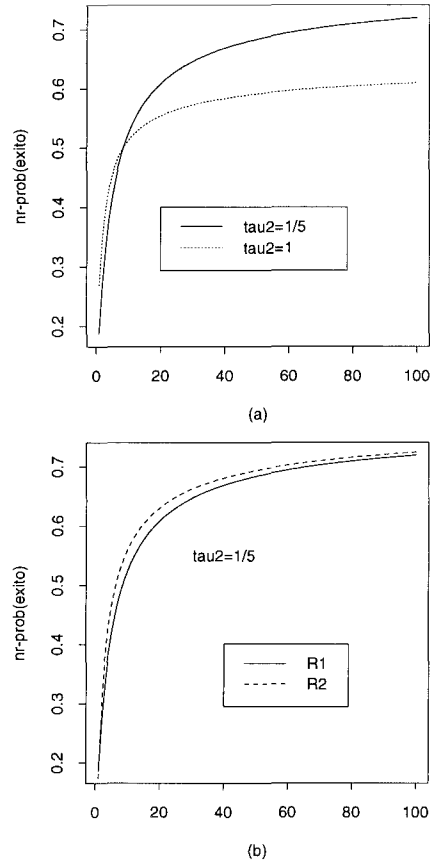


Figura 1. Reproducción de Conclusiones. Éxito $\equiv$ rechazo de $H_0^r : \theta_r \leq 0$. Probabilidad de éxito en función del tamaño de la réplica.

(1a) Rechazo con p-valor —R1—, para estudios próximos ($\tau^2 = 1/5$) y alejados ($\tau^2 = 1$).

(1b) Estudios próximos. Rechazos R1 —réplica aislada— y R2 —combinando original y réplica—.

3.3. Comparación de poblaciones

Un investigador podría estar interesado en repetir una experiencia estadística previa sobre cierta población de interés, distinta en principio a aquella sobre la que se muestreó en el estudio original, con la pretensión de comparar ambas poblaciones respecto del efecto estudiado. Este sería el contexto de las investigaciones basadas en estudiar la evolución (en el tiempo) de una determinada población, o su sensi-

bilidad ante modificaciones en las condiciones experimentales. Una discrepancia a priori grande entre las poblaciones muestreadas en los dos estudios vendrá dada por un valor «grande» para τ^2 (que representaba la variabilidad inicial de θ_o y θ_r).

Consideraremos a partir de ahora que tanto en ε_o como en la réplica, las inferencias se llevan a cabo (por parte del investigador) desde una perspectiva totalmente bayesiana. Así, el resultado del primer estudio (x_o) condicionará siempre la inferencia en la réplica (toda función de x_r dependerá implícitamente de x_o : $\gamma(x_r) \equiv \gamma(x_r|x_o)$).

Un objetivo natural que se plantea el investigador cuando replica en esta situación consiste en inferir sobre la diferencia de medias $\delta = \theta_r - \theta_o$. Una vez observada la réplica, la inferencia sobre las medias θ_o y θ_r se lleva a cabo con la distribución final:

$$p(\theta_o, \theta_r | x_o, x_r) \propto f(x_o, x_r | \theta_o, \theta_r) \cdot p(\theta_o, \theta_r),$$

con

$$p(\theta_o, \theta_r) = \int p(\theta_o, \theta_r | \mu) \cdot p(\mu) d\mu \propto p(\theta_r | \theta_o),$$

de donde se obtiene la distribución a posteriori para la diferencia δ :

$$(8) \quad p(\delta | x_o, x_r) = N(\delta | dQR, \gamma_o^2 Q + \gamma_r^2 Q^2 R),$$

con $d = x_r - x_o$,

$$Q = \frac{2\tau^2}{\gamma_o^2 + 2\tau^2}, \quad \text{y} \quad R = \frac{\gamma_o^2 + 2\tau^2}{\gamma_o^2 + \gamma_r^2 + 2\tau^2}.$$

Los objetivos en la réplica respecto a la inferencia a realizar sobre δ podrían ser:

- C1.** conseguir una estimación de δ lo suficientemente precisa,
- C2.** resolver el contraste de igualdad de efectos ($\theta_o = \theta_r$).

3.3.1. Precisión

Una forma de cuantificar la precisión de las inferencias a posteriori sobre la diferencia de medias consiste en utilizar la longitud de una región creíble para δ , obtenida a partir de su distribución final. Una réplica en las condiciones expuestas en 3.3 se considerará un éxito si dicha longitud, $L_\delta(x_r, n_r)$, resultara suficientemente pequeña (inferior a cierto S fijado por el investigador):

$$(9) \quad L_\delta(X_r, n_r) \leq S.$$

El experimentador elegiría para replicar aquel tamaño n_r que le proporcionara suficientes garantías de éxito:

$$Prob(\text{éxito}) = P^{X_r|x_o}[L_\delta(X_r, n_r) \leq S] \geq M.$$

En nuestro modelo, la longitud de la región creíble para δ , de contenido probabilístico $1 - \alpha$, resulta independiente de X_r :

$$L_\delta(X_r, n_r) \equiv L_\delta(n_r) = 2 \cdot z_{1-\alpha/2} \cdot \sqrt{v_\delta^2(n_r)},$$

donde $z_{1-\alpha/2}$ es el cuantil de la normal estándar y v_δ^2 la varianza de la distribución a posteriori de δ , que sólo depende del tamaño n_r de la réplica. Así, con probabilidad 1 el investigador podrá inferir sobre δ con una precisión dada, $1/v_\delta^2(n_r)$, siempre que replique con un tamaño n_r tal que:

$$v_\delta^2(n_r) \leq S^*,$$

para un S^* convenientemente pequeño.

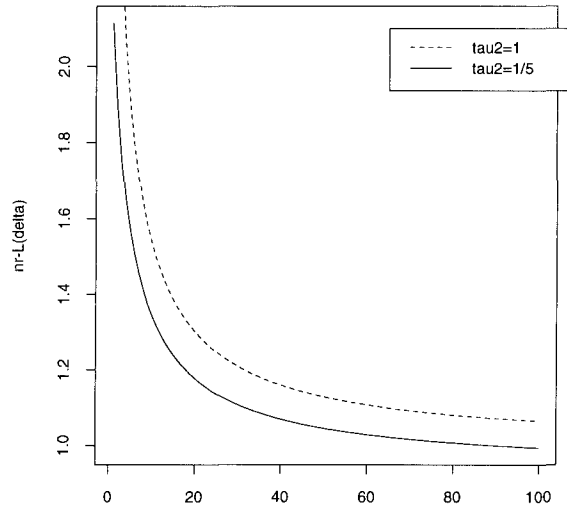


Figura 2. Comparación de Poblaciones. Éxito $\equiv$ Precisión al inferir sobre $\delta = \theta_r - \theta_o$. Longitud del intervalo creíble para δ (de probabilidad 0.95), L_δ , en función de n_r (estudios próximos y alejados).

En la Figura 2 hemos representado la longitud de la región creíble al 95% para δ , $L_\delta(n_r)$, en los casos $\tau^2 = 1$ y $\tau^2 = 1/5$. Una mayor proximidad inicial entre los estudios original y réplica repercute (como era de suponer) en una mayor precisión en la inferencia final sobre δ . En ambos supuestos iniciales el comportamiento es similar: un decrecimiento rápido de la longitud de la región creíble, seguido de una estabilización de su valor (a partir de $n_r = 50$) por debajo de 1.17. Ni una mayor

cercanía a priori entre las poblaciones, ni la incorporación de más observaciones cuando partamos de poblaciones supuestamente distintas, proporcionarán ya una ganancia apreciable de precisión en la inferencia sobre δ .

3.3.2. Contraste de igualdad

Una forma de inferir acerca de la discrepancia entre ambas poblaciones (original y réplica) es plantear el contraste sobre la igualdad de sus medias:

$$\begin{cases} H_0^\delta & : \delta = 0 \\ H_1^\delta & : \delta \neq 0, \end{cases}$$

donde $\delta = \theta_r - \theta_o$.

Desde una perspectiva clásica, suele considerarse «éxito» en un contraste, el rechazo de la hipótesis nula. Aquí llevaremos a cabo el análisis desde una perspectiva bayesiana y nuestro único objetivo será el de resolver el contraste, ya sea a favor de una o de otra hipótesis, con «suficiente» grado de confianza. El éxito en la réplica consistirá pues, en que con su resultado podamos discriminar fácilmente entre ambas hipótesis en base a sus probabilidades finales, o lo que es igual, en obtener un factor Bayes a favor de una de dichas hipótesis lo bastante grande.

El factor Bayes (el cociente entre los «odds» a posteriori y los «odds» a priori) es una herramienta bayesiana que cuantifica la evidencia a favor o en contra de las hipótesis del contraste. Puesto que no involucra explícitamente las probabilidades iniciales para cada hipótesis,

$$\frac{P(H_0)}{P(H_1)} \cdot B_{01} = \frac{P(H_0|\text{datos})}{P(H_1|\text{datos})},$$

puede interpretarse como un cociente de las verosimilitudes ponderadas:

$$\begin{aligned} B_{01}^\delta &= \frac{f(d|H_0^\delta)}{f(d|H_1^\delta)} \\ (10) \quad &= \frac{N(d|0, \gamma_o^2 + \gamma_r^2)}{N(d|0, \gamma_o^2 + \gamma_r^2 + 2\tau^2)}, \end{aligned}$$

donde:

$$d = x_r - x_o, \text{ y } f(d|H_1^\delta) = \int_{\delta \neq 0} f(d|\delta) \cdot p(\delta) d\delta.$$

Un factor Bayes B_{01} grande indica evidencia a favor de H_0 (y pequeño, a favor de la alternativa). Fijado el nivel de confianza ($B > 1$) que exigimos para concluir sobre el contraste planteado, éste se resolverá a dicho nivel cuando:

$$\begin{aligned} \text{se acepte } H_0 &\Leftrightarrow B_{01} > B \\ &\text{ó} \\ \text{se acepte } H_1 &\Leftrightarrow B_{10} > B \Leftrightarrow B_{01} \leq 1/B, \end{aligned}$$

siendo B_{10} el índice de evidencia a favor de la hipótesis alternativa.

Es decir, el **éxito** vendrá definido por obtener con la replicación:

$$(11) \quad B_{01}^{\delta} > B \quad \text{ó} \quad B_{01}^{\delta} \leq 1/B.$$

Al depender B_{01}^{δ} de X_r , para un n_r fijo la probabilidad de éxito se obtendrá integrando la distribución predictiva del resultado de la réplica sobre el rango de valores de X_r que satisfacen la condición [11], esto es, sobre la región:

$$R_B^{\delta}(X_r, n_r) = \{x_r : B_{01}^{\delta}(x_r, n_r) > B \text{ ó } B_{01}^{\delta}(x_r, n_r) \leq 1/B\}.$$

Obviamente, cuanto mayor sea el nivel de evidencia que exijamos para concluir a favor de una de las hipótesis, más reducida será la región de éxito y por tanto menor la probabilidad de resolver el contraste.

La planificación «óptima» de la réplica resultará del mínimo n_r para el que se obtengan las garantías de éxito (M) exigidas por el experimentador:

$$Prob(\text{éxito}) = \int_{R_B^{\delta}(x_r, n_r)} m(x_r | x_o) dx_r \geq M.$$

Al pretender concluir sobre un contraste con hipótesis nula **puntual**, es de esperar que en su resolución se aprecie una sensibilidad importante respecto de la información a priori utilizada. En efecto, si asumimos a priori una discrepancia «leve» entre las poblaciones original y réplica ($\tau^2 = 1/5$), serán precisos muchos datos (o muy «raros») para conseguir evidencia suficiente en contra de H_0 . La probabilidad de éxito en este caso se espera pequeña, como de hecho resulta (inferior a 0.7, aun con 300 datos). Con $\tau^2 = 1$, caso en el que ya a priori asumimos cierta discrepancia entre la población inicial y la replicada, en principio se espera poder discriminar mejor ambas poblaciones, al tener de partida más peso para la hipótesis «más grande» (la alternativa a la puntual). No son necesarios entonces muchos datos para lograr concluir sobre el contraste, aun en el caso más restrictivo con $B=4$ (ver Tabla 1). Cuando sólo se exige (para concluir sobre el contraste) el doble de evidencia a favor de alguna de las hipótesis ($B=2$), la probabilidad de éxito resulta mayor que 0.85 tan sólo duplicando el tamaño original en la réplica ($n_r = 30$).

Tabla 1. Comparación de Poblaciones. $Prob(B_{01}^\delta < 1/B \text{ ó } B_{01}^\delta > B)$

ÉXITO: Resolver $\theta_r = \theta_o$ versus $\theta_r \neq \theta_o$					
n_r	15	30	100	300	
$\tau^2 = 1/5$	0.3364	0.5735	0.6613	0.6846	
B=2					
$\tau^2 = 1$	0.8374	0.8645	0.8851	0.8914	
B=4					
$\tau^2 = 1$	0.6366	0.7633	0.8129	0.8262	

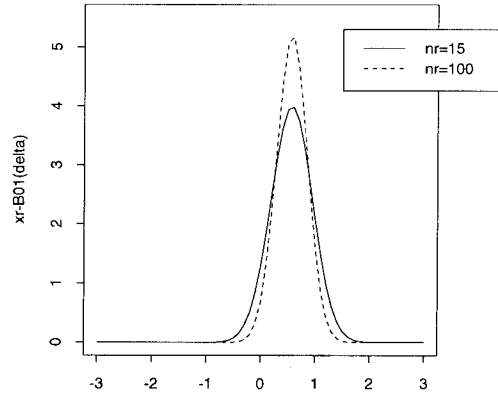
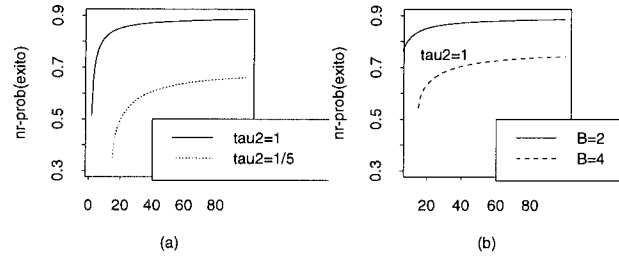


Figura 3. Comparación de Poblaciones. Éxito $\equiv$ resolver $H_0^\delta : \theta_r = \theta_o$ con el factor Bayes B_{01}^δ .

(3a) Éxito $\equiv B_{01}^\delta < 1/2$ ó $B_{01}^\delta > 2$. Probabilidad de éxito versus n_r para estudios próximos y alejados.

(3b) Éxito $\equiv B_{01}^\delta < 1/B$ ó $B_{01}^\delta > B$. Estudios alejados. Probabilidad de éxito versus n_r para B=2 y B=4.

(3c) Estudios alejados. $B_{01}^\delta(x_r)$ versus x_r para $n_r = 15$ y $n_r = 100$.

Dadas las restricciones de nuestro modelo y los valores asumidos a priori para los parámetros, la región de éxito, dependiente de n_r y B , resulta en ocasiones «imposible». Cuando τ^2 es pequeño, valores altos de B y bajos de n_r dan lugar a regiones imposibles (dicha condición de «posibilidad» está detallada en el apéndice 1). Es el caso de $\tau^2 = 1/5$ y $B = 4$.

En la Figura 3a está representada la probabilidad de éxito en función de n_r para ambas situaciones (considerablemente menor en $\tau^2 = 1/5$ que en $\tau^2 = 1$), definiendo la región de éxito con $B = 2$. En la Figura 3b apreciamos las diferencias en la amplitud de dicha región cuando tomamos $B = 2$ y $B = 4$ para resolver el contraste.

En la Figura 3c hemos representado la magnitud del factor Bayes en función del resultado a observar en la réplica, para dos tamaños distintos ($n_r = 15$ y $n_r = 100$). El rango de valores de X_r para los que se consigue evidencia a favor de H'_0 no tiene amplitud mayor a 2 ya en el caso más desfavorable (en el que la réplica cuenta con menos observaciones — $n_r = 15$ —), disminuyendo a medida que aumenta n_r . Más datos proporcionarán inferencias más precisas y permitirán discriminar mejor entre las hipótesis planteadas (un x_r próximo a x_o otorgará más evidencia a H'_0 cuando $n_r = 100$ que cuando $n_r = 15$).

4. DETECCIÓN DE SESGO CON LA RÉPLICA

Es frecuente encontrar primeros análisis llevados a cabo con un diseño deficiente, a veces incluso inexistente, que puede haber dado lugar a la introducción de sesgo en el efecto percibido en la muestra. La replicación permite repetir la experiencia controlando factores y/o variables que podrían haber quedado «suelos» en dicho primer estudio, reduciendo o eliminando así el sesgo. El éxito en este tipo de réplicas estará basado en la inferencia sobre el sesgo.

Una forma posible de modelizar el sesgo β en el experimento original (y su ausencia en ϵ_r) consiste en suponer para las medias poblacionales la relación:

$$\begin{aligned}\theta_o &= \mu + \beta \\ \theta_r &= \mu.\end{aligned}$$

El modelo jerárquico M3 incorpora la información a priori sobre el sesgo β a la modelización de las restantes variables.

4.1. Modelo M3

I. DATOS

Con muestras condicionalmente independientes, suponemos que la distribución de los estadísticos suficientes es:

$$(12) \quad f(x_o|\mu, \beta) = N(x_o|\mu + \beta, \gamma_o^2),$$

$$(13) \quad f(x_r|\mu) = N(x_r|\mu, \gamma_r^2),$$

con $\gamma_i^2 = \sigma^2/n_i$ ($i = o, r$), y σ^2 conocida ($=1$).

II. PARÁMETROS

Modelizamos con una distribución mínimo-informativa la información a priori sobre la media μ . Para el sesgo de ϵ_o asumimos una distribución inicial normal centrada en cero y con varianza conocida Δ^2 , mayor o menor en función de las sospechas iniciales sobre su existencia:

$$(14) \quad p(\mu) \propto \text{cte},$$

$$(15) \quad p(\beta) = N(\beta|0, \Delta^2).$$

En nuestra inferencia hemos considerado dos valores para Δ^2 relativamente distintos:

$\Delta^2 = 1$: Suponemos que el sesgo, en caso de existir, no puede ser demasiado grande.

$\Delta^2 = 5$: La información a priori sobre la existencia de sesgo en ϵ_o es débil; estaríamos dando probabilidad apreciable a magnitudes grandes para el sesgo.

El análisis de sensibilidad se obtendrá, en esta primera aproximación, de la comparación entre los resultados para estos dos valores. Con posterioridad (modelo M3A) se supondrá Δ^2 aleatorio.

Para calcular la probabilidad de éxito en la réplica utilizaremos, como ya hicimos en los otros tipos de réplicas, la distribución predictiva para su resultado, obtenida de la combinación de información:

$$(16) \quad \begin{aligned} m(x_r|x_o) &= \int f(x_r|\mu) \cdot p(\mu|x_o) d\mu \\ &= N(x_r|x_o, \gamma_o^2 + \gamma_r^2 + \Delta^2), \end{aligned}$$

donde

$$p(\mu|x_o) = \int p(\mu|\beta, x_o) \cdot p(\beta) d\beta,$$

dado que, en nuestro caso,

$$p(\beta|x_o) \propto f(x_o|\beta) \cdot p(\beta) \propto p(\beta).$$

4.2. Inferir sobre el sesgo

Proponemos la inferencia sobre el contraste de existencia de sesgo en el primer estudio:

$$\begin{cases} H_0^\beta & : \beta = 0 \\ H_1^\beta & : \beta \neq 0, \end{cases}$$

utilizando de nuevo el factor Bayes. Concluiremos sobre dicho contraste cuando obtengamos un valor para él lo suficientemente grande (a favor de H_0^β), o lo suficientemente pequeño (a favor de H_1^β).

Siendo:

$$(17) \quad B_{01}^\beta = \frac{N(x_r|x_o, \gamma_o^2 + \gamma_r^2)}{N(x_r|x_o, \gamma_o^2 + \gamma_r^2 + \Delta^2)},$$

la región de éxito vendrá dada (para un n_r fijo) por:

$$R_B^\beta(X_r, n_r) = \{x_r : B_{01}^\beta(x_r, n_r) > B \text{ ó } B_{01}^\beta(x_r, n_r) \leq 1/B\},$$

donde B representa el grado de confianza exigido para resolver el contraste.

Tabla 2. Inferencia sobre el Sesgo. $Prob(B_{01}^\beta \leq 1/B \text{ ó } B_{01}^\beta \geq B)$

ÉXITO: Resolver $\beta = 0$ versus $\beta \neq 0$.					
	n_r	15	30	100	300
B=2	$\Delta^2 = 1$	0.7417	0.7880	0.8220	0.8322
	$\Delta^2 = 5$	0.9081	0.9228	0.9342	0.9377
B=3	$\Delta^2 = 1$	#	0.6105	0.6890	0.7098
	$\Delta^2 = 5$	0.8490	0.8740	0.8931	0.8989
B=4	$\Delta^2 = 1$	#	#	#	#
	$\Delta^2 = 5$	0.8006	0.8352	0.8611	0.8689

La probabilidad de éxito para cada tamaño n_r se calcula integrando la predictiva $m(x_r|x_o)$ sobre dicha región, y de la condición [1] se obtendría el *mejor* tamaño con que replicar.

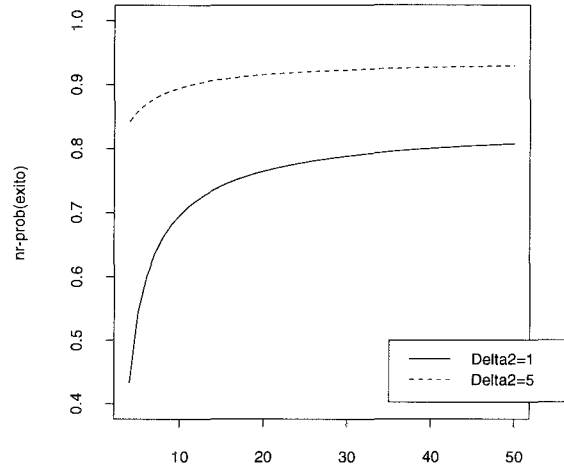


Figura 4. Detección de Sesgo. Éxito $\equiv$ resolver $H_0^\beta : \beta = 0$ con $B_{01}^\beta < 1/2$ ó $B_{01}^\beta > 2$. Probabilidad de éxito versus n_r , asumiendo distinta información inicial sobre el sesgo de ϵ_o .

En la Figura 4 está representada la probabilidad de éxito en función de n_r cuando el investigador está dispuesto a aceptar aquella hipótesis para la que obtenga el doble de evidencia ($B=2$). La diferencia entre las probabilidades de éxito con $\Delta^2 = 1$ y $\Delta^2 = 5$ es importante, si bien resultaba esperable (puesto que con $\Delta^2 = 5$ estamos contemplando la posibilidad de mayores sesgos, es razonable que esperemos poder detectarlo más fácilmente). Este efecto queda reflejado en la Tabla 2. Con $n_r = 30$ y $\Delta^2 = 1$, si bien para $B = 2$ la probabilidad de éxito está próxima a 0.8, para $B=3$ se queda sólo en 0.6. Para $\Delta^2 = 5$ la probabilidad de resolver el contraste con cualquiera de los tres tamaños es muy alta (con $B=4$, superior a 0.8). Valores pequeños de Δ^2 , niveles de confianza (B) altos y tamaños de la réplica (n_r) bajos dan lugar a regiones «imposibles» (ver apéndice 1). En la Tabla 2 el símbolo # designa incompatibilidades de ese tipo.

5. REPRODUCCIÓN Y COMPARACIÓN: τ^2 DESCONOCIDA

Como una extensión del modelo M12 que proponíamos para analizar la réplica cuando se pretendía reproducir las conclusiones de un estudio previo (sec. 3.2) o comparar la población muestreada en la réplica con la de cierto experimento original (sec. 3.3), modelizamos las relaciones entre los parámetros asumiendo un mayor desconocimiento inicial.

Conservando la estructura jerárquica de M12, las medias poblacionales θ_o y θ_r se asumen intercambiables, normales, con media μ y varianza τ^2 , ambas desconocidas. En el tercer nivel de dicho modelo será entonces necesario precisar la información inicial sobre el nuevo parámetro introducido (τ^2). Al tratarse de una medida de la distancia entre las poblaciones de ε_o y ε_r , se asume (es lo usual) una distribución a priori gamma invertida:

$$(18) \quad p(\tau^2) = \text{Gal}(\tau^2 | a + 1, ak).$$

Esta modelización es equivalente a asumir que, a priori y condicionando a μ , las medias poblacionales θ_o y θ_r son intercambiables (en lugar de independientes como en M12), con distribuciones marginales que son *t-Student* (las colas son algo mayores que las del modelo normal de M12):

$$p(\theta_i | \mu) = St\left(\theta_i | \mu, \frac{ak}{a-1}, 2(a+1)\right), \quad i = o, r.$$

Así k , el valor esperado para la variabilidad a priori de las medias, será un indicativo de la discrepancia asumida inicialmente entre la población de ε_o y la de la réplica. Al igual que en M12, consideraremos en nuestro ejemplo los valores $k = 1/5$ para una discrepancia a priori pequeña entre ambas poblaciones (S1), y $k = 1$ para cuando asumamos cierto «alejamiento» (S2).

El parámetro « a » da idea sobre la certidumbre de nuestra información inicial acerca de la distribución de las medias θ_o y θ_r . Un valor de $a = 4$ proporciona unas colas para la distribución Student a priori no demasiado altas (10 grados de libertad), de forma que el alejamiento de la modelización normal de M12 no resulte excesivo.

La modificación introducida en el modelo da lugar a una complicación considerable en el cálculo de la distribución predictiva, que se obtiene de una integral no analítica:

$$(19) \quad m(x_r | x_o) = \int m(x_r | x_o, \tau^2) \cdot p(\tau^2) d\tau^2,$$

donde $m(x_r|x_o, \tau^2)$ es la predictiva [5] resultante en M12:

$$m(x_r|x_o, \tau^2) = N(x_r|x_o, \gamma_o^2 + \gamma_r^2 + 2\tau^2).$$

(La expresión [19] de la predictiva refleja el hecho de que $p(\tau^2|x_o) = p(\tau^2)$).

Problemas de cálculo surgen también para las restantes densidades a posteriori y factores Bayes con los que se lleva a cabo el análisis. Utilizamos la simulación, la integración numérica y la aproximación Monte-Carlo para obtener estimaciones con qué evaluar la probabilidad de éxito al replicar. Las densidades de interés son aproximadas por estimadores *kernel* (ver Silverman, 1986) que suavizan los histogramas confeccionados con las simulaciones obtenidas.

Dada la complicación del cálculo, el análisis de la replicación lo reducimos considerando sólo tres valores de n_r : 15 (igual al tamaño de ϵ_n), 30 (doble) y 100 (considerablemente mayor).

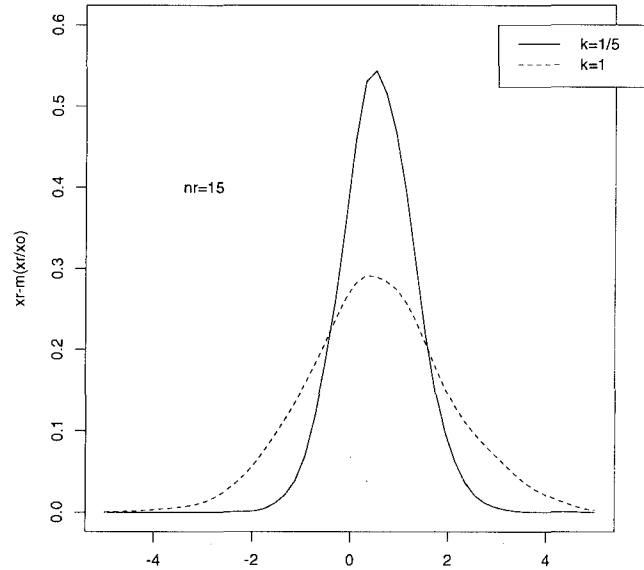


Figura 5. Predicción del resultado de la réplica con τ^2 desconocida. Estudios próximos: $k = 1/5$. Estudios alejados: $k = 1$.

En la Figura 5 hemos representado la aproximación que obtenemos de la densidad predictiva de X_r para cada uno de los tres tamaños muestrales de interés y las dos

posibilidades de k ($1/5$ y 1). Cuando la divergencia a priori entre las poblaciones es menor (en valor esperado, esto es, $k = 1/5$), la predictiva para X_r resulta más precisa (con menor varianza) y apuntada que en el caso en que inicialmente asumimos mayor discrepancia ($k = 1$). Dicha distribución aparece centrada alrededor de la observación del primer estudio (≈ 0.5). El apuntamiento y la precisión crecen con el tamaño de la réplica (n_r), como era de esperar.

Vayamos a cada uno de los supuestos que ya consideramos con M12.

5.1. Reproducción de conclusiones

El problema es idéntico al planteado en la sección 3.2, con sus dos tipos de éxito:

R1: $p_r \leq 0.05$,

R2: $P(\theta_r \leq 0 | x_o, x_r) \leq 0.05$.

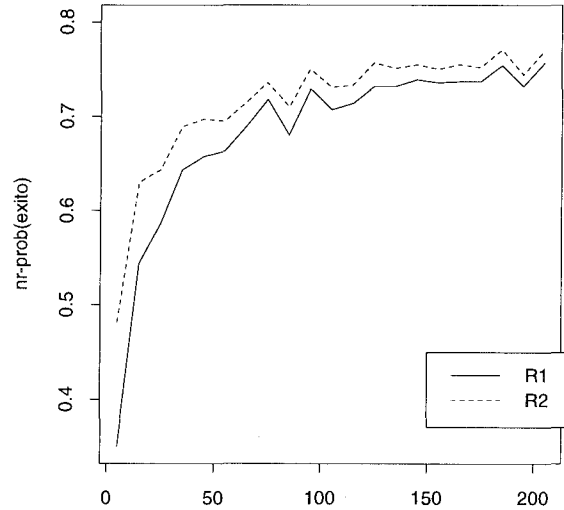


Figura 6. Reproducción de Conclusiones con τ^2 desconocida. Éxito $\equiv$ rechazo de H_0^r : $\theta_r \leq 0$. Estudios próximos. Probabilidad de éxito versus n_r . Rechazos R1 (con p-valor) y R2 (con probabilidad a posteriori para H_0^r).

Las dos formas de rechazo nos conducían a sendas regiones de éxito cuya probabilidad evaluábamos integrando respecto de la predictiva $m(x_r | x_o)$. Ahora, al no

disponer de su expresión analítica, estimamos las probabilidades de éxito mediante *conteos* a partir de simulaciones de la distribución predictiva. (Al tomar 2000 simulaciones, el error estándar del estimador resulta inferior a 0.01).

El comportamiento es similar al observado en el análisis de M12 (ver Figura 6). La probabilidad de rechazar combinando información (R2) resulta en todo momento, mayor que la de rechazar sólo con la réplica (R1), y la estabilización (por debajo de 0.8) es apreciable a partir de $n_r = 100$.

5.2. Comparación de poblaciones

Como en la sección 3.3, se pretende inferir sobre la diferencia de medias $\delta = \theta_r - \theta_o$, para lo que se proponen dos tipos de éxito:

- C1. conseguir una estimación lo suficientemente precisa;
- C2. resolver el contraste $H_0^\delta : \delta = 0$ con suficiente grado de confianza.

Asumimos para nuestro ejemplo $k = 1$ (poblaciones distantes «a priori»).

5.2.1. Precisión

Al no disponer de una expresión analítica para la distribución a posteriori de la diferencia de medias δ :

$$(20) \quad p(\delta|x_o, x_r) = \int p(\delta|x_o, x_r, \tau^2) \cdot p(\tau^2|x_o, x_r) d\tau^2,$$

donde

$$p(\tau^2|x_o, x_r) \propto m(x_r|x_o, \tau^2) \cdot p(\tau^2),$$

no tenemos una expresión «cerrada» para la varianza a posteriori de δ , $v_\delta^2(X_r, n_r)$, en función de la cual podíamos concluir sobre la precisión que proporcionaba la réplica para inferir sobre la diferencia de medias. Sin embargo, dado que dicha varianza es aleatoria (depende de la variable no observada X_r), fijando n_r obtenemos una muestra «aproximada» de su distribución predictiva (por estimación Monte-Carlo — ver apéndice 2—). Con la densidad predictiva de la varianza a posteriori se estima la probabilidad de éxito, definido este por:

$$v_\delta^2(X_r, n_r) \leq S^*,$$

evaluando el área encerrada a la izquierda de S^* , o simplemente como ya hacíamos con R1 y R2, utilizando conteos de las estimaciones de la muestra por debajo de S^* . La «mejor» planificación de la réplica se obtendrá de aquel tamaño n_r que proporcione suficientes garantías de éxito.

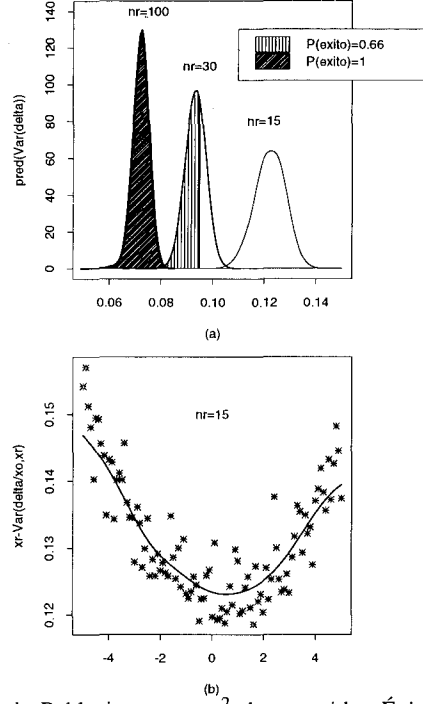


Figura 7. Comparación de Poblaciones con τ^2 desconocida. Éxito $\equiv$ Precisión al inferir sobre $\delta = \theta_r - \theta_o$. Estudios alejados.

(7a) Éxito: $\text{Var}(\delta|x_o, X_r) \leq 0.095$. Predicción de $\text{Var}(\delta|x_o, X_r)$ y probabilidad de éxito para $n_r = 15, 30$ y 100 .

(7b) Aproximación MC de $\text{Var}(\delta|x_o, x_r)$ versus x_r .

En la Figura 7a representamos conjuntamente las densidades predictivas de la varianza a posteriori de δ (aproximación kernel) para los tres tamaños de interés. La influencia del tamaño de la réplica es importante: las tres distribuciones representadas ($n_r = 15$, $n_r = 30$ y $n_r = 100$) no se solapan apenas. Cuanto mayor es n_r , mayor es el apuntamiento de esta distribución y más próxima a cero queda su media. Tomando $S^* = 0.095$, la probabilidad de éxito (área sombreada) planificando con $n_r = 15$ es nula; con $n_r = 30$ resulta ya de 0.66, y con $n_r = 100$ es máxima (todos los valores «probables» de la varianza a posteriori quedan por debajo de 0.095).

Comparamos en la Figura 7b la magnitud de la varianza $v_\delta^2(x_r, 15)$ (su aproximación) con el resultado de la réplica (x_r), ajustando una curva a la representación de los puntos. Esta varianza resulta mínima cuando el valor x_r de la réplica está próximo al observado en el primer estudio, creciendo a medida que éste es más «raro» respecto

de x_o . De lo observado con otros valores de n_r , cuanto mayor es el tamaño de la réplica mayor resulta la precisión con que se infiere sobre δ (menor es la magnitud de su varianza a posteriori), si bien el comportamiento respecto de x_r es similar.

5.2.2. Contraste de igualdad

Al utilizar el factor Bayes de la diferencia de medias B_{01}^δ para concluir sobre el contraste de igualdad de poblaciones, de nuevo nos encontramos con dificultades analíticas:

$$B_{01}^\delta(x_r, n_r) = \frac{N(x_r|x_o, \gamma_o^2 + \gamma_r^2)}{\int m(x_r|x_o, \tau^2) \cdot p(\tau^2) d\tau^2}.$$

Fijando el tamaño muestral en la réplica (n_r) obtenemos una muestra de la distribución predictiva para B_{01}^δ utilizando integración numérica. Puesto que el éxito lo definíamos [11] por resolver el contraste a favor de una u otra hipótesis a un nivel de confianza B, la probabilidad de éxito es estimada mediante conteos a partir de la muestra obtenida.

Representamos en la Figura 8a la distribución predictiva para B_{01}^δ considerando $n_r = 15$ y 100. El comportamiento es similar para ambos tamaños. La densidad predictiva del factor Bayes resulta bimodal, con la moda máxima alrededor de cero (evidencias en contra de H_0^δ) y la menor (evidencias a favor de la hipótesis puntual) desplazándose hacia la derecha a medida que crece el tamaño de la réplica. Con n_r aumenta la probabilidad de concluir **con mayor confianza** sobre el contraste, esto es, de discriminar mejor entre las hipótesis. En la gráfica sombreamos la región complementaria a la región de éxito, definido éste por conseguir para alguna de las hipótesis un grado de evidencia cuatro veces superior al conseguido para su alternativa. Con B=4, una réplica de 15 observaciones da una probabilidad muy baja a lograr concluir sobre el contraste; ya con 100 observaciones es posible discriminar entre las dos hipótesis del contraste. El efecto del tamaño muestral sobre la probabilidad de éxito es más apreciable cuando el grado de exigencia para resolver el contraste (B) es mayor (ver Tabla 3).

Tabla 3. Comparación de Poblaciones. $Prob(B_{01}^\delta \geq B \text{ ó } B_{01}^\delta \leq 1/B)$

ÉXITO: Resolver $\theta_r = \theta_o$ versus $\theta_r \neq \theta_o$			
n_r	15	30	100
B=2	0.797	0.8525	0.879
B=4	0.596	0.712	0.7935

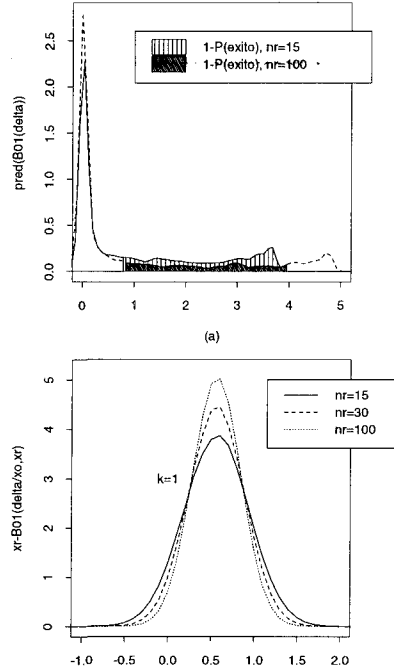


Figura 8. Comparación de Poblaciones con τ^2 desconocida. Éxito $\equiv$ resolver $H_0^\delta : \delta = 0$ con B_{01}^δ . Estudios alejados.

(8a) Éxito $\equiv B_{01}^\delta < 1/4$ ó $B_{01}^\delta > 4$. Predicción de B_{01}^δ y probabilidad de éxito para $n_r = 15$ y 100.

(8b) $B_{01}^\delta(x_r)$ versus x_r para $n_r = 15, 30$ y 100.

Hemos representado también la magnitud del factor Bayes para la diferencia de medias en función del valor observado para X_r (Figura 8b). Para valores de X_r próximos a x_0 (evidencias a favor de H_0^δ), el factor Bayes B_{01}^δ crece con n_r . Valores de X_r alejados de x_0 inducen mayores evidencias a favor de la alternativa y en contra de la igualdad de poblaciones (B_{01}^δ crece) cuando el tamaño de la réplica es más grande.

6. DETECCIÓN DE SESGO: δ^2 DESCONOCIDA

Buscando también robustecer nuestros resultados en el análisis de la réplica cuando el objetivo de ésta consiste en detectar sesgo en un estudio previo (o conseguir evidencia en contra de su existencia), añadimos a la modelización jerárquica M3, un nivel más.

6.1. Modelo M3A

En dicha modelización asumíamos para el sesgo β una distribución a priori normal, centrada en cero y con varianza Δ^2 conocida. Ahora consideramos esta varianza desconocida, por lo que surge un tercer nivel en la jerarquía destinado a la distribución a priori para Δ^2 (una gamma invertida):

$$(21) \quad p(\Delta^2) = \text{Gal}(\Delta^2 | b+1, bk).$$

Esto es equivalente a asumir como distribución inicial para el sesgo β , una distribución *t-Student*:

$$(22) \quad p(\beta) = \text{St}\left(\beta | 0, \frac{bk}{b+1}, 2(b+1)\right).$$

Siendo b el parámetro que da la amplitud de las colas en la distribución [22], tomamos $b = 1$ (que le da 4 g.l. $\equiv$ colas altas) asumiendo así muy poca certidumbre a priori sobre la existencia de sesgo en ϵ_o . Como valores esperados para la varianza inicial del sesgo consideramos, siguiendo M3, $k = 1$ y $k = 5$, que informan sobre su posible magnitud.

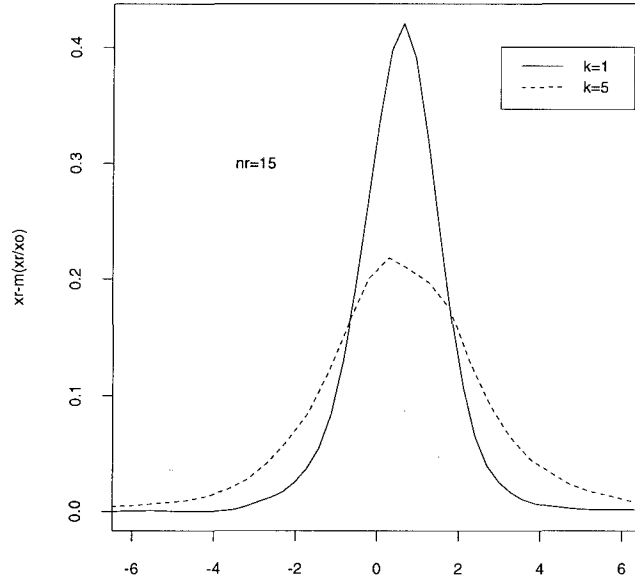


Figura 9. Predicción del resultado de la réplica asumiendo sesgo en ϵ_o y Δ^2 desconocida. Información precisa sobre el sesgo: $k = 1$. Información dispersa sobre β : $k = 5$.

La distribución predictiva que obtenemos ahora para llevar a cabo la planificación de la réplica resulta:

$$(23) \quad m(x_r|x_o) = \int m(x_r|x_o, \Delta^2) \cdot p(\Delta^2) d\Delta^2,$$

donde $m(x_r|x_o, \Delta^2)$ es la predictiva [16] obtenida en M3.

Al igual que en el modelo M12A, utilizamos la integración numérica para resolver las integrales no analíticas (caso del factor Bayes, que involucra integraciones unidimensionales) y la simulación y aproximación kernel para obtener las distribuciones de interés.

En la Figura 9 aproximamos la densidad predictiva de X_r con $n_r = 15$ para los dos valores de k , que afectan (como cabía esperar) al apuntamiento de la distribución: la variabilidad de la predicción de X_r resulta mayor cuanto menos sabemos a priori sobre el sesgo en ϵ_o (mayor k).

6.2. Inferir sobre el sesgo

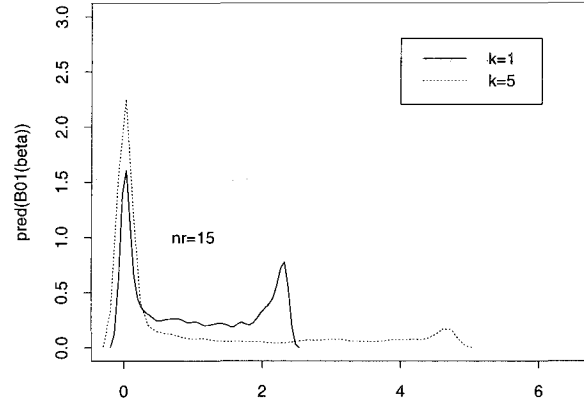
Análogamente a como hacíamos en la sección 4.2, pretendemos resolver el contraste sobre la existencia de sesgo en el primer estudio con el factor Bayes a favor de la hipótesis nula de sesgo cero:

$$B_{01}^B(x_r, n_r) = \frac{N(x_r|x_o, \gamma_o^2 + \gamma_r^2)}{m(x_r|x_o)}.$$

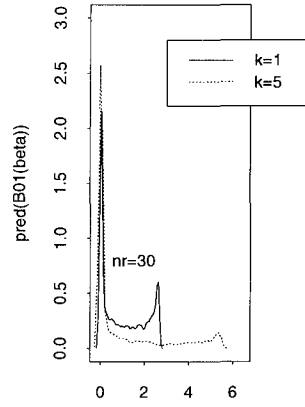
Fijado el índice de evidencia exigido para concluir, B , conseguiremos éxito cuando concluyamos a favor de alguna de las hipótesis:

$$\text{ÉXITO} \Leftrightarrow B_{01}^B(x_r, n_r) \geq B \text{ ó } B_{01}^B(x_r, n_r) \leq 1/B.$$

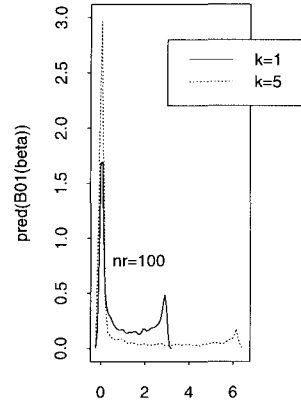
En la Figura 10 representamos las aproximaciones conseguidas para la densidad predictiva de $B_{01}^B(X_r)$ con $n_r = 15, 30$ y 100 . El comportamiento es siempre el mismo: la distribución de $B_{01}^B(X_r)$ es bimodal, más variable cuanto mayor es el desconocimiento inicial sobre el sesgo (k grande). A mayor k más se apunta la densidad hacia la izquierda (reconocimiento de sesgo) y se achata y desplaza hacia la derecha (negación de sesgo). Cuanto menos sepamos inicialmente sobre el sesgo y mayor hayamos de asumir por tanto su variabilidad inicial, al dar cabida (con probabilidad apreciable) a valores grandes para el sesgo, más probabilidad tendremos de resolver el contraste a un nivel de confianza (B) alto, ya sea a favor de una u otra alternativa. Si la información inicial es más precisa ($k = 1$) sobre un sesgo muy pequeño (en torno a 0), el contraste será más difícil de resolver para un buen nivel de confianza, como se observa en la gráfica.



(a)



(b)



(c)

Figura 10. Detección de sesgo con Δ^2 desconocida. Predicción de B_{01}^β con sesgo disperso ($k=5$) y preciso ($k=1$) a priori.

(10a) $n_r = 15$.

(10b) $n_r = 30$.

(10c) $n_r = 100$.

Según se aprecia en la Figura 11, para k fijo el efecto del tamaño n_r es siempre el mismo: con n_r aumenta la variabilidad del factor Bayes y la magnitud de su moda más oriental, es decir, aumenta la probabilidad de conseguir evidencias altas a favor de la hipótesis nula de inexistencia de sesgo.

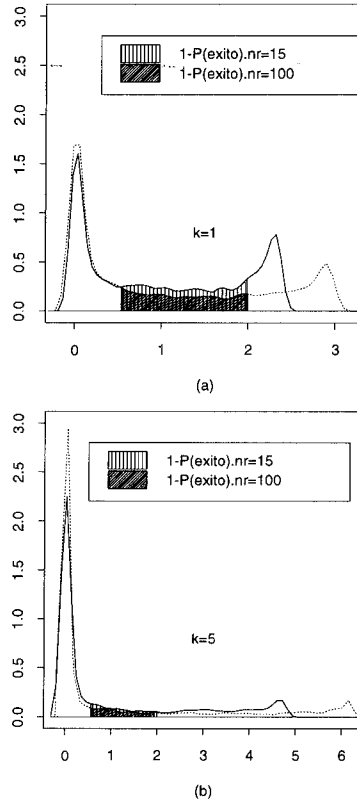


Figura 11. Detección de sesgo con Δ^2 desconocida. Éxito $\equiv$ resolver $H_0^\beta : \beta = 0$ con $B_{01}^\beta < 1/2$ ó $B_{01}^\beta > 2$. Predicción de B_{01}^β y probabilidad de éxito para $n_r = 15$ y $n_r = 100$.

(11a) Información a priori precisa sobre β .

(11b) Información a priori dispersa sobre β .

Tabla 4. Inferencia sobre Sesgo en ϵ_0 . $Prob(B_{01}^\beta \leq 1/B \text{ ó } B_{01}^\beta \geq B)$

ÉXITO: Resolver $\beta = 0$ versus $\beta \neq 0$.

15	30	100	n_r	15	30	100
0.6493	0.6953	0.7559	B=2	0.8804	0.8764	0.9119
0.3692	0.4072	0.4462	k=1 B=3 k=5	0.7984	0.8074	0.8599
0.3417	0.3812	0.4202	B=4	0.7133	0.7454	0.8094

Fijado el grado de confianza B con el que vaya a resolverse el contraste, el tamaño de la réplica afecta más a la probabilidad de éxito cuanto mayor es la precisión inicial (k menor), como puede verse en la Tabla 4.

7. CONCLUSIONES Y EXTENSIONES

En este artículo hemos tratado de mostrar que es posible sistematizar, estudiar y planificar la replicación de una experiencia estadística previa antes de llevarla a cabo, en base a los distintos objetivos que pueden justificarla. Como objetivos sólo hemos considerado tres (los que suponíamos más comunes), y algunos métodos inferenciales, pero confiamos en haber transmitido el mensaje de que el análisis de la réplica es viable en cualquier situación (un estudio mucho más exhaustivo puede encontrarse en Mayoral, 1995).

Los modelos jerárquicos bayesianos son una herramienta flexible para cuantificar las relaciones entre los experimentos originales y las réplicas. La distribución predictiva permite cuantificar la probabilidad de que la réplica tenga éxito y abordar con ella el diseño de la misma.

En este trabajo hemos utilizado modelos normales relativamente sencillos, con el fin de que los detalles técnicos no interfirieran con las ideas planteadas sobre el análisis de la réplica. Obviamente, en la práctica surgirán modelos mucho más complejos en los que habrá de recurrirse a otras técnicas de aproximación e incluso a otras soluciones estadísticas para predecir sobre la réplica.

Son muchas las posibilidades de la replicación en la resolución de problemas importantes, y muy diversos los objetivos con que un investigador puede justificar su práctica.

Si bien el punto de partida para la aplicación de nuestro análisis ha sido un supuesto de réplica ofrecido por Tversky y Kahnemann, son muy frecuentes las replications en Sociología, Psicología y Comunicación (véase una amplia recopilación de réplicas en estas disciplinas en Neuliep 1990). Asimismo, en Medicina, Agricultura, etc, son muy comunes las replications con fines diversos: probar una determinada droga/tratamiento sobre poblaciones distintas para comparar comportamientos, ratificar las conclusiones que sobre una droga/tratamiento algún investigador extrajo de un estudio experimental anterior, cuantificar el sesgo que las deficiencias experimentales de un estudio previo introdujeron en las conclusiones inferenciales, etc. La literatura científica está llena de ejemplos concretos de replications.

La utilización estadística eficaz de la réplica y la información contenida en estudios previos similares resultará esencial en la investigación científica. Ese es nuestro

objetivo de cara al futuro: explotar desde el punto de vista estadístico las posibilidades que ofrece la replicación de estudios en el desarrollo de la Ciencia.

APÉNDICE 1

Definido en los modelos M12 y M3 el factor Bayes para resolver un contraste con hipótesis nula puntual (igualdad de poblaciones o existencia de sesgo, respectivamente) por un cociente de densidades normales (verosimilitudes ponderadas),

$$B_{01} = \frac{N(d|0, h)}{N(d|0, h + \eta)},$$

donde:

$$\begin{aligned} * \eta &= \begin{cases} 2\tau^2 & \text{en M12} \\ \Delta^2 & \text{en M3,} \end{cases} \\ * d &= x_r - x_o, \\ * h &= \gamma_o^2 + \gamma_r^2, \end{aligned}$$

se llega a una simplificación de la forma:

$$B_{01} = \sqrt{1 + \frac{\eta}{h}} \cdot \exp \left\{ -\frac{d^2 \cdot \eta}{2h(h + \eta)} \right\}.$$

Así, al definir el éxito por conseguir un factor Bayes lo suficientemente grande ($B_{01} > B$) o lo suficientemente pequeño ($B_{01} < 1/B$), la región de éxito para resolver el contraste con un grado de confianza $B(> 1)$ vendrá dada por el conjunto de valores de X_r que satisfacen cualquiera de las dos condiciones:

$$(24) \quad (x_r - x_o)^2 \leq -2 \frac{h(h + \eta)}{\eta} \cdot \text{Log} \left\{ B \left(1 + \frac{\eta}{h} \right)^{-1/2} \right\}$$

$$(25) \quad (x_r - x_o)^2 \geq 2 \frac{h(h + \eta)}{\eta} \cdot \text{Log} \left\{ B \left(1 + \frac{\eta}{h} \right)^{1/2} \right\}.$$

Teniendo en cuenta los signos de cada uno de los términos de las desigualdades, para que dicha región resulte «**posible**» (una región coherente), será preciso que se verifique la condición:

$$B < \sqrt{1 + \frac{\eta}{h}}.$$

Los valores iniciales para las varianzas muestrales (γ_o^2 y γ_r^2) impondrán las restricciones sobre la coherencia de las regiones de éxito definidas en función del «grado de confianza» (B) que haya fijado el experimentador para resolver el contraste.

En nuestro ejemplo, las restricciones nos llevan a que para valores grandes de B , las regiones a que dan lugar valores pequeños de η y de n_r no resultan coherentes.

APÉNDICE 2

Fijado el tamaño para replicar, n_r , a partir de las simulaciones $\{x_r^i\}_{i=1}^N$ de la distribución predictiva para el resultado de la réplica, se pretende generar una muestra (de idéntico tamaño) de simulaciones de la distribución predictiva para la varianza a posteriori $v_\delta^2(x_r, n_r)$ de la diferencia de medias $\delta = \theta_r - \theta_o$.

Dado que la distribución a posteriori para δ se obtiene de:

$$(26) \quad p(\delta|x_o, x_r) = \int p(\delta|x_o, x_r, \tau^2) \cdot p(\tau^2|x_o, x_r) d\tau^2,$$

fijados x_r y n_r aproximamos la varianza a posteriori utilizando el método de Monte-Carlo:

A partir de simulaciones $\{\delta_j\}_{j=1}^M$ de la distribución a posteriori $p(\delta|x_o, x_r)$, estimamos la varianza $v_\delta^2(x_r, n_r)$ mediante:

$$\hat{v}_\delta^2(x_r, n_r) = \frac{1}{M-1} \sum_{j=1}^M (\delta_j - \bar{\delta})^2,$$

donde $\bar{\delta} = \frac{1}{M} \sum_{j=1}^M \delta_j$ es la media de las simulaciones.

Sustituyendo x_r por cada uno de los valores de las simulaciones $\{x_r^i\}_{i=1}^N$, obtendremos una muestra (aproximada) de la distribución predictiva para la varianza a posteriori de la diferencia de medias.

Las simulaciones $\{\delta_j\}_{j=1}^M$ según $p(\delta|x_o, x_r)$ se generan utilizando el método de composición a partir de simulaciones $\{\tau_j^2\}_{j=1}^M$ de la distribución a posteriori $p(\tau^2|x_o, x_r)$. Dichas simulaciones τ_j^2 se obtienen por el método de aceptación-rechazo, dada la forma de su densidad:

$$p(\tau^2|x_o, x_r^i) \propto m(x_r^i|x_o, \tau^2) \cdot p(\tau^2).$$

Referencias en Tanner, 1992.

Nota. El número de simulaciones que hemos utilizado en todo momento es $N = M = 2000$.

Nota. La implementación y representación se han llevado a cabo utilizando programación en C, el paquete de utilidades Mathematica y S-Plus.

REFERENCIAS

- [1] **Mayoral, A.M.** (1995). *La replicación de experimentos. Valoración del éxito*. Tesis de Licenciatura. Universitat de València.
- [2] **Neuliep, J.W.** (Ed.) (1990). *Handbook of Replication Research in the Behavioral and Social Sciences*. Select Press.
- [3] **Silverman, B.** (1986). *Density Estimation for Statistics and Data Analysis*. Chapman and Hall, London.
- [4] **Tanner, M.A.** (1992). «Tools for Statistical Inference. Observed Data and Data Augmentation Methods». *Lecture Notes in Statistics*, **67**. Eds. J. Berger *et al.* Springer Verlag, Berlin.
- [5] **Utts, Jessica** (1991) «Replication and Meta-Analysis in Parapsychology». *Statistical Science* 1991, **6**, **4**, 363–403.

ENGLISH SUMMARY

BAYESIAN DESIGN AND ANALYSIS OF REPLICATIONS IN SCIENTIFIC EXPERIMENTATION

M.J. BAYARRI*

M.A. MARTÍNEZ*

Universitat de València

Replication of experiments is a widespread activity in practice. The motivations for replicating are diverse; usually experimenters wish to verify some previous conclusions, but other motivations are also common, as achieving more precision, validating an assumed model, bias detection, studying the change on inferences when some conditions of the population under study changes, etc. Since, previous to the replication, the results of the original experiment are indeed available, it seems sensible to look for an optimal design of the replication based on all relevant information. We provide the minimal sample size such that there is high probability for the experimenter to be successful in her/his replication. We use Bayesian hierarchical models to quantify the relationship between the original and replicated experiments. The probability of success, key for the design of the replication, is derived from appropriate (Bayesian) predictive distributions.

Keywords: Replicación, modelo jerárquico bayesiano, factor Bayes, simulación, aproximación Monte-Carlo.

*This work has been partially supported by the Consellería de Cultura, Educació i Ciència de la Generalitat Valenciana, research grant GV-1081/93.

*M.J. Bayarri y M.A. Martínez. Dpto. Estadística e I.O. Universitat de València. Dr. Moliner, 50. 46100 BURJASSOT. València.

–Received decembre 1995.

–Accepted june 1996.

Replication of experiments is a most usual practice in scientific experimentation. However, and in spite of its widespread use, there does not seem to exist a systematization of almost anything related to it. Important questions that usually are not systematically addressed nor answered include, among others, the following: what goals are meant to be achieved? What does a successful replication mean? Should an isolated or a meta-analysis be performed? how should the replications be designed? How should the relations among the original and replicated experiments be incorporated?

The problem has been recognized in some areas and an entire meeting was devoted to it (Neuliep, 1990). However, nor even the exact meaning of a successful replication is free from controversies. In this paper, we systematize some of the goals that seem present when replicating. We also quantify what a «successful replication» might mean, and, based on the information provided by the original experiment, we quantify the probability that the replication will be a success. This is a most natural tool for design. In particular, we provide the needed sample size in the replication so that the probability of success is high enough.

Tversky and Kahnemann (as reported in Utts, 1990) performed an empirical study in which a number of scientists were asked about what a surprising t -value would be. The result was somehow paradoxical in that what was perceived as a failure in replicating the significant result could indeed be viewed as providing a stronger evidence against the null. We use simplified version of this experience to exemplify all of the calculations in the paper.

The paper is organized in six sections, of which Section 1 is devoted to an Introduction. In Section 2 we state the problem and the scenarios we shall be treating in the rest of the paper. In particular, we assume that it is wished to replicate an original experiment (ϵ_o) in which data X_o have been observed. The parameter of interest in ϵ_o is denoted by θ_o . The replicated experiment, data, and parameter are denoted by ϵ_r , X_r and θ_r respectively. We only consider in the paper three basic scenarios:

- S1: The main goal of ϵ_r is to reproduce the findings in ϵ_o . A most usual setup for ϵ_o is that of a hypothesis testing about θ_o and we assume that a similar testing will be performed in ϵ_r .
- S2: The main goal of ϵ_r is to study the similarities or dissimilarities of the replicated population when compared with the original one. A useful is that of estimation of θ_o in ϵ_o the purpose of ϵ_r being to infer about the discrepancy $\delta = \theta_o - \theta_r$ between the population parameters.
- S3: Some bias is suspected to have incurred in the original experiment. A replication is carried out with the main goal of detecting such bias.

Many different goals for replicating as well as methods for the statistical analysis of both experiments other than those treated here can be found in Mayoral (1995).

The flexibility of Bayesian hierarchical models allows for easy modeling of the relationships (exchangeability, dependency on covariates, temporal or spatial relations, etc.) among the original and replicated populations that should be taken into account when replicating. We use a three levels hierarchical model. The first level models the conditional distributions of the observable variables. The second, models the relationship among the parameters indexing both, original and replicated, models. The third level gives the prior distribution of the unspecified hyperparameters on the second and first levels. Bayesian machinery is used to develop the needed posterior and predictive distributions.

Depending on the goal of replication, we shall have different meanings for what a successful replication is. All of them, however, can be expressed by identifying (different) subsets of the sample space of the replication such that if X_r lies on such a subset the experimenter would consider to have succeeded when replicating. The probability of success (in its different meanings) can thus be easily computed from the completely specified predictive distribution for X_r , and a sample size for the replication can then be chosen such that this probability is high enough.

Section 3 is entitled «Reproduction and comparisons» and is devoted to analyzing a simple model to deal with scenarios S1 and S2 above. Specifically, X_o and X_r are assumed to be conditionally independent normals with means θ_o and θ_r respectively. These are taken to be exchangeable, which is modeled as a normal with unknown mean (which is given a non-informative, flat prior on the third level) and known variance, which is then given different values to reflect different opinions about the degree of similarity of both population means. Optimal sample sizes are derived, under several conditions, for scenario S1, both, when the replicated data is going to be analyzed by itself, and when it is going to be combined with the results of the original experiment. When we are in scenario S2, we might wish to compare both populations by simply inferring about δ , the difference of means, in which case the replications would be considered successful if the inferences are precise enough. On the other hand, the goal might actually be to directly test the equality of both means, in which case «success» would mean to solve the testing (in one or another direction) with enough confidence. We compute optimal sample sizes for the different goals, scenarios and prior opinions.

Section 4 is devoted to replication for bias detection. A different model is used in this Section, in which an unknown bias, β , is added to the mean, μ , of the original experiment, while the replication is assumed free from bias. μ is then given an non-informative prior while the bias is assumed normal with mean 0 and different variances, which reflect the possible sizes of the suspected bias. A most usual statistical analysis for bias detection consists in testing for 0 bias, and we provide

optimal sample sizes for the replication such that, with high probability, the testing is expected to be resolved with a small enough Bayes factor against one of the two hypotheses.

Section 5 is basically a robustification of the analysis in Section 3. Here, the variance of the population means is assumed unknown and is given an inverse gamma distribution in the third stage. This basically amounts to assessing a student t distribution to the means instead of the normal one used in Section 3, thus making the model more robust and more sensitive to data. This flexibility does not come without a price, and the resulting model is notably more cumbersome to analyze than the one in Section 3. The analyses in Section 3 are carried out resorting to Monte-Carlo methods to perform most integrations.

Section 6 generalizes Section 4 by allowing the variance of the bias to be unknown and then be given an inverse gamma distribution on the third stage. Again this is equivalent to assess a t distribution for the bias, which is most appropriate. Optimal sample sizes for the replication are provided assuming different prior distributions.

ANÁLISIS FACTORIAL DE TABLAS MIXTAS: NUEVAS EQUIVALENCIAS ENTRE ACP NORMADO Y ACM

M. ISABEL LANDALUCE CALVO*

Universidad del País Vasco

En este trabajo se pone de manifiesto que es posible el Análisis Factorial de tablas mixtas sin modificar la naturaleza de ninguno de los dos conjuntos, cualitativo y cuantitativo, que las integran. Se propone codificar de manera apropiada las indicadoras de cada variable cualitativa tratando de respetar, en la medida de lo posible, la estructura inicial de esta última y posteriormente aplicar un Análisis en Componentes Principales (ACP) Normado al conjunto de variables. Los factores obtenidos para el grupo de variables nominales serán iguales a los factores resultantes de un Análisis de Correspondencias Múltiples (ACM) de la Tabla Disyuntiva Completa (TDC).

Factorial Analysis of Mixed Tables: New Equivalences between Weighted PCA and MCA

Keywords: Análisis en Componentes Principales Ponderado, Análisis de Correspondencias Múltiples, Tablas mixtas, Variables indicadoras, Ponderación

Clasificación AMS: 62-07, 62H25

*M. Isabel Landaluce Calvo. Departamento de Economía Aplicada III. (Econometría y Estadística) Facultad de CC.EE. y Empresariales. Universidad del País Vasco. Avda. Lehendakari Aguirre, 83. 48015 BILBAO. e-mail:il@alcib.bs.ehu.es

– Article rebut l'abril de 1996.

– Acceptat el gener de 1997.

1. INTRODUCCIÓN

En el campo del Análisis Factorial no es infrecuente encontrarse con tablas de datos que recogen variables cuantitativas y cualitativas conjuntamente, esto es, tablas mixtas. La práctica habitual cuando se analizan este tipo de tablas transforma unas u otras para conseguir un conjunto de variables homogéneas. La técnica más corriente consiste en codificar las variables cuantitativas: se divide el intervalo de R en el que dichas variables toman valores en subintervalos, convirtiéndolas en variables nominales y se aplica al conjunto resultante un Análisis Factorial de Correspondencias. Este método, no obstante, plantea, por una parte, problemas de codificación, elección de la partición, etc, y, por otra parte, conlleva una pérdida de información. Pero, sobre todo pone de manifiesto el hecho de que la variable numérica desaparece para ser reemplazada por el conjunto de todas las funciones correspondientes a la variable codificada.

En este trabajo se propone un tratamiento de tablas mixtas a través de un Análisis en Componentes Principales en el que las variables nominales van a ser debidamente ponderadas.

2. ACP NORMADO Y ACM APLICADOS A VARIABLES CUALITATIVAS

Una variable cualitativa $\vec{x}_k$ se puede considerar como una partición del conjunto I de individuos. Esta variable está representada por el conjunto Q_k de las variables indicadoras de las clases de esta partición o por el subespacio de R^I que engendran. Este subespacio tiene como dimensión el número de modalidades, ya que las variables indicadoras correspondientes a una misma variable son ortogonales entre sí. Es, por tanto, el subespacio de funciones numéricas que toman el mismo valor para los individuos que han elegido la misma modalidad. Así consideradas, un conjunto de K variables cualitativas, con Q variables indicadoras en total, está codificado bajo la forma de una tabla disyuntiva completa, a la que tradicionalmente se aplica un Análisis de Correspondencias Múltiples.

Se puede comprobar que los resultados obtenidos al aplicar un ACM son equivalentes a los obtenidos a partir de un ACP normado de la tabla disyuntiva completa cuando se han ponderado las variables indicadoras a través de la proporción de individuos que no las han elegido. Esta equivalencia la demuestran Brigitte Escofier y Jérôme Pagès, (1990, Cap. 7), siguiendo los siguientes razonamientos:

1. En ACM como consecuencia de la transformación de las columnas, de la métrica en R^I (proporcional a la métrica identidad) y de los pesos de los elementos, las modalidades de las variables poseen las siguientes propiedades cuando se las considera respecto al origen:

- Las modalidades de una misma variable son ortogonales entre sí. La transformación en perfiles no cambia su dirección.
- Todas las modalidades, de peso $\frac{I_q}{IK}$, siendo I_q el número de individuos que han elegido la modalidad q y K el número de variables cualitativas, tienen la misma inercia respecto al origen:

Inercia de la modalidad q respecto al origen =

$$= \frac{I_q}{IK} \sum_i I \left(\frac{y_{iq}}{I_q} \right)^2 = \frac{1}{K}$$

Siendo y_{iq} el término general de la TDC.

2. Se construye una nube de variables indicadoras con las mismas propiedades inerciales que la nube de modalidades en ACM. Para su posterior tratamiento mediante ACP normado se las considera divididas por su desviación típica pero no centradas. Con este fin, se asigna a cada variable indicadora el peso $\frac{(I-I_q)}{I}$. La nube, así definida, posee las siguientes propiedades inerciales comparables a las mencionadas en el punto anterior:

- La métrica del espacio R^I es también la métrica identidad, excepto por el coeficiente $1/I$. La dirección de las variables indicadoras no se modifica por la división entre su desviación típica.
- Toda variable indicadora posee la misma inercia respecto al origen:

Inercia de la variable indicadora q respecto al origen =

$$= \frac{I-I_q}{I} \sum_i \frac{1}{I} \frac{y_{iq}^2}{I_q \frac{(I-I_q)}{I^2}} = 1$$

3. Existe equivalencia entre las operaciones de centrado del ACM y del ACP cuando se efectúan sobre las nubes definidas anteriormente.

En ACP, por una parte, el centrado de las variables se interpreta, en el espacio R^I , como una proyección de la nube de variables sobre el hiperplano ortogonal a la primera bisectriz.

En ACM, por otra parte, considerado como un AFC de la tabla disyuntiva completa, la nube de las variables indicadoras está centrada en otro sentido: el origen está situado en el centro de gravedad de la nube de modalidades N_Q , ($\sum_k Q_k = Q$). Esta nube, en ACM, presenta las siguientes propiedades:

- El centro de gravedad está situado sobre la primera bisectriz, se confunde con el perfil de la marginal sobre I , estando caracterizado por un perfil perfectamente plano.
- Está contenida en un hiperplano ortogonal a la primera bisectriz. Debido al carácter disyuntivo de la TDC, los vectores que unen el origen a las modalidades de una misma variable son ortogonales entre sí. El conjunto de las modalidades de las variables cualitativas engendran diferentes subespacios que, debido al carácter completo de la TDC, tienen una dirección común: la que une el origen con el centro de gravedad de la nube. Esta dirección se elimina con el centrado. En consecuencia, el centrado en ACM se interpreta como en ACP: una proyección sobre un hiperplano ortogonal a la primera bisectriz.

3. NUEVAS EQUIVALENCIAS ENTRE ACP NORMADO Y ACM

Un estudio profundo y detallado de la equivalencia entre ambos métodos bajo la ponderación señalada nos ha conducido a encontrar relaciones más concretas entre los elementos que intervienen en los análisis, relaciones que se resumen en los puntos siguientes.

3.1. Relación entre las matrices diagonalizadas

Las matrices analizadas en ACM y ACP normado de la TDC son equivalentes, excepto por un coeficiente, el inverso del número de variables cualitativas analizadas.

1. La matriz que se analiza en ACP normado de las variables indicadoras es la matriz de correlación entre las mismas, R . Siendo q_1 y q_2 dos variables indicadoras correspondientes a la misma variable cualitativa, I_{q_1} y I_{q_2} el número de individuos que han elegido la modalidad q_1 y la modalidad q_2 respectivamente, el coeficiente de correlación entre estas dos indicadoras se puede expresar de la siguiente manera:

$$r_{q_1, q_2} = -\sqrt{\frac{I_{q_1}}{(I - I_{q_1})}} \sqrt{\frac{I_{q_2}}{(I - I_{q_2})}}$$

Se comprueba que el coeficiente de correlación entre las variables indicadoras de una misma variable cualitativa va a tener siempre signo negativo.

El coeficiente de correlación entre dos variables indicadoras q y h , correspondientes a dos variables cualitativas, siendo y_{iq} y y_{ih} términos generales de la TDC, se puede expresar como sigue:

$$r_{q,h} = \frac{I \sum_{i=1}^I y_{iq} y_{ih} - I_q I_h}{\sqrt{I_q(I - I_q)} \sqrt{I_h(I - I_h)}}$$

Teniendo en cuenta que el término $\sum_{i=1}^I y_{iq} y_{ih}$ es el número de individuos que han elegido la modalidad q y la modalidad h al mismo tiempo, se comprueba que este coeficiente será positivo para el par de variables indicadoras que hayan sido elegidas mayoritariamente por los mismos individuos y negativo para las que hayan sido elegidas mayoritariamente por diferentes individuos.

2. La matriz que se diagonaliza en ACM es $M^{1/2} X^t D X M^{1/2}$ (al no ser $X^t D X M$ una matriz simétrica) (1990, Cap.4), siendo D y M las matrices diagonales de pesos de individuos y variables, respectivamente. Los términos generales de esta matriz se pueden expresar de la siguiente manera:

- El producto de una modalidad q por sí misma será:

$$\sum_{i=1}^I \left(\frac{I y_{iq}}{I_q} - 1 \right)^2 \frac{1}{I} \frac{I_q}{IK} = \frac{1}{K} \left(\frac{I - I_q}{I} \right)$$

- El producto de dos modalidades q_1 y q_2 de la misma variable será:

$$\begin{aligned} \sum_{i=1}^I \left(\frac{I y_{iq_1}}{I_{q_1}} - 1 \right) \left(\frac{I y_{iq_2}}{I_{q_2}} - 1 \right) \frac{1}{I} \sqrt{\frac{I_{q_1}}{IK}} \sqrt{\frac{I_{q_2}}{IK}} = \\ = - \frac{1}{K} \sqrt{\frac{I_{q_1}}{I}} \sqrt{\frac{I_{q_2}}{I}} \end{aligned}$$

- El producto de dos modalidades q y h pertenecientes a dos variables cualitativas será:

$$\begin{aligned} \sum_{i=1}^I \left(\frac{I y_{iq}}{I_q} - 1 \right) \left(\frac{I y_{ih}}{I_h} - 1 \right) \frac{1}{I} \sqrt{\frac{I_q}{IK}} \sqrt{\frac{I_h}{IK}} = \\ = \frac{1}{K} \frac{I \sum_{i=1}^I y_{iq} y_{ih} - I_q I_h}{I \sqrt{I_q} \sqrt{I_h}} \end{aligned}$$

3. Por otra parte, al igual que en el ACM, en ACP normado de las variables indicadoras al recibir éstas diferente ponderación, $\frac{I - I_q}{I}$, la matriz a diagonalizar, RM (considerando que todos los individuos poseen el mismo peso), no es simétrica. Se procede de igual manera que en aquel análisis y la matriz que se diagonaliza realmente es: $M^{1/2} R M^{1/2}$, esto es, cada elemento correspondiente de la matriz

R queda multiplicado por la raíz cuadrada de la ponderación asignada a cada una de las variables indicadoras que intervienen en su cálculo.

Se comprueba, por tanto, que las matrices diagonalizadas en los análisis contrastados son equivalentes, excepto por el coeficiente K , número de variables cualitativas.

3.2. Relación entre distancias

Las distancias definidas en ACP y en ACM aún siendo diferentes guardan una estrecha relación:

1. Con respecto a la distancia entre individuos, siendo i y l dos individuos cualesquiera, en ACP normado de la TDC ponderada se tiene:

$$\begin{aligned}
 d^2(i, l) &= \sum_{q=1}^Q m_q (y_{iq} - y_{lq})^2 = \\
 &= 0, \text{ cuando los individuos han elegido las mismas} \\
 &\quad \text{modalidades} \\
 &= \text{un valor que crece con el número de modalidades que} \\
 &\quad \text{difieren entre los individuos.}
 \end{aligned}$$

La presencia de una modalidad elegida por pocos individuos, modalidad rara, aleja a sus poseedores del resto de individuos.

En ACM, por otra parte, se tiene:

$$\begin{aligned}
 d^2(i, l) &= \sum_{q=1}^Q \frac{IK}{I_q} \left(\frac{y_{iq}}{K} - \frac{y_{lq}}{K} \right)^2 = \\
 &= \frac{I}{K} \sum_q \frac{1}{I_q} (y_{iq} - y_{lq})^2
 \end{aligned}$$

Como $(y_{iq} - y_{lq})^2$ vale 0 o 1, esta distancia crece con el número de modalidades que difieren entre los individuos. La presencia de una modalidad rara aleja a sus poseedores de los demás individuos.

2. Con respecto a la distancia entre dos modalidades, q y h , hay que recordar que en ACP se estudia la relación entre las variables indicadoras a través de su coeficiente de correlación. Por tanto, se reproduce aquí el resultado obtenido anteriormente:

$$r_{q,h} = \frac{I \sum_{i=1}^I y_{iq} y_{ih} - I_q I_h}{\sqrt{I_q(I - I_q)} \sqrt{I_h(I - I_h)}}$$

Se observa que la relación entre dos modalidades aumenta con el número de individuos que las han elegido a la vez, es decir, su distancia disminuye (hay que recordar que en ACP existe la siguiente equivalencia: $d^2(q, h) = 2(1 - r_{q,h})$).

En ACM, por otra parte, se tiene:

$$\begin{aligned} d^2(q, h) &= \sum_i I \left(\frac{y_{iq}}{I_q} - \frac{y_{ih}}{I_h} \right)^2 = \\ &= \frac{I}{I_q I_h} (I_q + I_h - 2 \sum_i y_{iq} y_{ih}) \end{aligned}$$

Se observa que esta distancia decrece con el número de individuos que han elegido las dos modalidades a la vez.

3.3. Otras relaciones

En este apartado quedan reflejadas otras relaciones existentes entre los dos análisis.

1. En ACM la inercia de una modalidad q respecto al centro de gravedad vale: $\frac{1}{K} (1 - \frac{I_q}{I})$. Al sumar las inercias de todas las modalidades se obtiene que la inercia total de la nube estudiada vale: $\frac{Q}{K} - 1$.

En ACP normado la inercia de una modalidad q respecto al centro de gravedad vale: $1 - \frac{I_q}{I}$. Al sumar las inercias de todas las modalidades se obtiene que la inercia total de la nube estudiada vale: $Q - K$.

Se alcanza, de nuevo, un resultado de gran interés: las inercias de los dos análisis son iguales, excepto por el coeficiente $\frac{1}{K}$.

2. En la representación simultánea obtenida a través de un ACP, un individuo i está situado próximo a las variables en las que posee mayoritariamente valores más altos que la media. Tal y como está definida la matriz X en el análisis

de las variables indicadoras ponderadas, un individuo estará próximo de las indicadoras que ha elegido.

En ACM un individuo i está situado, excepto por el coeficiente $1/\sqrt{\lambda}$, en el baricentro de las modalidades que ha elegido. Una modalidad q está situada, excepto por $1/\sqrt{\lambda}$, en el baricentro de los individuos que la poseen.

3. En ACM, por una parte, se sabe que las modalidades con pocos efectivos pueden contribuir mucho a la formación de los factores. En ACP, por otra parte, tal y como se ha definido la ponderación de las indicadoras, $\frac{I - I_q}{I}$, las modalidades raras tienen más peso que las modalidades con elevado número de efectivos.

En conclusión, se puede afirmar, por tanto, que un ACP normado de las variables indicadoras ponderadas, correspondientes a las modalidades de las variables cualitativas, conduce a los mismos factores sobre I que un ACM.

4. BIBLIOGRAFÍA

- [1] **B. Escofier** and **J. Pagès**. (1986). «Le Traitement des Variables Qualitatives et Tableaux Mixtes Par Analyse Factorielle Multiple». *Data Analysis and Informatics*, **IV(2)**, 179–191.
- [2] **B. Escofier** and **J. Pagès**. (1990). *Analyses Factorielles Simples et Multiples: Objectifs, Méthodes et Interprétation*. Dunod, Paris, 2ème edition.

ENGLISH SUMMARY

FACTORIAL ANALYSIS OF MIXED TABLES: NEW EQUIVALENCES BETWEEN WEIGHTED PCA AND MCA

M. ISABEL LANDALUCE CALVO*

Universidad del País Vasco

The Factorial Analysis of Mixed Tables, when the nature of the numerical and categorical variables remains unchanged, is possible because the Weighted Principal Components Analysis (WPCA) of qualitative indicator variables is equivalent to the Multiple Correspondence Analysis (MCA). The adequate weighting, i.e. the proportion of individuals that has not chosen the correspondent modality, allows us to analyze the Disjunctive Complete Table through a WPCA. In this way, we obtain the same factors as in a MCA.

Keywords: Weighted Principal Components Analysis, Multiple Correspondence Analysis, Mixed Tables, Indicator Variables, Weighting

AMS Classification: 62-07, 62H25

*M. Isabel Landaluce Calvo. Departamento de Economía Aplicada III. (Econometría y Estadística) Facultad de CC.EE. y Empresariales. Universidad del País Vasco. Avda. Lehendakari Aguirre, 83. 48015 BILBAO. e-mail:il@alcib.bs.ehu.es

–Received april 1996.

–Accepted january 1997.

The Factorial Analysis of Mixed Tables, when the nature of the numerical and categorical variables remains unchanged, is possible because the Weighted Principal Components Analysis (WPCA) of qualitative indicator variables is equivalent to the Multiple Correspondence Analysis (MCA). The adequate weighting, i.e. the proportion of individuals that has not chosen the correspondent modality, allows us to analyze the Disjunctive Complete Table through a WPCA. In this way, we obtain the same factors as in a MCA.

Escofier & Pagès (1990, Ch. 7) established this equivalence for a very general setting. In a detailed study of the relationship between these two methods, we have found new and more specific equivalences. These equivalences are as follows:

1. The analyzed matrices in MCA and WPCA of the Disjunctive Complete Table are equivalent, except for a constant given by the inverse of the number of qualitative variables. That is, with this exception, these analyses have the same inertia.
2. The distances between different elements are highly related in these two methods, in spite of not being the same.
3. In the simultaneous representation that a PCA provides an individual is projected close to those variables with values higher than the mean. In the analysis of the weighted indicator variables an individual is located next to the chosen modalities. In MCA an individual is located, except for a constant, on the baricentre of the chosen modalities. A modality is located, apart from the same constant, on the baricentre of the individuals that have chosen this modality.
4. We know that the modalities that have been rarely chosen in an MCA can have enough contribution to the formation of the factors. In addition, for PCA and due to the way the weights are defined, the indicators for the rare modalities have more weight than the modalities more frequently chosen.

As a conclusion, it is possible to carry out the Factorial Analysis of Mixed Tables without changing the nature of the numerical and categorical variables. It is not necessary to transform the numerical variables but we need to put same adequate weights to the indicators of the modalities associated with the categorical variables.

ESTIMADORES DE RAZÓN: UNA REVISIÓN

JOSÉ A. MAYOR GALLEGO*

Universidad de Sevilla

En este trabajo de revisión exponemos los principios básicos de la utilización de información adicional en la estimación de parámetros definidos sobre poblaciones finitas, mediante el empleo de estimadores de razón. Aunque la introducción de este tipo de estimadores se realizó inicialmente desde un punto de vista «heurístico», basado en la existencia de relaciones de proporcionalidad directa «aproximada» entre la variable de estudio y otras variables más controladas, su estudio detallado se desarrolla a partir del enfoque de modelos de superpoblación. La problemática principal que presentan este tipo de estimaciones es la existencia de sesgo, lo que obliga a utilizar diseños muestrales específicos o a modificar las expresiones de los estimadores con el fin de obtener estrategias insesgadas o con sesgo reducido, pero manteniendo la simplicidad en la estimación del error de muestreo.

Ratio estimators: a review

Keywords: Muestreo, poblaciones finitas, estimadores de razón.

Clasificación AMS: 62D05

*José A. Mayor Gallego. Dpto. Estadística e Investigación Operativa. Universidad de Sevilla. C/Tarfia s/n. 41012 SEVILLA.

—Article rebut el maig de 1996.

—Acceptat l'abril de 1997.

1. INTRODUCCIÓN

La utilización de información auxiliar es un recurso muy extendido en los diversos ámbitos del muestreo en poblaciones finitas, siendo su principal objetivo la obtención de estimaciones más acuradas.

En términos generales, podemos afirmar que la información auxiliar, generalmente suministrada por una o más **variables auxiliares**, conocidas o controladas al menos en cierto grado, puede ser aplicada en la fase de muestreo, en la fase de estimación o en ambas.

Así, los diseños PPS y PPS, es decir, con probabilidades de inclusión o de selección de elementos, proporcionales al tamaño, utilizan probabilidades de selección de elementos, que están afectadas por los valores de una variable auxiliar, X . También, en el muestreo estratificado, se emplean variables auxiliares en cuestiones tales como la afijación y la definición de estratos.

En este trabajo revisaremos diferentes formas de construir estimadores, con una estructura matemática de tipo **fraccional**, y que utilicen en la forma más adecuada, dicha información auxiliar, buscando la obtención de buenas estimaciones. Por supuesto, ello no va en detrimento de utilizar estos estimadores en combinación con diseños muestrales que también incorporen información auxiliar como los mencionados anteriormente. Una clasificación pormenorizada de éstas y otras formas de empleo de la información auxiliar puede verse en Hedayat y Sinha (1991).

Este tipo de estimadores de **razón** resultan muy apropiados cuando se presenta una relación aproximada de proporcionalidad directa entre la variable de estudio y otras variables auxiliares. Estas relaciones aparecen con frecuencia en situaciones reales. Por ejemplo, para estimar el contenido total de azúcares de un gran cargamento de naranjas, podemos utilizar la proporcionalidad existente entre el peso del fruto y el de azúcares que contiene, empleando el muestreo para estimar el factor de proporcionalidad. O para estimar el total de automóviles en una población, podemos tener en cuenta la proporcionalidad aproximada, existente entre éstos y el número de habitantes. De esta forma, la existencia de este tipo de relaciones puede ayudarnos a obtener estimaciones más precisas, aunque como veremos, con la contrapartida de la aparición de sesgos en las mismas.

Así, en la sección 2, iniciaremos el desarrollo de estas cuestiones estudiando lo que denominamos **soluciones heurísticas**, basadas en considerar el parámetro a estimar como una derivación de una expresión más compleja, sobre la cual se sustituyen ciertas cantidades poblacionales por muestrales. Un estudio posterior permitirá calibrar si hemos obtenido un estimador adecuado, y encontrar las mejores condiciones para su aplicación.

En la sección 3., enfocaremos nuestro estudio bajo la perspectiva de los modelos de superpoblación, desarrollando la solución general obtenida para el caso de diseño muestral aleatorio simple, $MAS(N, n)$, y para el caso de diseños PPS, es decir, con probabilidades de inclusión proporcionales al tamaño.

En la sección 4., se estudian varias estrategias insesgadas, basadas en el diseño de Lahiri-Midzuno, y en el diseño $MAS(N, n)$ combinado con estimadores especiales como el de Hartley-Ross y el de Mickey. Estas estrategias presentan estimadores de la varianza complicado, por lo que, en esta misma sección, introducimos varias estrategias cuasi-insesgadas, con estimadores de la varianza más simples.

La sección 5., se dedica al estudio del estimador de razón multivariante, y la sección 6., a las distintas formas de combinar el estimador de razón con la estructura de estratos.

Finalmente, en la sección 7. estudiamos propiedades de optimalidad del estimador de razón, bajo el modelo de superpoblación de proporcionalidad directa, exponiendo los resultados clásicos sobre optimalidad de ciertas estrategias basadas en muestreo intencional.

En lo que sigue, denotaremos por U a la población finita bajo estudio, siendo sus elementos,

$$U = \{1, 2, 3, \dots, N\}$$

También denotaremos por Y una variable genérica, definida sobre U que asocia a cada elemento, $i \in U$, un número real, Y_i .

Supondremos que las estimaciones se realizan a partir de una muestra, m , obtenida de la población mediante un determinado diseño muestral. En particular, emplearemos frecuentemente el diseño muestral formado por todas las muestras posibles (en el sentido de subconjuntos) de tamaño n , con distribución de probabilidad uniforme sobre las mismas, al que denominamos **diseño muestral aleatorio simple**, ya mencionado previamente, y que denotaremos $MAS(N, n)$.

2. SOLUCIONES HEURÍSTICAS

Si suponemos que el parámetro a estimar es la media poblacional de la variable Y ,

$$\bar{Y} = \frac{1}{N} \sum_{i \in U} Y_i$$

- y que X es una variable auxiliar perfectamente conocida para todos los elementos de U , podemos considerar la expresión,

$$\bar{Y} = \frac{\bar{Y}}{\bar{X}} \bar{X} = R \bar{X}$$

a partir de la cual podemos definir, **heurísticamente**, el estimador,

$$\hat{\bar{Y}}_R = \hat{R} \bar{X}$$

La forma final del estimador dependerá del diseño muestral utilizado. Por ejemplo, si la muestra se obtiene mediante diseño muestral aleatorio simple, MAS(N, n), podemos estimar R por la razón de medias muestrales, obteniendo,

$$\hat{\bar{Y}}_R = \frac{\bar{y}}{\bar{x}} \bar{X}$$

Obviamente, la eficiencia de este estimador depende en gran medida de la relación existente entre las variables Y y X , siendo el caso más favorable aquel en el que existe una relación aproximada de proporcionalidad entre la variable de estudio y la variable auxiliar. Gráficamente ello significa una nube de puntos concentrada en las proximidades de una línea recta que pasa por el origen. Véase la Figura 1.

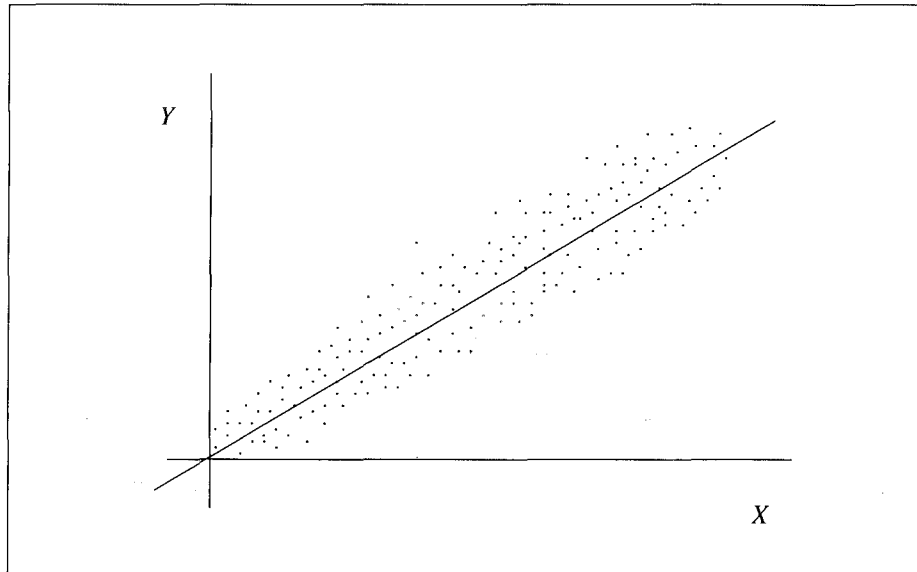


Figura 1. Relación de proporcionalidad aproximada entre dos variables.

El caso ideal, aunque utópico, se daría cuando la proporcionalidad entre las variables es exacta. En tal caso $V[\widehat{Y}_R] = 0$, y la estimación coincide con el verdadero valor. En general, y suponiendo diseño MAS(N, n), basta aplicar las expresiones usuales para la varianza aproximada de la razón, y su estimación, véase Fernández y Mayor (1995), para obtener inmediatamente las conocidas expresiones,

$$V[\widehat{Y}_R] \approx \frac{1-f}{n} (S_y^2 + R^2 S_x^2 - 2RS_{xy})$$

$$\widehat{V}[\widehat{Y}_R] = \frac{1-f}{n} (s_y^2 + \widehat{R}^2 s_x^2 - 2\widehat{R}s_{xy})$$

donde S_y^2 , S_x^2 , S_{xy} denotan cuasivarianzas y cuasicovarianza poblacionales, y s_y^2 , s_x^2 , s_{xy} las correspondientes muestrales. También denotamos $f = n/N$.

Es interesante observar que la solución obtenida heurísticamente volverá a aparecer al aplicar el enfoque predictivo, por lo que pospondremos para más adelante un estudio en profundidad de esta solución, en lo que se refiere al tratamiento del sesgo y del error cuadrático medio.

También es necesario resaltar que, en el enfoque heurístico se está empleando, aunque no *a priori* o de una forma manifiesta, la existencia de una relación aproximada de proporcionalidad directa, entre las variables Y y X , y en este sentido, hemos de recordar los trabajos pioneros de algunos científicos, en el campo de la demografía, que han utilizado esta idea.

Así, son clásicos los estudios del inglés John Graunt sobre la estimación del número de habitantes de Londres. Sus resultados se publicaron en el famoso trabajo *Natural and Political Observations made upon the Bills of Mortality*, aparecido en 1662. En este estudio, Graunt investigó un conjunto de familias pertenecientes a determinadas parroquias de la ciudad de Londres, donde los registros resultaban fiables, y observó que había un promedio de tres fallecimientos anuales en 11 familias, siendo la cantidad total de fallecimientos por año en esta ciudad de aproximadamente de 13000. De este forma, Graunt concluyó que el número de familias era de 48000, y suponiendo un tamaño medio familiar de 8, estimó en 384000 el número de habitantes de la ciudad.

Como puede observarse, en estos cálculos está implícito un modelo de proporcionalidad entre el número de fallecimientos y el de familias y también entre el número de éstas y el de habitantes. También hay que notar como Graunt no realizó ningún estudio adicional encaminado a cuantificar los posibles errores cometidos. Véanse Chang (1976) y Hald (1990) para un estudio pormenorizado del trabajo de Graunt.

Otro precedente histórico de gran relevancia lo constituyen los estudios sobre la población de Francia, llevados a cabo por Laplace, y cuyos primeros resultados

aparecieron en 1786. Sus métodos de muestreo y estimación fueron similares a los empleados por Graunt, pero el científico francés vio la necesidad de tener en cuenta de alguna forma la precisión de los resultados obtenidos, tanto en su control, seleccionando una muestra de calidad, como en su medición.

A partir de una muestra, Laplace estimó la población total del país utilizando una estimación de razón, empleando los nacimientos ocurridos el año precente como variable auxiliar. Adicionalmente, calculó la distribución de la diferencia entre el verdadero valor y el estimado, aproximando esta distribución por una normal. Sus métodos y resultados aparecieron en la clásica obra *Théorie Analytique des Probabilités*, publicada en 1812. Véase Cochran (1978) y Chang (1976).

3. MODELO DE SUPERPOBLACIÓN DE PROPORCIONALIDAD DIRECTA

El enfoque considerado en el apartado anterior, que denominamos **heurístico**, en cierto modo se contrapone, metodológicamente, al que vamos a considerar ahora, más formal, y denominado **enfoque predictivo**. Este enfoque se basa en suponer un «modelo» o relación funcional entre la variable de estudio, Y , y la variable auxiliar, X , de la forma $Y = f(X)$. Si $f(\cdot)$ fuera conocida completamente, el conocimiento de X nos llevaría al de Y y por tanto al de cualquier parámetro $\theta(Y)$.

Usualmente, $f(\cdot)$ es desconocida y su determinación sólo puede realizarse de un modo **aproximado**, a partir del conocimiento de la variable X , y de la información suministrada por el estadístico,

$$\{(X_i, Y_i) | i \in m\}$$

Con dicha información, buscaremos la función $\hat{f}(\cdot)$ que «mejor» explique la relación observada y que denominaremos **función predictora**.

En este sentido, es muy importante realizar estudios exploratorios de los datos muestrales, por ejemplo dibujando la **nube de puntos**, que proporcionen indicios sobre las pautas que relacionan X con Y . Véase Fernández y Mayor (1995).

El siguiente paso será estimar $\theta(Y)$ mediante $\theta(\hat{Y})$, siendo $\hat{Y}_i$, $i \in U$ los valores aproximados a los verdaderos valores Y_i , $i \in U$, proporcionados por la función predictora $\hat{f}(\cdot)$, es decir,

$$\hat{Y}_i = \hat{f}(X_i) \quad i \in U$$

Por ejemplo, para un parámetro lineal del tipo $\theta(Y) = \sum_{i \in U} a_i Y_i$ se tendrá,

$$\theta(Y) = \sum_{i \in U} a_i Y_i = \sum_{i \in U} a_i \hat{Y}_i + \sum_{i \in U} a_i (Y_i - \hat{Y}_i)$$

cuyo primer sumando es conocido, y el segundo es desconocido, siendo función del error de estimación de $f(\cdot)$. Si es posible despreciar este segundo sumando, entonces podemos dar como estimador,

$$\hat{\theta}_1(Y) = \sum_{i \in U} a_i \hat{Y}_i$$

en otro caso, podemos estimarlo, por ejemplo mediante el estimador de Horvitz-Thompson, obteniendo el estimador alternativo,

$$\hat{\theta}_2 = \sum_{i \in U} a_i \hat{Y}_i + \sum_{i \in m} a_i \frac{Y_i - \hat{Y}_i}{\pi_i}$$

Es importante observar que en el **enfoque predictivo** se combinan dos procesos estadísticos, el **ajuste** y la **estimación**, dependiendo la bondad de las estimaciones de numerosos factores entre los que destacamos la habilidad en la conjunción de ambos procesos, la bondad del ajuste realizado y la estructura de la población en lo que atañe a las variables involucradas.

Con el fin de obtener el estimador de razón a partir del enfoque predictivo, vamos a suponer que la población, en relación a la variable de estudio, Y , y la variable auxiliar, X , posee el siguiente modelo de superpoblación, de **proporcionalidad directa**,

$Y_i = \beta X_i + \epsilon_i$ $E_s[\epsilon_i] = 0$ $V_s[\epsilon_i] = \sigma^2 v(X_i)$ $E_s[\epsilon_i \epsilon_j] = 0, \quad i \neq j$

siendo $v(\cdot)$ una función conocida que marca la estructura de la varianza. Adicionalmente, supondremos que X toma únicamente valores no negativos.

Notemos que si fuera posible observar la totalidad de los valores $\{(X_i, Y_i) | i \in U\}$, como en el caso de un censo, podríamos obtener una estimación de β basándonos en el teorema de Gauss-Markov generalizado dado por C.R. Rao (1965), mediante la resolución del siguiente problema de minimización,

$$\min_{\beta} \sum_{i \in U} \frac{(Y_i - \beta X_i)^2}{\sigma^2 v(X_i)}$$

lo que proporcionaría,

$$\hat{\beta} = \frac{\sum_{i \in U} Y_i X_i / v(X_i)}{\sum_{i \in U} X_i^2 / v(X_i)}$$

Pero como sólo disponemos de la información suministrada por la muestra, utilizaremos el método de estimación propuesto por Kish y Frankel (1974) y Fuller (1975), consistente en reemplazar la suma poblacional por una estimación muestral, y más concretamente, si es la de Horvitz-Thompson, resolviendo el problema,

$$\min_{\beta} \sum_{i \in m} \frac{(Y_i - \beta X_i)^2}{\sigma^2 v(X_i) \pi_i}$$

Para simplificar la solución del mismo, tomaremos $v(x) = x$, lo que representa una situación muy general, en la cual la varianza en la superpoblación aumenta proporcionalmente al valor de la variable auxiliar (que ha de ser no negativa). Con esta hipótesis, obtenemos, sin más que derivar, la siguiente ecuación normal,

$$\sum_{i \in m} \frac{Y_i - \hat{\beta} X_i}{\pi_i} = 0$$

y la siguiente estimación de β ,

$$\hat{\beta} = \frac{\sum_{i \in m} Y_i / \pi_i}{\sum_{i \in m} X_i / \pi_i}$$

Si ahora empleamos el estimador $\hat{\theta}_2$ para estimar la media poblacional, obtendremos,

$$\hat{\theta}_2 = \sum_{i \in U} \hat{Y}_i / N + \sum_{i \in m} \frac{Y_i - \hat{Y}_i}{N \pi_i} = \sum_{i \in U} \hat{\beta} X_i / N + \sum_{i \in m} \frac{Y_i - \hat{\beta} X_i}{N \pi_i} = \hat{\beta} \bar{X} \triangleq \hat{Y}_R$$

donde el segundo sumando se ha anulado por la ecuación normal. Así pues, hemos obtenido el siguiente estimador de la media poblacional,

$$\hat{Y}_R = \frac{\sum_{i \in m} Y_i / \pi_i}{\sum_{i \in m} X_i / \pi_i} \bar{X}$$

caso particular del de Sánchez-Crespo (1980).

El estudio de la varianza de $\widehat{Y}_R$ dependerá del diseño muestral que se emplee, y en general, se puede realizar por los métodos usuales basados en aproximación lineal.

A continuación particularizaremos la solución obtenida, para el diseño muestral aleatorio simple, MAS(N, n), y para los diseños Π PS, esto es, con probabilidades de inclusión de los elementos proporcionales a su valor de la variable auxiliar.

3.1. Diseño MAS(N, n)

Al ser las probabilidades de inclusión de primer orden para este diseño muestral, $\pi_i = n/N$, $\forall i \in U$, obtenemos el estimador,

$$\widehat{Y}_R = \frac{\bar{y}}{\bar{x}} \bar{X}$$

que coincide con el obtenido heurísticamente, y obviamente es, en general sesgado, ya que,

$$B[\widehat{Y}_R] = E[\widehat{Y}_R] - \bar{Y} = \bar{X} E\left[\frac{\bar{y}}{\bar{x}}\right] - E[\bar{y}] = E[\bar{x}] E\left[\frac{\bar{y}}{\bar{x}}\right] - E[\bar{y}] = -\text{Cov}\left[\bar{x}, \frac{\bar{y}}{\bar{x}}\right]$$

Con objeto de realizar un estudio cuantitativo de este sesgo, así como del error cuadrático medio y de la varianza, construiremos los valores,

$$\delta\bar{y} = \frac{\bar{y} - \bar{Y}}{\bar{Y}} \quad \delta\bar{x} = \frac{\bar{x} - \bar{X}}{\bar{X}}$$

que nos servirán, adoptando la línea expuesta en David y Sukhatme (1974), para definir las siguientes cantidades,

$$B_k[\widehat{Y}_R] = \bar{Y} E\left[\sum_{i=0}^{2k-1} (-\delta\bar{x})^i (\delta\bar{y} - \delta\bar{x})\right]$$

$$\text{ECM}_k[\widehat{Y}_R] = \bar{Y}^2 E\left[(\delta\bar{y} - \delta\bar{x})^2 \sum_{i=0}^{2k-2} (i+1)(-\delta\bar{x})^i\right]$$

con $k \geq 1$ entero. Estas cantidades serán utilizadas como aproximaciones del sesgo y del error cuadrático medio, y con respecto a sus órdenes de aproximación se verifica el resultado que exponemos a continuación.

Teorema 1 *Bajo diseño muestral MAS(N, n), y suponiendo que la media muestral de X verifica $\bar{x} \geq x_0 > 0$, se tiene para las cantidades $B_k[\widehat{Y}_R]$ y $\text{ECM}_k[\widehat{Y}_R]$,*

$$\begin{aligned} \left| B[\widehat{Y}_R] - B_k[\widehat{Y}_R] \right| &\leq O(n^{-(k+1)}) \\ \left| \text{ECM}[\widehat{Y}_R] - \text{ECM}_k[\widehat{Y}_R] \right| &\leq O(n^{-(k+1)}) \end{aligned}$$

Para un estudio pormenorizado de este resultado, y su demostración, véanse David y Sukhatme (1974) y Sukhatme *et al.* (1984).

Tomando ahora $k = 1$, y teniendo en cuenta que entre el error cuadrático medio y la varianza existe la relación,

$$\text{ECM}[\widehat{Y}_R] = V[\widehat{Y}_R] + \left(B[\widehat{Y}_R] \right)^2$$

se obtienen fácilmente las siguientes aproximaciones.

Teorema 2

$$\begin{aligned} B[\widehat{Y}_R] &= \frac{1-f}{n} \left(\frac{\bar{Y}}{\bar{X}^2} S_x^2 - \frac{1}{\bar{X}} S_{xy} \right) + O(n^{-2}) = O(n^{-1}) \\ \text{ECM}[\widehat{Y}_R] &= \frac{1-f}{n} \left(S_y^2 + \frac{\bar{Y}^2}{\bar{X}^2} S_x^2 - 2 \frac{\bar{Y}}{\bar{X}} S_{xy} \right) + O(n^{-2}) = O(n^{-1}) \\ V[\widehat{Y}_R] &= \frac{1-f}{n} \left(S_y^2 + \frac{\bar{Y}^2}{\bar{X}^2} S_x^2 - 2 \frac{\bar{Y}}{\bar{X}} S_{xy} \right) + O(n^{-2}) = O(n^{-1}) \end{aligned}$$

El resultado anterior es importante por varias razones. Por una parte, nos dice que el sesgo puede ser reducido incrementando el tamaño muestral.

Por otra parte, proporciona una expresión del sesgo, aproximada hasta el orden $O(n^{-2})$. Similares consideraciones se derivan para el error cuadrático medio y la varianza.

Como una aplicación interesante, hemos realizado una comprobación empírica del comportamiento del sesgo, utilizando una población construida artificialmente, EXP1000, con $N = 1000$ elementos, para los cuales las variables Y y X están ligadas por la relación,

$$Y_i = 1000 + 10 \times X_i + 3 \times \varepsilon_i, \quad i = 1, \dots, 1000$$

Los valores de X_i han sido generados de una distribución exponencial de media $\mu = 200$ y los de ϵ_i se han obtenido restando 50 a los valores generados a partir de una distribución exponencial de media 50.

Para cada uno de los valores $n = 10, 15, \dots, 100$ hemos simulado 500 veces un muestreo aleatorio simple, MAS(1000, n), y para cada valor de n hemos tabulado el **término guía** del sesgo,

$$B_1 = B_1[\widehat{Y}_R] = \frac{1-f}{n} \left(\frac{\bar{Y}}{\bar{X}^2} S_x^2 - \frac{1}{\bar{X}} S_{xy} \right)$$

así como $\widehat{B} = \widehat{B}[\widehat{Y}_R]$ calculado promediando $\widehat{Y}_R - \bar{Y}$ sobre las 500 simulaciones. También hemos tabulado la diferencia, en valor absoluto, entre ambas cantidades. De esta forma hemos obtenido los resultados siguientes,

n	B_1	$\widehat{B}$	$ \widehat{B} - B_1 $
10	94.367097	79.384970	14.982127
15	62.593663	67.305330	4.711667
20	46.706946	50.874577	4.167631
25	37.174917	43.594316	6.419399
30	30.820230	26.349465	4.470765
35	26.281168	22.712932	3.568236
40	22.876871	20.871180	2.005691
45	20.229085	20.171039	0.058046
50	18.110856	12.119815	5.991041
55	16.377760	20.709838	4.332078
60	14.933513	18.269774	3.336261
65	13.711458	18.450887	4.739429
70	12.663982	14.513552	1.849570
75	11.756170	13.513885	1.757715
80	10.961834	12.005794	1.043960
85	10.260949	11.873348	1.612399
90	9.637941	10.489407	0.851466
95	9.080512	10.106495	1.025983
100	8.578827	9.283448	0.704621

Como puede verse, el comportamiento del sesgo estimado coincide con la expresión teórica, en lo que se refiere a su dependencia del tamaño muestral, n .

También la rápida disminución observada para $|\hat{B}[\hat{Y}_R] - B_1[\hat{Y}_R]|$ parece concordar con el orden $O(n^{-2})$ hallado teóricamente. Una gráfica de estos valores se muestra en la Figura 2.

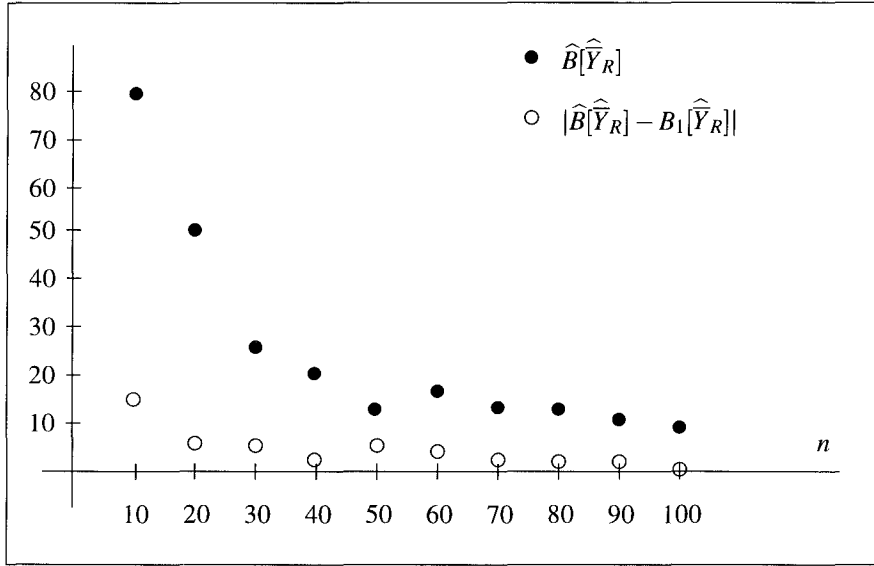


Figura 2. Sesgo estimado y su diferencia con el término orden $O(n^{-1})$.

Con respecto a la comparación entre el estimador de razón, $\hat{Y}_R$ y el estimador $\hat{\bar{Y}} = \bar{y}$, que no emplea información auxiliar, se tiene,

Teorema 3 Si el tamaño muestral es lo suficientemente grande para despreciar los términos de orden $O(n^{-2})$, entonces el estimador de razón, $\hat{Y}_R$ es más eficiente que $\hat{\bar{Y}} = \bar{y}$ si el coeficiente de correlación lineal, ρ verifica,

$$\rho > \frac{1}{2} \frac{CV_x}{CV_y}$$

donde $CV_y = S_y/\bar{Y}$ denota el cuasicoeficiente de variación de Y , y análogo para X .

Resultado que se obtiene inmediatamente sin más que resolver en ρ la inecuación,

$$\frac{1-f}{n} \left(S_y^2 + \frac{\bar{Y}^2}{\bar{X}^2} S_x^2 - 2 \frac{\bar{Y}}{\bar{X}} S_{xy} \right) < \frac{1-f}{n} S_y^2$$

Observemos que la anterior condición para ρ , suele verificarse cuando existe una elevada correlación entre las variables, con lo cual obtendremos mejores resultados con el estimador de razón que con la media muestral simple. Por supuesto, ello presenta el inconveniente del sesgo, aunque este puede ser reducido por algunos procedimientos como el aumento del tamaño muestral, la aplicación de técnicas de tipo jackknife o la estimación del sesgo, siendo referencias fundamentales para el estudio de estas cuestiones, además de las citadas, Hansen, Hurwitz y Madow (1953), Cochran (1993) y Hedayat y Sinha (1991).

También es interesante la comparación con el estimador de regresión de la media,

$$\hat{Y}_{RG} = \bar{y} + \frac{s_{xy}}{s_x^2} (\bar{X} - \bar{x})$$

En este sentido, se verifica que si, como en el anterior resultado, despreciamos los términos de orden $O(n^{-2})$, entonces siempre es más eficiente el estimador de regresión que el de razón, es decir, $V[\hat{Y}_{RG}] \leq V[\hat{Y}_R]$. Véase Hedayat y Sinha (1991).

Finalmente añadiremos que tanto el sesgo, como el error cuadrático medio y la varianza de $\hat{Y}_R$, pueden ser estimados mediante,

$$\begin{aligned} \widehat{B}[\hat{Y}_R] &= \frac{1-f}{n} \left(\frac{\bar{y}}{\bar{x}^2} s_x^2 - \frac{1}{\bar{x}} s_{xy} \right) \\ \widehat{ECM}[\hat{Y}_R] &= \frac{1-f}{n} \left(s_y^2 + \frac{\bar{y}^2}{\bar{x}^2} s_x^2 - 2 \frac{\bar{y}}{\bar{x}} s_{xy} \right) \\ \widehat{V}[\hat{Y}_R] &= \frac{1-f}{n} \left(s_y^2 + \frac{\bar{y}^2}{\bar{x}^2} s_x^2 - 2 \frac{\bar{y}}{\bar{x}} s_{xy} \right) \end{aligned}$$

3.2. Diseños IIIPS

Para estos diseños, suponiendo tamaño de muestra fijo, n , se tiene $\pi_i = n X_i / T(X)$, con lo cual, el estimador de razón adopta la forma,

$$\hat{Y}_R = \frac{1}{n} \sum_{i \in m} \frac{Y_i}{X_i} \bar{X}$$

Es interesante observar que dicho estimador coincide con el estimador de Horvitz-Thompson de la media, empleando X como variable auxiliar y probabilidades de inclusión proporcionales al tamaño, siendo pues la estimación insesgada. El estudio de su varianza se realizará por los métodos usuales. Por ejemplo, aplicando las expresiones de Yates-Grundy-Sen, obtendremos,

$$V[\widehat{Y}_R] = -\frac{\bar{X}^2}{2n^2} \sum_{i,j \in U} (\pi_{ij} - \pi_i \pi_j) \left(\frac{Y_i}{X_i} - \frac{Y_j}{X_j} \right)^2$$

$$\widehat{V}[\widehat{Y}_R] = -\frac{\bar{X}^2}{2n^2} \sum_{i,j \in m} \frac{(\pi_{ij} - \pi_i \pi_j)}{\pi_{ij}} \left(\frac{Y_i}{X_i} - \frac{Y_j}{X_j} \right)^2$$

Así, la varianza disminuye cuanto más ajustada es la relación de proporcionalidad entre las variables Y y X .

4. ESTRATEGIAS INSESGADAS Y CUASI-INSESGADAS

En este apartado estudiaremos algunas combinaciones de diseño y estimador que proporcionan estimaciones insesgadas o cuasi-insesgadas, es decir, con sesgo de orden $O(n^{-2})$. Estas estrategias se basan, tanto en las particularidades del diseño muestral, como en modificaciones introducidas en la expresión del estimador.

4.1. Estrategia insesgada basada en el esquema de Lahiri-Midzuno

Lahiri (1951) y Midzuno (1952), independientemente, han descrito el esquema de muestreo consistente en la selección de un elemento, i , con probabilidad proporcional a su tamaño, es decir, $p_i = X_i/T(X)$, $\forall i \in U$; y la selección de $n-1$ elementos adicionales mediante diseño MAS($N-1, n-1$) en $U - \{i\}$. La importancia de este esquema radica en el siguiente resultado,

Teorema 4 *Bajo el diseño muestral originado por el esquema de Lahiri-Midzuno se verifica que $\widehat{Y}_R = (\bar{y}/\bar{x})\bar{X}$ es un estimador insesgado de $\bar{Y}$.*

Demostración: Estudiemos en primer lugar el diseño muestral resultante. Su espacio muestral está formado por todas las muestras de tamaño n . Para calcular la probabilidad, $p(m)$, de una muestra del diseño, m , consideremos $m^* = (j_1, j_2, \dots, j_n)$, muestra ordenada con los mismos elementos de m . Se tiene entonces,

$$p(m^*) = \frac{X_{j_1}}{T(X)} \frac{1}{(n-1)! \binom{N-1}{n-1}}$$

y para la muestra m será,

$$p(m) = \sum_{m^* \sim m} p(m^*) = \frac{\sum_{i \in m} X_i}{T(X)} \frac{1}{\binom{N-1}{n-1}}$$

donde la notación $m^* \sim m$ expresa que la suma se extiende a todas las muestras ordenadas, con los mismos elementos que la muestra m . Calculemos ahora la esperanza de $\widehat{Y}_R$,

$$\begin{aligned} E[(\bar{y}/\bar{x})\bar{X}] &= \sum_{m \in M} p(m)(\bar{y}(m)/\bar{x}(m))\bar{X} = \frac{1}{T(X)} \frac{1}{\binom{N-1}{n-1}} \sum_{m \in M} n\bar{x}(m)(\bar{y}(m)/\bar{x}(m))\bar{X} \\ &= \frac{1}{N} \frac{1}{\binom{N-1}{n-1}} \sum_{m \in M} \sum_{i \in m} Y_i = \frac{1}{N} \frac{1}{\binom{N-1}{n-1}} T(Y) \binom{N-1}{n-1} = \bar{Y} \end{aligned}$$

■

Para estudiar la varianza de esta estimación, se pueden seguir varios caminos. Uno de ellos se basa en la técnica de linealización, tomando los términos lineales del desarrollo de Taylor de $\widehat{Y}_R = (\bar{y}/\bar{x})\bar{X}$ en un entorno de $(\bar{Y}, \bar{X})$,

$$\widehat{Y}_R = \frac{\bar{y}}{\bar{x}} \bar{X} \approx \bar{Y} + (\bar{y} - \bar{Y}) - \frac{\bar{Y}}{\bar{X}}(\bar{x} - \bar{X}) = \bar{Y} + (\bar{y} - R\bar{x}) = \bar{Y} + \bar{z}$$

donde $Z_i = Y_i - R X_i$, $i \in m$. Se tiene pues,

$$\begin{aligned} V[\widehat{Y}_R] &\approx V[\bar{z}] = -\frac{1}{2n^2} \sum_{i,j \in U} (\pi_{ij} - \pi_i \pi_j)(Z_i - Z_j)^2 \\ &= \sum_{i,j \in U} c_{ij}(Z_i - Z_j)^2 \end{aligned}$$

siendo,

$$c_{ij} = -\frac{1}{2n^2}(\pi_{ij} - \pi_i \pi_j)$$

Y si tenemos en cuenta que para el diseño que estamos considerando se verifica, véase Fernández y Mayor (1995),

$$\begin{aligned} \pi_i &= \frac{N-n}{N-1} \frac{X_i}{T(X)} + \frac{n-1}{N-1} \\ \pi_{ij} &= \frac{n-1}{N-1} \left[\frac{N-n}{N-2} \frac{(X_i + X_j)}{T(X)} + \frac{n-2}{N-2} \right] \quad i \neq j \end{aligned}$$

sustituyendo y desarrollando, se obtiene,

$$c_{ij} = \frac{N-n}{2n^2(N-1)^2} \left[(N-n) \frac{X_i X_j}{(T(X))^2} + \frac{(n-1)(1 - X_i/T(X) - X_j/T(X))}{N-2} \right] \quad i \neq j$$

De forma similar, empleando la expresión de Yates-Grundy-Sen para la varianza estimada, se puede obtener una estimación de $V[\widehat{Y}_R]$.

Otra línea diferente para el estudio de la varianza, sin emplear técnicas aproximadas, es la introducida por Rao y Vijayan (1977). Estos autores consideran la siguiente forma cuadrática para expresar la varianza, obtenida mediante un cálculo directo,

$$V[\widehat{Y}_R] = \sum_{i \in U} \alpha_{ii} Y_i^2 + \sum_{\substack{i, j \in U \\ i \neq j}} \alpha_{ij} Y_i Y_j$$

donde,

$$\begin{aligned} \alpha_{ii} &= \frac{\bar{X}}{n^2 \binom{N}{n}} \left(\sum_{m \ni i} \frac{1}{\bar{x}(m)} \right) - \frac{1}{N^2} \quad \forall i \in U \\ \alpha_{ij} &= \frac{\bar{X}}{n^2 \binom{N}{n}} \left(\sum_{m \ni i, j} \frac{1}{\bar{x}(m)} \right) - \frac{1}{N^2} \quad \forall i \neq j \in U \end{aligned}$$

Y aplicando los resultados clásicos sobre estimación insesgada de formas cuadráticas (véase por ejemplo Hedayat y Sinha (1991)), obtenemos directamente el siguiente teorema,

Teorema 5 *Todo estimador insesgado y no negativo de la anterior varianza, $V[\widehat{Y}_R]$, adopta necesariamente la forma,*

$$\widehat{V}[\widehat{Y}_R] = -\frac{1}{2} \sum_{i, j \in m} \alpha_{ij}(m) X_i X_j \left(\frac{Y_i}{X_i} - \frac{Y_j}{X_j} \right)^2$$

verificando los coeficientes la condición de insesgades,

$$\sum_{m \ni i, j} \alpha_{ij}(m) p(m) = \alpha_{ij} \quad \forall i \neq j \in U$$

Como posibles elecciones de $\alpha_{ij}(m)$, Rao y Vijayan (1977) sugieren las siguientes,

$$(a) \quad \alpha_{ij}(m) = \alpha_{ij} / \pi_{ij} \quad i, j \in m$$

$$(b) \alpha_{ij}(m) = \frac{\bar{X}^2}{n^2 \bar{x}^2(m)} - \frac{\bar{X}(N-1)}{\bar{x}(m)Nn(n-1)} \quad i, j \in m$$

En Rao (1966, 1972), Lanke (1974), Rao y Vijayan (1977) y Hedayat y Sinha (1991) pueden encontrarse resultados adicionales relacionados con esta línea.

4.2. Estrategia insesgada basada en el estimador de Hartley-Ross

Esta estrategia utiliza como diseño muestral el aleatorio simple, MAS(N, n), en combinación con un estimador especial definido por Hartley y Ross (1954). Para construir este estimador, consideremos previamente el siguiente estimador, **heurístico**, de la media poblacional,

$$\hat{\bar{Y}}'_R = \frac{1}{n} \sum_{i \in m} \frac{Y_i}{X_i} \bar{X} = \bar{z} \bar{X}$$

donde $Z_i = Y_i/X_i$. Este estimador es sesgado ya que,

$$\begin{aligned} B[\hat{\bar{Y}}'_R] &= E[\bar{z} \bar{X}] - \bar{Y} = \bar{Z} \bar{X} - \bar{Y} \\ &= \bar{Z} \bar{X} - \bar{Z} \bar{X} = -\frac{N-1}{N} S_{zx} \end{aligned}$$

siendo entonces $[-(N-1)s_{zx}/N]$ un estimador insesgado de dicho sesgo.

Podemos entonces, sin más que restar la estimación insesgada del sesgo, construir el siguiente estimador insesgado de **Hartley y Ross** para la media poblacional,

$$\begin{aligned} \hat{\bar{Y}}_{HR} &= \bar{z} \bar{X} + \frac{N-1}{N} s_{zx} \\ &= \bar{z} \bar{X} + \frac{n(N-1)}{N(n-1)} (\bar{y} - \bar{z} \bar{x}) \end{aligned}$$

La varianza de este estimador, así como su estimación, adoptan formas muy complicadas, lo que explica que el estimador de Hartley-Ross no se haya popularizado. Para su obtención, Robson (1957) ha empleado el formalismo de «polykays» multivariantes, obteniendo una expresión de $V[\hat{\bar{Y}}_{HR}]$ en función de medias simétricas poblacionales, a partir de la cual, podemos obtener un estimador insesgado sin más que sustituir éstas por las correspondientes medias simétricas muestrales.

Es posible realizar una simplificación suponiendo que la población es infinita, en cuyo caso se obtiene, empleando la notación usual,

$$\begin{aligned}\lim_{N \rightarrow \infty} V[\widehat{Y}_{HR}] &= \frac{1}{n} \left(\sigma_y^2 + \bar{Z}^2 \sigma_x^2 - 2\bar{Z}\sigma_{xy} \right) + \frac{1}{n(n-1)} \left(\sigma_z^2 \sigma_x^2 + \sigma_{xz}^2 \right) \\ &\approx \frac{1}{n} \left(\sigma_y^2 + \bar{Z}^2 \sigma_x^2 - 2\bar{Z}\sigma_{xy} \right)\end{aligned}$$

aproximación obtenida también, independientemente, por Goodman y Hartley (1958). Estos autores proporcionan además el siguiente resultado acerca de la comparación del estimador de Hartley y Ross y el estimador de razón usual, $\widehat{Y}_R = (\bar{y}/\bar{x})\bar{X}$,

Teorema 6 *Si se ignoran los términos de orden $O(n^{-2})$ y superior, el estimador $\widehat{Y}_{HR}$ es más eficiente que el estimador $\widehat{Y}_R = (\bar{y}/\bar{x})\bar{X}$ si y sólo si el coeficiente de regresión, β , entre las variables Y y X está más próximo a $\bar{Z}$ que $R = \bar{Y}/\bar{X}$. En caso de ser $\bar{Z} = R$, ambos estimadores son igualmente eficientes.*

Observemos que la condición propuesta en este teorema no es fácil de verificar en la práctica, véase Sukhatme *et al.* (1984). Otras aportaciones relacionadas pueden también consultarse en Ruiz y Santos (1989), y en Sahoo y Ruiz (1994).

4.3. Estrategia insesgada basada en el estimador de Mickey

Una metodología diferente, propuesta por Mickey (1959), se fundamenta en la partición al azar de una muestra aleatoria simple, m , de tamaño n , en k grupos, cada uno con $l = n/k$ elementos. Denotando por $m_1, m_2, \dots, m_k$ a dichos grupos, se definen entonces las cantidades,

$$\bar{y}(m - m_j), \quad \bar{x}(m - m_j) \quad j = 1, \dots, k$$

es decir, las medias de Y y X calculadas sobre las submuestras obtenidas omitiendo en m , sucesivamente, los grupos $m_1, m_2, \dots, m_k$.

A partir de las cantidades anteriores, definimos,

$$\bar{y}_R^{(j)} = \frac{\bar{y}(m - m_j)}{\bar{x}(m - m_j)} \bar{X} \quad j = 1, \dots, k$$

$$\bar{y}_M^{(j)} = \bar{y}_R^{(j)} + \frac{k(N - n + l)}{N} \left[\bar{y} - \bar{y}_R^{(j)} \frac{\bar{x}}{\bar{X}} \right]$$

que nos sirven para construir el estimador de Mickey,

$$\widehat{Y}_M = \frac{1}{k} \sum_{j=1}^k \bar{y}_M^{(j)}$$

cuya importancia radica en el siguiente resultado,

Teorema 7 $\widehat{Y}_M$ es un estimador insesgado de la media poblacional.

Demostración: Para probarlo, es suficiente demostrar que para cada j , $\bar{y}_M^{(j)}$ es insesgado respecto a $\bar{Y}$, para ello expresamos dichas cantidades en la forma,

$$\begin{aligned} \bar{y}_M^{(j)} &= \frac{N - (n - l)}{N} \left[\bar{y}(m_j) - \frac{\bar{y}(m - m_j)}{\bar{x}(m - m_j)} \left(\bar{x}(m_j) - \frac{N\bar{X} - (n - l)\bar{x}(m - m_j)}{N - (n - l)} \right) \right] \\ &\quad + \frac{n - l}{N} \bar{y}(m - m_j) \end{aligned}$$

Por las propiedades del diseño muestral MAS(N, n), sabemos que si $(n - l)$ unidades son seleccionadas al azar en m , las l unidades restantes forman una muestra aleatoria seleccionada mediante diseño MAS($N - (n - l), l$).

Así pues, descomponiendo la esperanza en dos fases, la primera fase, E_1 , de obtención de l elementos de entre $N - (n - l)$, supuesto que $(n - l)$ unidades ya han sido seleccionadas; y la segunda fase, E_2 , de selección de $(n - l)$ elementos a partir de N , obtenemos,

$$\begin{aligned} E[\bar{y}_M^{(j)}] &= E_2 E_1[\bar{y}_M^{(j)}] \\ &= E_2 \left[\frac{N - (n - l)}{N} \left(\frac{N\bar{Y} - (n - l)\bar{y}(m - m_j)}{N - (n - l)} \right) + \frac{n - l}{N} \bar{y}(m - m_j) \right] \\ &= E_2 [\bar{Y}] = \bar{Y} \end{aligned}$$

■

4.4. Estrategias cuasi-insesgadas

Ya hemos visto que el estimador de razón, $\widehat{Y}_R = (\bar{y}/\bar{x})\bar{X}$, en combinación con el diseño muestral aleatorio simple, da lugar a una estrategia sesgada, en el sentido de

que la estimación lo es, siendo el sesgo de orden $O(n^{-1})$. Denominaremos **estrategias cuasi-insesgadas** a aquellas que siendo sesgadas, proporcionan estimaciones con sesgo de orden $O(n^{-2})$ o inferior.

Las estrategias de este tipo, que vamos a considerar, están basadas en el diseño muestral aleatorio simple, $MAS(N, n)$, en combinación con estimadores especiales, usualmente contruidos a partir del estimador $\widehat{Y}_R = (\bar{y}/\bar{x})\bar{X}$, de forma tal que el sesgo se reduzca en cierto orden de aproximación.

Así, el primer estimador que estudiamos se obtiene aplicando una técnica de tipo **jackknife**, al estimador $\widehat{Y}_R = (\bar{y}/\bar{x})\bar{X}$. Para ello, utilizaremos las cantidades introducidas para construir el estimador de Mickey, y definiremos a partir de las mismas el siguiente estimador,

$$\widehat{Y}_J = \frac{N-n+l}{N} k \widehat{Y}_R - \frac{N-n}{N} \frac{k-1}{k} \sum_{j=1}^k \bar{y}_R^{(j)}$$

Observemos que si N es muy elevado, de forma que,

$$(N-n+l)/N \approx 1$$

$$(N-n)/N \approx 1$$

obtenemos la versión simplificada,

$$\widehat{Y}_J' = k \widehat{Y}_R - \frac{k-1}{k} \sum_{j=1}^k \bar{y}_R^{(j)}$$

que coincide con la forma de jackknife usual, tomando $k = n$, y por tanto $l = 1$. Véase Quenouille (1956).

Con respecto a los anteriores estimadores, se tiene el siguiente resultado sobre su sesgo y error cuadrático medio, cuya demostración sigue la misma línea que la realizada para el estimador de Mickey.

Teorema 8 *El sesgo y el error cuadrático medio del estimador $\widehat{Y}_J$ verifican,*

$$B[\widehat{Y}_J] = O(n^{-2})$$

$$ECM[\widehat{Y}_J] = \frac{1-f}{n} \left(S_y^2 + \frac{\bar{Y}^2}{\bar{X}^2} S_x^2 - 2 \frac{\bar{Y}}{\bar{X}} S_{xy} \right) + O(n^{-2})$$

y análogo para $\widehat{Y}_J'$.

Este resultado nos dice que, en aproximación de primer orden, los estimadores $\widehat{Y}_J$ e $\widehat{Y}_R$ son igualmente eficientes que $\widehat{Y}_R = (\bar{y}/\bar{x})\bar{X}$. Véase Sukhatme *et al.* (1984).

Otras estrategias cuasi-insesgadas, basadas en diseño MAS(N, n), son las definidas por los siguientes estimadores,

- Estimador de De Pascual (1961),

$$\widehat{Y}_{DP} = \widehat{Y}_R + \frac{\bar{y} - \bar{z}\bar{x}}{n-1} \quad \text{con } Z_i = \frac{Y_i}{X_i}$$

- Estimador de Beale (1962),

$$\widehat{Y}_B = \widehat{Y}_R \left[1 + \left(\frac{1}{n} - \frac{1}{N} \right) \frac{s_{xy}}{\bar{x}\bar{y}} \right] \left[1 + \left(\frac{1}{n} - \frac{1}{N} \right) \frac{s_x^2}{\bar{x}^2} \right]^{-1}$$

- Estimador de Tin (1965),

$$\widehat{Y}_T = \widehat{Y}_R \left[1 - \left(\frac{1}{n} - \frac{1}{N} \right) \left(\frac{s_x^2}{\bar{x}^2} - \frac{s_{xy}}{\bar{x}\bar{y}} \right) \right]$$

Todos estos estimadores poseen un sesgo de orden $O(n^{-2})$, y con respecto a su eficiencia, se tiene el siguiente resultado,

Teorema 9 *Los errores cuadráticos medios de los estimadores de Tin, Beale y De Pascual verifican,*

$$\text{ECM}[\widehat{Y}_T] = \frac{1-f}{n} \left(S_y^2 + \frac{\bar{Y}^2}{\bar{X}^2} S_x^2 - 2 \frac{\bar{Y}}{\bar{X}} S_{xy} \right) + O(n^{-2})$$

$$\text{ECM}[\widehat{Y}_B] = \frac{1-f}{n} \left(S_y^2 + \frac{\bar{Y}^2}{\bar{X}^2} S_x^2 - 2 \frac{\bar{Y}}{\bar{X}} S_{xy} \right) + O(n^{-2})$$

$$\text{ECM}[\widehat{Y}_{DP}] = \frac{1-f}{n} \left(S_y^2 + \frac{\bar{Y}^2}{\bar{X}^2} S_x^2 - 2 \frac{\bar{Y}}{\bar{X}} S_{xy} \right) + O(n^{-2})$$

Este teorema asegura que, en aproximación de primer orden, los estimadores de Tin, Beale y De Pascual son igualmente eficientes que el estimador de razón $\widehat{Y}_R = (\bar{y}/\bar{x})\bar{X}$, y que el estimador $\widehat{Y}_J$.

En las referencias citadas, y también en Williams (1961), Rao (1967), Rao y Beegle (1967) y Rao y Rao (1971) se pueden encontrar comparaciones, tanto analíticas como empíricas, entre estos estimadores. Algunas conclusiones que se desprenden de estos estudios comparativos son,

1. Bajo el modelo de proporcionalidad directa considerado en el enfoque predictivo, con función guía de la varianza $v(X_i) = X_i$, y tamaño muestral no muy reducido, es preferible el empleo del estimador clásico de razón, $\widehat{Y}_R$.
2. El sesgo de $\widehat{Y}_T$ es pequeño, y su error cuadrático medio es menor que para el resto de los estimadores, salvo para $\widehat{Y}_R$ con $v(X_i) = X_i$.
3. El estimador de Beale, $\widehat{Y}_B$, no difiere sustancialmente de $\widehat{Y}_T$ salvo que n sea muy pequeño.
4. Si lo importante es la reducción del sesgo, y no necesariamente de error cuadrático medio, $\widehat{Y}_J$ y $\widehat{Y}_M$ han de ser preferidos al resto de los estimadores.
5. El estimador de Hartley-Ross, $\widehat{Y}_{HR}$, puede no ser apropiado en ciertas circunstancias, bajo el modelo de proporcionalidad directa considerado. Hartley y Ross (1954), proponen para su estimador una modificación basada en la partición en grupos de la muestra.

Además de las referencias citadas, véanse también Sukhatme *et al.* (1984), Rao (1988) y Hedayat y Sinha (1991).

5. ESTIMADOR DE RAZÓN MULTIVARIANTE

El estimador de razón, $\widehat{Y}_R = (\bar{y}/\bar{x})\bar{X}$, bajo diseño muestral aleatorio simple, se generaliza, de manera natural, al siguiente **estimador de razón multivariante**, propuesto por Olkin (1958),

$$\widehat{Y}_{MR} = \bar{y} \sum_{i=1}^p w_i \frac{\bar{X}_i}{\bar{x}_i} = \sum_{i=1}^p w_i \widehat{Y}_{Ri} \quad \text{con } \widehat{Y}_{Ri} = \frac{\bar{y}}{\bar{x}_i} \bar{X}_i$$

siendo $X_{i1}, X_{i2}, \dots, X_{iN}$; $i = 1, \dots, p$, p variables auxiliares conocidas, relacionadas con la variable de estudio, Y , con medias poblacionales respectivas $\bar{X}_i$, $i = 1, \dots, p$, y medias muestrales $\bar{x}_i$, $i = 1, \dots, p$. Las cantidades w_i son unos **pesos** que se determinarán en relación a la eficiencia del estimador.

Al ser combinación lineal de estimadores sesgados, $\widehat{Y}_{MR}$ también lo será, y su sesgo estará influenciado por la elección de los pesos. En este sentido, se tiene el siguiente teorema, cuya demostración se basa en los resultados para el sesgo del estimador de razón univariante bajo diseño muestral MAS(N, n).

Teorema 10 *Una condición necesaria y suficiente para que el estimador $\widehat{Y}_{MR}$ posea sesgo de orden $O(n^{-1})$ es que $\sum_{i=1}^P w_i = 1$, siendo en tal caso,*

$$B[\widehat{Y}_{MR}] = \frac{1-f}{n} \sum_{i=1}^P w_i \left(\frac{\bar{Y}}{\bar{X}_i} S_{xi}^2 - \frac{1}{\bar{X}_i} S_{yxi} \right) + O(n^{-2})$$

donde S_{xi}^2 denota la cuasivarianza poblacional de $X_{i1}, \dots, X_{iN}$, y S_{yxi} la correspondiente cuasicovarianza.

En lo que sigue, supondremos ya que $\sum_{i=1}^P w_i = 1$, con objeto de controlar el sesgo.

Para el estudio del error cuadrático medio utilizaremos la siguiente expresión, que se obtiene mediante un cálculo directo,

$$E \left[(\widehat{Y}_{Ri} - \bar{Y})(\widehat{Y}_{Rj} - \bar{Y}) \right] = \frac{1-f}{n} \left(S_y^2 - \frac{\bar{Y}}{\bar{X}_i} S_{yxi} - \frac{\bar{Y}}{\bar{X}_j} S_{yxj} + \frac{\bar{Y}^2}{\bar{X}_i \bar{X}_j} S_{xixj} \right) + O(n^{-2})$$

verificándose el siguiente resultado,

Teorema 11 *El error cuadrático medio del estimador multivariante de razón, $\widehat{Y}_{MR}$, es de orden $O(n^{-1})$, siendo además,*

$$ECM[\widehat{Y}_{MR}] = \frac{1-f}{n} \sum_{i,j=1}^P w_i w_j \left(S_y^2 - \frac{\bar{Y}}{\bar{X}_i} S_{yxi} - \frac{\bar{Y}}{\bar{X}_j} S_{yxj} + \frac{\bar{Y}^2}{\bar{X}_i \bar{X}_j} S_{xixj} \right) + O(n^{-2})$$

Demostración: Al ser $\sum_{i=1}^P w_i = 1$, podemos escribir,

$$\begin{aligned} ECM[\widehat{Y}_{MR}] &= E \left[\left(\sum_{i=1}^P w_i \widehat{Y}_{Ri} - \bar{Y} \right)^2 \right] = E \left[\left(\sum_{i=1}^P w_i (\widehat{Y}_{Ri} - \bar{Y}) \right)^2 \right] \\ &= E \left[\sum_{i=1}^P w_i^2 (\widehat{Y}_{Ri} - \bar{Y})^2 + \sum_{i \neq j=1}^P w_i w_j (\widehat{Y}_{Ri} - \bar{Y}) (\widehat{Y}_{Rj} - \bar{Y}) \right] \\ &= \sum_{i=1}^P w_i^2 ECM[\widehat{Y}_{Ri}] + \sum_{i \neq j=1}^P w_i w_j E \left[(\widehat{Y}_{Ri} - \bar{Y}) (\widehat{Y}_{Rj} - \bar{Y}) \right] \end{aligned}$$

Y basta aplicar los resultados para el error cuadrático medio del estimador de razón univariante, así como la expresión de $E[(\widehat{Y}_{Ri} - \bar{Y})(\widehat{Y}_{Rj} - \bar{Y})]$, para obtener el resultado propuesto. ■

Planteamos ahora el problema de calcular los pesos de forma que el error sea lo menor posible. Para ello trabajaremos con la aproximación de primer orden del error cuadrático medio, es decir,

$$\text{ECM}_1[\widehat{Y}_{MR}] = \frac{1-f}{n} \sum_{i,j=1}^p w_i w_j \left(S_y^2 - \frac{\bar{Y}}{\bar{X}_i} S_{yxi} - \frac{\bar{Y}}{\bar{X}_j} S_{yxj} + \frac{\bar{Y}^2}{\bar{X}_i \bar{X}_j} S_{xixj} \right)$$

que se puede expresar como $\text{ECM}_1[\widehat{Y}_{MR}] = wAw'$ siendo $w = (w_1, \dots, w_p)$, y $A = (a_{ij})$ la matriz de dimensión $p \times p$ definida como,

$$a_{ij} = \frac{1-f}{n} \left(S_y^2 - \frac{\bar{Y}}{\bar{X}_i} S_{yxi} - \frac{\bar{Y}}{\bar{X}_j} S_{yxj} + \frac{\bar{Y}^2}{\bar{X}_i \bar{X}_j} S_{xixj} \right) \quad i, j = 1, \dots, p$$

Observemos que A es simétrica y definida positiva, luego, por la desigualdad de Cauchy generalizada, se tiene,

$$(ab')^2 \leq (aAa')(bA^{-1}b') \quad \forall a, b \in R^p$$

donde la igualdad se da si y sólo si $aA = \alpha b$ siendo α un escalar no nulo. En particular, tomando $a = w$, y $b = e = (1, \dots, 1)$, tendremos,

$$1 = (we')^2 \leq (wAw')(eA^{-1}e')$$

con lo que (wAw') toma el valor mínimo si y sólo si $wA = \alpha e$, es decir, $w = \alpha eA^{-1}$, y a partir de la condición de normalidad de los pesos, se obtiene el siguiente vector de pesos **óptimos**,

$$w_{\text{opt}} = \frac{eA^{-1}}{eA^{-1}e'}$$

y con esta elección, el término guía, es decir, de orden $O(n^{-1})$ del error cuadrático medio resulta ser,

$$\text{ECM}_1[\widehat{Y}_{MR}] = \frac{1}{eA^{-1}e'}$$

Además de la fundamental referencia de Olkin, pueden consultarse Sukhatme *et al.* (1984), donde se realiza un tratamiento exhaustivo del caso $p = 2$, y Cochran (1993).

Hemos de observar también que si la correlación entre la variable Y y las variables auxiliares es positiva para algunas y negativas para otras, se puede plantear un estimador multivariante mixto basado en una combinación lineal de estimadores de razón y producto, propuesto por Rao y Mudholkar (1967), tomando como base las propiedades del estimador de producto de Murthy (1964).

6. ESTIMADOR DE RAZÓN Y ESTRATIFICACIÓN

En este apartado estudiaremos las formas que adopta el estimador usual de razón cuando se utiliza muestreo estratificado, y más concretamente, diseño muestral aleatorio simple estratificado, $MASE(N, n)$, consistente en realizar diferentes muestreos aleatorios simples en cada uno de los estratos, siendo $N = (N_1, \dots, N_L)$ un vector cuyas componentes son los tamaños de los diferentes estratos, y lo mismo para $n = (n_1, \dots, n_L)$, compuesto por los tamaños de muestra en cada estrato.

Dependiendo de la metodología aplicada, obtendremos dos tipos de estimación, **separada** y **combinada**.

6.1. Estimación separada de razón

Teniendo en cuenta la descomposición usual de la media poblacional como combinación lineal de medias en cada estrato,

$$\bar{Y} = \sum_{h=1}^L W_h \bar{Y}_h$$

podemos estimar las medias de cada estrato empleando estimadores de razón, obteniendo el estimador separado,

$$\hat{\bar{Y}}_{RS} = \sum_{h=1}^L W_h \hat{\bar{Y}}_{hR}$$

Obviamente, las propiedades de este estimador dependerán de los estimadores de razón utilizados en los distintos estratos. En particular podrá ser sesgado, insesgado o cuasi-insesgado, y su eficiencia, medida por ejemplo en términos del error cuadrático medio, estará influenciada por la eficiencia de las estimaciones en cada estrato, ya que, por la independencia de las extracciones, se tendrá,

$$ECM[\hat{\bar{Y}}_{RS}] = \sum_{h=1}^L W_h^2 ECM[\hat{\bar{Y}}_{hR}]$$

En particular, si empleamos los estimadores de razón usuales para el muestreo aleatorio,

$$\widehat{Y}_{hR} = (\bar{y}_h / \bar{x}_h) \bar{X}_h = \widehat{R}_h \bar{X}_h \quad h = 1, \dots, L$$

obtendremos,

$$\widehat{Y}_{RS} = \sum_{h=1}^L W_h \frac{\bar{y}_h}{\bar{x}_h} \bar{X}_h$$

que será sesgado, siendo su sesgo,

$$\begin{aligned} B[\widehat{Y}_{RS}] &= E \left[\sum_{h=1}^L W_h \frac{\bar{y}_h}{\bar{x}_h} \bar{X}_h \right] - \sum_{h=1}^L W_h \bar{Y}_h \\ &= E \left[\sum_{h=1}^L W_h \left(\frac{\bar{y}_h}{\bar{x}_h} \bar{X}_h - \bar{Y}_h \right) \right] = \sum_{h=1}^L W_h \left(E[\widehat{R}_h] E[\bar{x}_h] - E[\widehat{R}_h \bar{x}_h] \right) \\ &= - \sum_{h=1}^L W_h \text{Cov}[\widehat{R}_h, \bar{x}_h] \end{aligned}$$

y suponiendo para los coeficientes de variación de $\bar{x}_h$ la acotación común $\text{CV}[\bar{x}_h] = \sigma[\bar{x}_h] / E[\bar{x}_h] \leq C_0$, $\forall h$, obtenemos,

$$\left| B[\widehat{Y}_{RS}] \right| \leq \sum_{h=1}^L W_h \sigma[\widehat{Y}_{hR}] \text{CV}[\bar{x}_h] \leq C_0 \sum_{h=1}^L W_h \sigma[\widehat{Y}_{hR}]$$

que proporciona la siguiente cota,

$$\frac{\left| B[\widehat{Y}_{RS}] \right|}{\sigma[\widehat{Y}_{RS}]} \leq C_0 \sqrt{L}$$

Podemos pues afirmar que si el número de estratos es elevado, el sesgo de la estimación puede llegar a ser apreciable. Ello sugiere el empleo de estimadores insesgados o cuasi-insesgados, sobre todo si los tamaños de la muestra en los diferentes estratos no son muy elevados.

Con respecto al error cuadrático medio de la estimación anterior, tendremos, en primer orden de aproximación,

$$\text{ECM}_1[\widehat{Y}_{RS}] = \sum_{h=1}^L W_h^2 \left(\frac{1}{n_h} - \frac{1}{N_h} \right) (S_{yh}^2 + R_h^2 S_{xh}^2 - 2R_h S_{xyh})$$

6.2. Estimación combinada de razón

Otra posibilidad para reducir el sesgo es emplear el estimador combinado de razón, propuesto por Hansen, Hurwitz y Gurney (1946), y definido como,

$$\widehat{Y}_{RC} = \frac{\sum_{h=1}^L W_h \bar{y}_h}{\sum_{h=1}^L W_h \bar{x}_h} \bar{X} = \frac{\bar{y}_{\text{est}}}{\bar{x}_{\text{est}}} \bar{X} = \widehat{R}_{\text{est}} \bar{X}$$

Este estimador también es sesgado, siendo ahora el sesgo,

$$\begin{aligned} B[\widehat{Y}_{RC}] &= E[\widehat{R}_{\text{est}}] \bar{X} - \bar{Y} \\ &= E[\widehat{R}_{\text{est}}] E[\bar{x}_{\text{est}}] - E[\widehat{R}_{\text{est}} \bar{x}_{\text{est}}] \\ &= -\text{Cov}[\widehat{R}_{\text{est}}, \bar{x}_{\text{est}}] \end{aligned}$$

de donde se obtiene,

$$\frac{|B[\widehat{Y}_{RC}]|}{\sigma[\widehat{Y}_{RC}]} \leq \text{CV}[\bar{x}_{\text{est}}]$$

es decir, si el coeficiente de variación de $\bar{x}_{\text{est}}$ es pequeño, el sesgo puede ser despreciable.

Para el error cuadrático medio del estimador combinado tendremos, en primer orden de aproximación,

$$\text{ECM}_1[\widehat{Y}_{RC}] = \sum_{h=1}^L W_h^2 \left(\frac{1}{n_h} - \frac{1}{N_h} \right) (S_{yh}^2 + R^2 S_{xh}^2 - 2R S_{xyh})$$

Es interesante comparar este error cuadrático medio con el correspondiente del estimador separado. De esta forma, es inmediato obtener,

$$\begin{aligned} \text{ECM}_1[\widehat{Y}_{RC}] - \text{ECM}_1[\widehat{Y}_{RS}] &= \sum_{h=1}^L W_h^2 \left(\frac{1}{n_h} - \frac{1}{N_h} \right) (S_{xh}^2 (R - R_h)^2 \\ &\quad + 2(R - R_h)(R_h S_{xh}^2 - S_{xyh})) \end{aligned}$$

Como puede verse, la diferencia depende de la variabilidad de las razones, R_h , en los diferentes estratos, y también de las cantidades $(R_h S_{xh}^2 - S_{xyh})$. Estas últimas serán, usualmente, pequeñas, pudiendo ser despreciables si el ajuste al modelo de proporcionalidad directa es muy satisfactorio en todos los estratos. Así pues, el estimador separado resulta ser más eficiente que el combinado, a menos que las razones R_h presenten muy poca variabilidad a lo largo de los estratos. Véanse Sukhatme *et al.* (1984), Cochran (1993) y para ejemplos numéricos, Fernández y Mayor (1995).

7. OPTIMALIDAD DEL ESTIMADOR DE RAZÓN

Ya hemos visto como, en general, los estimadores de razón no son insesgados **en el diseño**, esto es, no tienen por qué cumplir,

$$E_d [\widehat{Y}_R] = \sum_{m \in M} p(m) \widehat{Y}_R(m) = \bar{Y}$$

lo que, en principio, los descarta como candidatos a ser considerados óptimos. No obstante, sí es posible considerar, bajo un determinado modelo de superpoblación, la denominada insesgadez **en el modelo**, según la definición que exponemos a continuación.

Definición Dado un diseño muestral, $d = (M, p(\cdot))$, y un modelo de superpoblación, S , diremos que el estimador de $\theta(Y)$, $\widehat{\theta}$, es insesgado en el modelo S , si cumple,

$$E_s [\widehat{\theta} - \theta(Y)] = 0 \quad \forall m \in M$$

Por ejemplo, si consideramos el modelo de proporcionalidad directa ya empleado para obtener el estimador de razón genérico,

$$Y_i = \beta X_i + \varepsilon_i$$

$$E_s[\varepsilon_i] = 0$$

$$V_s[\varepsilon_i] = \sigma^2 v(X_i)$$

$$E_s[\varepsilon_i \varepsilon_j] = 0, \quad i \neq j$$

se tendrá, con $v(x) = x$,

$$E_s \left[\widehat{Y}_R - \bar{Y} \right] = E_s \left[\frac{\sum_{i \in m} Y_i / \pi_i}{\sum_{i \in m} X_i / \pi_i} \bar{X} - \bar{Y} \right] = \frac{\sum_{i \in m} \beta X_i / \pi_i}{\sum_{i \in m} X_i / \pi_i} \bar{X} - \beta \bar{X} = 0$$

es decir, el estimador de razón genérico es insesgado en el modelo, aunque, en general, no lo sea en el diseño. Es obvio que este nuevo concepto de insesgadez se puede considerar más débil que el usual, en el sentido de que está supeditado a que el modelo de superpoblación sea el adecuado. Observemos también que, al igual que un estimador puede ser insesgado en el modelo pero no en el diseño, es posible que sea insesgado en el diseño pero no en el modelo. Por ejemplo, en el diseño MAS(N, n), la media muestral, $\bar{y}$ es un estimador insesgado (en el diseño) de $\bar{Y}$, sin embargo, para el modelo de proporcionalidad directa se tiene,

$$E_s[\bar{y} - \bar{Y}] = \beta \bar{x} - \beta \bar{X} = \beta(\bar{x} - \bar{X})$$

que, en general, es distinto de cero.

Bajo este enfoque, que podemos denominar **dependiente del modelo** es posible obtener ciertas propiedades de optimalidad sobre el estimador de razón, que estudiaremos a continuación, siguiendo la línea iniciada en Royall (1970).

Para simplificar la notación, supondremos que el parámetro a estimar es el total poblacional, $T(Y) = \sum_{i \in U} Y_i$, y que el estimador, $\widehat{T}(Y)$, se va a buscar en la clase de los estimadores lineales e insesgados en el modelo.

Observemos que, dada una muestra, m , perteneciente al diseño muestral que estemos empleando, $\widehat{T}(Y)$ admite la descomposición,

$$\widehat{T}(Y) = \sum_{i \in m} Y_i + \frac{\widehat{T}(Y) - \sum_{i \in m} Y_i}{\sum_{i \in U-m} X_i} \sum_{i \in U-m} X_i = \sum_{i \in m} Y_i + \widehat{\beta} \sum_{i \in U-m} X_i$$

donde la notación introducida se justifica al ser $E_s[\widehat{\beta}] = \beta$, lo que se comprueba fácilmente. Además, al ser $\widehat{T}(Y)$ lineal en Y_i , $i \in m$, también lo será $\widehat{\beta}$, es decir,

$$\widehat{\beta} = \sum_{i \in m} \alpha_i Y_i$$

Dado un diseño muestral, $d = (M, p(\cdot))$, podemos utilizar la siguiente cantidad,

$$\text{ECM}[\widehat{T}(Y), d] = E_s \text{ECM}_d[\widehat{T}(Y)] = E_s \left[\sum_{m \in M} \left(\widehat{T}(m, Y) - T(Y) \right)^2 p(m) \right]$$

como medida de la eficiencia del estimador $\widehat{T}(Y)$. Esta medida involucra tanto el efecto del diseño muestral, como la influencia del modelo de superpoblación, y en relación a la misma, se tiene el siguiente resultado,

Lema Para un diseño muestral cualquiera, $d = (M, p(\cdot))$, y los estimadores del total,

$$\begin{aligned}\hat{T}'(Y) &= \sum_{i \in m} Y_i + \hat{\beta}' \sum_{i \in U-m} X_i \\ \hat{T}''(Y) &= \sum_{i \in m} Y_i + \hat{\beta}'' \sum_{i \in U-m} X_i\end{aligned}$$

se verifica que si,

$$E_s \left[\left(\hat{\beta}' - \beta \right)^2 \right] \leq E_s \left[\left(\hat{\beta}'' - \beta \right)^2 \right] \quad \forall m \in M$$

entonces,

$$\text{ECM} \left[\hat{T}'(Y), d \right] \leq \text{ECM} \left[\hat{T}''(Y), d \right]$$

Demostración: Teniendo en cuenta el modelo de superpoblación, podemos realizar el siguiente desarrollo,

$$\begin{aligned}\text{ECM} \left[\hat{T}'(Y), d \right] &= E_s \left[\sum_{m \in M} \left(\hat{T}(m, Y) - T(Y) \right)^2 p(m) \right] \\ &= \sum_{m \in M} E_s \left[\left(\hat{T}(m, Y) - T(Y) \right)^2 \right] p(m) \\ &= \sum_{m \in M} E_s \left[\left(\sum_{i \in m} Y_i + \hat{\beta}' \sum_{i \in U-m} X_i - \sum_{i \in m} Y_i - \sum_{i \in U-m} Y_i \right)^2 \right] p(m) \\ &= \sum_{m \in M} E_s \left[\left(\hat{\beta}' \sum_{i \in U-m} X_i - \beta \sum_{i \in U-m} X_i - \sum_{i \in U-m} \varepsilon_i \right)^2 \right] p(m) \\ &= \sum_{m \in M} \left(\left(\sum_{i \in U-m} X_i \right)^2 E_s \left[\left(\hat{\beta}' - \beta \right)^2 \right] + \sigma^2 \sum_{i \in U-m} v(X_i) \right) p(m)\end{aligned}$$

de donde se deduce inmediatamente el resultado propuesto. ■

Con respecto a la optimalidad del estimador de razón, se verifica el siguiente resultado fundamental, en cuya demostración emplearemos el anterior lema.

Teorema 12 Dado un diseño muestral cualquiera, $d = (M, p(\cdot))$, sea,

$$\hat{\beta}^* = \frac{\sum_{i \in m} X_i Y_i / v(X_i)}{\sum_{i \in m} X_i^2 / v(X_i)}$$

y el estimador del total,

$$\hat{T}^*(Y) = \sum_{i \in m} Y_i + \hat{\beta}^* \sum_{i \in U-m} X_i$$

Entonces se verifica que para cualquier estimador del total, $\hat{T}(Y)$, lineal e insesgado en el modelo,

$$\text{ECM}[\hat{T}^*(Y), d] \leq \text{ECM}[\hat{T}(Y), d]$$

siendo además,

$$\text{ECM}[\hat{T}^*(Y), d] = \sigma^2 \sum_{m \in M} \left(\frac{(\sum_{i \in U-m} X_i)^2}{\sum_{i \in m} X_i^2 / v(X_i)} + \sum_{i \in U-m} v(X_i) \right) p(m)$$

Demostración: Para la demostración, nos basaremos en el lema anterior. Sea pues m una muestra del diseño, y $\hat{\beta}$ lineal en Y_i , $i \in m$, esto es,

$$\hat{\beta} = \sum_{i \in m} \alpha_i Y_i$$

A partir de la condición $E_s[\hat{\beta}] = \beta$ se obtiene la restricción,

$$\sum_{i \in m} \alpha_i X_i = 1$$

Por otra parte,

$$\begin{aligned} E_s \left[(\hat{\beta} - \beta)^2 \right] &= E_s \left[\left(\sum_{i \in m} \alpha_i (Y_i - \beta X_i) \right)^2 \right] \\ &= E_s \left[\sum_{i, j \in m} \alpha_i \alpha_j (Y_i - \beta X_i)(Y_j - \beta X_j) \right] \\ &= E_s \left[\sum_{i \in m} \alpha_i^2 (Y_i - \beta X_i)^2 \right] = \sigma^2 \sum_{i \in m} \alpha_i^2 v(X_i) \end{aligned}$$

que, con la restricción ya obtenida anteriormente,

$$\sum_{i \in m} \alpha_i X_i = 1$$

alcanza el mínimo para los valores,

$$\alpha_i = \frac{X_i / v(X_i)}{\sum_{i \in m} X_i^2 / v(X_i)} \quad i \in m$$

lo que proporciona precisamente $\hat{\beta}^*$. Hemos demostrado pues que $\forall m \in M$,

$$E_s \left[\left(\hat{\beta}^* - \beta \right)^2 \right] \leq E_s \left[\left(\hat{\beta} - \beta \right)^2 \right] \quad \forall \hat{\beta}$$

y basta tener en cuenta el anterior lema para obtener el resultado de optimalidad enunciado.

Finalmente, la expresión para $\text{ECM}[\hat{T}^*(Y), d]$ se obtiene mediante un cálculo directo. ■

Observemos que para el caso $v(x) = x$, ya considerado en el enfoque predictivo aplicado al principio del tema, se obtiene como estimador óptimo,

$$\hat{T}^*(Y) = \sum_{i \in m} Y_i + \frac{\sum_{i \in m} Y_i X_i / X_i}{\sum_{i \in m} X_i^2 / X_i} \sum_{i \in U-m} X_i = \frac{\bar{y}}{\bar{x}} T(X)$$

Este resultado no está en contradicción con las buenas propiedades que presenta el estimador genérico de razón obtenido en la sección 3. de esta revisión,

$$\hat{T}_R(Y) = \frac{\sum_{i \in m} Y_i / \pi_i}{\sum_{i \in m} X_i / \pi_i} T(X)$$

en lo que respecta a la estimación del error. Por otra parte, hay que tener en cuenta que el estimador óptimo ha sido buscado en una clase muy especial de estimadores lineales, **insesgados para el modelo**, con los inconvenientes de dependencia del mismo que ello supone.

Observemos también que para el caso $v(x) = x$, el error cuadrático medio resulta ser,

$$\begin{aligned} \text{ECM}[\hat{T}^*(Y), d] &= \sigma^2 \sum_{m \in M} \left(\frac{(\sum_{i \in U-m} X_i)^2}{\sum_{i \in m} X_i^2 / X_i} + \sum_{i \in U-m} X_i \right) p(m) \\ &= \sigma^2 T(X) E_d \left[\frac{\sum_{i \in U-m} X_i}{\sum_{i \in m} X_i} \right] \end{aligned}$$

Y si denotamos por m^* la muestra formada por los elementos de la población, tal que,

$$\sum_{i \in m^*} X_i = \max_{m \in M} \left\{ \sum_{i \in m} X_i \right\}$$

el anterior error cuadrático medio será mínimo para el siguiente diseño muestral **intencional**,

$$d^* = (\{m^*\}, p(m^*) = 1)$$

es decir, un diseño con una única muestra, m^* , que es seleccionada con probabilidad uno. Hemos obtenido pues el siguiente resultado,

Teorema 13 *Bajo el modelo de superpoblación de proporcionalidad directa, con función guía de la varianza $v(x) = x$, la estrategia muestral $(d^*, (\bar{y}/\bar{x})T(X))$ es óptima para la estimación de $T(Y)$.*

Notemos que el resultado anterior sigue siendo válido si solamente exigimos que $v(x)$ sea no decreciente, y $v(x)/x^2$ no creciente.

A pesar de su importancia teórica, estos resultados, como afirman Cassel, Särndal y Wretman (1977), están en conflicto con uno de los principios más extendidos de la Estadística como es el de la aleatorización. Por otra parte, la estrategia óptima anterior es incompatible con el cálculo o la estimación del error.

Como fuentes importantes para profundizar en estas cuestiones, citaremos Royall (1970), Royall y Herson (1973a, 1973b), Royall y Eberhardt (1975), Cassel, Särndal y Wretman (1977), Bellhouse (1984), Sukhatme *et al.* (1984) y Chaudhuri y Vos (1988).

Observemos finalmente que es posible plantear el problema complementario de hallar las probabilidades de inclusión óptimas para que el estimador de razón genérico,

$$\hat{T}_R(Y) = \frac{\sum_{i \in m} Y_i / \pi_i}{\sum_{i \in m} X_i / \pi_i} T(X)$$

sea óptimo bajo el modelo de superpoblación de proporcionalidad directa. Véase Särndal, Swensson y Wretman (1992) para un estudio del mismo.

8. REFERENCIAS

- [1] **Azorín, F. y Sánchez-Crespo, J.L.** (1986). *Métodos y Aplicaciones del Muestreo*. Alianza Universidad Textos. Madrid.

- [2] **Beale, E.M.L.** (1962). «Some use of computers in operational research». *Industrielle Organization*, **31**, 27–28.
- [3] **Bellhouse, D.R.** (1984). «A review of optimal designs in survey sampling». *The Canadian Journal of Statistics*, **12**, 53–65.
- [4] **Cassel, C., Särndal, C. y Wretman, J.** (1977). *Foundations of Inference in Survey Sampling*. Wiley. New York.
- [5] **Cochran, W.G.** (1978). «Laplace's Ratio Estimator». *Contributions to Survey Sampling and Applied Statistics*. H.A. David (ed.). Academic Press. New York.
- [6] **Cochran, W.G.** (1993). *Técnicas de Muestreo*. Décima reimpresión. CECSA. México.
- [7] **Chang, W.C.** (1976). «Statistical theories and sampling practice». *On the History of Statistics and Probability*. D.B. Owen (ed.). Dekker. New York.
- [8] **Chaudhuri, A. y Vos, J.** (1988). *Unified Theory and Strategies of Survey Sampling*. North Holland. Amsterdam.
- [9] **David, I.P. y Sukhatme, B.V.** (1974). «On the bias and mean square error of the ratio estimator». *J. Amer. Statist. Assoc.*, **69**, 464–466.
- [10] **De Pascual, N.** (1961). «Unbiased ratio estimators in stratified sampling». *J. Amer. Statist. Assoc.*, **56**, 70–87.
- [11] **Fernández, F.R. y Mayor, J.A.** (1995). *Muestreo en Poblaciones Finitas: Curso Básico*. E.U.B. Barcelona. (También en P.P.U., (1994)).
- [12] **Fuller, W.A.** (1975). «Regression analysis for sample survey». *Sankhyā*, **C37**, 117–132.
- [13] **Goodman, L.A. y Hartley, H.O.** (1958). «The precision of unbiased ratio-type estimators». *J. Amer. Statist. Assoc.*, **53**, 491–508.
- [14] **Hald, A.** (1990). *A History of Probability and Statistics and Their Applications before 1750*. Wiley. New York.
- [15] **Hansen, M.H., Hurwitz, W.N. y Gurney, M.** (1946). «Problems and methods of the sample survey of business». *J. Amer. Statist. Assoc.*, **41**, 173–189.
- [16] **Hansen, M.H., Hurwitz, W.N. y Madow, W.G.** (1953). *Sample Survey Methods and Theory. Vol. I y II*. Wiley. New York.
- [17] **Hartley, H.O. y Ross, A.** (1954). «Unbiased ratio estimators». *Nature*, **174**, 270–271.
- [18] **Hedayat, A.S. y Sinha, B.K.** (1991). *Design and Inference in Finite Population Sampling*. Wiley. New York.
- [19] **Kish, L. y Frankel, M.R.** (1974). «Inference from complex samples». *J. Roy. Statist. Soc.*, **B36**, 1–37.

- [20] **Lahiri, D.B.** (1951). «A method of sample selection providing unbiased ratio estimates». *Bulletin of the International Statistical Institute*, **33**, 133–140.
- [21] **Lanke, J.** (1974). «On nonnegative variance estimators in survey sampling». *Sankhyā*, **C36**, 33–42.
- [22] **Mickey, M.R.** (1959). «Some finite population unbiased ratio and regression estimators». *J. Amer. Statist. Assoc.*, **54**, 594–612.
- [23] **Midzuno, H.** (1952). «On the sampling system with probability proportionate to sum of sizes». *Annals of the Institute of Statistical Mathematics*, **3**, 99–107.
- [24] **Murthy, M.N.** (1964). «Product method of estimation». *Sankhyā*, **A26**, 69–74.
- [25] **Olkin, I.** (1958). «Multivariate ratio estimation for finite populations». *Biometrika*, **45**, 154–165.
- [26] **Quenouille, M.H.** (1956). «Notes on bias in estimation». *Biometrika*, **43**, 353–360.
- [27] **Rao, C.R.** (1965). *Linear Statistical Inference and its Applications*. Wiley. New York.
- [28] **Rao, J.N.K.** (1967). «The precision of Mickey's unbiased ratio estimator». *Biometrika*, **54**, 321–324.
- [29] **Rao, J.N.K. y Beegle, L.D.** (1977). «A Monte Carlo study of some ratio estimators». *Sankhyā*, **B29**, 47–56.
- [30] **Rao, J.N.K. y Vijayan, K.** (1977). «On estimating the variance in sampling with probability proportional to aggregate size». *J. Amer. Statist. Assoc.*, **72**, 579–584.
- [31] **Rao, P.S.R.S.** (1971). «Small sample results for ratio estimators». *Biometrika*, **58**, 625–630.
- [32] **Rao, P.S.R.S.** (1988). «Ratio and regression estimators». *Handbook of Statistics 6. Sampling*. Krishnaiah y Rao, (Eds.). North Holland. Amsterdam.
- [33] **Rao, P.S.R.S. y Mudholkar, G.S.** (1967). «Generalized multivariate estimator for the mean of finite populations». *J. Amer. Statist. Assoc.*, **62**, 1009–1012.
- [34] **Rao, T.J.** (1966). «On the variance of the ratio estimator for Midzuno-Sen sampling scheme». *Metrika*, **10**, 89–91.
- [35] **Rao, T.J.** (1972). «On the variance of the ratio estimator». *Metrika*, **18**, 209–215.
- [36] **Robson, D.S.** (1957). «Application of multivariate polykays to the theory of unbiased ratio-type estimators». *J. Amer. Statist. Assoc.*, **52**, 511–522.
- [37] **Royall, R.M.** (1970). «On finite population sampling theory under certain linear regression models». *Biometrika*, **57**, 377–387.

- [38] **Royall, R.M. y Herson, J.** (1973a). «Robust estimation in finite populations I». *J. Amer. Statist. Assoc.*, **68**, 880–890.
- [39] **Royall, R.M. y Herson, J.** (1973b). «Robust estimation in finite populations II». *J. Amer. Statist. Assoc.*, **68**, 890–894.
- [40] **Royall, R.M. y Eberhardt, K.R.** (1975). «Variance estimates for the ratio estimator». *Sankhyā*, **C37**, 43–52.
- [41] **Royall, R.M. y Cumberland, W.G.** (1981). «An empirical study of the ratio estimator and estimators of its variance». *J. Amer. Statist. Assoc.*, **76**, 66–77.
- [42] **Ruiz, M. y Santos, J.** (1989). «Unbiased mean-of-the-ratios estimators». *Statistica*, **49**, 617–622.
- [43] **Sahoo, L.N. y Ruiz, M.** (1994). «Unbiased estimators using auxiliary information in sample surveys: a review». *Revista de la Academia de Ciencias de Zaragoza*, **49**, 137–146.
- [44] **Sánchez-Crespo, J.L.** (1980). *Curso Intensivo de Muestreo en Poblaciones Finitas*. 2ª edición. Instituto Nacional de Estadística. Madrid.
- [45] **Särndal, C.** (1984). *Inférence Statistique et Analyse des Données sous des Plans d'Échantillonnage Complexes*. Presses de l'Université de Montréal.
- [46] **Särndal, C., Swensson, B. y Wretman, J.** (1992). *Model Assisted Survey Sampling*. Springer-Verlag. New York.
- [47] **Singh, P. y Srivastava, A.K.** (1980). «Sampling schemes providing unbiased regression estimators». *Biometrika*, **67**, 205–209.
- [48] **Sukhatme, P.V., Sukhatme, B.V., Sukhatme, S. y Asok, C.** (1984). *Sampling Theory of Surveys Applications*. Tercera edición. Iowa State University Press. Ames. Iowa.
- [49] **Tin, M.** (1965). «Comparison of some ratio estimators». *J. Amer. Statist. Assoc.*, **60**, 294–307.
- [50] **Williams, W.H.** (1961). «Generating unbiased ratio and regression estimators». *Biometrics*, **17**, 267–274.

ENGLISH SUMMARY

RATIO ESTIMATORS: A REVIEW

JOSÉ A. MAYOR GALLEGO*

Universidad de Sevilla

In this review we expose the basic principles relating to the use of auxiliary information, in order to estimate linear parameters defined over finite populations, by means of ratio type estimators. The construction of these estimators is firstly carried out by means of an heuristic approach based on the existence of a direct proportionality relation between the study variable and the auxiliary variable, but a more formal study is carried out under a superpulation model approach. The main problem of these estimators is the existence of bias, and in order to reduce it, we have to use special sampling designs or to modify the structure of the estimators to obtain unbiased or almost unbiased strategies.

Keywords: Sampling, finite populations, ratio-type estimators

AMS Classification: 62D05

*José A. Mayor Gallego. Dpto. Estadística e Investigación Operativa. Universidad de Sevilla. C/Tarfia s/n. 41012 SEVILLA.

–Received may 1996.

–Accepted april 1997.

A way to integrate the auxiliary information, in order to increase the accuracy of the estimates over a finite population is to use estimators with a rational mathematical structure, involving the study variable, Y , and an auxiliary variable, X , completely known.

In this paper, we **review** this class of estimators, emphasizing the importance of the proportionality between the X and Y variables, in order to obtain good estimates, but controlling the bias.

We are going to suppose that we are estimating the linear parameter,

$$\theta(Y) = \sum_{i \in U} a_i Y_i$$

over a finite population, $U = \{1, 2, \dots, N\}$, by means of a sample, s , chosen from a sampling design, d .

Thus we study firstly the **heuristic approach**, based on the following factorizing, (for the population mean),

$$\bar{Y} = \frac{\bar{Y}}{\bar{X}} \bar{X} = R \bar{X}$$

then, we can define the estimator,

$$\widehat{\bar{Y}}_R = \widehat{R} \bar{X}$$

The final form of this estimator depends on the sampling design. If we use the simple random sampling, $\text{SRS}(N, n)$, we can estimate R using the sample means ratio, that is to say,

$$\widehat{\bar{Y}}_R = \frac{\bar{y}}{\bar{x}} \bar{X}$$

The efficiency of this estimator depends on the relation between the study variable, Y , and the auxiliary variable, X . It has a good behaviour if there is an approximate relation, $Y \approx \alpha X$, that is to say, a direct proportionality. Using approximate techniques, we obtain the following expressions for the variance and its estimation,

$$V[\widehat{\bar{Y}}_R] \approx \frac{1-f}{n} (S_y^2 + R^2 S_x^2 - 2RS_{xy})$$

$$\widehat{V}[\widehat{\bar{Y}}_R] = \frac{1-f}{n} (s_y^2 + \widehat{R}^2 s_x^2 - 2\widehat{R}s_{xy})$$

where S_y^2 , S_x^2 , S_{xy} are the population quasivariances and quasicovariance of the Y values, and s_y^2 , s_x^2 , s_{xy} the corresponding over the sample. Also, we denote $f = n/N$.

An alternative and more formal approach is the **predictive approach**, based on the existence of the following superpopulation model,

$$\begin{aligned} Y_i &= \beta X_i + \varepsilon_i \\ E_s[\varepsilon_i] &= 0 \\ V_s[\varepsilon_i] &= \sigma^2 v(X_i) \\ E_s[\varepsilon_i \varepsilon_j] &= 0, \quad i \neq j \end{aligned}$$

where $v(\cdot)$ is a known function. We suppose that X only takes positive values. Under this model, and if $v(x) = x$, we obtain the estimator,

$$\widehat{Y}_R = \frac{\sum_{i \in s} Y_i / \pi_i}{\sum_{i \in s} X_i / \pi_i} \bar{X}$$

particular case of Sánchez-Crespo (1980).

The properties of this estimator depends on the sampling design. If the sample, s , is obtained using a SRS(N, n), then, the estimator is,

$$\widehat{Y}_R = \frac{\bar{y}}{\bar{x}} \bar{X}$$

This estimator is biased, since,

$$B[\widehat{Y}_R] = E[\widehat{Y}_R] - \bar{Y} = \bar{X} E \left[\frac{\bar{y}}{\bar{x}} \right] - E[\bar{y}] = E[\bar{x}] E \left[\frac{\bar{y}}{\bar{x}} \right] - E[\bar{y}] = -\text{Cov} \left[\bar{x}, \frac{\bar{y}}{\bar{x}} \right]$$

and defining the quantities,

$$\begin{aligned} B_k[\widehat{Y}_R] &= \bar{Y} E \left[\sum_{i=0}^{2k-1} (-\delta \bar{x})^i (\delta \bar{y} - \delta \bar{x}) \right] \\ \text{ECM}_k[\widehat{Y}_R] &= \bar{Y}^2 E \left[(\delta \bar{y} - \delta \bar{x})^2 \sum_{i=0}^{2k-2} (i+1) (-\delta \bar{x})^i \right] \end{aligned}$$

we have the following result, about the bias and the mean square error,

$$\begin{aligned} \left| B[\widehat{Y}_R] - B_k[\widehat{Y}_R] \right| &\leq O(n^{-(k+1)}) \\ \left| \text{ECM}[\widehat{Y}_R] - \text{ECM}_k[\widehat{Y}_R] \right| &\leq O(n^{-(k+1)}) \end{aligned}$$

If $k = 1$, we obtain the first order approximations,

$$\begin{aligned} B[\widehat{Y}_R] &= \frac{1-f}{n} \left(\frac{\bar{Y}}{\bar{X}^2} S_x^2 - \frac{1}{\bar{X}} S_{xy} \right) + O(n^{-2}) = O(n^{-1}) \\ \text{ECM}[\widehat{Y}_R] &= \frac{1-f}{n} \left(S_y^2 + \frac{\bar{Y}^2}{\bar{X}^2} S_x^2 - 2 \frac{\bar{Y}}{\bar{X}} S_{xy} \right) + O(n^{-2}) = O(n^{-1}) \\ V[\widehat{Y}_R] &= \frac{1-f}{n} \left(S_y^2 + \frac{\bar{Y}^2}{\bar{X}^2} S_x^2 - 2 \frac{\bar{Y}}{\bar{X}} S_{xy} \right) + O(n^{-2}) = O(n^{-1}) \end{aligned}$$

Thus, the bias and the mean square error can be controlled, increasing the sample size.

If we use a sampling design with first order inclusion probabilities proportional to size, X , we obtain the unbiased estimation,

$$\widehat{Y}_R = \frac{1}{n} \sum_{i \in s} \frac{Y_i}{X_i} \bar{X}$$

with variance (for fixed sample size),

$$\begin{aligned} V[\widehat{Y}_R] &= -\frac{\bar{X}^2}{2n^2} \sum_{i,j \in U} (\pi_{ij} - \pi_i \pi_j) \left(\frac{Y_i}{X_i} - \frac{Y_j}{X_j} \right)^2 \\ \widehat{V}[\widehat{Y}_R] &= -\frac{\bar{X}^2}{2n^2} \sum_{i,j \in s} \frac{(\pi_{ij} - \pi_i \pi_j)}{\pi_{ij}} \left(\frac{Y_i}{X_i} - \frac{Y_j}{X_j} \right)^2 \end{aligned}$$

that is to say, the variance decreases if the Y and X variables are proportional or nearly proportional.

Other unbiased strategies are obtained combining the ratio estimator,

$$\widehat{Y}_R = \frac{\bar{y}}{\bar{x}} \bar{X}$$

with the sampling design generated by the Lahiri-Midzuno scheme.

The Harley-Ross estimator,

$$\widehat{Y}_{HR} = \bar{z} \bar{X} + \frac{n(N-1)}{N(n-1)} (\bar{y} - \bar{z} \bar{x}), \quad Z_i = Y_i/X_i$$

with the SRS(N, n) design.

The Mickey's estimator with the SRS(N, n) design, and others related by Ruiz and Santos (1989).

Also, is possible to obtain almost unbiased strategies, that is to say, with bias $O(1/n^2)$, for example,

- The De Pascual's estimator,

$$\widehat{Y}_{DP} = \widehat{Y}_R + \frac{\bar{y} - \bar{z}\bar{x}}{n-1} \quad \text{with } Z_i = \frac{Y_i}{X_i}$$

- The Beale's estimator,

$$\widehat{Y}_B = \widehat{Y}_R \left[1 + \left(\frac{1}{n} - \frac{1}{N} \right) \frac{s_{xy}}{\bar{x}\bar{y}} \right] \left[1 + \left(\frac{1}{n} - \frac{1}{N} \right) \frac{s_x^2}{\bar{x}^2} \right]^{-1}$$

- The Tin's estimator,

$$\widehat{Y}_T = \widehat{Y}_R \left[1 - \left(\frac{1}{n} - \frac{1}{N} \right) \left(\frac{s_x^2}{\bar{x}^2} - \frac{s_{xy}}{\bar{x}\bar{y}} \right) \right]$$

in combination with the SRS(N, n) design. And also, the estimator obtained applying the jackknife technique to the ratio estimator, $\widehat{Y}_R = (\bar{y}/\bar{x})\bar{X}$.

To finish this review, we study the multivariate ratio estimator, the ratio estimator in stratified sampling and the optimality of the ratio estimator.

Investigació Operativa

HEURÍSTICO PARA LOS PROBLEMAS DE RUTAS CON CARGA Y DESCARGA EN SISTEMAS LIFO

JOAQUÍN A. PACHECO*

E.U.E. Empresariales de Burgos

En este trabajo se propone un algoritmo heurístico para el «Problema de Carga y Descarga (PDP) con un solo vehículo sin restricciones de capacidad en sistemas de descarga LIFO» —es decir, en cada momento sólo se puede descargar la última mercancía que ha entrado en el vehículo de entre todas las que se encuentran en él—. Este algoritmo es una extensión y adaptación del método de Or para el Problema del Viajante (TSP) que sirve también para matrices asimétricas. Con este heurístico se consiguen resolver problemas de gran tamaño en un tiempo de computación razonable en ordenadores personales, con una desviación del óptimo muy pequeña.

Heuristic for the pick-up and delivery routing problems in lifo unloading systems

Keywords: Carga y descarga con sistema LIFO; intercambios de Or; chequeo de factibilidad.

*Joaquín A. Pacheco. Dpto. Matemáticas. E.U.E. Empresariales de Burgos. Fco. Vitoria s/n. 09006 Burgos.

—Article rebut el març de 1995.

—Acceptat el juliol de 1996.

1. INTRODUCCIÓN

El Problema de Carga y Descarga (PDP) puede ser descrito de la forma siguiente: Sea un conjunto de clientes $N = \{2, 3, \dots, n\}$; cada cliente r requiere que le sea transportada una mercancía desde un origen r^+ a un destino r^- . Para ello se dispone de m vehículos que parten de una ciudad inicial 1, a la que regresan al final del trayecto. Se trata de diseñar rutas de vehículos con distancia total a recorrer mínima. Se va a denotar $N^+ = \{r^+ / r = 2, \dots, n\}$, el conjunto de puntos de carga y $N^- = \{r^- / r = 2, \dots, n\}$, el conjunto de puntos de descarga, y $q(r)$, $r = 2, \dots, n$, la mercancía de cada cliente r . Al caso especial en el que las demandas requeridas son todas iguales se le denomina el *Dial-A-Ride-Problem* (DARP), que tiene lugar en el transporte de viajeros, transporte escolar, etc.

El PDP y el PDP con Ventanas de Tiempo (PDPTW) son problemas de rutas muy estudiados, especialmente a partir de la segunda mitad de la década de los ochenta. Un repaso a los principales métodos de solución a estos problemas, tanto exactos como heurísticos, se puede encontrar en el trabajo de Pacheco (1994).

Entre los algoritmos exactos destacan inicialmente el de Kalantari y otros (1985), para el PDP con un solo vehículo sin restricciones de capacidad, que proponen una serie de modificaciones en el algoritmo de Little y otros (1963), para el TSP, con el fin de impedir la selección de arcos incompatibles con las relaciones de precedencia *origen-destino* (r^+ anterior a r^-); Psaraftis, (1983), usa programación dinámica para resolver el DARP con Ventanas de Tiempo (DARPTW) con un solo vehículo. Los estados son de la forma $(j, k_2, \dots, k_n)$, donde j es el vértice actualmente visitado, y cada k_i puede tener tres valores diferentes, que indican el *estatus* de la mercancía de cada cliente: -1 , no ha sido cargada; 0 , ha sido cargada pero no descargada; y $+1$, ha sido descargada; Desrosiers y otros (1986), adaptan este algoritmo de $2n$ estados, a PDPTW con un solo vehículo. Dumas (1985), y Desrosiers y otros (1987), presentan un algoritmo para el PDPTW planteándolo como un problema de Partición de Conjuntos y adaptando las técnicas de solución de éste.

En cuanto a los algoritmos heurísticos Jaw y otros (1986), describen un algoritmo de inserción para una variante del DARPTW. El uso de las técnicas descritas en el párrafo anterior también pueden usarse como subrutinas en otras técnicas heurísticas para problemas con varios vehículos, como en el caso de las del tipo *Cluster 1^o – Ruta 2^a*. Ejemplos de este tipo de aproximaciones los tenemos en los trabajos de Desrochers y otros, (1991).

Ahora bien, considérese en el planteamiento del PDP con un vehículo la siguiente restricción: En cada momento sólo puede ser descargado la última mercancía que ha sido cargada, de entre las que se encuentren en el vehículo. Este problema, que se va a denominar como PDP con un vehículo con sistema de descarga LIFO (último en

entrar primero en salir), aparece en muchas situaciones en el mundo del Transporte y la Industria: transporte de determinados gases industriales, distribución de pales en la industria del automóvil, o, en general, en aquellas actividades de reparto en las que se quiere un tiempo controlado para la descarga y se contempla la situación de la mercancía en el vehículo.

Existen referencias de trabajos que tratan modelos que también tienen en cuenta la situación de la carga en el vehículo, principalmente el VRPB. (Vehicle Routing Problem with Backhauls), como Casco y otros (1988), Deif y otros (1984), Goetschalkx y otros (1986), y Golden y otros (1985). Sin embargo, el problema tal como se ha definido en el párrafo anterior, sólo tiene referencias en los trabajos de Pacheco (1994), y Pacheco y otros (1994). En éstos se propone una estrategia para desarrollar algoritmos exactos a este problema, consistente en, a partir de un algoritmo Branch & Bound para el TSP, introducir una serie de modificaciones que den lugar a un algoritmo exacto para el problema anterior. Sin embargo, se ha observado que para problemas de mayor tamaño a los utilizados en estos trabajos (10 o más clientes), el tiempo de computación en ordenadores personales deja de ser razonable, cuando se utilizan estos algoritmos u otros algoritmos exactos (como los basados en métodos de programación dinámica).

En este trabajo, se propone un heurístico capaz de resolver problemas de gran tamaño en un tiempo de computación moderado, desviándose del óptimo sólo en un pequeño número de casos y en una cantidad mínima. Esta técnica es una extensión y adaptación a este problema del algoritmo de Or para el TSP (1976). El algoritmo de Or es una variante de los conocidos algoritmos r -óptimos desarrollados por Lin (1965), para el TSP simétrico. En el caso del algoritmo de Or se toma $r = 3$ y se reduce la búsqueda de los 3-intercambios a aquellos que supongan una recolocación de cadenas de 1, 2 o 3 elementos entre otros dos consecutivos en la ruta actual. La ventaja de este método es que el número de operaciones que se realiza es de $\theta(m^2)$, frente al algoritmo 3-óptimo original que utiliza $\theta(m^3)$ (siendo m el número de puntos que intervienen en el problema). Además, como se verá mas adelante, el algoritmo de Or se puede adaptar fácilmente a problemas asimétricos. La eficacia de este algoritmo ha sido contrastada recientemente en el trabajo de Nurmi (1991). Adaptaciones de este algoritmo al TSP con ventanas de tiempo (TSPTW) se pueden encontrar en el trabajo de Salomon y otros (1988).

El trabajo se estructura de la siguiente forma: en la sección 2 se explica la idea básica del algoritmo que se quiere desarrollar; en la sección 3 se estudian las condiciones para que un 3-intercambio, en este caso recolocación, sea factible; en la sección 4 se estudian las situaciones posibles dependiendo de los elementos que componen la cadena que se quiere recolocar; en la sección 5 se describe el algoritmo —desarrollado a partir de los resultados descritos en los apartados anteriores—; en la sección 6 se chequea la eficacia de este algoritmo a partir de la resolución de una

serie de problemas simulados; en la sección 7 se expondrán las conclusiones que se extraen de los resultados obtenidos en la sección anterior.

Igual que en trabajos anteriores —Pacheco (1994) y Pacheco y otros (1994)— se supondrá que la capacidad del vehículo es mayor que la carga total del conjunto de clientes; el caso general se estudiará posteriormente.

2. IDEA BÁSICA

Or (1976) propone restringir la búsqueda de intercambios a los 3-intercambios en los que cadenas de uno, dos o tres clientes consecutivos son recolocadas entre otros dos clientes (ver figura 1). Esto hace reducir el número de 3-intercambios a considerar a una cantidad de orden $\theta(n^2)$. Nótese que con estos intercambios no se cambia el sentido de los diferentes tramos.

En nuestro caso seguiremos la misma idea, pero no el límite del tamaño de la cadena a recolocar; además debemos chequear la factibilidad de cada posible recolocación respecto a las restricciones del problema.

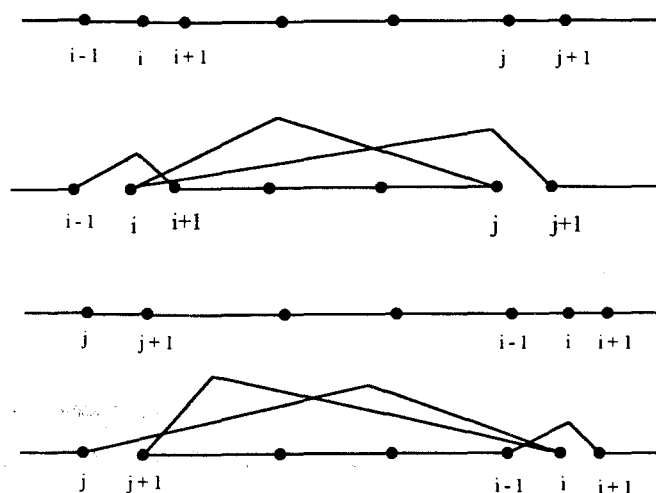


Figura 1. Posible recolocación del elemento i hacia adelante y hacia atrás entre j y $j+1$.

En el caso general de recolocación de cadenas de k elementos consecutivos, comenzando en la posición i , entre dos puntos que ocupan las posiciones j y $j+1$, como

muestra la figura 2, supone suprimir los arcos $(i-1, i)$, $(i+k-1, i+k)$, $(j, j+1)$ e incorporar los arcos $(i-1, i+k)$, (j, i) y $(i+k-1, j+1)$.

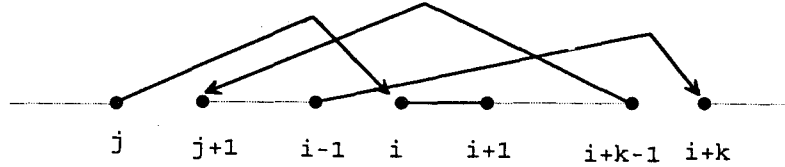


Figura 2. Recolocación de una cadena de k elementos.

En nuestro caso la ruta contiene $2 \cdot n$ puntos: el origen 1, $n-1$ puntos de carga, $n-1$ puntos de descarga, y además, por comodidad, se añade al final el punto 1 que se interpreta como la vuelta al origen. Se definen las siguientes variables: k , el tamaño de la cadena a recolocar; i el punto inicial de la cadena; y j el punto tras el que recolocamos la cadena.

Para cada número de elementos k de la cadena, los posibles valores que puede tomar i varían desde 2, hasta $2 \cdot n - k$. Para cada valor de i y cada valor de k , los posibles valores que puede tomar j , punto a partir del cual insertamos la cadena, son $j = 1, \dots, i-2$; (recolocación hacia atrás) y $j = i+k, \dots, 2 \cdot n - 1$ (recolocación hacia adelante). Por comodidad, se añade a las rutas el punto $2 \cdot n$ que se interpreta como la vuelta al origen 1.

El algoritmo actúa básicamente de la forma siguiente: a partir de una ruta inicial examina todos los intercambios o recolocaciones posibles que sean factibles. Para cada uno de estos intercambios factibles calcula la variable *ganancia* definida como la suma de las distancias de los arcos que se suprimen menos la suma de las distancias de los arcos que se añaden. Se realiza el intercambio correspondiente al mayor valor de *ganancia*, y se vuelve a repetir el proceso con la nueva ruta. El algoritmo termina cuando el mayor valor de *ganancia* no sea positivo.

El chequeo de la factibilidad de cada intercambio, es decir de la nueva ruta a la que éste da lugar, puede suponer un tiempo de computación excesivo. Si para cada recolocación posible (definida por i, k y j) se chequea su factibilidad, el número de operaciones en cada iteración sería de $\Theta(n^4)$. A continuación vamos a proponer un método que rápidamente examine las recolocaciones factibles para cada cadena, es decir, para cada valor de i y de k , determinar directamente los valores de j factibles; de esta forma que el número de operaciones sea de $\Theta(n^3)$.

3. INTERCAMBIOS FACTIBLES

Para examinar la factibilidad de un intercambio se deben considerar las dos restricciones, o precedencias obligatorias, siguientes del problema:

1. Para cada cliente r cada punto de carga r^+ debe ser anterior a cada punto de descarga r^- .
2. Sólo se puede descargar la última mercancía que ha sido recogida, es decir, cuando haya dos mercancías en el vehículo, se ha de descargar antes la última en entrar. En otras palabras: no se pueden dar situaciones de este tipo en una ruta

$$1 - \dots - r^+ - \dots - s^+ - \dots - s^- - \dots - r^- - \dots - 1.$$

Cuando se estudia la recolocación de una cadena concreta (determinada por los valores de i y k definidas anteriormente), se ha de examinar en que lugares se puede insertar, es decir qué valores de j van a dar lugar a rutas que cumplan las dos restricciones anteriores. Para estudiar los valores de j que cumplan la primera restricción, basta con fijarse por una parte en las posiciones de los puntos de carga correspondientes a los puntos de descarga que están en la cadena; y por otra parte en las posiciones de los puntos de descarga correspondientes a los puntos de carga que están en la cadena. Como se verá de forma explícita más adelante, dichas posiciones marcarán el mínimo y el máximo valor que puede tomar j .

El estudio de los valores que cumplan la segunda restricción es más complicado. Para ello considérese C una cadena de esta ruta. Definimos el siguiente conjunto:

M = Conjunto de clientes cuyos puntos de carga y descarga no pertenecen ninguno a C ;

para todo cliente r , cumpliendo que o bien r^+ o bien r^- , pero sólo uno, pertenece a C se definen:

$A(r)$ = Subconjunto de M de clientes s tales que s^- es anterior a r^+ en la ruta actual;

$$s^+ - \dots - s^- - \dots - r^+ - \dots - r^-$$

$B(r)$ = Subconjunto de M de clientes s tales que s^+ es anterior a r^+ y s^- es posterior a r^- en la ruta actual;

$$s^+ - \dots - r^+ - \dots - r^- - \dots - s^-$$

$C(r)$ = Subconjunto de M de clientes s tales que s^+ es posterior a r^+ y s^- es anterior a r^- en la ruta actual;

$$r^+ - \dots - s^+ - \dots - s^- - \dots - r^-$$

$D(r)$ = Subconjunto de M de clientes s tales que s^+ es posterior a r^- en la ruta actual

$$r^+ - \dots - r^- - \dots - s^+ - \dots - s^-.$$

$\forall t \in N^+ \cup N^-$ se define $\text{orden}(t)$ como el número de orden o posición que ocupa el elemento t en la ruta actual. Se tiene el siguiente

Lema 1 Sea j la variable que indica la posición donde se recoloca C , según se ha definido anteriormente:

- a) $\forall s \in A(r)$, C no puede ser recolocado entre s^+ y s^- , es decir, j no puede tomar valores comprendidos entre $\text{orden}(s^+)$ y $\text{orden}(s^-) - 1$, ambos incluidos;
- b) Sean $m1 = \max\{\text{orden}(s^+)/s \in B(r)\}$ y $m2 = \min\{\text{orden}(s^-)/s \in B(r)\}$, j no puede tomar valores comprendidos entre 1 y $m1 - 1$ ambos incluidos, ni entre $m2$ y $2 \cdot n - 1$ ambos inclusive.
- c) $\forall s \in C(r)$, j no puede tomar valores comprendidos entre $\text{orden}(s^+)$ y $\text{orden}(s^-) - 1$, ambos incluidos;
- d) $\forall s \in D(r)$, j no puede tomar valores comprendidos entre $\text{orden}(s^+)$ y $\text{orden}(s^-) - 1$, ambos incluidos.

Demostración: Si j tomara alguno de los valores anteriormente señalados, se tendría, para algún cliente s , una de las dos posiciones relativas siguientes entre los puntos de carga y descarga de r y s en la nueva ruta obtenida:

$$s^+ - \dots - r^+ - \dots - s^- - \dots - r^-$$

o bien

$$r^+ - \dots - s^+ - \dots - r^- - \dots - s^-.$$

■

4. SITUACIONES POSIBLES DEL CASO GENERAL

Se define $r(t)$ el elemento de la ruta actual que ocupa la posición t ; por consiguiente, según lo comentado anteriormente, que $r(1) = r(2n) = 1$. Sea la cadena C determinada por los valores de i y k :

$$r(i) - r(i+1) - \dots - r(i+k-1).$$

Definimos los siguientes conjuntos de clientes:

$$A^+ = \{r/r^+ \in C, r^- \notin C\} \quad \text{y} \quad A^- = \{r/r^- \in C, r^+ \notin C\};$$

se definen también:

$$\text{límite_inferior} = \max\{\text{orden}(j^+)/j \in A^-\}, \quad \text{si } A^- = \emptyset \quad \Leftrightarrow \quad \text{límite_inferior} = 1;$$

$$\text{límite_superior} = \min\{\text{orden}(j^-)/j \in A^+\}, \quad \text{si } A^+ = \emptyset \quad \Leftrightarrow \quad \text{límite_superior} = 2n.$$

Sean además $j1$ el elemento de A^+ verificando $\text{orden}(j1^-) = \text{límite_superior}$ si $A^+ \neq \emptyset$; y $j2$ el elemento de A^- verificando $\text{orden}(j2^+) = \text{límite_inferior}$ si $A^- \neq \emptyset$; se tiene el siguiente

Lema 2 *Sea j la variable que indica la posición donde se recoloca C , según se ha definido anteriormente, se tiene que j sólo puede tomar los valores comprendidos entre límite_inferior y $\text{límite_superior}-1$, ambos inclusive.*

Demostración: Trivial, ya que si j tomara otros valores en la ruta resultante se violaría, para al menos un cliente, la restricción 1 del apartado 3, es decir, algún punto de descarga sería anterior a su correspondiente de carga.

Por otra parte, según los elementos de A^- y A^+ se tienen los siguientes 4 casos:

1. $A^+ = \emptyset, \quad A^- = \emptyset$:

La cadena está compuesta por pares completos de puntos de carga descarga, como por ejemplo, $r^+r^-s^+s^-$ o bien $r^+s^+s^-r^- \dots$; en este caso ninguna recolocación puede romper ninguna precedencia obligatoria señalada en la sección anterior.

2. $A^+ \neq \emptyset, \quad A^- = \emptyset$:

$$1 - \dots - r(i) - r(i+1) - \dots - r(i+k-1) - \dots - j1^- - \dots - 1$$

Los clientes que se van a considerar, según las posiciones relativas de sus correspondientes puntos de carga y descarga, para estudiar las posibles recolocaciones son las siguientes:

- a) $1 - \dots - s^+ - \dots - s^- - \dots - r(i) - r(i+1) - \dots - r(i+k-1) - \dots - j1^- - \dots - 1$
 Conjunto de clientes s verificando que $\text{orden}(s^-) < i$; que va a coincidir con $A(j1)$, (ver apartado anterior);
- b) $1 - \dots - s^+ - \dots - r(i) - r(i+1) - \dots - r(i+k-1) - \dots - j1^- - \dots - s^- - \dots - 1$
 Conjunto de clientes s , verificando que $\text{orden}(s^+) < i$ y $\text{orden}(s^-) > \text{orden}(j1^-)$; es decir $B(j1)$;
- c) $1 - \dots - r(i) - r(i+1) - \dots - r(i+k-1) - \dots - s^+ - \dots - s^- - \dots - j1^- - \dots - 1$
 Conjunto de clientes s verificando que $\text{orden}(s^+) > i+k-1$ y $\text{orden}(s^-) < \text{orden}(j1^-)$, es decir $C(j1)$.

A efectos del estudio de recolocaciones factibles no se considera el conjunto de clientes s verificando que $\text{orden}(s^+) > \text{orden}(j1^-)$, (es decir $D(j1)$), pues la cadena C no puede ser insertada en una posición posterior a la de $j1^-$ (Lema 2).

3. $A^- \neq \emptyset$, $A^+ = \emptyset$:

$$1 - \dots - j2^+ - \dots - r(i) - r(i+1) - \dots - r(i+k-1) - \dots - 1$$

Los conjuntos de clientes a considerar son los siguientes:

- a) $1 - \dots - s^+ - \dots - j2^+ - \dots - r(i) - r(i+1) - \dots - r(i+k-1) - \dots - s^- - \dots - 1$
 Conjunto de clientes s verificando que $\text{orden}(s^+) < \text{orden}(j2^+)$ y $\text{orden}(s^-) > i+k-1$, es decir $B(j2)$;
- b) $1 - \dots - j2^+ - \dots - s^+ - \dots - s^- - \dots - r(i) - r(i+1) - \dots - r(i+k-1) - \dots - 1$
 Conjunto de clientes s verificando que $\text{orden}(s^+) > \text{orden}(j2^+)$ y $\text{orden}(s^-) < i$, es decir $C(j2)$;
- c) $1 - \dots - j2^+ - \dots - r(i) - r(i+1) - \dots - r(i+k-1) - \dots - s^+ - \dots - s^- - \dots - 1$
 Conjunto de clientes s verificando que $\text{orden}(s^+) > i+k-1$, es decir $D(j2)$.

Al igual que en el caso anterior no se considera el conjunto $A(j2)$.

4. $A^+, A^- \neq \emptyset$:

$$1 - \dots - j2^+ - \dots - r(i) - r(i+1) - \dots - r(i+k-1) - \dots - j1^- - \dots - 1$$

Sólo se consideran los dos siguientes conjuntos:

a) $1 - \dots - j2^+ - \dots - s^+ - \dots - s^- - \dots - r(i) - r(i+1) - \dots - r(i+k-1) - \dots - j1^- - \dots - 1$

Conjunto de clientes s verificando que $\text{orden}(s^+) > \text{orden}(j2^+)$ y $\text{orden}(s^-) < i$, es decir $C(j2)$;

b) $1 - \dots - j2^+ - \dots - r(i) - r(i+1) - \dots - r(i+k-1) - \dots - s^+ - \dots - s^- - \dots - j1^- - \dots - 1$

Conjunto de clientes s verificando que $\text{orden}(s^+) > i+k-1$ y $\text{orden}(s^-) < \text{orden}(j1^+)$, es decir $C(j2)$;

Obviamente, no se van a considerar otros conjuntos de clientes.

■

5. DESCRIPCIÓN DEL ALGORITMO EN SEUDOCÓDIGO

A continuación se describe el algoritmo que sigue la idea básica descrita en el apartado 2, y con la determinación de intercambios factibles según se estudia en los apartados 4 y 5. Se tienen los siguientes parámetros de entrada:

n : valor que determina los conjuntos N^+ y N^- según la definición del apartado;

t : matriz de distancias;

y los siguientes parámetros por variable:

ruta: vector que registra la solución en cada momento, es decir, $\text{ruta}[i]$ indica el punto que ocupa la posición i en la ruta actual;

costetotal: distancia total de la solución;

además se hacen uso de las siguientes variables:

gmax: Indica el valor máximo de *ganancia* (definida en el apartado 2) en cada iteración;

imax: Indica el lugar donde comienza la cadena correspondiente a *gmax*;
nele: indica el número elementos de dicha cadena;
jmax: indica el lugar donde se recoloca la cadena correspondiente a *gmax*;
orden: vector que indica la posición de cada punto en la ruta actual;
valido: vector lógico que indica, para cada cadena, en qué posiciones puede ser recolocada de forma factible.

Asimismo, siguiendo las definiciones tomadas hasta ahora, *i* es la variable auxiliar que indica la posición donde comienza la cadena *C*, *k* el número de elementos que la componen, y *j* la posición donde va a ser recolocada; el algoritmo actúa de la forma siguiente:

Hacer ruta como la ruta inicial y coste total como la distancia de ésta.

Repetir

hacer *gmax*:=0;

Para *k*:=1 hasta *n* – 1 hacer:

Para *i*:=2 hasta $2 \cdot n - k - 1$ hacer

inicio

Para *j*:=1 hasta $2n - 1$ hacer *valido*[*j*]:=TRUE;

Determinar los conjuntos A^+ y A^- (según la definición del apartado 4).

Si $A^- \neq \emptyset$ entonces

inicio

Hacer *límite_inferior* = $\max\{\text{orden}[j^+]/j \in A^-\}$;

Determinar $j2 \in A^- / \text{orden}[j2^+] = \text{límite_inferior}$

fin

si no: Hacer *límite_inferior* = 1;

Si $A^+ \neq \emptyset$ entonces

inicio

Hacer *límite_superior* = $\min\{\text{orden}[j^-]/j \in A^+\}$;

Determinar $j1 \in A^+ / \text{orden}[j1^-] = \text{límite_superior}$

fin

si no: Hacer $\text{límite_superior} = 2n$;

Para $j:=1$ hasta $\text{límite_inferior}-1$ hacer $\text{valido}[j]:=FALSE$;

Para $j:=\text{límite_superior}$ hasta $2n-1$ hacer $\text{valido}[j]:=FALSE$;

Si $A^+ \neq \emptyset$, $A^- = \emptyset$ entonces:

inicio

Determinar conjuntos $A(j1), B(j1), C(j1)$ (según la definición del apartado 3)

$\forall s \in A(j1) \cup C(j1)$ hacer:

Para $j:=\text{orden}[s^+]$ hasta $\text{orden}[s^-]-1$ hacer $\text{valido}[j]:=FALSE$;

Si $B(j1) \neq \emptyset$ hacer

inicio

Determinar $m1 = \max\{\text{orden}[s^+] / s \in B(j1)\}$;

Determinar $m2 = \min\{\text{orden}[s^-] / s \in B(j1)\}$;

Para $j:=1$ hasta $m1-2$ y de $m2$ hasta $2 \cdot n-1$ hacer $\text{valido}[j]:=FALSE$

fin;

fin;

Si $A^+ = \emptyset$, $A^- \neq \emptyset$ entonces:

inicio

Determinar conjuntos $B(j2), C(j2), D(j2)$; hacer:

$\forall s \in C(j2) \cup D(j2)$ hacer:

Para $j:=\text{orden}[s^+]$ hasta $\text{orden}[s^-]-1$ hacer $\text{valido}[j]:=FALSE$;

Si $B(j2) \neq \emptyset$ hacer

inicio

Determinar $m1 = \max\{\text{orden}[s^+]/s \in B(j2)\}$;

Determinar $m2 = \min\{\text{orden}[s^-]/s \in B(j2)\}$;

Para $j:=1$ hasta $m1 - 2$ y de $m2$ hasta $2 \cdot n - 1$ hacer $\text{valido}[j]:=FALSE$

fin;

fin;

Si $A^+, A^- \neq \emptyset$ entonces:

inicio

Determinar conjuntos $C(j1), C(j2)$;

$\forall x \in C(j1) \cup C(j2)$ hacer:

Para $j:=\text{orden}[s^+]$ hasta $\text{orden}[s^-] - 1$ hacer $\text{valido}[j]:=FALSE$;

fin;

Para $j:=1$ hasta $i - 2$, y $j:=i + k$ hasta $2 \cdot n - 1$ hacer:

Si $\text{valido}[j]$ entonces:

inicio

Calcular el valor de *ganancia* (según se definió el apartado 2);

Si $\text{ganancia} > gmax$ entonces:

inicio

$gmax:=ganancia$;

$imax:=i$;

$kmax:=k;$

$jmax:=j;$

fin

fin;

fin;

Si $gmax > 0$ entonces:

inicio

Sustituir los arcos $(imax-1, imax)$, $(imax+kmax-1, imax+kmax)$, $(jmax, jmax+1)$ por los arcos $(imax-1, imax+kmax)$, $(jmax, imax)$ y $(imax+kmax-1, jmax+1)$ y registrar el resultado en *ruta*;

Hacer $costetotal:=costetotal-gmax$

fin

hasta que $gmax = 0$.

Existen copias del algoritmo programado en PASCAL a disposición de las personas interesadas.

6. RESULTADOS COMPUTACIONALES

Para examinar la eficacia del algoritmo anteriormente descrito se han realizado cuatro tipos de pruebas: el primero tiene como objeto reflejar la eficacia del método desarrollado para determinar el ahorro en el tiempo de computación de los intercambios factibles, el segundo tiene como objeto comprobar la desviación del óptimo de la solución obtenida; el tercero determinar en qué medida la incorporación de este heurístico en los algoritmos exactos descritos por Pacheco, (1994), y Pacheco y otros (1994), para este problema hace reducir el tiempo de computación de éstos; y el cuarto comprobar la evolución de los tiempos de computación de este algoritmo en función del tamaño del problema.

Tanto la implementación de los algoritmos que se han utilizado, como la programación de los diferentes problemas que han servido de prueba, se han hecho utilizando el compilador Borland Pascal (ver. 7.0). El equipo técnico usado es un ordenador

personal tipo PC AT i 486dx2 a 50 Mhz para las tres últimas pruebas, y un ordenador Pentium a 100 Mhz. para la primera. (Existen copias de estos programas a disposición de las personas interesadas.)

6.1. Chequeo de intercambios factibles. Reducción de tiempo

Como se ha comentado en el apartado 2, el chequeo de la factibilidad de cada intercambio uno a uno puede emplear mucho tiempo de computación. El objeto de los apartados 3 y 4 es desarrollar un método para determinar de forma más *rápida* qué intercambios son factibles. Para establecer la eficacia del método propuesto, se han simulado 20 matrices de distancia, para cada tamaño de problema: 5, 10, 15, 20, 25 y 30 clientes. Los problemas así generados se han resuelto utilizando el algoritmo con este método (tal como se describe en el apartado 5) y chequeando los intercambios uno a uno. La forma de generar distancias es asignar a cada punto del problema dos coordenadas x e y , cuyos valores son generados aleatoriamente con distribución uniforme entre 0 y 100. La distancia entre cada par de puntos se define como la distancia euclídea correspondiente.

A continuación se muestra un cuadro que resume los resultados obtenidos: tiempos de computación para ambas formas y reducción porcentual:

N. clientes	5	10	15	20	25	30
Uno_a_Uno	0'023	0'247	2'199	9'617	34'594	85'415
Método	< 0'001	0'021	0'083	0'199	0'464	0'823
Reducción	95'65	91'49	96'22	97'93	98'37	99'03

6.2. Desviación del óptimo

Para estudiar la proximidad al óptimo de las soluciones obtenidas por este heurístico se han simulado matrices de distancias de dos formas diferentes. En la primera de ellas las distancias son generadas aleatoriamente tomando valores enteros con la misma probabilidad entre 0 y 99. La segunda forma de generar distancias coincide con la utilizada en el apartado 6.1. En ambos casos se han generado 20 matrices para cada tamaño del problema: 5, 6, 7, 8 y 9 clientes (correspondientes a 11, 13, 15, 17 y 19 puntos).

A continuación, para cada tipo de matriz, se muestra un cuadro que resume los resultados obtenidos: En estos cuadros se muestra, para cada número de clientes, la desviación porcentual del óptimo (media, máxima y mínima), así como el número de casos en los que se alcanzó.

6.2.1. Primer tipo de matriz

Para obtener la ruta inicial, (Paso 1 descrito en el apartado 5) se han generado aleatoriamente 50 rutas factibles y se ha tomado aquella con menos distancia total. Los resultados son los siguientes:

Número de clientes	Desviación porcentual del óptimo			Casos en los que se alcanza
	Mínima	Media	Máxima	
5	0	0	0	20
6	0	0'011	0'104	16
7	0	0'015	0'240	16
8	0	0'028	0'350	18
9	0	0'014	0'255	17

6.2.2. Segundo tipo de matriz

En este caso, para la ruta inicial se ha elegido aquella con menos distancias entre 10 rutas factibles generadas aleatoriamente y las obtenidas utilizando la adaptación a este problema de los algoritmos de inserción más cercana y más lejana para el TSP. (en el trabajo de Golden y otros (1980), se hace una amplia recopilación y descripción de los principales algoritmos de inserción para el TSP, así como de otros algoritmos constructivos que pueden ser adaptados fácilmente a este problema). Los resultados son los siguientes:

Número de clientes	Desviación porcentual del óptimo			Casos en los que se alcanza
	Mínima	Media	Máxima	
5	0	0	0	20
6	0	0	0	20
7	0	0	0	20
8	0	0'015	0'302	19
9	0	0	0	20

6.3. Reducción en el tiempo de computación de los exactos

Como se ha mencionado anteriormente, en los trabajos de Pacheco (1994) y Pacheco y otros (1994), se desarrolla una estrategia para el diseño de algoritmos

exactos para este problema, a partir de métodos Branch & Bound para el TSP. En los algoritmos exactos para problemas de rutas se puede conseguir, en muchos casos, reducir el tiempo de computación mediante determinadas estrategias. Entre éstas destacamos: generar cotas superiores en cada vértice del árbol de búsqueda si el algoritmo es de tipo Branch & Bound, (ver Volgenant y Jonker, (1982). o Nurmi, (1991), más recientemente); obtener cotas inferiores más ajustadas al óptimo mediante métodos de penalización lagrangiana (ver Held & Karp, (1970) y (1971)), o métodos de Programación Dinámica (ver Christodides y otros (1981), etc. En nuestro caso, vamos a reducir el tiempo de computación realizando dos modificaciones en los algoritmos exactos anteriormente señalados; éstas son las siguientes:

1. Partir como solución inicial la obtenida por el heurístico descrito en este trabajo.
2. Cada vez que se llegue a una solución, utilizar dicho heurístico para mejorarla.

El objeto de estas modificaciones es obtener rápidamente cotas superiores ajustadas al óptimo para evitar exploraciones innecesarias, y reducir el tiempo de computación. Para verificar estos ahorros de tiempo se han generado 20 matrices de distancias, como en el caso del apartado 6.1., para cada tamaño del problema: 5, 6, 7, 8 y 9 clientes. Para obtener una solución inicial se empleó como ruta inicial aquella con menos distancia total entre 50 rutas factibles generadas aleatoriamente. Los resultados son los siguientes:

Número de clientes	Tiempo de Computación en segundos		Reducción Porcentual
	Algoritmo Original	Algoritmo Modificado	
5	0'288	0'215	25'347
6	1'396	0'630	54'871
7	17'205	11'362	33'961
8	59'644	44'558	25'293
9	474'427	238'822	49'661

6.4. Evolución de los tiempos de computación

Finalmente se ha realizado un tercer tipo de pruebas destinado a establecer la evolución de los tiempos de computación en ordenadores personales, del heurístico desarrollado en este trabajo, para problemas de mayor tamaño. Para ello se han

simulado 20 matrices de distancia como en el apartado 6.1.2. (distancias euclídeas), para cada tamaño del problema: 10 clientes, 20, ..., hasta 120 clientes (241 puntos). Como ruta inicial se ha elegido aquella con menos distancias entre 10 rutas factibles generadas aleatoriamente y las obtenidas por los algoritmos de inserción (ver apartado 6.1.2). Los resultados son los siguientes:

Número de clientes	Tiempo medio de Computación en segundos		
	Solución Inicial	Heurístico Propuesto	Total
10	0'120	0'035	0'155
20	1'282	0'321	1'603
30	5'029	1'516	6'545
40	14'891	3'212	18'103
50	35'403	8'295	43'698
60	71'609	18'775	90'384
70	125'081	24'044	149'125
80	210'159	39'886	250'045
90	339'721	57'293	397'014
100	520'928	101'775	622'703
110	761'427	130'193	891'620
120	1074'992	174'439	1249'431

7. CONCLUSIONES

A la vista de estos resultados se llegan a las siguientes conclusiones:

- El heurístico descrito ofrece soluciones muy poco desviadas del óptimo: en caso de las matrices de distancias del apartado 6.1.1. un máximo del 0'35%; en caso de matrices de distancias euclídeas la solución coincide con la óptima en 99 de los 100 casos; el único caso en el que no coincide se desvía el 0'302%. Por consiguiente, especialmente para este último tipo de matrices, la aproximación es realmente grande.
- El ofrecer soluciones cercanas al óptimo (cuando no coinciden con él), permite reducir el tiempo de computación de los algoritmos exactos para este problema cuando se combinan con este heurístico — tal como se hace en el apartado

6.3—. Esto se debe a que se parte de una cota superior en el vértice inicial muy ajustada al óptimo, con lo que se rechazan más exploraciones innecesarias.

- El tiempo de computación que emplea es insignificante si se le compara con el que emplean los exactos existentes. Esto hace que se puedan resolver problemas de mucho mayor tamaño. En el apartado 6.4 se resuelven problemas de hasta 241 puntos, pero la evolución de los tiempos de computación permite adivinar que se podrían resolver problemas de mayor tamaño, en un tiempo aceptable, con ordenadores o compiladores que soportaran matrices de distancia de mayor dimensión.
- El empleo de un tiempo de computación pequeño, se debe fundamentalmente al método utilizado para establecer los intercambios factibles en cada iteración, basado en los resultados teóricos desarrollados obtenidos en los apartados 3 y 4. Este método permite una reducción porcentual del tiempo de computación de más del 90% y que aumenta con el tamaño de los problemas.

Concluir finalmente que el algoritmo heurístico propuesto resulta extremadamente eficaz al ofrecer soluciones muy cercanas al óptimo —sobre todo con distancias euclídeas—, en un tiempo de computación mucho menor que los exactos, lo que permite poder resolver problemas de mayor tamaño.

8. REFERENCIAS Y BIBLIOGRAFÍA

- [1] **Christofides, N., Mingozi, A. & Toht, P.** (1981). «State-Space Relaxation Procedures for the Computation of Bounds to Routing Problems». *Networks*, **11**, 2, 145–164.
- [2] **Casco, D.O., Golden, B.L. & Wasi, E.A.** (1988). «Vehicle Routing with Backhauls: Models, Algorithms, and Case Studies». In *Vehicle Routing: Methods and Studies*, (Studies in Management Sciences and Systems, vol.16), eds: GOLDEN, B.L. and ASSAD, A.A., Nort-Holland, 7–46.
- [3] **Deif, I. & Bodin, L.** (1984). «Extension of the Clark and Wright Algorithm for Solving the Vehicle Routing Problem with Backhauling». *Proceedings of the Babson Conference on Software Uses in Transportation and Logistics Management* (A.E.KIDDER, edit.), 75–96. Babson Park, Massachussets.
- [4] **Desrochers, M. & Soumis, F.** (1991). «Implantation et Complexité des Techniques de Programmation Dynamique dans les Méthodes de Confection des Tournées et d'Horaires». *Recherche Operationnelle/Operations Research*, **25**, 3, 291–310.

- [5] **Desrosiers, J., Dumas, Y. & Soumis, F.** (1987). «Vehicle Routing Problem with Pickup, Delevery and Time Windows». *Cahiers du GERARD*, Ecole Des Hautes Etudes Commerciales de Montreal.
- [6] **Dumas, Y.** (1985). «Confection d'iteneraires de vehicules en vue du transport de plusieurs origines à plusieurs destinations». *Centre de Recherche sur le Transports*. Université de Montreal.
- [7] **Goetschalcks, M. & Horsley, C.** (1986). «The Vehicle Routing Problem with Backhauls». *Material Handling Researcsh Center*, Dept. of Industrial and Systems Engineering, Georgia Institute of Technology.
- [8] **Golden, B., Baker,E. , Alfaro, J. & Schaffer, J.** (1985). «The Vehicle Routing Problem with Backhauling: Two Approaches». *Twenty-First Annul Meeting of S.E. TIMS* (R.D.HAMMESFAHR, edit.), Myrtle Beach, SC, 90–92.
- [9] **Golden, B., Bodin, L. Doyle,T. & Stewart,W.Jr.** (1980). «Approximate Traveling Salesman Algorithms». *Operations Researsch*, **28, 3**, parte II, 694–711.
- [10] **Held, M. & Karp, R.M.** (1970). «The Traveling Salesman Problem and Minimum Spanning Trees». *Ops. Res.*, **18**, 1138–1162.
- [11] **Held, M. & Karp, R.M.** (1971). «The Traveling Salesman Problem and Minimum Spanning Trees: Part II». *Mathematical Programing I*, 6–25.
- [12] **Kalantari, B., Hill, A.V. & Arora, S.R.** (1985). «An Algorithm for the Traveling Salesman Problem with Pickup and Delevery Customers». *European Journal of Operational Research*, **22**, 377–386.
- [13] **Little, J., Murty, K., Sweeney, D. & Karel, C.** (1963). «An Algorithm for the Traveling Salesman Problem». *Operations Research*, **11 (5)**, 972–989.
- [14] **Lin, S.** (1965). «Computer Solutions to the Traveling Salesman Problem». *Bell Syst. Tech. Jou.*, **44**, 2245–2269.
- [15] **Nurmi, K.** (1991). «Traveling Salesman Problem Tools for Microcomputers». *Computers & Ops. Res.* **18, 8**, 741–749.
- [16] **O'Brien, S.K. & Nameroff, S.** (1993). *Turbo-Pascal 7: Manual de Referencia*. Osborne McGraw-Hill. Madrid.
- [17] **Or, I.** (1976). *Traveling Salesman Type Combinatorial Problems and their Relations to the Logistics of Blood Banking*. Ph. Thesis, Dpt. of Industrial Engineering and Management Sciences, Northwestern Univ.
- [18] **Pacheco, J.** (1994). *Problemas de Rutas con Ventanas de Tiempo*. Tesis Doctoral leída en el Dpto. de Estadística e I. Optva. de la Facultad de Matemáticas de la U. Complutense de Madrid. Mayo 1994.
- [19] **Pacheco, J., García, A. & Aragón, A.** (1994). *Problemas de Ruta con Carga y Descarga en Sistemas LIFO*. XXI Congreso de Estadística e I. Optva. celebrado en Calella en Abril de 1994.

- [20] **Psaraftis, H.** (1983). «An Exact Algorithm for the Single Vehicle Many-to-Many Dial-a-Ride Problem with Time Windows». *Transportation Sci.*, **17**, 351–360.
- [21] **Solomon, M., Baker, E. & Schaffer, J.** (1988a). «Vehicle Routing and Scheduling Problems with Time Window Constraints». In *Vehicle Routing: Methods and Studies*, (Studies in Management Sciences and Systems, vol.16), eds: GOLDEN, B.L. and ASSAD, A.A., Nort-Holland, 85–106.
- [22] **Volgenant, T. & Jonker, R.** (1982). «A Branch and Bound Algorithm for the Symmetric Traveling Salesman Problem based on the 1-tree Relaxation». *Eur. J. Ops. Res.*, **9**, 83–89.

ENGLISH SUMMARY

HEURISTIC FOR THE PICK-UP AND DELIVERY ROUTING PROBLEMS IN LIFO UNLOADING SYSTEMS

JOAQUÍN A. PACHECO*

E.U.E. Empresariales de Burgos

In this paper a heuristic is proposed to solve the «Pick-up and Delivery Problem (PDP) using only one vehicle with limitless capacity in LIFO unloading systems», i.e., only the last goods to be loaded onto the vehicle may be unloaded at one time. This algorithm is an extension an adaptation of the Or method for the Travelling Salesman Problem (TSP) which can also be used for asymmetrical instances. With this heuristic one can solve large scale problems in reasonable calculation times on PC with only a small deviation from the optimum.

Keywords: Pick-up & Delivery with LIFO systems, Or-exchanges, feasibility, checking.

*Joaquín A. Pacheco. Dpto. Matemáticas. E.U.E. Empresariales de Burgos. Fco. Vitoria s/n. 09006 Burgos.

–Received march 1995.

–Accepted july 1996.

The Pick-up and Delivery Problem (PDP) may be described in the following way: One has a set of clients $N = \{2, 3, \dots, n\}$; each client i requires that goods be transported from an origin i^+ to a destination i^- . With this objective in mind one has m vehicles which leave an initial city 1, which they then return to when their route is completed. The object is to design feasible vehicle routes with the shortest total distances. The set of pick-up points will be denoted by $N^+ = \{r^+ / r = 2, 3, \dots, n\}$; the set of delivery by $N^- = \{r^- / r = 2, 3, \dots, n\}$; and the loads of each client r by $q(r)$, $r = 2, 3, \dots, n$. The special case in which the orders required are all the same is known as Dial-A-Ride-Problem (DARP) which occurs in passenger transport, school transport, etc...

Now let us consider the PDP with one vehicle with the following restriction: At any one moment only the last goods loaded may be unloaded. This problem, which will be denoted PDP with one vehicle with a LIFO (last-in-first-out) unloading system, appears in many situations in the world of Transport and Industry: transport of certain industrial gases, distribution of pallets in the auto industry, or in general, in those delivery activities in which a controlled time is required for unloading.

In this paper a heuristic is proposed capable of solving large scale problems in moderate calculation times, deviating from the optimum in only a few cases and by a small amount. This technique is an extension and adaptation to this problem of the Or algorithm for the TSP (1976).

The paper is organized in the following way: section 2 explains the basic idea of the algorithm to be developed; section 3 studies the conditions necessary for a 3-interchange, in this case repositioning, to be feasible; section 4 studies the possible situations depending on the elements that make up the chain which is to be repositioned; section 5 describes the algorithm with the results from the previous sections; section 6 checks the efficiency of this algorithm starting from the solving of a series of the simulated problems; section 7 explains the conclusions drawn from the results obtained in the previous section. Furthermore, an appendix is added which describes the algorithm in more detail in PASCAL.

As in previous papers —Pacheco, (1994), and Pacheco *et al.* (1994) — it is assumed that the total load from all the clients; the general case will be studied later in other works.

CÁLCULO DE LA DISTRIBUCIÓN DEL TIEMPO DE VIDA DE COMPONENTES MEDIANTE AUTOPSIA EN SISTEMAS BINARIOS ADITIVOS, SERIE-PARALELO Y PARALELO-SERIE

FERMÍN MALLOR GIMÉNEZ*

CRISTINA AZCÁRATE CAMIO*

ANTONIO PÉREZ PRADOS*

Universidad Pública de Navarra

En este artículo se estudia el problema de determinar la función de distribución del tiempo de vida de las componentes de un sistema binario, a partir del conocimiento de las leyes que rigen el funcionamiento del sistema y del conjunto de componentes que causa su fallo (obtenida mediante autopsia del sistema en el momento de su deterioro).

Se presentan los resultados de Meilijson (1981) y Nowik (1990) que proponen un sistema de ecuaciones implícito para obtener estas distribuciones. Sin embargo, se observa que este sistema es de muy difícil resolución práctica, por lo que nosotros consideramos un método cuya utilización es más restringida pero más sencilla, y estudiamos su aplicación a sistemas binarios aditivos, serie-paralelo y paralelo-serie.

Determination of the lifetime distribution of the components from the autopsy in additive, series-paralell and paralell-series binary systems

Keywords: Sistema binario, distribución del tiempo de vida, autopsia

Clasificación AMS: 62G05, 62N05, 90B25

*Departamento de Estadística e Investigación Operativa. Universidad Pública de Navarra. Campus de Arrosadía, s/n. 31006 Pamplona.

—Article rebut el gener de 1996.

—Acceptat el novembre de 1996.

1. INTRODUCCIÓN

Un problema que se plantea en el análisis de fiabilidad de sistemas, es la determinación del tiempo de vida de sus componentes, a partir de la observación del sistema.

Consideraremos un sistema coherente binario (Barlow y Proschan (1981), Zacks (1992), Aggarwal (1993)) constituido por n componentes cuyos tiempos de vida son independientes y poseen el mismo rango de variación. En el tiempo $t = 0$ todas las componentes funcionan (y por tanto, el sistema también). La observación del sistema en cualquier instante de tiempo proporciona su estado, funcionamiento o fallo, pero no el estado de cada una de sus componentes. En el instante de fallo del sistema se realiza una autopsia, esto es, se «abre», y se observa el estado de las componentes. Las leyes que rigen el funcionamiento del sistema, dadas por su función estructura, proporcionan un conocimiento adicional del tiempo en que pudieron fallar las componentes estropeadas.

Utilizaremos la siguiente notación:

- Z , el tiempo de vida del sistema coherente binario;
- $x = (x_1, \dots, x_n)$, el vector de estados de las componentes; x_i toma valores 1 ó 0 según la componente funcione o no;
- $x(t) = (x_1(t), \dots, x_n(t))$, el vector de estados de las componentes en el tiempo t ;
- $\phi(x)$, la función estructura del sistema;
- $T_1, \dots, T_n$, los tiempos de vida de las n componentes del sistema;
- $F_1, \dots, F_n$, sus funciones de distribución, respectivamente;
- $I = \{i/X_i \leq Z\}$, el conjunto de componentes que se encuentran deterioradas en el instante de fallo del sistema;
- $\{I_1, \dots, I_m\}$, el conjunto de cortes minimales;
- W , la matriz de incidencia $m \times n$ corte minimal-componente, es decir, $W(i, j) = 1$ si $j \in I_i$ y $W(i, j) = 0$ en otro caso;
- LS , el conjunto de componentes que pertenecen a todos los cortes minimales. Obsérvese que cuando alguna componente de LS está en funcionamiento, entonces el sistema también funciona.

Un subconjunto A del conjunto de componentes $C = \{1, 2, \dots, n\}$ se llama fatal si puede ser observado con probabilidad positiva en la autopsia del sistema, es decir, el conjunto de componentes A es un corte que verifica $P\{I = A\} > 0$.

Observemos que en estos sistemas, todo corte minimal es un conjunto fatal. Sin embargo, es evidente que no todo corte es un conjunto fatal. Así, por ejemplo, en un sistema en serie, cualquier subconjunto formado por dos componentes es un corte, pero no es fatal.

Para cada conjunto fatal A denotamos con D_A la intersección de todos los cortes minimales que están contenidos en A . Por tanto, D_A es el conjunto de componentes que han podido fallar simultáneamente con el sistema. Es inmediato comprobar que un conjunto de componentes A es fatal si y sólo si $D_A \neq \emptyset$.

El problema de identificación que se plantea consiste en saber cuándo la información obtenida de la autopsia del sistema, esto es, el conocimiento del tiempo Z transcurrido hasta el fallo del sistema y del conjunto fatal A que lo ocasiona, junto con el conocimiento de la función estructura, permite determinar la distribución de los tiempos de vida de las componentes del sistema. Es decir, cuándo la distribución conjunta de Z e I determina unívocamente el vector de las funciones de distribución $F = (F_1, \dots, F_n)$.

Meilijson (1981) determina condiciones suficientes para la identificación de las funciones de distribución, que se recogen en el siguiente resultado:

Si las variables aleatorias $T_1, \dots, T_n$ no toman ningún valor con probabilidad estrictamente positiva, son independientes y poseen los mismos extremos esenciales, y si el rango de W es n , entonces la distribución conjunta de Z e I determina unívocamente la distribución de cada T_i .

La información obtenida de la autopsia del sistema es $\{Z = t, I = A\}$, que, en términos de los tiempos de vida de las componentes, es equivalente a saber que

$$\left\{ \max_{j \in A - D_A} T_j < t, \max_{j \in D_A} T_j = t, \min_{j \notin A} T_j > t \right\}$$

Por tanto, llamando $G_A(t) = P\{Z \leq t, I = A\}$ y en particular, para los cortes minimales $G_i(t) = P\{Z \leq t, I = I_i\}$, se verifica que:

$$G_A(t) = \int_0^t (\prod_{A-D_A} F_j(s)) (\prod_{A^c} (1 - F_j(s))) d(\prod_{D_A} F_j(s))$$

$$G_i(t) = \int_0^t (\prod_{I_i^c} (1 - F_j(s))) d(\prod_{I_i} F_j(s))$$

Operando sobre esta última expresión, se obtiene la siguiente relación:

$$F(t) = \exp\left\{V \log \int_0^t \exp\{-\bar{W} \log(1 - F(s))\} dG(s)\right\}$$

donde $\bar{W}(i, j) = 1 - W(i, j)$, $V = (W^t W)^{-1} W^t$ y $G = (G_1, \dots, G_m)$.

Meilijson demuestra que este sistema de ecuaciones implícito para las funciones de distribución F_i posee una solución única, y propone un método iterativo de tipo Newton-Kantorovich para su resolución. La descripción de este método puede consultarse en Kantorovich y Akilov (1982) y Yamamoto (1986).

Posteriormente, Nowik (1990) establece condiciones necesarias y suficientes para la identificación:

Sea Φ una función estructura coherente binaria. Supongamos que los tiempos de vida de las componentes T_i , $i \in C$, son independientes. Supongamos, además, que las F_i , $i \in C$, son mutuamente y absolutamente continuas y que cada una posee un único punto con probabilidad positiva en el ínfimo esencial (común). Entonces:

- i) Una condición necesaria y suficiente para la identificabilidad de todas las distribuciones F_i , $i \in C$, es que el conjunto LS contenga a lo sumo una componente.*
- ii) Las distribuciones F_i , $i \in C - LS$, y $\prod_{i \in LS} F_i$ son identificables.*

De un modo similar a como lo hizo Meilijson, en la demostración del teorema se obtiene un sistema implícito de ecuaciones para las funciones de distribución y se propone para su resolución el mismo método iterativo Newton-Kantorovich.

En el artículo de Antoine, Doss y Hollander (1993), se prueba una condición suficiente para la identificación, más general que la de Meilijson, suponiendo que las distribuciones son funciones analíticas. Se establece además una condición necesaria para la identificación de las componentes, que es también suficiente para componentes cuyas distribuciones pertenecen a una cierta familia paramétrica, que incluye la distribución exponencial, la distribución normal truncada positiva, y las distribuciones Gamma y Weibull con parámetros enteros.

Los resultados anteriores se fundamentan en propiedades de la matriz de incidencia corte minimal-componente o conjunto fatal-componente. Estas matrices, incluso para sistemas pequeños, pueden tener grandes dimensiones, lo que las hace intratables numéricamente.

Así, por ejemplo, un sistema k-out-of-n tiene $\binom{n}{k}$ cortes minimales distintos, por lo que un sistema pequeño con $n = 20$ y $k = 10$ tendría una matriz asociada corte minimal-componente de dimensión 184756×20 . Resulta, pues, impensable siquiera

plantearse en la práctica la resolución del sistema de ecuaciones para la identificación. Este hecho nos lleva a buscar métodos alternativos más sencillos para la identificación de las componentes del sistema.

Como señala Nowik, en algunas ocasiones, es posible determinar la función de distribución de ciertas componentes de forma directa, sin necesidad de resolver ningún sistema de ecuaciones.

Así, si para la componente r existe un conjunto fatal A que no contiene a r y $B = A \cup \{r\}$ es fatal, con $D_A = D_B$, entonces F_r es la única función no decreciente y continua por la derecha de entre todas las que son iguales a $\frac{1}{1 + \frac{dG_A}{dG_B}}$.

Existen sistemas que se rigen por una función estructura que permite la identificación directa de la función de distribución del tiempo de vida de algunas o de todas sus componentes. Es el caso de los sistemas aditivos (p, α) que se definen en la sección 2, en la que obtendremos las condiciones necesarias y suficientes para que este tipo de sistemas posea componentes identificables de forma directa, utilizando desigualdades lineales para variables binarias. Finalmente, se incluye, a modo de ejemplo, la aplicación de estas técnicas para la identificación de las componentes de un sistema aditivo (p, α) .

En la sección 3, estudiamos la identificabilidad de las componentes de los sistemas binarios serie-paralelo. Demostraremos que para estos sistemas todas las componentes son directamente identificables, excepto las pertenecientes a los módulos formados por una única componente.

En la sección 4, realizamos un estudio similar de los sistemas binarios paralelo-serie, comprobando que todas las componentes que pertenecen a un módulo con más de dos componentes se pueden identificar directamente.

2. SISTEMAS ADITIVOS (p, α)

2.1. Definición y propiedades

Definición 2.1 *Un sistema se dice binario aditivo (p, α) si su función estructura puede expresarse como:*

$$\phi(x) = 1_{\{\sum_{i=1}^n p_i x_i > 1 - \alpha\}}$$

donde $0 < p_i < 1$, $i = 1, \dots, n$; $\sum_{i=1}^n p_i = 1$; $0 < \alpha \leq 1$; y $1_{\{a\}}$ es la función indicador, que toma valor 1 si a es cierto y toma valor 0 si a es falso. Mediante p se denota el vector $(p_1, \dots, p_n)$ que llamaremos vector de pesos.

Se puede considerar que el valor de p_i mide la contribución de la componente i al funcionamiento del sistema, en el sentido de que para que este sistema falle es preciso que la suma de los valores p_i asociados a las componentes deterioradas alcance, al menos, el umbral α ; o, lo que es lo mismo, el sistema funciona mientras la suma de las cantidades p_i asociadas a las componentes en funcionamiento supera el nivel $1 - \alpha$.

Como ejemplo, consideremos un sistema productivo formado por un conjunto de n máquinas, cada una de las cuales puede estar en dos estados: funcionamiento o fallo. La máquina i opera con una tasa de producción λ_i si funciona, y 0 si no funciona. Supongamos que el sistema únicamente se considera operativo si el nivel global de producción supera cierto nivel P . Para este sistema, $\phi(x) = 1_{\{\sum_{i=1}^n \lambda_i x_i > P\}}$ es decir, se trata de un sistema aditivo (p, α) con $p_i = \frac{\lambda_i}{\sum_j \lambda_j}$ y $1 - \alpha = \frac{P}{\sum_j \lambda_j}$.

Denotamos por d_A a la diferencia $\sum_{i \in A} p_i - \alpha$.

Conjuntos fatales y cortes minimales

En un sistema binario aditivo (p, α) , un conjunto de componentes A es fatal si $d_A \geq 0$ y la intersección de todos los cortes minimales en A es distinta del vacío; es decir, $D_A \neq \emptyset$, que ocurre cuando $\exists j \in A$ cumpliendo $p_j > d_A$.

Un corte minimal es un conjunto I_k de componentes que cumple

$$d_{I_k} \geq 0 \quad \text{de forma que} \quad \forall j \in I_k, \quad p_j > d_{I_k}.$$

El siguiente resultado proporciona una caracterización de los cortes minimales:

Proposición 2.2 *Existe una aplicación biyectiva entre el conjunto de los cortes minimales $\{I_k\}$ y el conjunto de vectores $\{(Y_1, \dots, Y_n)\}$ que verifican las siguientes restricciones:*

$$(1) \quad \left. \begin{array}{ll} \sum_{i=1}^n p_i Y_i \geq \alpha & (a) \\ \sum_{i=1}^n p_i Y_i + (1 - p_j) Y_j < 1 + \alpha & \forall j = 1, \dots, n \quad (b) \\ Y_i = \{0, 1\} & \forall i = 1, \dots, n \end{array} \right\}$$

Demostración: Sea $(Y_1, \dots, Y_n)$ un punto factible para las restricciones anteriores e I_k el conjunto de componentes definido como $I_k = \{i/Y_i = 1\}$. Por la restricción (a) es claro que las componentes de I_k cuando fallan ocasionan el fallo del sistema.

Veamos ahora que el deterioro de las componentes de un subconjunto propio de I_k no puede causar el fallo del sistema, para lo cual es suficiente ver que $\sum_{i \in I_k} p_i - p_j < \alpha$ $\forall j \in I_k$:

$$\sum_{i \in I_k} p_i - p_j < \alpha \iff \sum_{i=1}^n p_i Y_i - p_j Y_j < \alpha \iff \sum_{i=1}^n p_i Y_i - p_j Y_j + Y_j - 1 < \alpha$$

que se corresponde con una de las restricciones de tipo (b); por tanto I_k es conjunto minimal.

Sea ahora I_k un corte minimal. Consideramos el vector binario $(Y_1, \dots, Y_n)$ definido como $Y_i = 1$ si la componente $i \in I_k$, $Y_i = 0$ si la componente $i \notin I_k$. Es inmediato, debido a que el fallo de las componentes que están en I_k ocasionan el fallo del sistema, que $(Y_1, \dots, Y_n)$ cumple (a). La restricción de (b) correspondiente a una componente j cuyo $Y_j = 1$, es decir, que pertenece a I_k , equivale a que la suma de los deterioros de las componentes pertenecientes a I_k , menos el de la componente j , sea menor que α , lo cual es cierto por la definición del conjunto I_k . Las restricciones de (b) correspondientes a componentes j cuyo $Y_j = 0$ se verifican trivialmente. ■

Casos particulares de sistemas aditivos (p, α)

- i) Si $\alpha \leq \min_i \{p_i\}$, entonces el sistema se comporta como un sistema en serie, ya que el deterioro de una cualquiera de sus componentes provoca el fallo del sistema.
- ii) Si $1 - \alpha < \min_i \{p_i\}$, entonces el sistema se comporta como un sistema en paralelo, ya que el fallo del sistema se produce únicamente cuando se estropean todas sus componentes.
- iii) Si $p_i = p = 1/n$, $\forall i = 1, \dots, n$, entonces el sistema se comporta como un sistema k-out-of-n, con $(k-1)/n < \alpha \leq k/n$.

2.2. Identificación directa de componentes

Lema 2.3 *Un conjunto fatal A posibilita la identificación directa de la función de distribución de la componente r si $r \notin A$ y $\forall i \in A$ se cumple*

- (I) $\sum_{j \neq i \in A} p_j \geq \alpha$, o bien
- (II) $\sum_{j \neq i \in A} p_j < \alpha - p_r$,

y al menos para una componente i se verifica esta última relación.

Demostración: Veamos que el conjunto fatal A verifica las condiciones de identificación directa enunciadas en la introducción, es decir, $B = A \cup \{r\}$ es fatal y se cumple que $D_A = D_B$.

Como $\sum_{i \in B} p_i = \sum_{i \in A} p_i + p_r$, entonces $j \in D_B$ si y sólo si $p_j > d_A + p_r$. Esta relación implica que $r \notin D_B$.

Además, dado que $j \in D_A$ si y sólo si $p_j > d_A$, $D_A = D_B$ si y sólo si se cumple que $\forall j \in A, p_j \notin (d_A, d_A + p_r]$, siendo $D_A = D_B \neq \emptyset$, si al menos una componente $j \in A$ verifica $p_j > d_A + p_r$.

Es decir, $\forall i \in A, p_i \leq d_A$ o $p_i > d_A + p_r$, y al menos una componente verifica la última desigualdad; lo cual equivale a que $\forall i \in A$ se cumple $\sum_{j \neq i \in A} p_j \geq \alpha$, (I), o bien $\sum_{j \neq i \in A} p_j < \alpha - p_r$, (II), y al menos para una componente i es válida esta última relación.

Notemos que B es fatal por ser $\sum_{i \in B} p_i \geq \sum_{i \in A} p_i \geq \alpha$ y $D_B \neq \emptyset$.

■

El siguiente teorema proporciona una caracterización de la identificabilidad de una componente en términos de un sistema de restricciones lineales para un conjunto de variables binarias.

Teorema 2.4 Dado un sistema aditivo (p, α) y con pesos de las componentes $p_1, \dots, p_n$, la componente r -ésima es identificable directamente cuando la siguiente región de factibilidad es no vacía:

$$(2) \quad \left. \begin{array}{ll} \sum_{j \neq r} p_j Y_j \geq \alpha & (a) \\ \sum_{j \neq i, r} p_j Y_j + Z_i \geq \alpha & \forall i = 1, \dots, n, \quad i \neq r \quad (b) \\ \sum_{j \neq i, r} p_j Y_j + p_r + (Z_i - 1) < \alpha & \forall i = 1, \dots, n, \quad i \neq r \quad (c) \\ \sum_i Z_i \geq 1 & (d) \\ Y_j = \{0, 1\} & \forall j = 1, \dots, n, \quad j \neq r \\ Z_i = \{0, 1\} & \forall i = 1, \dots, n, \quad i \neq r \end{array} \right\}$$

Demostración: Veamos que existe una aplicación biyectiva entre todos los conjuntos fatales A que permiten la identificación directa de la componente r y los vectores factibles $Y = (Y_1, \dots, Y_n)$ de la región (2) anterior.

Para ello, primeramente, veremos que, dado un conjunto fatal A que permite la identificación directa, los valores de las variables:

$$Y_j = 1 \text{ si } j \in A$$

$$Y_j = 0 \text{ si } j \notin A$$

$$Z_j = 0 \text{ si } j \notin A \text{ o si } j \in A \text{ con } p_j \leq d_A$$

$$Z_j = 1 \text{ si } j \in A \text{ y } p_j > d_A + p_r$$

proporcionan un punto factible de (2).

La restricción (a) se cumple por ser A un conjunto fatal.

Comprobamos el cumplimiento del bloque de restricciones (b). Dada una componente i ,

- si $i \notin A$, la restricción (b) se reduce a la restricción (a), y por tanto, se cumple.
- si $i \in A$, existen dos situaciones posibles:
 1. La componente i verifica la condición (I) del lema 2.3, en cuyo caso, $Z_i = 0$ y (b) equivale a (I).
 2. La componente i verifica (II) del lema 2.3, en cuyo caso, $Z_i = 1$ y (b) se satisface trivialmente.

Comprobamos ahora el cumplimiento de las restricciones (c). Para una componente i dada,

- si $i \notin A$, entonces $Z_i = 0$ y la restricción (c) es trivial.
- si $i \in A$, existen dos situaciones posibles:
 1. La componente i verifica la condición (I) del lema 2.3, entonces, $Z_i = 0$ y (c) es trivial.
 2. La componente i verifica (II) del lema 2.3, entonces $Z_i = 1$ y (c) equivale a (II).

La restricción (d) se cumple ya que, por el lema anterior, al menos una componente de A satisface la condición (II).

Veremos a continuación que a cada punto factible de (2) se le puede asociar un conjunto fatal A que permite la identificación directa de la componente r .

Definimos este conjunto A como $A = \{i/Y_i = 1\}$, que demostraremos que es fatal y que cumple las condiciones del lema 2.3.

Sea $i \in A$, es decir, $Y_i = 1$; entonces,

Si $Z_i = 0$, (b) equivale a la condición (I) del lema.

Si $Z_i = 1$, (c) equivale a la condición (II) del lema, (en este caso $i \in D_A$).

Notemos que si $Z_i = 1$, entonces también $Y_i = 1$, ya que si fuera $Y_i = 0$, entonces (a) y (c) serían contradictorias.

La restricción (d) asegura que al menos una componente $i \in A$ verifica $Z_i = 1$, es decir, $D_A \neq \emptyset$, que junto con (a) implican que A es fatal.

■

Corolario 2.5 *Si la componente k es directamente identificable, entonces también lo son las componentes i con $p_i \leq p_k$.*

Demostración: Veamos que si la componente k es directamente identificable, también lo es la componente $k - 1$.

Para ello, basta observar que si para todo conjunto fatal A que permite la identificación directa de k se tiene $\{k - 1\} \notin A$, entonces A permite también la identificación directa de $k - 1$. En otro caso, el conjunto fatal $E = \{k\} \cup A - \{k - 1\}$ permite identificar directamente dicha componente.

■

Corolario 2.6 *En un sistema aditivo (p, α) con todos los pesos de las componentes iguales a p , excepto una componente q con peso $3p$, es posible la identificación de todas las componentes.*

Demostración: La identificación de la distribución de una componente r cualquiera con peso p se puede realizar utilizando el conjunto A de componentes que contiene a la componente q y otras k componentes distintas de r , con $k + 1 = \max\{m \text{ entero tal que } mp \leq \alpha\}$. Definiendo $Y_i = 1$ para las componentes en A e $Y_i = 0$ para el resto, $Z_q = 1$ y $Z_i = 0 \forall i \neq q$, es inmediato comprobar por sustitución que se trata de un punto factible para el sistema (2) del teorema anterior.

Una vez identificadas las distribuciones de todas las componentes menos la de q , ésta puede obtenerse de forma directa despejándola de la expresión para la distribución

del tiempo en que se observa como conjunto fatal el conjunto A definido anteriormente:

$$dG_A(t) \doteq \left(\prod_{j \neq q \in A} F_j(t) \right) \left(\prod_{A^c} (1 - F_j(t)) \right) d(F_q(t))$$

■

2.3. Ejemplos prácticos

Ejemplo 1 Cálculo de los conjuntos fatales que permiten la identificación directa

Consideremos un sistema aditivo (p, α) con 10 componentes.

Los pesos de estas componentes son : $p_1 = 0.025$, $p_2 = 0.025$, $p_3 = 0.05$, $p_4 = 0.05$, $p_5 = 0.05$, $p_6 = 0.1$, $p_7 = 0.1$, $p_8 = 0.15$, $p_9 = 0.15$, $p_{10} = 0.3$. Supongamos que el sistema falla cuando la suma de los pesos de las componentes estropeadas es de al menos 0.5.

Para cada componente se plantea el correspondiente sistema de restricciones (2) y se busca, si existe, un punto factible. La localización de un punto factible puede realizarse resolviendo un problema de optimización binaria que tiene por restricciones el sistema (2) y por función objetivo la función nula, con lo que todo punto factible es solución óptima.

La identificación directa de la función de distribución de cada una de las diez componentes tiene asociado un sistema (2) con, a lo sumo, 18 variables binarias y 20 restricciones. La optimización de la función nula sobre la región factible (2) proporciona los siguientes conjuntos A fatales que permiten la identificación:

Componente	Conjunto fatal A	D_A	$\sum_{i \in A} p_i$
1	{ 8, 9, 10 }	{ 8, 9, 10 }	0.6
2	{ 8, 9, 10 }	{ 8, 9, 10 }	0.6
3	{ 6, 8, 9, 10 }	{ 10 }	0.7
4	{ 6, 8, 9, 10 }	{ 10 }	0.7
5	{ 6, 8, 9, 10 }	{ 10 }	0.7
6	{ 3, 8, 9, 10 }	{ 10 }	0.65
7	{ 3, 8, 9, 10 }	{ 10 }	0.65
8	{ 3, 4, 6, 7, 10 }	{ 10 }	0.6
9	{ 3, 4, 6, 7, 10 }	{ 10 }	0.6

Además, el conjunto fatal $A = \{ 6, 8, 9, 10 \}$ permite la identificación de la función de distribución de la componente 10:

$$dG_A(t) = (1 - F_1(t))(1 - F_2(t))(1 - F_3(t))(1 - F_4(t))(1 - F_5(t))(1 - F_7(t))F_6(t)F_8(t)F_9(t)dF_{10}(t)$$

$$F_{10}(t) = \int_0^t \frac{dG_A(s)}{(1 - F_1(s))(1 - F_2(s))(1 - F_3(s))(1 - F_4(s))(1 - F_5(s))(1 - F_7(s))F_6(s)F_8(s)F_9(s)}$$

Sin embargo, si cambiamos el nivel de fallo α de 0.5 a 0.66, entonces únicamente son identificables directamente las componentes 1, 2, 3, 4, 5, 6 y 7, pero no así las componentes 8 y 9, cuyos sistemas asociados son incompatibles.

Componente	Conjunto fatal A	D_A	$\sum_{i \in A} p_i$
1	$\{ 6, 8, 9, 10 \}$	$\{ 6, 8, 9, 10 \}$	0.7
2	$\{ 6, 8, 9, 10 \}$	$\{ 6, 8, 9, 10 \}$	0.7
3	$\{ 6, 8, 9, 10 \}$	$\{ 6, 8, 9, 10 \}$	0.7
4	$\{ 6, 8, 9, 10 \}$	$\{ 6, 8, 9, 10 \}$	0.7
5	$\{ 6, 8, 9, 10 \}$	$\{ 6, 8, 9, 10 \}$	0.7
6	$\{ 3, 4, 5, 7, 8, 9, 10 \}$	$\{ 10 \}$	0.85
7	$\{ 3, 4, 5, 7, 8, 9, 10 \}$	$\{ 10 \}$	0.85

Ejemplo 2 Identificación del tiempo de vida de las componentes

Consideremos un sistema aditivo con $\alpha = 0.7$ y cuatro componentes, cuyos pesos son $p_1 = 0.1$, $p_2 = 0.1$, $p_3 = 0.2$, $p_4 = 0.6$. Se conoce para los conjuntos fatales

$$A_1 = \{2, 3, 4\}, A_2 = \{1, 3, 4\}, A_3 = \{1, 2, 4\} \quad \text{y} \quad B = \{1, 2, 3, 4\},$$

sus funciones $G_{A_1}(t)$, $G_{A_2}(t)$, $G_{A_3}(t)$ y $G_B(t)$ de modo que:

$$G_{A_1}(t) = (-3e^{-2t} + 3e^{-4t} - e^{-6t} + 1)/6$$

$$G_{A_2}(t) = G_{A_3}(t) = (-20e^{-3t} + 15e^{-4t} + 12e^{-5t} - 10e^{-6t} + 3)/60$$

$$G_B(t) = (-30e^{-t} + 15e^{-2t} + 20e^{-3t} - 15e^{-4t} - 6e^{-5t} + 5e^{-6t} + 11)/30$$

Por tanto,

$$\begin{aligned}dG_{A_1}(t) &= (e^{-2t} - 2e^{-4t} + e^{-6t})dt \\dG_{A_2}(t) = dG_{A_3}(t) &= (e^{-3t} - e^{-4t} - e^{-5t} + e^{-6t})dt \\dG_B(t) &= (e^{-t} - e^{-2t} - 2e^{-3t} + 2e^{-4t} + e^{-5t} - e^{-6t})dt\end{aligned}$$

- *Identificación de $F_1(t)$*

Los conjuntos de componentes A_1 y B permiten la identificación directa de $F_1(t)$,

$$F_1(t) = \frac{1}{1 + \frac{dG_{A_1}(t)}{dG_B(t)}} = 1 - e^{-t}$$

que corresponde a una variable exponencial de media 1.

- *Identificación de $F_2(t)$ y $F_3(t)$*

Los conjuntos de componentes A_2 , B y A_3 , B permiten la identificación directa de $F_2(t)$ y $F_3(t)$, respectivamente.

$$F_2(t) = F_3(t) = \frac{1}{1 + \frac{dG_{A_2}(t)}{dG_B(t)}} = 1 - e^{-2t}$$

que corresponde a una variable exponencial de media 1/2.

- *Identificación de $F_4(t)$*

Una vez conocidas las funciones $F_1(t)$, $F_2(t)$ y $F_3(t)$, la función $F_4(t)$ puede calcularse a partir de cualquier conjunto fatal que contenga a la componente 4; por ejemplo, en este caso, el conjunto B :

$$dG_B(t) = F_1(t)F_2(t)F_3(t)dF_4(t)$$

Por tanto,

$$F_4(t) = \int_0^t \frac{dG_B(s)}{F_1(s)F_2(s)F_3(s)} = \int_0^t e^{-s} ds = 1 - e^{-t}$$

que corresponde a una variable exponencial de media 1.

3. SISTEMAS BINARIOS SERIE-PARALELO

3.1. Definición y propiedades

Sea un sistema binario con n componentes distribuidas en N subconjuntos disjuntos de componentes que denominaremos módulos, colocados en serie, a los que denotamos $M_1, \dots, M_N$. Cada uno de estos módulos está formado por $\#M_i$, $i = 1, \dots, N$, componentes dispuestas en paralelo. El sistema resultante se denomina sistema binario serie-paralelo.

Denotando por x_j^i el estado de la componente j -ésima del módulo i , la función estructura de un sistema binario serie-paralelo viene dada por:

$$\phi(x) = \min_{i=\{1,\dots,N\}} \{ \max_{j \in \{1,\dots,\#M_i\}} x_j^i \}$$

Conjunto fatales y cortes minimales

Los únicos cortes minimales de un sistema binario serie-paralelo son los conjuntos M_i con $i = 1, \dots, N$.

Un conjunto de componentes A es fatal si existe un único $j \in \{1, \dots, N\}$ tal que $M_j \subset A$. Siendo, por tanto, $D_A = M_j$.

El número de conjuntos fatales en un sistema binario serie-paralelo es

$$\sum_{j=1}^N \prod_{i \neq j} (2^{\#M_i} - 1)$$

que se demuestra observando que, fijado el módulo M_j que está contenido en el conjunto fatal A , existen $\prod_{i \neq j} (2^{\#M_i} - 1)$ formas distintas válidas para el subconjunto $A - M_j$.

3.2. Identificación directa de componentes

En el teorema siguiente, se determinan las componentes que pueden identificarse de forma directa.

Teorema 3.1 *En un sistema binario serie-paralelo, todas las componentes son directamente identificables, salvo las pertenecientes a los módulos formados por una única componente.*

Demostración: Sea una componente r cualquiera de un módulo M_i , para $i \in \{1, \dots, N\}$, cumpliendo $\#M_i > 1$. Consideramos los conjuntos A y B siguientes:

$$A = M_j \text{ con } j \neq i \in \{1, \dots, N\} \quad \text{y} \quad B = M_j \cup \{r\}.$$

A y B son conjuntos fatales con $D_A = D_B = M_j$, por lo que la componente r es directamente identificable.

Sin embargo, si r es la única componente de un módulo, el fallo de esta componente ocasiona el fallo del sistema, por lo que si $r \in A$, entonces $D_A = \{r\}$, y no es posible su identificación directa. ■

Corolario 3.2 *En un sistema binario serie-paralelo, una condición necesaria y suficiente para que un conjunto fatal permita identificar directamente la componente r es que exista una componente distinta de r en su mismo módulo que se encuentre funcionando.*

Observación 3.3 El número de conjuntos fatales que permiten identificar una componente $r \in M_j$ es:

$$(2^{\#M_i} - 2) \sum_{j=1}^N \prod_{\substack{h \neq j, i \\ j \neq i}} (2^{\#M_h} - 1)$$

que se obtiene comprobando que el número de posibles formas en las que pueden encontrarse los $N - 1$ módulos que no contienen a r es, $\sum_{j=1}^N \prod_{\substack{h \neq j, i \\ j \neq i}} (2^{\#M_h} - 1)$. Además, por el corolario 3.2, en el módulo i , debe existir al menos una componente distinta de r funcionando, por tanto, las $\#M_i - 1$ componentes distintas de r , pueden encontrarse en $(2^{\#M_i-1} - 1)$ formas diferentes, y, en cada una de éstas, la componente r puede estar en funcionamiento o en estado de fallo. □

4. SISTEMAS BINARIOS PARALELO-SERIE

4.1. Definición y propiedades

Un sistema binario paralelo-serie está formado por N módulos $M_1, \dots, M_N$, dispuestos en paralelo, de forma que las $\#M_i$ componentes que constituyen el módulo

M_i están colocadas en serie. La función estructura de estos sistemas es:

$$\phi(x) = \max_{i=\{1,\dots,N\}} \left\{ \min_{j \in \{1,\dots,\#M_i\}} x_j^i \right\}$$

Conjuntos fatales y cortes minimales

En un sistema binario paralelo-serie, los cortes minimales están formados por una componente de cada módulo: $I_h = \{j_1, j_2, \dots, j_N\}$, con $j_k \in M_k$. Por tanto, el número de cortes minimales es $\prod_{k=1}^N \#M_k$.

Un conjunto de componentes A es fatal si existe al menos un módulo M_j con una única componente fallada.

El número de conjuntos fatales es:

$$\sum_{i=1}^N (\#M_i) \prod_{j=1}^{i-1} (2^{\#M_j} - 1 - \#M_j) \prod_{j=i+1}^N (2^{\#M_j} - 1)$$

Contamos el número total de conjuntos fatales condicionando con el módulo del sistema de menor índice que tiene una única componente deteriorada, (módulo i). Así, los $i-1$ primeros módulos M_j tienen que tener al menos dos componentes deterioradas, y por tanto, $(2^{\#M_j} - 1 - \#M_j)$ configuraciones para cada uno de ellos, y los $N-i$ últimos módulos M_j tienen al menos una componente en estado cero, es decir, $(2^{\#M_j} - 1)$ posibilidades de elección en cada módulo. Sumando para todos los posibles valores de i , se tiene el número anterior.

4.2. Identificación directa de componentes

Analizamos a continuación las componentes identificables directamente en estos sistemas.

Teorema 4.1 *En un sistema binario paralelo-serie, todas las componentes son directamente identificables, excepto las pertenecientes a un módulo con dos o menos componentes.*

Demostración: Sean r, u, v componentes de un módulo M_k , y j_i una componente cualquiera del módulo M_i , $i \neq k$.

Entonces, el conjunto fatal $A = \{j_i/i = 1, \dots, N; i \neq k\} \cup \{u, v\}$, y el conjunto fatal $B = A \cup \{r\}$ permiten la identificación de la componente r , puesto que verifican $D_A = D_B = \{j_i/i = 1, \dots, N; i \neq k\}$.

Si $r, s \in M_k$, con $\#M_k = 2$, basta observar que el funcionamiento o fallo de la componente r implica que la componente s pertenezca o no al subconjunto intersección de todos los cortes minimales contenidos en el conjunto fatal, por lo que no existen dos conjuntos fatales A y B cumpliendo las hipótesis de identificación. Si $r \in M_k$, con $\#M_k = 1$, obviamente $r \in D_A$, para cualquier conjunto fatal A , por lo que tampoco es identificable. ■

Corolario 4.2 *En un sistema binario paralelo-serie, una condición necesaria y suficiente para que un conjunto fatal permita identificar directamente la componente r es que existan dos componentes distintas de r en su mismo módulo que hayan fallado.*

Observación 4.3 El número de conjuntos fatales que permiten identificar una componente $r \in M_h$ es:

$$2(2^{\#M_h-1} - \#M_h) \sum_{\substack{i=1 \\ i \neq h}}^N \#M_i \prod_{\substack{j=1 \\ j \neq h}}^{i-1} (2^{\#M_j} - 1 - \#M_j) \prod_{\substack{j=i+1 \\ j \neq h}}^N (2^{\#M_j} - 1)$$

Razonando como en resultados anteriores, basta tener en cuenta que las $\#M_h - 1$ componentes distintas de r del módulo h , pueden encontrarse de $(2^{\#M_h-1} - \#M_h)$ formas diferentes, y que la componente r , puede estar en cualquier estado: 0 ó 1. □

5. OBSERVACIONES FINALES

1. El modelo de autopsia tratado en este artículo fue introducido por Meilijson (1981), y supone una generalización del modelo *Competing Risks* clásico, que considera únicamente sistemas en serie. Los principales resultados de este modelo clásico pueden encontrarse en Nadas (1970), Moeschberger (1974), Gail (1975), Tsiatis (1975) y Prentice y otros (1978).
2. El problema de identificación de funciones de distribución mediante autopsia en sistemas multiestado, (estos sistemas son tratados, por ejemplo, en El-Newehi y Proschan (1984), y Avent (1993)), ha sido estudiado por Costa Bueno (1988). Este autor extiende directamente los resultados de Meilijson para el caso binario. Para estos sistemas multiestado también es posible aplicar el concepto de identificación directa de un modo similar al presentado aquí para sistemas binarios.

3. El estudio de este problema de identificación desde un punto de vista estadístico fue abordado por primera vez por Watelet (1990), que propuso un estimador no paramétrico reemplazando las funciones de distribución teóricas por empíricas, en el sistema de ecuaciones implícito propuesto por Meilijson, aunque no consiguió obtener resultados teóricos de carácter general.
4. En el trabajo de Meilijson (1994), suponiendo funciones de distribución pertenecientes a familias paramétricas con buenas propiedades, se estima, mediante máxima verosimilitud, el conjunto de parámetros desconocidos. En este trabajo, además, se propone un modelo de autopsia más general, permitiendo conocer el tiempo exacto de fallo de algunas componentes, y de otras en ocasiones, dependiendo del estado de las anteriores.

AGRADECIMIENTOS

Los autores desean agradecer las indicaciones y sugerencias formuladas por dos expertos anónimos que sin duda han contribuido a mejorar este trabajo.

BIBLIOGRAFÍA

- [1] **Aggarwal, K.K.** (1993). *Reliability Engineering*. Kluwer Academic Publishers.
- [2] **Antoine, R., H. Doss y M. Hollander** (1993). «On the identifiability in the autopsy model of reliability theory». *Journal of Applied Probability*, **30**, 913–930.
- [3] **Avent, T.** (1992). *Reliability and Risk Analysis*. Elsevier Applied Science.
- [4] **Barlow, R.E. y F. Proschan** (1981). *Statistical Theory of Reliability and Life Testing*. To Begin with, Silver Spring, Ms.
- [5] **Costa Bueno, V.** (1988). «A note on the component lifetime estimation of a multistate monotone system through the system lifetime». *Advances Applied Probability*, **20**, 686–689.
- [6] **El-Newehi, E. and F. Proschan** (1984). «Degradable systems: a survey of multistate coherent systems». *Commun. Statist. - Theory Math.*, **13**, 405–432.
- [7] **Gail, M.** (1975). «A review and critique of some models used in competing risk analysis». *Biometrics*, **31**, 209–222.
- [8] **Kantorovich, L.V. y G.P. Akilov** (1982). *Functional Analysis*. Pergamon Press.
- [9] **Meilijson, I.** (1981). «Estimation of the lifetime distribution of the parts from the autopsy statistics of the machine». *Journal of Applied Probability*, **18**, 829–838.

- [10] **Meilijson, I.** (1994). «Competing risks on coherent reliability systems: estimation in the parametric case». *Journal of the American Statistical Association*, **89**, 1459–1464.
- [11] **Moeschberger, M.L.** (1974). «Life tests under dependent causes of failure». *Technometrics*, **16**, 39–47.
- [12] **Nadas, A.** (1970). «On estimating the distribution of a random vector when only the smallest coordinate is observable». *Technometrics*, **12**, 923–924.
- [13] **Nowik, S.** (1990). «Identifiability Problems in Coherent Systems». *Journal of Applied Probability*, **28**, 862–872.
- [14] **Prentice, R.L.; Kalbfleisch, J.D.; Peterson, A.V.; Farewell, V.T. and N.E. Breslow** (1978). «The analysis of failure times in the presence of competing risks». *Biometrics*, **34**, 541–554.
- [15] **Tsiatis, A.** (1975). «A nonidentifiability aspect in the problem of competing risks». *Proceedings of the National Academy of Science USA*, **72**, 20–22.
- [16] **Yamamoto, T.** (1986). «A method for finding sharp error bounds for Newton's method under the Kantorovich assumption». *Numer. Math.*, **49**, 203–220.
- [17] **Watelet, L.F.** (1990). «Nonparametric estimation of component life distributions in Meilijson's competing risk model». Ph.D. Thesis, University of Washington.
- [18] **Zacks, S.** (1992), *Introduction to Reliability Analysis. Probability Models and Statistical Models*. Springer-Verlag.

ENGLISH SUMMARY

DETERMINATION OF THE LIFETIME DISTRIBUTION OF THE COMPONENTS FROM THE AUTOPSY IN ADDITIVE, SERIES-PARALLEL AND PARALLEL-SERIES BINARY SYSTEMS

FERMÍN MALLOR GIMÉNEZ*

CRISTINA AZCÁRATE CAMIO*

ANTONIO PÉREZ PRADOS*

Universidad Pública de Navarra

In this paper, we analyze the question of determining the lifetime distributions of the binary system components, from the knowledge of the system functioning rules and from the joint distribution of the lifetime system and the component set that caused the death of the system.

We present the results obtained by some authors, who propose a system of implicit equations for the computation of those distributions. However, we observe that it has a difficult resolution in practice; therefore, we consider a method that can be used in fewer situations than the former, but with a rather simple computation. We study its application in the additive, series-parallel and parallel-series binary systems.

Keywords: Binary system, lifetime distribution, autopsy

AMS Classification: 62G05, 62N05, 90B25

*Departamento de Estadística e Investigación Operativa. Universidad Pública de Navarra. Campus de Arrosadía, s/n. 31006 Pamplona.

–Received january 1996.

–Accepted november 1996.

1. INTRODUCTION

The determination of the component lifetime distributions from the observation of the system is an analyzing problem in system reliability.

Let us consider a coherent binary system (Barlow and Proschan (1981), Zacks (1992), Aggarwal (1993)) composed by n components which act independently. At $t = 0$, every component and the system are working. We observe the system until it fails. The observed data obtained from the autopsy of the system consist both of the failure time of the system and of the set of components that are dead at the instant of the system failure (fatal set), without knowing the exact time at which each component has failed. The structure function of the system provides an additional information about the instant in which the failed components could have died.

From these autopsy data, the goal is to estimate the lifetime distribution functions of the components. In Meilijson (1981), Nowik (1990) and Antoine et al (1993), necessary and sufficient conditions for the identifiability, on some assumptions, are given. These results are based on the properties of the incidence matrix cut set-component or fatal set-component. These matrices are of large size, even for simple systems, and they are therefore complex from the point of view of numerical calculus.

Sometimes the distribution function of certain components can be directly determined, without solving any complex system of equations. For instance, it is possible for component r when there exists a fatal set A not containing r and if the set $B = A \cup \{r\}$ is also fatal, and both, A and B , have the same subset of components that may be failed simultaneously with the system.

The structure function of some systems allows the direct identification of the components. In this paper we analyze the direct identifiability of the (p, α) additive binary systems, the series-parallel, and the parallel-series binary systems.

2. (p, α) ADDITIVE BINARY SYSTEMS

In this section (p, α) additive binary system are introduced as follow:

Definition 2.1 A binary system is called (p, α) additive system if its function structure can be expressed as follows:

$$\Phi(\vec{x}) = 1_{\{\sum_{i=1}^n p_i x_i > 1 - \alpha\}}$$

where $0 < p_i < 1$, $i = 1, \dots, n$; $\sum_{i=1}^n p_i = 1$; $0 < \alpha \leq 1$; and $1_{\{a\}}$ is the indicator function,

which takes value 1 if a is true, and 0 if a is false. Let $p = (p_1, \dots, p_n)$ be the weight vector.

In this section minimal cut sets and fatal sets are analyzed and the directly identifiable components are characterized by means of the existence of feasible points in a set of linear constraints with binary variables.

Theorem 2.4 *Let a (p, α) additive binary system with weights $p_1, \dots, p_n$; any component r is directly identifiable if the following region is no empty:*

$$(3) \quad \left. \begin{array}{ll} \sum_{j \neq r} p_j Y_j \geq \alpha & (a) \\ \sum_{j \neq i, r} p_j Y_j + Z_i \geq \alpha & \forall i = 1, \dots, n, \quad i \neq r \quad (b) \\ \sum_{j \neq i, r} p_j Y_j + p_r + (Z_i - 1) < \alpha & \forall i = 1, \dots, n, \quad i \neq r \quad (c) \\ \sum_i Z_i \geq 1 & (d) \\ Y_j = \{0, 1\} & \forall j = 1, \dots, n, \quad j \neq r \\ Z_i = \{0, 1\} & \forall i = 1, \dots, n, \quad i \neq r \end{array} \right\}$$

Observe that if the r component can be directly identified, then so can the i component with $p_i \leq p_r$.

We conclude this section with two illustrative examples.

3. SERIES-PARALLEL SYSTEMS

A series-parallel system consists of n components distributed in K disjoint subsets, that we say modules, series-arranged, $M_1, \dots, M_K$. Each of these modules has $\#M_i$ components, ($i = 1, \dots, K$), in parallel.

The next theorem provides the directly identifiable components of these systems.

Theorem 3.1 *In a series-parallel binary system, all the components can be directly identified, except those belonging to a module that consists of a single component.*

4. PARALLEL-SERIES BINARY SYSTEMS

A parallel-series binary system consists of K parallel-arranged modules, such that $\forall i = 1, \dots, K$ the $\#M_i$ components of the module M_i being series-arranged.

We shall provide the components that can be directly identified in a parallel-series binary system.

Theorem 4.1 *In a parallel-series binary system, all the components can be directly identified, except those belonging to a module with less than three components.*

A TABU SEARCH ALGORITHM TO SCHEDULE UNIVERSITY EXAMINATIONS*

RAMÓN ÁLVAREZ-VALDÉS*

ENRIC CRESPO**

JOSÉ M. TAMARIT*

Universitat de València

Scheduling examinations in a large university is an increasingly complex problem, due to its size, the growing flexibility of students' curricula and the interest in including a wide set of objectives and constraints. In this paper we present a new algorithm for this problem and its application to a university in Spain.

A combination of heuristics, based on Tabu Search, first finds a solution in which no student has two exams simultaneously and then improves it by evenly spacing the exams in the examination period. The algorithm has been imbedded into a package to be used by Faculty administrators and proposes solutions to be considered by the parties involved: administration, departments and students. Computational results, comparing our results with the timetables actually used, are presented.

Keywords: Timetabling, scheduling, heuristics, tabu search.

* Ramón Álvarez-Valdés and José M. Tamarit. Department of Statistics and Operations Research. University of València. Dr. Moliner, 50. 46100 BURJASSOT (València).

** Enric Crespo. Department of Financial and Mathematical Economy.

* This work was subsidised in part by the Conselleria d'Educació i Ciència, Generalitat Valenciana, GV 2212/94

– Article rebut el juliol de 1996.

– Acceptat el febrer de 1997.

1. INTRODUCTION

The examination scheduling problem consists of assigning exams to periods and classrooms in such a way that all exams are scheduled within a given time interval and a number of different objectives are satisfied. Each university offers a large list of courses every term and allows students a great deal of flexibility in their course selection. The result is a large and complex scheduling problem. Although the main objective is that no student has two exams simultaneously, there are many other secondary objectives, such as students' convenience, classroom availabilities and the special requirements of some exams.

The problem has attracted the attention of researchers over the past 30 years. Though most papers are tailored to the characteristics of authors' universities, two trends can be pointed out: the increasing interest in considering a wide set of objectives and side constraints, and the use of new and powerful heuristics to cope with the growing size and complexity of the problems.

The survey by Carter(1986) covers the work carried out before 1986. Most of it used some modified colouring heuristics —Desroches *et al.* (1978), Mehta(1981), White *et al.* (1979)—. Two different approaches were the algorithm HORHEC of Laporte and Desroches(1984), an assignment procedure with limited backtracking, and the paper by Romero(1982), in which the computer assists in the negotiation process between administration, departments and students.

All the mentioned papers had avoiding or minimizing *first order conflicts* as their primary interest, that is, students having two exams simultaneously. However, some of them started to consider some secondary objectives, such as minimizing *back-to-back conflicts*, that is students having exams in consecutive periods, or distributing the exams evenly in the examination periods. In the papers appearing after 1986 this trend has become more pronounced, with multiphase procedures in which each objective is taken into account at each step in the process —Johnson(1990), Lofti *et al.* (1991), Thompson *et al.* (1993)— and papers in which, assuming that a solution with a minimum number of first order conflicts is known, the objective is minimizing back-to-back conflicts —Arani *et al.* (1988), Balakrishnan *et al.* (1992)—.

In recent years, examination scheduling problems have been solved by using some new heuristic procedures that have been shown to be very useful in other related problems. Thompson and Dowsland(1993) and Johnson(1990) have used Simulated Annealing. Hertz(1991) and Clark(1993) have developed Tabu Search algorithms. Clark addresses a problem quite similar to ours, though he does not consider many special characteristics we have included (forbidden periods for some exams, precedences, changes in the availability of classrooms,...). His algorithm is based on an

intensification procedure and diversification is reduced to a dynamic variation of the length of the tabu list.

Still, an emerging trend that may intensify is the evolution of solution procedures into packages that can be used directly by the final users, preferably on personal computers. A good example is EXAMINE, developed by Carter *et al.* (1994) from the algorithm of Laporte and Desroches(1984).

The purpose of this paper is to describe a new algorithm for examination scheduling and its application to a large university in Spain. In the same line of the above mentioned research on this problem, our algorithm combines several heuristics based on tabu search and aims to obtain not only a solution without simultaneous exams, but the best distribution of exams among the periods for all the students. These objectives are basically shared by all universities, but they are considered here according to the priorities and specifications established by our university in València, Spain.

The description of the problem, its constraints and the hierarchy of objectives are presented in Section 2. Section 3 describes our algorithmic approach and Section 4 the implementation and computational results. Finally, in Section 5 we draw some conclusions and outline future lines of research.

2. DATA, OBJECTIVES AND CONSTRAINTS

The University of València has 65,000 students divided into four main Areas: Social Sciences, Health Sciences, Humanities and Basic and Technical Sciences. Each Area has several Faculties or Schools that are responsible for the academic organization of classes and examinations. However, the students in a Faculty usually take some optional courses from other Faculties in their Area. Moreover, the classrooms are shared by Faculties. Therefore, examinations must be planned at Area level, making the task much more complex.

The necessary data for exam scheduling are:

- *Students*

Each student is registered on a set of compulsory and optional courses. The number of courses is variable, but most students have to sit 4 or 5 exams each term, many of them with a theoretical and a practical part.

From the actual registration file we can obtain the conflict matrix, that is, the number of students common to each pair of exams. This matrix has many non-zero elements and not only in the expected square submatrices corresponding to some course streams of typical students.

- *Periods*

There is a prescribed exam interval, usually three weeks each term. Due to the length of exams, with theoretical and practical parts, it has been established that each day has only one examination period. That gives us a total of 15 periods (18 if all Saturdays are used).

- *Courses*

Some courses do not have examinations. Semester courses usually have one exam and annual courses require one or more partial exams and a final examination. Therefore, from the list of courses we build the list of exams with their characteristics.

Some exams have a restricted set of allowed periods or are directly preassigned to a period. Sometimes, a pair of exams are related in such a way that one must precede the other in at least a given number of periods, for instance, oral and written language exams.

In some cases, the exam needs a special type of classroom, such as a laboratory or computer room. The size of the required classroom depends on the course enrollment and on a parameter, *the attendance factor*, which is the estimated percentage of registered students actually attending the exam. This parameter is inputted by the user, who may decide, for instance, a 100% attendance for a partial exam and a 70% for a final exam if the final is compulsory only for students below a certain mark in the partials.

- *Classrooms*

At each period we have an available set of classrooms with their type and exam capacities, usually a half or a third of the seats, depending on classroom structure.

In the data description above the constraints have implicitly appeared: The exam interval must be strictly respected, preassignments and forbidden periods for exams must also be satisfied, as well as the precedence relations.

The objectives are basically three and there is a clear hierarchy among them:

1. No student should have two exams simultaneously. These first order conflicts have to be avoided. In fact, this condition could be considered a constraint, but we cannot be sure of finding a feasible solution if we enforce it because of the multiple combinations of students' registrations and the reduced number of periods. Therefore, we put it as the main objective and we will try to get a solution without first order conflicts.
2. For each student, his/her exams should be as evenly separated as possible along the exam interval. This idea extends the usual objective of avoiding consecutive

exams, back-to-back conflicts, and conditions of the type «Students should not write x or more examinations within any y consecutive periods» —Carter et al.(1994)—. In an examination interval of 15 periods, students with 6 exams cannot expect an average distance between exams greater than 2 periods, but that should not be applied to students with 2 or 3 exams who could obtain more distant exam dates.

This idea of evenly scattering the exams forms part of the non-written tradition of our university, especially in those faculties in which the examination time-tabling has been done by hand after a negotiation process with teachers and students. Now, the complexity of the process makes it difficult to maintain, but an acceptable automatic substitute has to include this objective.

3. Classroom capacities should be respected. Clearly, this is a constraint, but not a hard one. The inclusion of some students above classroom capacity is not physically impossible, though it may not be desirable. If this condition is included as a constraint, we may lose some good solutions with reduced capacity overflow that could be considered acceptable by the parties involved.

As we mentioned above, a tradition of negotiation is being substituted with or assisted by computer systems. So far, in the Area of Basic and Technical Sciences, with four Faculties and 6,000 students, the scheduling process is done by hand in two stages. First, each Faculty builds its timetable in a negotiation between administration, departments and students. Then, the Faculties in the Area put together their proposals and look for a common classroom assignment, modifying examination dates if necessary. This procedure has serious drawbacks. In the first step, it cannot assure minimizing first order conflicts. In particular, it cannot adequately consider the courses attended by students from several Faculties. In the second step, it is very difficult to fit the timetables into the available classrooms, and changes in the dates deteriorate the quality of the solutions. Therefore, all parties involved are well aware of the need for changing this lengthy and unsatisfactory process by a more efficient procedure. However, from their experience, they will not accept solutions they do not consider at least as good as those manually obtained. Hence, timetables not only have to *be good* according to some agreed cost function, but they have to *look good* to the users, according to their knowledge of the situation and some characteristics that cannot be modelled but are clearly perceived by them just by looking at the solutions.

3. SOLUTION METHOD

The process of obtaining a good solution, according to the mentioned objectives, is divided into two phases. First, a solution without first order conflicts is built, whenever it exists; then it is improved with consideration to overall student convenience.

3.1. Building a feasible solution

An initial solution is built by randomly assigning exams to periods. The assignment for each exam is randomly chosen from the allowed periods for the exam that satisfy the precedence relations.

This process is repeated several times to obtain different initial solutions. They will evolve separately to produce several final solutions to be presented to the people involved in the decision.

These initial solutions have first order conflicts and do not satisfy classroom availability. For each of them we start a tabu search procedure, a general heuristic method for guiding the search to obtain good solutions in complex solution spaces. A general introduction to Tabu Search may be found in Glover *et al.* (1993). Though these techniques have been developed and extended in very sophisticated ways in recent years, the basic aspects of the procedure may be described as follows:

1. Construct an initial solution s in the space of solutions X

$s^* = s$ (s^* = best solution found so far)

$k = 1$ (k = number of iterations)

2. **While** $k < \text{maximum number of iterations}$ **do**

$k = k + 1$

Generate a set of solutions $V^* \subseteq N(s, k)$ (set of neighbours of s that are not tabu or for which the aspiration criterion applies)

Choose the best s' in V^*

$s = s'$

If $f(s') < f(s^*)$, then $s^* = s'$

end while

The space of solutions X in which we move is the set of timetables satisfying the allowed periods for exams and the precedence relations.

The objective function combines the number of first order conflicts with the excess of students over classroom capacities. For each solution s we define

$$f(s) = \sum_{i=1}^N (p \sum_{j,k \in E_i} x_{jk} + r_i)$$

where N is the number of periods, E_i the set of exams assigned to period i , r_i the room shortage at period i , x_{jk} the number of common students for exams j and k , and p the relative weight of conflicts with respect to overall room shortage.

A solution $s' \in X$ is a neighbour of solution $s \in X$ if it can be obtained from s by moving an examination to another period. To decide which neighbour s' we should move to from the current solution s , we do not explore the whole neighbourhood because that would be too costly. Instead, we select the *most conflictive period*, the period of biggest contribution to the objective function, and from among the exams in it we choose the *most conflictive examination*. We generate and evaluate all its possible moves and make the best one, as long as it is not tabu. If a move is tabu but it improves the best current solution, it is made in spite of its tabu status, applying the aspiration criterion. If all the moves are tabu and the aspiration criterion does not apply, we select the second most conflictive exam and repeat the procedure. The process is repeated as many times as necessary until a possible move is found and made.

For each move the tabu list keeps the examination that is moved and the period from which it comes. A move involving an exam and period in the tabu list is considered tabu. The aspiration criterion allows a move to be made, in spite of its tabu status, if the new solution is better than the best solution found so far.

This iterative procedure ends when we reach a conflict-free timetable. Nevertheless, sometimes the problem does not have a zero conflict solution, or even if it has, our algorithm is not guaranteed to find it. In these cases, the process stops when a limit of iterations is reached, then returning a timetable in which some students have more than one examination in a period.

3.2. Improving the solution

Once we get a feasible solution, or at least a solution with a minimum number of conflicts, we start to consider other objectives, such as a distribution of exams in the periods maximizing the distance between exams for the students. Also, the capacity of classrooms is now more important than it was in the previous phase.

The objective function of this phase reflects this new situation. For each solution s we define

$$f(s) = \sum_{i=1}^N \left(\sum_{j=1}^N (p_{|i-j|} \sum_{k \in E_i} \sum_{l \in E_j} x_{kl}) + w r_i \right)$$

where N is the number of periods, E_i the set of exams assigned to period i , r_i the excess of room capacity at period i , x_{kl} the number of common students of exams k and l , and $p_{|i-j|}$ the penalty associated with a pair of exams scheduled at a distance

of $|i - j|$ periods. The parameter w reflects the relative weight of room shortage with respect to conflicts.

The objective function includes conflicts of distance zero, that is, first order conflicts. Allowing these conflicts to appear makes the search more flexible. A high penalty p_0 will guide the process to solutions with a minimum number of these conflicts.

To improve the solution we have developed two procedures: *permutation of exam lists*, defined as the sets of exams assigned to the same period, and *interchange of exams*. Both procedures are complementary. The permutation of exam lists considers the list of exams in a period as a fixed set and moves the whole list from one period to another. Obviously, this procedure cannot produce first order conflicts and may improve the quality of the solution very fast. However, it is a very constrained process, because the limitations of exams (non-allowed periods, precedence relations,...) are limitations of their list. On the other hand, the interchange of exams takes the exams one at a time and considers moves or swaps. That allows us to study many more alternatives, though the improvement of the objective function is slower.

Interchange of exams

The procedure for interchanging exams is similar in spirit to that used to obtain a feasible solution but here we define a more flexible move. At each iteration we select a period t (*the most conflictive period*) and in it an exam e (*the most conflictive exam*). Then we consider its possible moves and evaluate them. If the exam e is going to a period t' and it does not produce first order conflicts, the change in the objective function depends on the distance to other exams and the room shortage. But if it produces some new first order conflicts, a high value of p_0 will produce a dramatic increase in the objective function and this move will never be made. In these cases, we study the possibility of moving some exams $e', e'', \dots$ from period t' to period t to avoid some of these conflicts. Hence, instead of a single move, we make an *interchange*. This move is more complex, but allows us to study new alternatives to the current solution, reaching regions of the space of solutions that would not have been visited otherwise.

The permutation of exam lists

In this procedure, we consider changing the whole set of exams assigned in a given period to another one. If in the objective function described above we do not consider the weight of room shortage, the objective function could be written as

$$f(s) = \sum_{i=1}^N \sum_{j=1}^N p_{ij} C_{l(i)l(j)}$$

where $C_{l(i)l(j)}$ is the number of students having exams in both periods i and j and p_{ij} is the associated penalty. This is the objective function of the Quadratic Assignment Problem (QAP). Hence, the problem of finding the best permutation of lists may be viewed as a QAP. By doing that we do not get an easier problem, but we may take advantage of the recent work on it. Probably the most successful heuristic approach at this moment is the Tabu Search, as can be seen in the papers of Taillard(1991) and Skorin-Kapov(1994).

We have developed an algorithm that closely follows the work of the above mentioned authors. Nevertheless, we maintain our objective function because the availability of classrooms may vary and changes in room shortage must be reflected in the quality of moves. At each iteration we explore all the possible moves, discard those which are unfeasible and evaluate the rest (note that the reduced number of periods allows us such a complete exploration). Then we make the best feasible permutation, according to considerations about the tabu list and aspiration criterion similar to those described previously.

Iterative scheme

The two procedures are linked in an iterative scheme in which the solution obtained in one of them is the starting point for the other, until the solution cannot be improved further. This scheme may be viewed as the alternance of *intensification* and *diversification* phases of Tabu Search method. In the procedure of interchange of exams, local changes try to get a better solution in an intensification phase. When this phase cannot improve the solution, the procedure of list permutations deeply changes the solution in a diversification phase that moves the search to completely different regions of the solution space. The two phases follow the basic scheme of Section 3.1, with the same objective function and the types of moves (and corresponding neighbourhoods) defined above. The procedure switches from one phase to the other after a given number of iterations.

3.3. Assigning classrooms to examinations

The final phase of the procedure consists in assigning classrooms to examinations. This phase is by no means trivial in our problem because of the interaction of three factors: first, the number of available seats is quite reduced compared with the number of students; second, the solution obtained in the preceding phases uses the global number of seats and compares it with the total number of students having an exam at a given period, but a mechanic translation of the exam list to the classroom will produce a large number of exams having to share a classroom and that should be avoided as far as possible. Third, in our system each day is a period because the exams are very long and it is not desirable for the students to schedule two exams on the same day. However, each day has two sessions, morning and afternoon. Therefore,

the number of seats is counted twice, unless otherwise stated. The problem now is to assign exams to rooms and sessions in such a way that every exam is assigned to only one session, using the minimum number of rooms and without sharing the rooms with any other exam. Moreover, if in the timetable there are still some first order conflicts, the corresponding exams should be assigned to different sessions, to partially solve the problem of some students having two exams simultaneously. The solution, though undesirable, would be physically possible.

In this procedure the user must input two parameters: *the minimum number factor*, the minimum number of students a course must have to require a classroom (the exams of courses with very few students may take place in some other places, like some departmental room) and *the residual factor*, which is the percentage of remaining students for which a new classroom is not needed when we assign more than one room to a large exam.

For this phase we have developed an assignment algorithm that considers the list of examinations ordered by size and first assigns each of them to a session, trying to solve, if necessary, the first order conflicts, and then to a set of classrooms. For each exam, knowing its size and the remaining available classrooms, the algorithm selects rooms minimizing their number and adjusting their overall capacity to exam size. Likewise it tries to minimize the number of invigilators and the number of classrooms shared by more than one examination.

4. IMPLEMENTATION AND COMPUTATIONAL RESULTS

The algorithms of the preceding section have been imbedded in a package to allow the user to input data and parameters, obtain several solutions, print them and modify the initial conditions and values of parameters to correct them or study new alternatives. The package has three basic and one auxiliary modules:

1. *Input module*, in which the user, through a set of menus and selection lists, edits the data about courses and rooms, sets the values of the parameters and gets information about student registration.
2. *Solution module*, including the algorithms described above, that produces several solutions from which the users may choose the most convenient.
3. *Classroom assignment module*, that produces for each solution the final list in which each examination appears, assigned to a day, a session and a set of classrooms.

This module may also produce a classroom assignment for timetables provided by the user.

4. *Evaluation module*, an auxiliary module that evaluates a given solution according to the current objective function. This module is internally used by the algorithms of the Solution module, but may be used separately to study and compare other solutions obtained by other methods. In our case, coming from a system in which the timetable is done by hand after a negotiation process, it is especially important to compare the solutions in order to shed light on the deficiencies of the system and convince the users of the advantages of the new solutions.

The programs have been written in C language and the package has been designed to run on personal computers, the usual equipment of Faculty administrators.

To assess the performance of our algorithms, we have made two sets of tests. First, a small example, involving only the examinations of the Faculty of Mathematics, with 1,200 students, in which the difficulty comes from the complexity of students' registrations, making it very hard to find a timetable without first order conflicts.

We include four instances of the problem, corresponding to the year 1992-93. In that year all courses were annual, with partial examinations in February and June and final examinations in July. Therefore, they needed a timetable for February and another for June-July. The last one admitted two possibilities: making two separate timetables for June and July or building one common timetable.

The second test involves the whole area of Basic and Technical Sciences, with 6,000 students where together with its large size, the difficulty lies in the reduced number of available rooms and the existence of courses attended by students from different Faculties. In the new academic system we have some annual courses and a majority of semester courses. We have solved the problem for the end of the academic year 1994/95, with partial and final examinations for annual courses and final exams for spring semester courses, and also for final exams of the fall term in February 1996. In June 1995 the examination interval was of four weeks for partial examinations and five weeks for final ones. In February 1996 the exams interval was of three weeks.

The values of parameters were set after a series of trials. In the first phase of building a feasible solution we made $p = 100$, giving much more weight to first order conflicts than to room shortage. In the second phase of improving the solutions the parameters were $p_0 = 3000$, $p_1 = 100$, $p_2 = 20$, $p_3 = 5$, $p_4 = 3$, $p_5 = 1$, $p_k = 0$, $k > 5$ and the cost of room shortage $w = 40$. These values reflect the main objective of getting solutions without first order conflicts and the need for adjusting the examination timetables to the available classrooms.

The length of tabu lists were fixed at $\sqrt{n}$, where n is the size of each part of the problem. In the procedures in which the move consists of changing an exam from one period to other, the size is the number of exams (ne) times the number

of periods (np). In the procedure of list permutations, the size is $np(np - 1)$. The number of iterations was set to 1000 in the moving-exams procedure and 500 in the list-permutation procedure. The tabu lists have a circular structure and at each iteration the oldest element on the list is replaced by the new one associated with the move.

Table 1 shows the data of the six tests (number of periods, number of exams, the total number of examinations to be done by the students and the density of conflicts matrix) and the comparison between the actual solutions obtained by the current process and those proposed by our algorithms, with respect to first order conflicts, room shortage, conflicts of distance 1 (*back-to-back conflicts*) and conflicts of distance 2.

Table 1. *Comparison of actual and proposed solutions*

Date	Per.	Ex.	Ex.-	Mat.		F.or.	Room	C.d.1	C.d.2
		N.	Stud.	dens.		conf.	short.		
feb93	17	40	4853	0.50	actual	63	208	235	202
					alg	0	0	147	426
jun93	17	40	4853	0.50	actual	21	0	293	460
					alg	0	0	150	418
jul93	17	40	4853	0.50	actual	27	30	336	480
					alg	0	0	136	474
jj93	34	80	9706	0.50	actual	48	0*	629	946
					alg4	0	0	335	991
					joint10	1	0	439	1056
jun95	30	107	22011	0.15	actual	203	179	3076	1828
					alg	0	0	474	2753
feb96	18	132	26994	0.15	actual	267	0	3833	3860
					alg	3	0	1606	4563

The block *jj93* contains several solutions to the joint problem of partial and final examinations in 1993. The first row, *actual*, is the actual solution, obtained by putting together the solutions of *jun93* and *jul93*, with a minor modification made by the users: a final examination was moved backwards into the periods corresponding to partial exams. That improved the solution, mainly in the use of classrooms. In the second row, *alg4*, we put together our solutions for *jun93* and *jul93*. In this solution, the

minimum distance from partial and final examinations of a course is 4 periods, which is considered unacceptable for teachers and students. Therefore, the third row, *joint10* shows the solution imposing a minimum distance of 10 periods.

For all instances our algorithm obtained much better solutions than those manually built. The new academic system, which is much more flexible and contains more optional courses, plus the growing number of students, will make the problem harder and the difference between manual and automatic timetables larger, as already happens in last instances *jun95* and *feb96*.

The solutions were obtained on a personal computer with a Pentium processor at 100 Mhz. For the large problem *jun95*, it took 5 seconds to obtain the first feasible solution and 30 minutes to get the final solution. **Table 2** justifies the computational effort of the improving phase by comparing the initial feasible solutions with the final solutions obtained by asking the algorithm to obtain six alternative solutions. The last two columns show the number of times the improving phase calls on the exam interchange and list permutation procedures. Though some solutions may seem better than others, according to the criteria of the table (for instance, solution 1 looks better than solution 2), there may be some other reasons for the users to adopt one solution and discard the rest.

Table 2. *Comparison of initial and final solutions*

First order conflicts	Conflicts of distance 1	Conflicts of distance 2	Exam interchange	List permutation
0	4753	4809	-	-
0	474	2753	3	3
0	5432	4405	-	-
0	669	2828	7	7
0	5551	4578	-	-
0	726	2503	4	4
0	3684	5327	-	-
1	468	2911	3	3
0	4535	3164	-	-
0	614	2717	4	3
0	5758	3497	-	-
0	616	2355	2	2

In the trials, some other combinations of parameters, tabu list lengths and number of iterations produced better results for some of the test instances. Nevertheless, the values mentioned above produced consistently good results for the whole set of problems.

5. CONCLUSIONS AND FUTURE RESEARCH

We have developed a package to solve the examination timetabling problem for our university. Many of its characteristics are common to other universities and the procedures could be adapted and used in other places.

The use of Tabu Search algorithms has allowed us to obtain high quality solutions in very short computing times, in spite of the size of the problem and the complexity of data and objectives.

The solutions obtained in the test were presented to Faculty administrators in the Basic and Technical Sciences Area. The quality of our solutions and the serious problems faced in the process of obtaining the actual solution have convinced them of the convenience of adopting our package as a decision support system to assist the construction of future timetables.

The next step in our work will be the use of our procedure in the other Areas of the University. We are also planning to imbed these algorithms in a general academic management system we are currently developing.

6. REFERENCES

- [1] **Arani T., Karwan M. and Lofti V.** (1988). «A Lagrangian relaxation approach to solve the second phase of the exam scheduling problem». *Eur. J. Opt Res.*, **34**, 372–383.
- [2] **Balakrishnan N., Lucena A. and Wong R.T.** (1992). «Scheduling examinations to reduce second-order conflicts». *Computers Ops Res.*, **19**, 353–361.
- [3] **Carter M.W.** (1986). «A survey of practical applications of examination time-tabling problems». *Mgmt Sci.*, **34**, 193–202.
- [4] **Carter M.W., Laporte G. and Chinneck J.W.** (1994). «A General Examination Scheduling System». *Interfaces*, **24**, 109–120.
- [5] **Clark D.** (1993). «Exam scheduling by Tabu Search». *Australian Soc. Ops Res. Bulletin*, **12**, 5–9.

- [6] **Desroches S., Laporte G. and Rousseau J.M.** (1978). «HOREX: A computer program for the construction of examination schedules». *INFOR*, **16**, 294–298.
- [7] **Glover F., Taillard E. and de Werra D.** (1993). «A user's guide to Tabu Search». *Annals of Ops Res.*, **41**, 3–28.
- [8] **Hertz A.** (1991). «Tabu search for large scale timetabling problems». *Eur. J. Opl Res.*, **54**, 39–47.
- [9] **Johnson D.** (1990). «Timetabling university examinations». *J. Opl Res. Soc.*, **41**, 39–47.
- [10] **Laporte G. and Desroches S.** (1984). «Examination timetabling by computer». *Computers Ops Res.*, **11**, 351–360.
- [11] **Lofti V. and Cervený R.** (1991). «A Final-exam-scheduling Package». *J. Opl Res. Soc.*, **42**, 205–216.
- [12] **Mehra N.K.** (1981). «The application of graph coloring methods to an examination scheduling problem». *Interfaces*, **11**, 57–64.
- [13] **Romero B.P.** (1982). «Examination scheduling in a large engineering school: A computer assisted participative approach». *Interfaces*, **12**, 17–24.
- [14] **Skorin-Kapov J.** (1994). «Extensions of a Tabu Search adaptation to the Quadratic Assignment Problem». *Computers Ops Res.*, **21**, 855–865.
- [15] **Taillard E.** (1991). «Robust taboo search for the quadratic assignment problem». *Parallel Computing*, **17**, 443–445.
- [16] **Thompson J. and Dowsland K.** (1993). «Multi-objective university examination scheduling». *Working paper EBMS/1993/12*.
- [17] **White G.M. and Chan P.W.** (1979). «Towards the construction of optimal examination schedules». *INFOR*, **17**, 219–229.

Estadística Oficial:

Recerca estadística promoguda per Eurostat

CONFIDENTIALITY IN THE EUROPEAN STATISTICAL SYSTEM

PHOTIS NANOPOULOS*

Eurostat

Confidentiality and data protection are the counterparts of the obligation for response and the corner stone for confidence building between the national statistical services and the respondents. On this common principle the various national statistical systems among the 15 (and other) European states have been established.

There is a great concern at national and international level on statistical confidentiality. The respect of privacy has led national authorities to develop ad-hoc legislation which, inter alia, prevents national statistical institutes transmitting confidential data to Eurostat. The effect on the quality and completeness of European statistics has been disastrous. Because of this important problem Eurostat undertook a series of actions some of which are as followed:

In the legal framework, a regulation was proposed to the Council of the Union which sets out the conditions for data transmission from national administrations to Eurostat. The regulation was adopted in 1990. On the 17th February 1997 the Council approved a new regulation which completes and enlarges that previously on statistical confidentiality.

Eurostat launched in parallel international seminars on confidentiality which brought together statisticians, academics and other officials and set the milestones for international co-operation in that area. The first seminar was organised in 1992 in Dublin in co-operation with the ISI (International Statistical Institute). The next one was held in 1994 in Luxembourg and this third in Slovenia in October 1996.

*Photis Nanopoulos. Direction A: Systèmes d'information statistique; recherche et analyse des données; coopération technique avec les pays Phare et Tacis. Eurostat. L-2920 Luxembourg.

—Article rebut l'octubre de 1996.

—Acceptat el febrer de 1997.

Methodological work also was encouraged. Through the fourth European framework program on research and development, financial support was given to a multinational team for the development of a software to control statistical disclosure of both micro and tabular data (Waal, Willenborg, 1995).

Eurostat made an inventory of existing methods and started in 93 to explore new techniques in the subject of the microaggregation.

The efforts in this direction will continue in a more intensive way. It is planned to launch other projects for the statistical disclosure control, especially in the field of microaggregation, to explore cognitive aspects of confidentiality and generally to promote research aimed at permitting access to the data without jeopardising confidentiality.

PROTECTING MICRO-DATA BY MICRO-AGGREGATION: THE EXPERIENCE IN EUROSTAT

DANIEL DEFAYS*

Eurostat

A natural strategy to protect the confidentiality of individual data is to aggregate them at the lowest possible level. Some studies realised in Eurostat on this topic will be presented: properties of classifications in clusters of fixed sizes, micro-aggregation as a generic method to protect the confidentiality of individual data, application to the Community Innovation Survey. The work performed in Eurostat will be put in line with other projects conducted at European level on the topic of statistical confidentiality.

Keywords: Confidentiality, clustering, micro-aggregation.

*Daniel Defays. Recherche et développement, méthodes et analyses des données. Eurostat. L-2920 Luxembourg.

—Article rebut l'octubre de 1996.

—Acceptat el febrer de 1997.

1. STATISTICAL CONFIDENTIALITY WITHIN EUROSTAT

There is a great concern at national and international level on statistical confidentiality. How to provide statistical information of high quality without disclosing confidential data? The respect of privacy has led national authorities to develop ad-hoc legislation which among others prevented national statistical institutes to transmit confidential data to Eurostat. The effect on the quality and completeness of European Statistics was disastrous. In most cases, it was impossible to provide European aggregates because some national data were missing. Facing this serious problem, Eurostat initiated beginning of the nineties, different actions.

- In the legal sphere, it proposed to the Council of the Union a regulation which sets the condition through the Commission for data transmission from national administrations to Eurostat. The regulation was adopted in 1990.
- Eurostat launched in parallel international seminars on confidentiality which brought together statisticians, academics and other officials and set the milestones for international co-operation in that area. The first seminar was organised in 1992 in Dublin in co-operation with the ISI. The next one was held in 1994 in Luxembourg and in 1996 it will take place in Bled (Slovenia).
- Eurostat organisation was updated in order to take into account the new constraint coming from the transmission of confidential data by the Member States. Administrative principle and procedures were established and technical measures taken in order to control the respect of those principles and procedures.
- Methodological work was also encouraged. Through the fourth European framework programme on research and development, financial support was given to a multinational team for the development of a software to control statistical disclosure of both microdata and tubular data (Waal, Willenborg, 1995).
- Eurostat made an inventory of existing methods and started in 93 to explore new techniques to protect micro-data by using micro-aggregates. The development of these techniques is the subject of this paper.

2. THE LEGISLATIVE FRAMEWORK

As already written, statistical confidentiality at the European Union level is governed by a Council Regulation (Euratom/EEC, 1588/90). This regulation was a response to a triple need:

- need of safeguards against possibilities of abusing data;
- need to allow the Member States of the Union to adopt a less restrictive data transmission policy regarding Eurostat;
- need for more and better Community statistics.

It makes it possible to remove national legal obstacles to the flow of statistical data from the national authorities to Eurostat. In that regulation, there is no European definition of confidentiality; confidential statistical data are defined as data declared confidential by the Member States in line with national legislation or practices. In most cases, precise rules exist for aggregated data: any cell of a table either containing data relative to less than three units, or dominated by one (or in some countries two) unit should not be disclosed.

For individual data, national policies are more different. In some countries, any transmission of micro-data is forbidden, in other ones only the research community has access under specific condition; part of micro-data can be publicly accessible in some countries.. Under this regulation, the use of confidential data transmitted to Eurostat is very limited. Dissemination of micro-data to scientists for statistical purposes is not allowed.

The investigation of Eurostat in the field of micro-aggregation was a response to the conflicting needs to give a maximum of information to the scientific community without disclosing confidential data.

3. BRIEF HISTORY OF MICRO-AGGREGATION AT EUROSTAT

The first idea was to start from the definition of confidentiality for tabular data: only cells with at least a minimum number of units and no dominance are not confidential. Applied to micro-data, this meant aggregation of units three by three, the units to be aggregated being as similar as possible in order to avoid dominance by one of them. Therefore we made the proposal to replace individual data by averages of small aggregates, which can play the role of fictive individuals on which data analysis could be performed (Defays, Nanopoulos, 1993). The problem was then to partition the whole population in cluster of fixed sizes. The averages of the optimal clusters could be transmitted. It was shown that if the primary data are points in R^p , the optimal partitioning (minimisation of within-group variances is characterised by the following property: every pair of clusters is separated by a hyperplane perpendicular to the line joining their barycenters.

To find the optimal partitioning seems to be a difficult problem. He have proposed to improve an algorithm proposed by Uri Hanani to cluster a set of points under an

equal groups constraint (Hanani, 1979). The problem was presented as a particular case of multicriteria dynamic clustering. Unfortunately there is no guarantee that that method will always reach the optimum. Another problem is the quality of the micro-aggregated data. It is easy to see that if the variables are not all correlated, the groups will be heterogeneous and the method will transform drastically some of the variables.

The difficulties to get the exact solution and the fact that in some cases the data can be strongly perturbed led us to investigate alternative methods, but in the same spirit.

4. A NEW MICRO-AGGREGATION TECHNIQUE

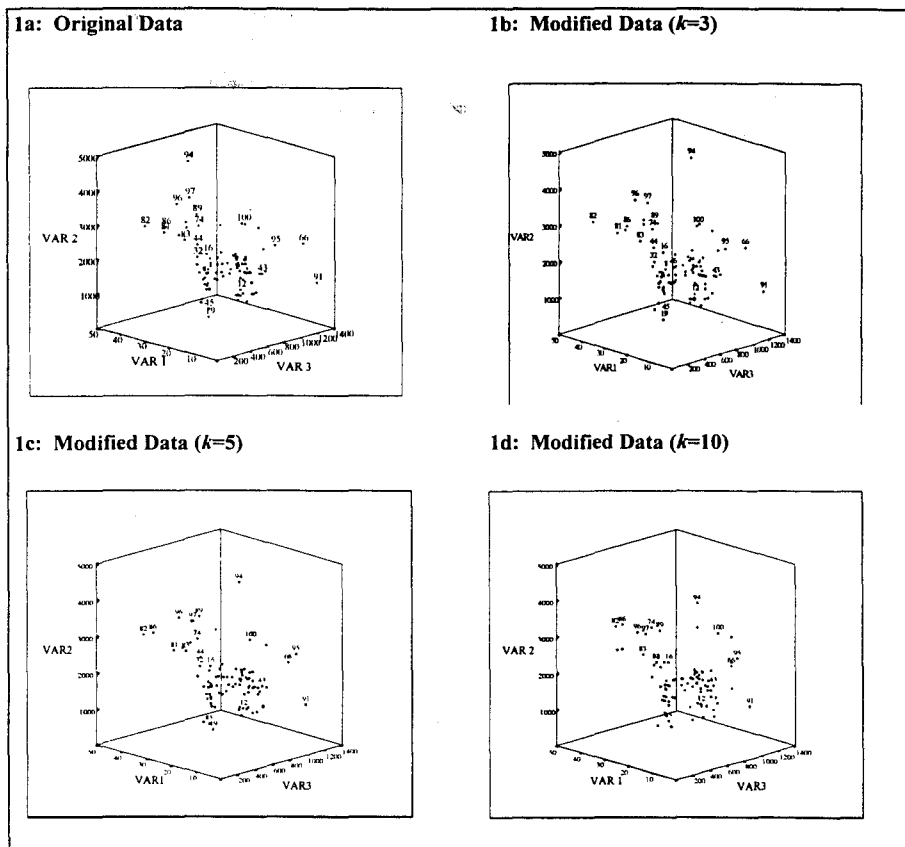
In order to avoid the loss of information caused by aggregation in clusters of fixed size of the original units, it has been proposed that the different unidimensional variables be aggregated separately, by ranking the values assumed by these variables and by an aggregation in fixed size groups of contiguous values.

The basic idea comes an application developed in the U.S. Internal Revenue Service (Strudler).

To illustrate the point let's take data for a set of 100 fictitious units covering three variables and aggregate into groups of 3, 5, and 10 to show the likely structural changes under different group sizes.

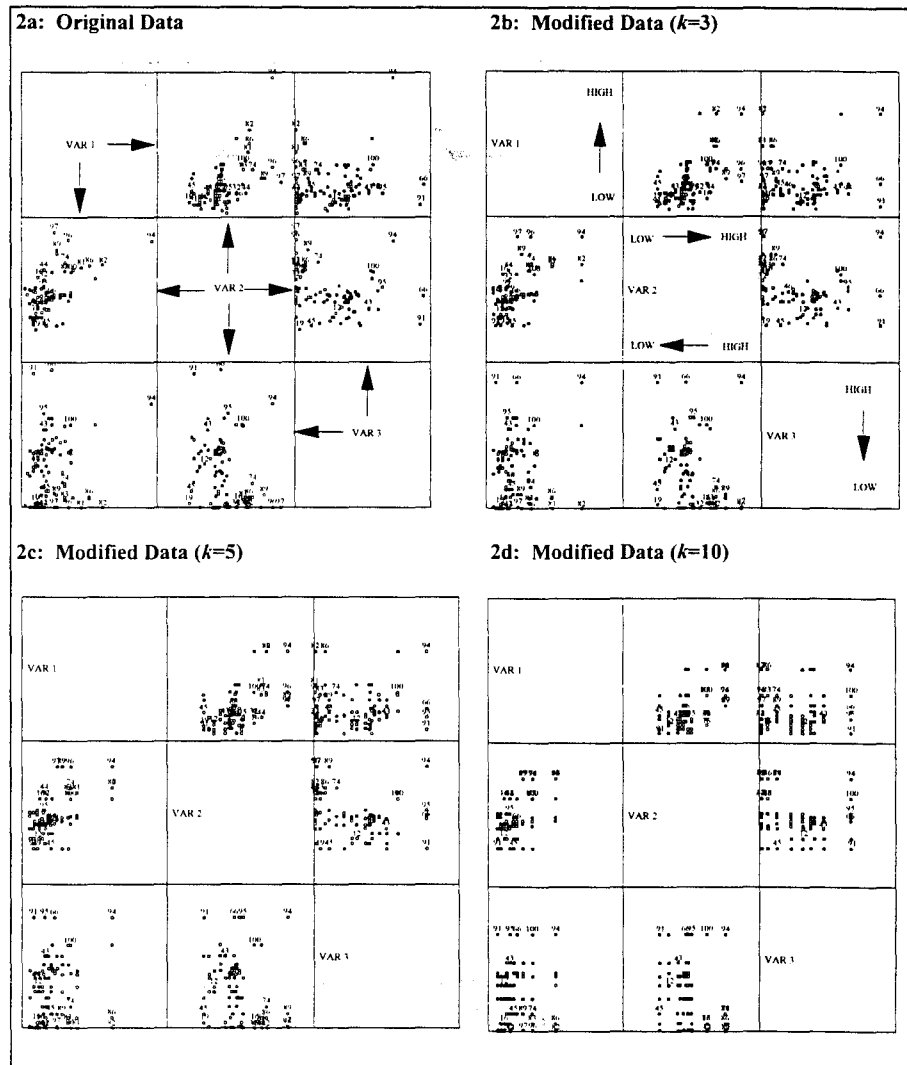
In a first step, the units are sorted in ascending (or descending) order of variable 1 and grouped k by k (where k in our case was 3, 5 and 10). The original variable 1 value for each unit is then replaced by the average for variable 1 of the corresponding group. In next step the units are again sorted, but by variable 2 this time. Groups of k are formed and the original variable 2 values are replaced by the averages of the corresponding groups. This procedure—sorting, grouping, replacement with average values—is repeated for the third variable, and a new file is created consisting of 100 surrogate observations. Figures 1a to 1d show the degree of perturbation to be expected under different group sizes.

Set in this three dimensional space it is hard to tell apart the different figures, and as such this is the essence of the method. The method acts more on outlying observations while leaving the majority of the data structure intact; this property is both interesting and useful from a statistical and confidentiality viewpoint. The masking of extreme values is a prerequisite of any method purporting to safeguard confidentiality, yet a method which destroys data structure also destroys the statistical properties of the data. The method proposed both decreases risk of disclosure and maintains relationships.



Figures 1a-1d: Data transformation under different size classes

A better picture of the data transformations can be seen by using two dimensional plots as shown below with the arrows in figure 2a indicating the corresponding axis. It is easier to focus on the three plots in the top right hand quadrant of each figure. The three plots in the bottom left hand quadrant are exact mirror images of the former plots and care needs to be taken in their interpretation since the axis are also mirror transformations. The axis increase vertically upwards and horizontally to the right and decrease in value vertically downwards and horizontally to the left, see figure 2b.



Figures 2a-2d: Data transformation matrices

The figures above show the «grid» structure imposed by the method which becomes more pronounced as k is increased. This grid structure provides a guarantee

of confidentiality by creating observations with identical values on a single given dimension.

Tables 1 and 2 below present a range of statistics for the three variables in our example. The statistics summarise what can be seen from the figures above especially the reduction in the variance as the number of k is increased, however the method is mean invariant and the degree to which the summary statistics are altered is minimal. The figures for correlations are also stable and near the original.

Table 1. *Summary Statistics*

Original Statistics				
Variable	Mean	Std Dev	Minimum	Maximum
VAR 1	10.24	6.83	1.00	48.00
VAR 2	2338.70	671.64	1127.00	4500.00
VAR 3	375.48	306.49	.00	1338.00
Modified Statistics ($k = 3$)				
Variable	Mean	Std Dev	Minimum	Maximum
VAR 1	10.24	6.62	2.25	35.00
VAR 2	2338.70	669.50	1223.50	4307.00
VAR 3	375.48	305.28	.00	1216.67
Modified Statistics ($k = 5$)				
Variable	Mean	Std Dev	Minimum	Maximum
VAR 1	10.24	6.52	2.40	31.00
VAR 2	2338.70	667.36	1243.20	4115.40
VAR 3	375.48	302.78	.00	1085.20
Modified Statistics ($k = 10$)				
Variable	Mean	Std Dev	Minimum	Maximum
VAR 1	10.24	6.21	3.10	25.00
VAR 2	2338.70	654.27	1332.5	3753.30
VAR 3	375.48	298.61	.10	948.50

Table 2. *Correlations Matrices*

Original Correlations			
	VAR 1	VAR 2	VAR 3
VAR 1	1.0000	.5770	.1641
VAR 2	.5770	1.0000	-.2344
VAR 3	.1641	-.2344	1.0000
Modified Correlations ($k = 3$)			
	VAR 1	VAR 2	VAR 3
VAR 1	1.0000	.5516	.1544
VAR 2	.5516	1.0000	-.2183
VAR 3	.1544	-.2183	1.0000
Modified Correlations ($k = 5$)			
	VAR 1	VAR 2	VAR 3
VAR 1	1.0000	.5521	.1213
VAR 2	.5521	1.0000	-.2274
VAR 3	.1213	-.2274	1.0000
Modified Correlations ($k = 10$)			
	VAR 1	VAR 2	VAR 3
VAR 1	1.0000	.5516	.1226
VAR 2	.5516	1.0000	-.2374
VAR 3	.1226	-.2374	1.0000

5. VARIANTS OF THE METHOD

Some generalisations of this method were then proposed to tackle other types of variables and to generalise the replacement of original values by averages.

a) Segmentation of the set of variables

The micro-aggregation-method presented above-has been referred to in the reference documents as «micro-aggregation method by individual ranking». This term underlines the necessarily separate treatment of the different individual variables which results in separate classifications and aggregations of units as illustrated above.

But what is regarded as an individual variable in this context? A set of p variables can be treated as a single multivariate variable, resulting in a single grouping, or as separate p variables, resulting in p groupings each.

More generally, the initial vector of variables can be segmented into a number of variables, multivariate or univariate, which we called segments. Each segment is treated separately.

b) Characterisation of the groups

For each segment formed, units are regrouped and within each group, the values of the units are replaced, when the variable is quantitative, by an average. This again can be generalised. First, if the variables are not numeric, one can use other central values than the average, for instance, the median or the mode. Whatever is the choice of the central value to be used, the method will cause a diminution of the original dispersion of the variable. The variance of the microaggregated values will always be lower than the original variance for instance. To counter balance this effect, one has proposed to replace in each group the original values by new ones, taken from a distribution which as close as possible to the original one (in the group). With that logic, the values of a cluster are not identical anymore, but the original distribution is replaced by another one with similar characteristics.

c) Definition of the groups

In the original methods, groups all have the same size. One could imagine to use different size groups according to the type of variable, its sensitivity. Even for a given variable, one could below and above a given threshold use different size constraints. For instance, for company statistics, small enterprises could be grouped three by three whereas larger enterprises could be aggregated in larger groups.

d) Measure of the homogeneity of a group

When the segments are defined by one dimensional variables, the notion of homogeneity of the groups is easy to define. In the multivariate case, the concept of underlying similarity in the definition of groups leaves room for interesting variants; homogeneity can be measured in different ways: within group-variances, entropy or measure based on any type of distance.

6. APPLICATION TO THE COMMUNITY INNOVATION SURVEY

The generalised method was applied to the processing of the Community Innovation Survey (CIS). This survey combined both quantitative and qualitative data, key data on the enterprise which might permit indirect identification and more neutral subjective assessments, simple questions or questions with a more complex structure. The objective was to put at the disposal of research teams working on behalf of the Commission on Innovation topics a maximum of information from the CIS. Given the restriction put on Eurostat dissemination policy by the above mentioned regulation, we were not allowed to disclose to the scientists the original data. We thus decided to use micro-aggregates. The 50.000 company data transmitted to Eurostat were micro-aggregated by country and sector of activity. The quality of the protection of the individual data brought by the method was then checked by the countries and the perturbed data were sent to a limited number of contractors which analysed them. Results of these analyses were reported upon during an international conference on Innovation and its measurement held in Luxembourg in May 1996. Eurostat services have started to check on the original data the robustness of the results established with the micro-aggregated data. The conclusions reached so far are very encouraging. The distortions introduced by the method do not seem to affect the structure of the variables and their relations.

7. CONCLUSIONS

The conflicts between the needs for access to statistical information and demands for confidentiality will not disappear in the coming years. More powerful methods than the existing ones will have to be developed to protect the privacy and at the same time to disclose as much information as possible. Micro-aggregation is an attempt in that direction, unfortunately an empirical one. We have established the need for an such a transformation of the data and its feasibility in a specific case. More evidence of its robustness will have to be given. Properties and conditions of applications of its variants will have to be explored. This challenge is part of the mission of the official statisticians which is to provide the user with a high-quality statistical information service. Eurostat is determined to contribute to the advancement of ideas and techniques which will tackle these issues.

ACKNOWLEDGEMENT. The work presented in this paper is a collective effort at Eurostat level in which the author with many other colleagues was involved.

REFERENCES

- [1] **Anwar, M. N.** (1993). «Micro-aggregation - The small Aggregates Methods». *Internal report*. Luxembourg, Eurostat.
- [2] **Defays, D. & Anwar, M. N.** (1995). «Micro-aggregation: A Generic Method». *Proceedings of the 94 International Seminar on Statistical Confidentiality*. Luxembourg. Office for Official Publications of the European Communities, 1995.
- [3] **Defays, D. & Nanopoulos, Ph.** (1993). «The Small Aggregates Method». *Proceedings of the 92 Symposium on «Design and Analysis of Longitudinal Surveys»*.. Ottawa, Statistics Canada.
- [4] **Hanani, U.** (1979). «Multicriteria dynamic clustering». *Research report*, **358**, IRIA, Rocquencourt.
- [5] **Strudler, M., Lock, H. & Scheuren, F.** «Protection of Taxpayer Confidentiality with respect to the Tax Model». *Washington, U.S. Internal Revenue Service*.
- [6] **Waal, A.G. & Willenborg, L.C.R.J.** (1995). «Development of ARGUS: Past, Present, Future». *Proceedings of the 94 International Seminar on Statistical Confidentiality*. Luxembourg, Office for Official Publications of the European Communities.

STATISTICAL DATABASES: THE REFERENCE ENVIRONMENT AND THREE LAYERS PROPOSED BY EUROSTAT

ROGER DUBOIS*

Eurostat

The functions of the Eurostat information system are divided into four sectors which correspond to the various stages in the processing of data from their collection to their diffusion:

- *Production: collection, validation and storage of the data and meta-data;*
- *Storage of the reference data (acceptance of the information);*
- *Use of the reference data (visibility/security and find/deliver);*
- *Diffusion.*

*The system of acquisition and validation represents for Eurostat a major workload given the diversity of the partners and of the different data processing situations encountered. The upshot of the diversity and change in both the production and dissemination environments is the need to de-couple them. If there are n production systems and m dissemination systems, then there will be a need for $n*m$ interfaces between them. This will lead an explosion of work to long term to have the data available to end-users. If a reference layer is inserted to de-couple the production and dissemination environments, the interface problem reduces to $n+m$ and therefore becomes much manageable. Once the interfaces to and from the reference layer are defined, new tools, products or domains in either environment are easily added. By de-coupling production from dissemination, a reference layer enables producers the freedom that they require to choose the most appropriate tools and techniques to do their job. Producers also gain a common interface into the reference area, rather than the many different interfaces that they have to cope with at present. By providing well documented data in a standard form, the reference layer provides disseminators with better raw material for their products. This enables them to provide better and more diverse products in a more timely and customer-driven, as desired in Eurostat's mission statement.*

Keywords: Architecture, reference environment, statistical information system.

*Roger Dubois. Administrator Unit A3. Information Data-Shop. Eurostat. L-2920 Luxembourg.

–Article rebut l'octubre de 1996.

–Acceptat el maig de 1997.

1. FUNCTIONAL OBJECTIVES OF THE INFORMATION SYSTEM

The functions of the Eurostat information system are divided into four sectors which correspond to the various stages in the processing of data from their collection to their diffusion:

- production: collection, validation and storage of the data and meta-data;
- storage of the reference data (acceptance of the information);
- use of the reference data (visibility/security and find/deliver);
- diffusion.

Co-operation between these four specialised sectors involves generalised functions of cataloguing and control (manipulation of data/resources management).

In terms of their nature, the functions can be classified into six categories:

- acquisition and validation of the data into six categories;
- management of the production data;
- manipulation of the data and meta-data;
- external presentation;
- associated services.

The system of acquisition and validation represents for Eurostat a major workload given the diversity of the partners and of the different data processing situations encountered. The objective for the development of this system is to minimise the workload required despite the increase in demand for statistical data. The use of EDI techniques (exchange of computerised data) is the direction to be followed in an attempt to eliminate the varieties of magnetic media, of format and of interpretation in the collection of data.

The management of the production data has to allow the migration of the various types of objects which combine to form Eurostat's information systems. The Nomenclatures are called to play an essential role in the creation, the homogenisation and then the documentation of the statistical data; they constitute de facto the heart of the information system.

The manipulation of the data must, in the first place, support the process of data transfer between the various environments. A number of services appears which allow the establishment of groups of users and of their associated privileges; certain meta-data become control rules in the information system.

The diffusion must consist of quality products which are well-known and are based on a platform which is easy to acquire and use (databases, statistical documents and publications, extractions from databases, magnetic media, CD-ROM, etc.).

The last two systems are activities of a general service nature which are necessary for any data-processing architecture, while emphasising particularly two very important activities for Eurostat: the STRINGS project for electronic publishing and the security of EDP applications.

2. TYPOLOGY OF THE DATA

Statistical data can be divided into two interdependent subjects:

- Statistical observations which correspond to all quantitative (and qualitative) measures associated with economic phenomena. This group is primarily composed of numerical values associated with information on their temporal evolution.

The principal types of projects which make it possible to organise and to structure this set of data are:

- vectors and matrices of numbers;
 - time series;
 - tables with one or more dimensions.
- All the statistical meta-data relating to the characterisation of economic phenomena and of statistical observations.

An item of statistical meta-data is characteristically associated with one or a collection of statistical observations so as to specify the significance of the related data.

The principal types of objects which characterise meta-data are:

- The Nomenclatures and the dimensions, which allow the organisation, the classification and the association of economic phenomena and statistical observations.
- The properties and documentation, which allow the characterisation and the documentation of the phenomena and the observations.

Models: which define the processes (of derivation or of validation) on the statistical observations.

It is very difficult to give a single definition of the meta-data concept and its representation for the following reasons:

- The meta-data serves as a link between a concrete entity, the statistical observation, and an abstract entity the economic phenomenon associated with the observation.

Note: the temporal evolution of these meta-data is the complexity factor.

- The meta-data concept may be used either «to classify» or «to seek» a subsets of observations. The hypotheses which underlie these two activities can be different and give rise to specific meta-data.
- In addition to the «link» aspect, meta-data characterise the «properties» of a subset of information.

For example, the meta-data associated with a time series are:

- ⇒ the economic phenomenon whose evolution is defined by the series;
- ⇒ the frequency, the source, the type of series (flow, stock, etc.), the unit, the order of magnitude, etc.;
- ⇒ one or more comments or descriptive texts on one or other aspect of the series.

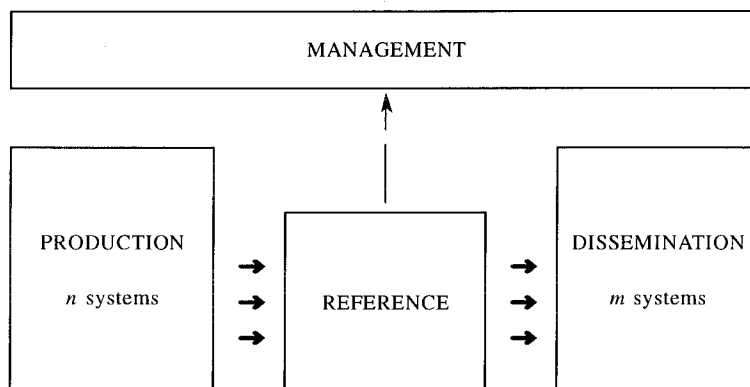
3. ARCHITECTURAL SOLUTION: A REFERENCE LAYER

3.1. Architecture

The upshot of the diversity and change in both the production and dissemination environments is the need to de-couple them. If there are n production systems and m dissemination systems, then there will be a need for $n*m$ interfaces between them. This will lead to an explosion of work to long time to become available to end-users. This is not in line with Eurostat's stated objective of timely provision of data.

If a reference layer is inserted to de-couple the production and dissemination environments, the interface problem reduces to $n + m$ and therefore becomes much more manageable. Once the interfaces to and from the reference layer are defined, new tools, products or domains in either environment are easily added. This enables

Eurostat to manage the evolution and change inherent in the European statistical environment, because changes to one part of either production or dissemination will not have detrimental impacts to any other part.



3.2. Potential benefits

A. For producers

By de-coupling production from dissemination, a reference layer enables producers the freedom that they require to choose the most appropriate tools and techniques to do their job. It provides well documented data to producers when they are using other domains. It also helps them to understand their own domains when they look at old data in their own domain, not only because their own memories may not be completely accurate, but because the movement of producers within the office means that expertise is not always available.

Producers also gain a common interface into the reference area, rather than the many different interfaces that they have to cope with at present. They have more mobility for other jobs with Eurostat, because the common working environment will be familiar to them in a new post.

B. For disseminators

By providing well documented data in a standard form, the reference layer provides disseminators with better raw material for their products. This enables them to provide better and more diverse products in a more timely and customer-driven manner, as desired in Eurostat's mission statement.

C. For customers

While not directly involved in the reference layer, the reference layer should provide customers with data that is:

- Visible. They will know what is available.
- Understandable. It will be documented to an agreed minimum standard.
- Accessible. They will be able to get the data more easily.
- Comparable. the standardisation of meta-data will help to enable cross-domain comparison.
- Timely. The data should be available more quickly.
- Accurate. The authorisation of data moving out of production should ensure its high quality.
- Appropriate. Feedback on usage by users should help Eurostat provide what users need.

D. For managers

A reference layer assists managers because it can provide:

- Central security and access control to Eurostat data. This ensures that data is not left on insecure and vulnerable PC platforms, and can be safeguarded as the vital resource that it is.
- Harmonisation of nomenclatures and code lists leading to less duplication of effort in capture and translation of meta-data. Producers will be able to re-use previous work by others. There will also be increased comparability of data between different domains because of harmonised definitions of their dimensions.
- Flexibility of choice in production tools and methods and in dissemination products.
- Stability of Eurostat data and meta-data over time.
- Co-ordinated market feedback to enable Eurostat to disseminate what their customers want and need. This information will come from both the reference and dissemination environments.

4. TYPOLOGY OF THE PROCESSES

The processes corresponding to one or to a number of functionalities, must occur in one of the defined environments: production, reference, diffusion.

An EDP application is seen as an ordered set of processes; the latter need not necessarily be on the same computer, but can be on two (or several) machines of the same or different platforms.

The distribution of the processing within an application over the different platforms is a difficult decision for which some rules have to be fixed:

- ⇒ general rules applicable throughout the European Union;
- ⇒ volume considerations concerning the data and the population of potential users;
- ⇒ security considerations requiring a more or less protected environment, access control functions, backup and recovery, archiving etc.;
- ⇒ the availability of suitable software on the different platforms.

4.1. Production

The production layer divides into several subjects according to the type of objects managed; precise description of the objects is needed and depends on the construction of standard applications giving a framework of utilisation to Eurostat's various producers. Eurostat has selected, on this basis, FAME for series and ACUMEN for multidimensional objects.

4.2. Reference

The «acceptance of information» poses the problem of dynamic definitions and the automated control of the passage between production and the reference. It depends on the production structures and on the complexity of the rules to be implemented. In view of the volume of existing information, only a completely automated procedure is viable.

The concept of quality of managed information is vital for Eurostat: the separation into two environments must preserve the wealth of existing information, i.e. allow more thorough interrogation of the set of information available if that is necessary.

The management of reference is concerned with a homogenous solution allowing uniform access to the set of managed information.

The visibility and security layer must ensure control of the activities of users privileges and of data confidentiality. It is important to guarantee that no leaks of information cross this layer of software.

The «FIND and DELIVER» layer allows a non-specialist user easy navigation around Eurostat s information and a rapid and exhaustive reply to his questions.

4.3. Diffusion

Eurostat disseminates statistical information by means of 3 types of products:

- ⇒ Publications or statistical reports, which need to mix texts, graphs and tables, using desk-top publishing products.
- ⇒ Databases.
- ⇒ Other electronic diffusion services (download, CD-ROM, etc.).

REDACTION AUTOMATIQUE DE COMMENTAIRES A PARTIR DE BASES DE DONNÉES STATISTIQUES

Dr. JEAN LOUIS ROOS*

Institut National de la Statistique et des Etudes Economiques

ALIEN is a knowledge based system that modelizes the economist's behaviour. It is used at the French Statistical Institute and at Eurostat. It can be seen as an assistant that help economists to write economic diagnosis in natural language as well as a human economist could do.

Alien is very useful for analysis in very large data base and/or for regular diagnosis. It study the shape of all series in an application, the instantaneous relationship and determine what is important and what is not.

ALIEN is used under a user friendly software (DARWIN) developed for Windows 95. DARWIN build text generation with ALIEN, but also construct graphs and table, help the user to analyse the data base and the semantic of an application. Last, DARWIN build a nice layout of the mixed comment, graph and table with WINWORD or an equivalent software.

Text generation from statistical data

Keywords: Text Generation, data Base, Artificial Intelligence, Automatic economic diagnosis

*Dr. Jean Louis Roos. Institut National de la Statistique et des Etudes Economiques. Avenue Albert Einstein BP 26000, 13791 Aix En Provence Cedex 3. Tel 42 16 29 59. Fax 42 16 29 00. roos@cil3a.insee.atlas.fr

—Article rebut l'octubre de 1996.

—Acceptat el febrer de 1997.

Les bases de données statistiques sont de plus en plus fréquemment de très grandes tailles. Lorsqu'un statisticien doit rédiger des commentaires réguliers sur les données contenues dans de telles bases, son travail peut être particulièrement long et difficile car il devra examiner une masse considérable de chiffres. Qui plus est, les données possèdent une sémantique qui n'apparaît jamais (ou que très rarement) dans les bases, donc le commentaire que le statisticien fera dépendra en grande partie de ses propres connaissances. Dès lors, tout changement du responsable chargé d'une étude statistique régulière aura des conséquences parfois dramatiques sur la qualité des textes d'analyse. Curieusement, les fabricants de logiciels ont privilégié le développement d'outils puissants d'interrogation, mais rien n'a été fait en ce qui concerne l'interprétation des données contenues dans de telles bases.

A notre avis, et dès maintenant, les statisticiens ont besoin d'un outil qui doit répondre à trois préoccupations :

◆ **Disposer d'un assistant**

Régulièrement, un statisticien doit pouvoir demander à son assistant un texte correctement rédigé qui analyse l'information dans une base de données. Ce texte doit pouvoir s'accompagner de tableaux et de graphiques. Il doit être aisément modifiable à travers un traitement de texte. L'assistant doit être suffisamment efficace pour ne rien oublier d'important et, évidemment, ses analyses doivent être « stables » dans le temps, et totalement indépendantes de facteurs subjectifs extérieurs.

◆ **Pouvoir traiter des commentaires de masse**

Il est fréquent de ne pas réaliser des commentaires à partir de la totalité de l'information dont on dispose. Ceci simplement parce que la masse d'information à traiter est trop importante ; c'est le cas lors d'analyses sectorielles, par communes, par pays, etc. Un être humain ne peut faire face — parfois simplement pour des raisons de temps — à la rédaction de dizaines, voire de centaines de diagnostics. Les difficultés seront encore plus lourdes si les rédactions doivent s'accompagner de graphiques et être mises en page.

◆ **Disposer de références**

On appelle références (guideline) la définition, dans des domaines divers, de recommandations de « bonnes pratiques ». Une bonne pratique est ce qu'il faut faire, dire, écrire face à une situation donnée. Ces « guidelines » ont été surtout développés en médecine, à partir de travaux de formalisation de diagnostics ou de traitements sur lesquels des consensus existaient. En simplifiant, une référence permet de décrire une situation et un traitement de celle-ci formalisés et acceptés par une majorité des praticiens. Ce consensus, certes très relatif, fournit une réponse efficace et argumentée face à une maladie donnée.

Pour un statisticien, il est tout aussi important de disposer de telles références associées à chaque base sur laquelle il travaille. Sinon, les risques d'erreurs dans l'interprétation des données peuvent être importants. Dans le cadre de travaux spécifiques de formalisation de l'analyse conjoncturelle, l'**Insee** a commencé à développer un logiciel répondant à ces préoccupations. **Eurostat** s'y est très vite associé et a assuré une grande partie du financement. Des utilisations sont maintenant en cours tant à l'**Insee** qu'à **Eurostat**.

Nous allons présenter ici successivement le logiciel de génération de texte: **Alien**, et son environnement de fonctionnement sous Windows: **Darwin**. L'ensemble se présente comme un assistant «informatique» pour le statisticien. L

1. LA GENERATION DE TEXTE: ALIEN

Alien est un logiciel qui formalise —en réalité qui modélise —les mécanismes qui amènent un statisticien à «rédiger» un diagnostic à partir d'un ensemble de données quantitatives. Cette formalisation repose sur l'observation suivante: la plupart des textes d'analyse simples ont un aspect «stéréotypés» —et ce quelle que soit la langue dans laquelle ils ont été écrits— il existe donc des schémas d'analyse très répétitifs, et il est alors possible de les modéliser.

1.1. Les bases theoriques

Les fondements d'**Alien** sont assez simples: le statisticien ne manipule plus des séries mais des indicateurs; ceux-ci sont des objets incorporant une sémantique. Deuxièmement, le statisticien utilise des formes de discours stéréotypés. Au sein de ces formes, des dictionnaires permettront la génération de texte et le changement de langues.

1.1.1. Les indicateurs

Alien repose sur l'idée que l'économiste ne manipule pas des séries statistiques, mais plutôt un ensemble de connaissances sur celles-ci, y compris de connaissances sur le vocabulaire associé à chaque série. Ces connaissances ont deux origines: au départ, et pour une application donnée, le système possède une connaissance minimale standard; cette connaissance est ensuite développée par un processus d'apprentissage qui repose sur un mécanisme d'inférence.

Alien range la connaissance utile dans des «objets» que l'on nomme les «indicateurs», et qui englobent, après une étape d'apprentissage, toute l'information indispensable pour la rédaction d'un diagnostic.

Les connaissances «ex ante» seront, par exemple, le nom des séries à analyser, différents codages sur la sémantique possible de la série (par exemple: la série peut-elle être interprété en variation? est-ce un flux ou est-ce un stock?,...), mais aussi le vocabulaire, parfois le «jargon», le plus courant associé à l'application.

Les connaissances «ex post» résultent d'opérations de calculs et permettent d'interpréter les valeurs en niveau et en variation pour les périodes à étudier. Ces «jugements» sont liés aux valeurs passées des séries.

Le plus important, pour obtenir des commentaires de qualité, est de bien coder certaines valeurs ex-ante. C'est un travail délicat, mais identique à celui que fournit l'être humain lorsqu'il rédige un commentaire; bien souvent nous manipulons des séries sans connaître toute leur signification, et toute erreur sur la sémantique de départ peut être ensuite une cause d'erreur dans l'analyse que l'on en fait.

1.1.2. Les discours stéréotypés

La richesse d'une langue est telle qu'il est illusoire, sans moyens importants, de définir une méthode qui permettrait d'écrire tous les discours possibles pour un diagnostic donné. Cependant l'usage montre que, lors de la rédaction de textes de commentaires statistiques, une certaine paresse, ou tout au moins l'habitude, fait que les discours utilisés sont fortement stéréotypés. En voici deux exemples choisis dans la presse:

Sales of new cars in western Europe increased by 1.2 per cent last month, with higher demand in Germany and France compensating for a heavy decline in sales in Italy. Sales in the first 11 months of the year, at 12.53m, were 1.3 per cent lower than in the corresponding period a year earlier.

Consumer optimism about the economy rose slightly in August from July. But the gain was small: from 77 in July to 77.3 in August.

Ceci étant, d'une part il existe de nombreuses variantes à de semblables discours, d'autre part il est possible de rédiger certains discours plus spécifiques dans des situations particulières. Qui plus est, le discours stéréotypé reste insuffisant pour faire une rédaction de qualité; dans un domaine particulier, par exemple la finance, la conjoncture, le commerce,... nous pouvons observer des «dialectes», c'est-à-dire un vocabulaire, parfois une forme lexicale particulière, qui est spécifique au sujet traité. **Alien** prendra en compte cette préoccupation.

1.1.3. Les dictionnaires

Les dictionnaires jouent plusieurs rôles dans la rédaction d'un texte: évidemment, ils fournissent avant tout le vocabulaire, mais ils indiquent aussi la plupart des règles

grammaticales courantes: les formes de pluriels par exemple ou encore les conjugaisons. C'est ce qui se passe avec les dictionnaires d'**Alien**. Et tout comme un être humain, **Alien** utilisera aussi des dictionnaires de synonymes, ou encore des dictionnaires de dialecte dont nous venons de parler.

1.2. Le système formel: une double arborescence

L'écriture est une activité humaine particulière qui implique de fréquent «retours en arrière», c'est aussi le cas d'une analyse statistique. De ce fait, il est difficile de construire un logiciel totalement déterministe. Alien se présente en pratique comme un système organisé autour d'une double arborescence d'informations, la rédaction étant traitée par une série de petits traitements organisés sur les noeuds terminaux et utilisant les dictionnaires. Le résultat est un logiciel indéterminisme «a priori» sur un discours précis.

1.2.1. L'arborescence des indicateurs

Un indicateur principal définit le sujet traité, il possède comme descendance des indicateurs *domaines* qui constitueront les paragraphes de textes. Chaque domaine est constitué d'indicateurs quantitatifs ou qualitatifs dont certains possèdent des occurrences (par exemple la production peut être analysée par branche, par pays,...). Ces indicateurs forment les discours sur chaque série à commenter. La hiérarchie sujet, domaine, indicateur final peut être comparée à celle d'un modèle économétrique où les domaines sont des blocs et les indicateurs sont des équations.

1.2.2. L'arborescence du contenu des indicateurs

Tous les indicateurs sont structurés de la même façon: ils possèdent des propriétés, et pour chacune d'elle, existe une valeur (parfois par défaut). Un indicateur possède environ deux cents propriétés. Il est possible, pour chaque application, de rajouter des propriétés originales.

1.2.3. Les dictionnaires

Ils forment une arborescence élémentaire à partir de points d'entrées. Les dictionnaires sont des objets particuliers qui regroupent le vocabulaire commun aux indicateurs, mais aussi la majeure partie des règles grammaticales associées et des conjugaisons. Le contenu des dictionnaires est défini par un index qui est un type abstrait particulier. On trouvera les traitements suivants:

- ◆ Les conjugaisons: l'entrée est un verbe, ou un groupe verbal, à l'infinitif. Sont conjugués: le présent, le passé, et le conditionnel.

- ◆ Les accords: à chaque mot, ou groupe de mots, correspond le même ensemble au masculin singulier, féminin singulier, masculin pluriel, féminin pluriel, neutre singulier, neutre pluriel (ces deux cas sont pour d'autres langues), ainsi que deux cas spécifiques (par exemple en français l'élision avec h).
- ◆ Les synonymes: six synonymes sont donnés pour chaque mot ou expression.
- ◆ Le vocabulaire complémentaire: il est constitué de mots isolés, principalement invariables: prépositions, conjonctions, adverbess, etc.

1.3. Les solutions informatiques

1.3.1. Les indicateurs

L'objet que manipule le statisticien est un indicateur, c'est-à-dire une structure dans laquelle se trouve rangée toute l'information indispensable pour interpréter une (ou plusieurs) série(s). En pratique, une grande partie de l'information étant partageable, un indicateur ne porte que l'information qui lui est spécifique. Le partage de l'information se fait ensuite par en grande partie par héritage - ce qui implique l'organisation hiérarchique des indicateurs - mais aussi par accès à des dictionnaires. Il existe ainsi une information commune à tous les indicateurs.

Le contenu d'un indicateur, et donc l'information qu'il peut porter, est susceptible de modifications d'une application à une autre, ou encore en fonction de développement du système; les indicateurs sont donc construits à partir de prototypes, c'est-à-dire de modèles, aisément modifiables. Ces modèles définissent un objet organisé d'une certaine façon, avec des attributs et des valeurs. Ce sont donc des types abstraits ou des prototypes.

Le modèle «indicateur» décrit, sous une forme arborescente, ce que doit contenir un indicateur. Chaque noeud joue un rôle d'attribut. Ces attributs sont en nombre variables. A un attribut peut être associé un sous-arbre ou une valeur. Les valeurs associées dans les indicateurs instances sont aussi modifiables. Voici quelques-uns des principaux attributs existants:

- ◆ des types, permettant de classer chaque indicateur. A chaque groupe d'indicateurs sera associé des traitements, et parfois un discours possible particulier.
- ◆ des noms pour l'indicateur,
- ◆ des informations sur la périodicité de la, ou des séries, propre(s) à l'indicateur,
- ◆ des formats d'impression de valeurs,

- ◆ du vocabulaire particulier: verbes, adjectifs, adverbes, prépositions, et des règles grammaticales,
- ◆ des listes d'indicateurs occurrences,
- ◆ des formules de calcul expliquant comment est obtenue la série de base. C'est parfois une différence (un solde par exemple), parfois une somme de plusieurs séries, voire une somme pondérée. Il y a aussi les formules décrivant les variations, lesquelles peuvent être calculées de nombreuses façons et avec différents décalages; enfin, on trouvera quelquefois des formules d'accélération.
- ◆ des index de jugement, et des outils permettant de les construire,
- ◆ la description des périodes à analyser en terme de variation et éventuellement d'accélération (ou de décélération)
- ◆ les valeurs et les noms de la, ou des séries, propre(s) à l'indicateur,
- ◆ des programmes et des traitements.

Cette liste n'est pas exhaustive. Chaque utilisateur peut rajouter de l'information. En pratique, plus de deux cents attributs définissent ainsi la sémantique utile d'un indicateur.

1.3.2. Les traitements

Pour obtenir la rédaction d'un texte de diagnostic par un logiciel, il est indispensable de donner à celui-ci la sémantique de l'analyse à faire; il est cependant inconcevable de fournir la totalité de celle-ci; c'est le système lui-même qui doit la construire par un processus d'apprentissage. Le statisticien ne fournira que des informations minimales telles le nom et les valeurs des séries à analyser. Dans **Alien** trois étapes de traitements se succèdent; La première est la construction de l'arborescence des indicateurs et le remplissage de ceux-ci, ensuite se déroule une étape de calculs sur les valeurs, enfin a lieu une sélection de l'information pertinente.

1.3.2.1. La mécanique d'apprentissage

La sémantique nécessaire à la rédaction est construite automatiquement à partir d'un modèle standard d'indicateur et de particularités propres à chaque indicateur. Cette mécanique d'apprentissage est paramétrable et peut être complétée en substituant aux données construites automatiquement d'autres informations décrites manuellement (dans un fichier).

Informatiquement «parlant» l'apprentissage implique l'utilisation de trois fichiers de «connaissances» spécifiques au domaine traité: la description minimale des indicateurs (nom, types, liens avec les séries), les informations communes à tous, et enfin des compléments d'informations particulières (facultatives).

De nombreux choix a priori seront ici définis par le statisticien: Entre autres des choix sur les types d'indicateurs, les périodes analysables, et les paramètres de calculs d'index de jugements. Par défaut, c'est le système qui effectuera ces choix.

1.3.2.2. L'analyse numérique extensive

Une fois la base de connaissances construite, Alien réalise une analyse extensive des valeurs numériques. Ceci consiste à appliquer sur toutes les périodes de temps analysables —celles-ci ont été fixées lors de la phase d'apprentissage— des formules mathématiques prédéfinies. Il s'agit d'obtenir le maximum de données quantitatives pour chaque indicateur: les niveaux des valeurs, mais aussi les variations et parfois les accélérations. Pour chaque information, un jugement qualitatif sera porté et codifié sur une échelle de sept valeurs: de «très faible» à «très fort». Ceci constituera des index de jugement.

Dans une analyse manuelle, un statisticien portera généralement un jugement sur les valeurs qu'il observe. Ainsi une augmentation du chômage, des prix, ou encore de la production, sera qualifiée par un expert de faible, de moyenne ou de forte! Comment de tels qualificatifs se forment-ils?

En fait, il n'existe pas de règles absolues. Ainsi entre 1976 et 1978 l'Insee qualifiait ainsi diverses hausses de prix:

- ◆ en 1978: faible, avec 7,9%
- ◆ en 1976:
 - ☞ vive, avec 9,6%
 - ☞ modérée, avec 7%
 - ☞ quasi-stable, avec 2,7%

Chacun aura constaté que la qualification d'un chiffre se fait souvent dans un contexte historique, mais doit-on toujours faire référence au passé? ou doit-on au contraire avoir une référence absolue? Dans **Alien** différentes options sont possibles:

- ◆ les tranches peuvent être définies manuellement,

- ◆ les tranches peuvent être calculées automatiquement à partir de l'historique de la série (par défaut), Dans ce dernier cas, la «longueur» de la série est importante et influence les résultats.
- ◆ Si les tranches sont calculées automatiquement, l'utilisateur peut encore fixer certaines options: Alien définit une valeur centrale, correspondant à une situation moyenne. Pour une série en niveau ou en stock, c'est la moyenne; pour les séries en variation c'est 0 si on raisonne en différence, 1 en rapport, 100 en pourcentage. Mais l'utilisateur peut aussi fixer librement ces valeurs centrales.

Dans le calcul des jugements, les valeurs proches de cette moyenne feront partie de la tranche moyenne et plus on s'éloignera de celle-ci, plus on déterminera les tranches permettant de qualifier: assez faible, assez fort —faible, fort— très faible, très fort (ces adjectifs sont bien évidemment dépendants de la série traitée). L'éloignement de la valeur moyenne est aussi paramétrable.

1.3.3. La recherche de l'information pertinente

Une fois qu'un développement extensif sur les valeurs et les jugements a été réalisé, il n'est pas possible de faire une rédaction qui utiliserait tout ce matériau. Une sélection est indispensable car on dispose de trop d'information. En effet, un texte bien écrit va directement à l'essentiel. La troisième étape des traitements dans **Alien** réalise cette sélection. Seule l'information la plus «forte» est conservée et mise en avant à partir de règles et de choix prédéfinis, mais eux aussi modifiables. Le résultat de cette opération est un «vecteur narratif», c'est-à-dire une séquence d'informations qui définit le «quoi dire».

1.3.4. La génération et les structures narratives stéréotypées

Lorsque les indicateurs sont complètement renseignés et que le vecteur narratif est codé, il ne reste plus qu'à rédiger. Jusqu'ici, toutes les étapes étaient indépendantes de la langue. A partir de maintenant, ce ne sera plus le cas. On va traduire le vecteur narratif à l'aide de deux outils. Le premier est un ensemble de vecteurs d'ordres de rédaction. Il existe plusieurs vecteurs d'ordres, ou séquences d'ordres, ce sont les «schémas narratifs» ou «structures narratives». Le second outil est bien sûr constitué des dictionnaires.

Le changement de langue implique de changer de dictionnaires, mais aussi de modifier les vecteurs des schémas narratifs. Ainsi l'Ifo a réécrit entièrement ces schémas pour la langue allemande. Actuellement, pour les tests effectués en Anglais, en Italien et en Espagnol, on utilise les structures narratives de la langue française. Le résultat est donc imparfait.

Les structures narratives gèrent la plupart des problèmes courant de lexicalité:

- ◆ l'élision qui consiste dans la suppression d'une voyelle finale devant une autre voyelle ou un h muet. L'élision est gérée pour les articles, mais aussi devant des mots invariables, tels «alors que». L'une des difficultés consiste à gérer les h (muets ou non). Les mots débutants par des «h» sont donc codés spécifiquement.
- ◆ les accords et les conjugaisons sont traités par les dictionnaires et des fonctions associées. Une variable globale du système, **syntax**, identifie à tout moment le genre et le nombre actif; elle peut être «sauvegardée» pour introduire dans un discours des éléments ayant un genre et un nombre différent du discours central; lorsque l'on retourne à ce discours, le genre et le nombre qui étaient actifs retrouvent leur état antérieur.
- ◆ les indicateurs quantitatifs ont un discours qui intègre, de diverses façons, les valeurs qui interviennent dans l'analyse. Mais ces valeurs ne doivent généralement pas apparaître lorsqu'elles sont égales à 0.
- ◆ l'enchaînement des prépositions et des pronoms; on dit:
«la production augmente sauf dans la mécanique où elle diminue»
«Les carnets de commande sont stationnaires, sauf pour les carnets étrangers qui se sont dégarnis».
- ◆ l'enchaînement stylistique des propositions: il est souvent plus élégant de faire suivre une proposition contenant un verbe transitif par une autre contenant une forme intransitive, et réciproquement. Il existe aussi des enchaînements particuliers avec l'auxiliaire être. Alien assure une telle gestion en mémorisant, à tout moment, le sujet traité.
- ◆ l'enchaînement des discours entre indicateurs (par des nuances: cependant, mais ...).

1.3.5. Le concept d'application

Une application **Alien** est un ensemble de connaissances sur un sujet donné. C'est sur ce sujet que sera réalisée la génération de texte.

Une application est constituée tout simplement de fichiers de connaissances et d'une base de données. Ces connaissances sont les connaissances particulières au sujet traité, et non les connaissances générales comme les dictionnaires. Elles sont formées des éléments suivants:

- ◆ le squelette de l'application: c'est le descriptif minimal pour CHAQUE futur indicateur: on donnera ici les types de l'indicateur, son nom et la, ou les, série(s) à utiliser,
- ◆ le modèle d'indicateur renseigné: c'est un prototype d'indicateur, mais complètement garni avec des valeurs communes pour tous les indicateurs,
- ◆ le fichier de complément stylistique: il est facultatif et permet d'améliorer l'apprentissage automatique. Il contiendra les tournures et le vocabulaire propre à l'application.

Le travail le plus délicat est de renseigner le squelette, en particulier ce dernier doit contenir une description typologique des indicateurs, ce qui est parfois délicat à définir. Ainsi plusieurs types sont utilisés. A titre d'illustration, examinons le type principal (il existe aussi des types secondaires).

Ce type permet de rattacher la série de base de l'indicateur à une famille parmi plusieurs possibles:

- ◆ type 0 et 1: les valeurs de la série à analyser sont significatives d'un niveau, d'un stock, d'une quantité absolue, et ce niveau a une signification économique par lui-même. Par exemple, un effectif salarié, un stock de produits finis, un chiffre d'affaires. Dans le type 1 l'étude de la variation de ce niveau ne pourra être faite, car elle n'aurait pas de sens.
- ◆ type 2: les valeurs de la série à analyser sont significatives d'un niveau, mais celui-ci ne veut rien dire; c'est le cas d'un indice, comme l'indice des prix, par exemple. On ne peut donc le commenter. Par contre, la variation de la série, entre deux périodes, est significative.
- ◆ type 3: les valeurs de la série à analyser sont ici significatives d'une variation. Le niveau est inconnu. La variation, entre deux périodes, de la série implique une accélération (ou un ralentissement, ou encore d'une stabilité de la variation).

Le choix du type implique des commentaires, des analyses, parfois fort différentes. Il est possible d'obtenir des commentaires fort différents sur des séries en ne modifiant que le type principal dans le descriptif de l'application!

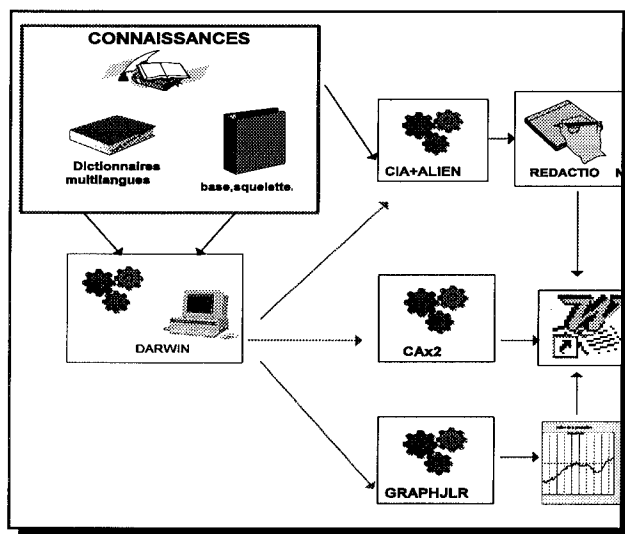
Si le type est mal choisi, le commentaire sera peut-être tout simplement faux! Mais ceci ne peut être reproché à Alien: un tel commentaire faux sera aussi bien fait par un être humain qui connaîtrait mal son sujet!

2. L'INTERFACE UTILISATEUR: DARWIN

Alien reste complexe à utiliser, aussi un environnement agréable d'utilisation, **Darwin**, a été réalisé. **Darwin** est un environnement de travail sous Windows. Il apporte une utilisation facile des textes générés à travers des éditeurs, un outil graphique de visualisation des séries (**Graphjlr**) et une mise en forme finale du texte, des graphiques et des tableaux par appel d'un traitement de texte, tel WINWORD. Cette mise en forme est automatique par l'intermédiaire du logiciel **Cax2**.

Darwin et ses logiciels associés sont écrits en langage C avec l'Api Windows 32 bits. Ils sont compatibles Win 95 et Windows 3.1.

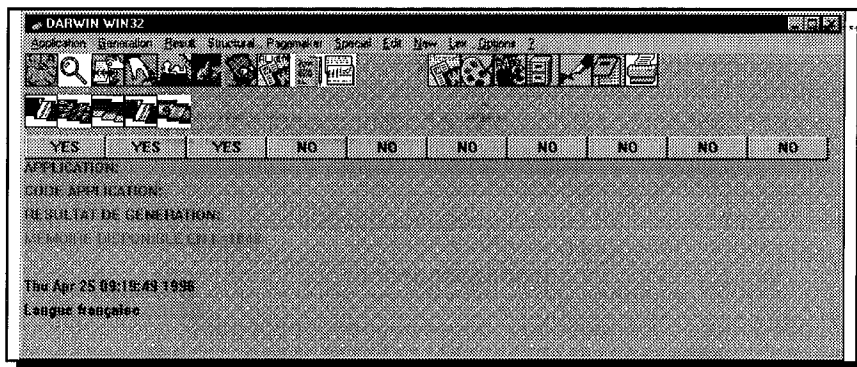
Outre un environnement de génération de texte, **Darwin** permet aussi de manipuler, de générer, de modifier les éléments d'une application; il autorise encore de choisir la langue de rédaction (cette partie est en cours de développement).



Dans l'avenir, **Darwin** automatisera une partie des tâches de construction des nouvelles applications, des dictionnaires, et peut-être même des structures narratives. Examinons les premières possibilités de cet outil.

2.1. Les capacités du logiciel

Avant tout Darwin permet de réaliser des générations de texte à partir d'Alien, mais il informe aussi l'utilisateur sur le contenu d'une application, il réalise la mise en page (sous Word) avec des graphiques; enfin, il permet facilement le changement de langue.

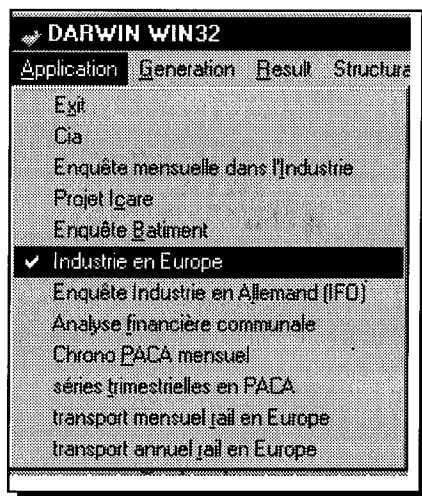


2.1.1. La génération

Réaliser une génération se fait en trois étapes: choisir une application, lancer la génération, examiner les résultats.

2.1.1.1. Le menu application

Toute utilisation implique d'avoir sélectionné une **Application** dans un menu dynamique:



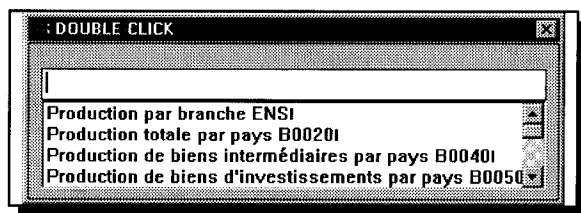
La sélection se fait sur le menu principal. Elle fait référence au contenu de fichiers d'initialisation qui définissent où se trouvent les applications.

2.1.1.2. Le menu génération

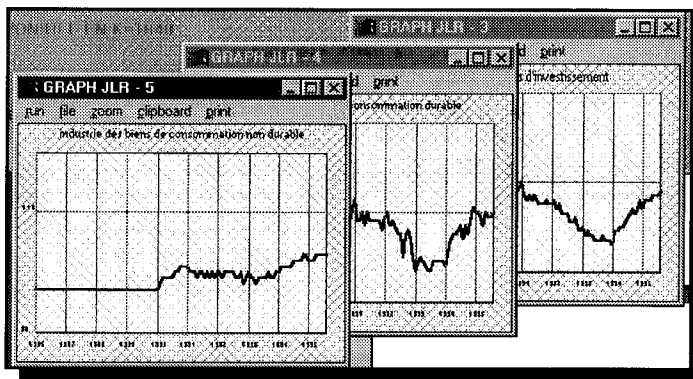
Il fait appel au logiciel d'intelligence artificielle Cia. La génération peut être, ou non, visible à l'écran.

2.1.1.3. Le menu résultat

La génération étant bien terminée, le menu **Result** permet de visionner les résultats. Par exemple à travers des éditeurs de texte, et donc de la modifier. Il est aussi possible de visionner les graphiques et les tableaux associés. Si la génération a été très importante, parfois une centaine de fichiers sont construit à raison d'un par paragraphe, Darwin permet à l'utilisateur de choisir le paragraphe qu'il désire visionner, ou dont il désire examiner les graphiques.



Les graphiques apparaissent superposés, avec un menu qui permet de nombreuses manipulations (zoom total ou partiel, sauvegardes...).

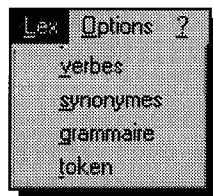


2.1.2. L'information

2.1.2.1. Les éditeurs simples

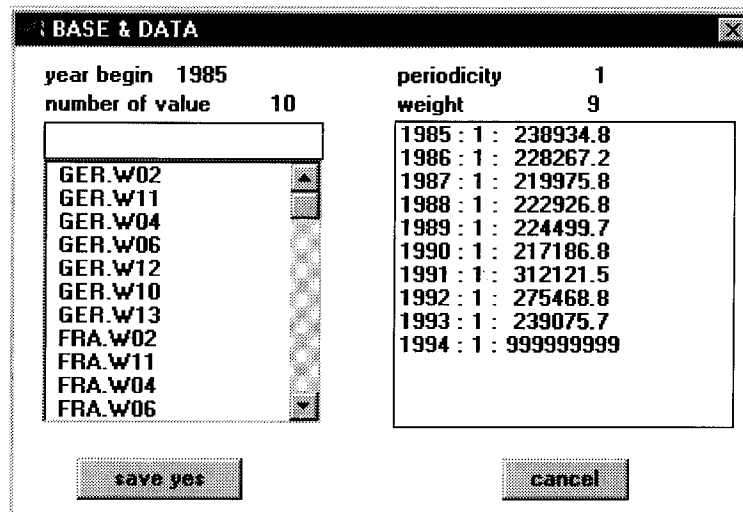
Un menu **Edit** fournit des éditeurs permettant de visionner rapidement et de modifier tous les fichiers indispensables pour une application.

Un menu lexical permet de visionner les dictionnaires dans le contexte d'une langue:



2.1.2.2. Les éditeurs structurels

Ces éditeurs particuliers autorisent une vision structurée des connaissances d'une application: Ainsi, par exemple, la base de données est visible par le nom des séries qu'**Alien** utilise; Et à partir d'un clic et d'un ascenseur, chaque série peut être explorée.



2.1.3. La mise en page

Un menu **Word** permet une mise en page automatique sous un autre traitement de texte de la génération issue d'**Alien**, et éventuellement modifiée par l'utilisateur. Un fond de page doit être décrit ainsi qu'une liste de commandes dans un fichier particulier. Les graphiques et tableaux peuvent être insérés. L'ensemble des opérations est géré par le logiciel **Cax2**. Les versions 2,6 et 7 de Word sont reconnues.

2.1.4. L'utilisation multilingue

Des drapeaux permettent de visionner les dictionnaires et de choisir une génération dans une langue particulière. En fait seul l'anglais est en cours de développement. L'italien et l'espagnol sont envisagés. Une version allemande partielle a été développée par l'Ifo.



3. LES UTILISATIONS

L'insee et Eurostat utilisent expérimentalement ce logiciel.

3.1. A l'Insee: la conjoncture

Alien est utilisé tous les mois pour l'enquête mensuelle dans l'industrie. Un commentaire de quelques pages est fabriqué, avec tableaux et graphiques (cf. annexe).

Simultanément, un projet de retour d'information aux entreprises du diagnostic conjoncturel sur leur secteur est maintenant opérationnel. A titre d'information Alien compose actuellement deux pages de texte pour 41 produits industriels en quelques secondes.

3.2. A Eurostat: les transports et l'industrie

Dans le cadre du projet de recherche SUPCOM, le logiciel est adapté et testé pour les résultats des statistiques sur les transports (données annuelles et mensuelles) et des données sectorielles par branche.

3.3. D'autres utilisations

L'Institut de recherche économique **Ifo** à Munich a une version en langue allemande de **Darwin**. Un projet en cours envisage avec d'autres partenaires une version plus complète intégrant aussi la langue hollandaise.

4. CONCLUSION

Darwin et **Alien** fonctionnent à ce jour, mais suscitent de nombreuses et parfois de violentes critiques. On peut classer les opposants à son utilisation en trois groupes:

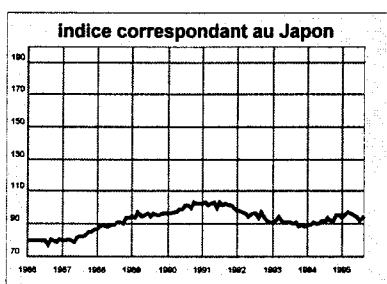
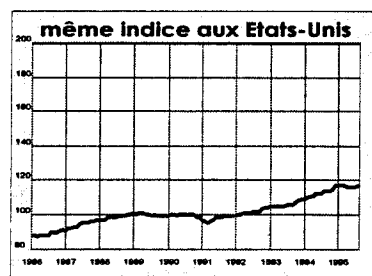
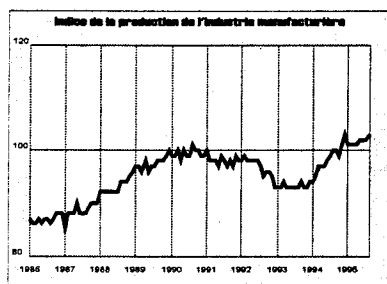
- ◆ ceux qui sont certains que l'analyse statistique faite par l'être humain ne peut, en aucune manière, être reproduite par une machine,
- ◆ ceux qui pensent que la reproductibilité sera possible dans le futur, mais qu'elle est encore trop complexe et trop mal connue. Donc, celle qui peut être faite actuellement est insuffisante.
- ◆ Bien sûr, tous ces opposants ont partiellement raison. Aucune machine ne pourra égaler une analyse humaine faite par un spécialiste de haut niveau. Par contre, il est vrai qu'il est possible de formaliser des commentaires simples, mais qu'une telle méthode reste complexe et difficile. Enfin le problème des gains de productivité que cela pourrait entraîner reste entier.

5. EXEMPLES

5.1. Exemple 1: Eurostat

Production de l'industrie manufacturière par pays B0200

Sur la période juin-août, l'indice de la production de l'industrie manufacturière dans la communauté européenne a continué d'augmenter faiblement de 0.8%, en moyenne mobile sur les trois derniers mois par rapport à la période mars-mai 1995. Alors que l'indice correspondant au Japon a subi encore un sensible recul de -2.4%. Le même indice aux Etats-Unis est resté stationnaire à 0.1%. En ce qui concerne les états membres de l'Union Européenne, et pour le même indice et la même périodicité, l'augmentation a été vérifiée particulièrement pour l'Irlande (5.1%) et pour l'Autriche (3.8%); en revanche pour le Danemark (-0.9%) l'indice de la production a diminué faiblement.

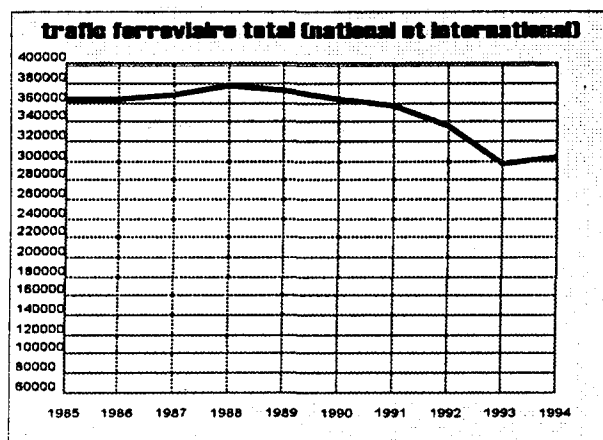


	juin 1995	juillet 1995	août 1995	septemb 1995	octobre 1995	novembr 1995
indice de la production de l'industrie manufacturière	101	101	102	102	102	103
indice correspondant au Japon	97	98	95	94	92	94
même indice aux Etats-Unis	117	118	116	116	116	117

5.2. Exemple 2: Eurostat

Trafic ferroviaire total ENS

Dans l'Union européenne en 1994, le trafic ferroviaire total (national et international) (sauf Allemagne de l'Ouest, Danemark, Espagne, Italie, Luxembourg) est resté très faible (303147.0 tonnes) mais a connu une sensible augmentation (1.7%). Cette augmentation a été notée pour la Belgique (10.2%), pour la France (5.9%) et pour les Pays-bas (6.5%); tandis que la situation est opposée principalement pour la Grèce (-58.%). Le trafic total par chemin de fer en tonnes kilomètres (excepté Allemagne de l'Ouest, Danemark, Espagne, Italie, Luxembourg) est demeuré faible (69133.00 tonnes) mais a connu une très nette augmentation (4.2%). Cette augmentation a été observée pour la Belgique (7.0%), pour la France (7.9%) et pour les Pays-bas (5.1%); à l'opposé en particulier pour la Grèce (-35.%) et pour la Grande-Bretagne (-5.7%) le trafic total par chemin de fer en tonnes kilomètres a diminué.

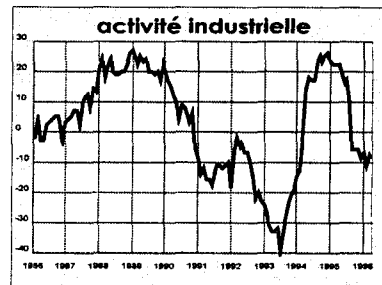
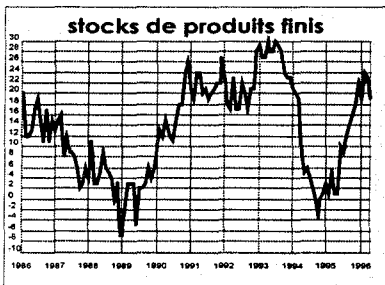
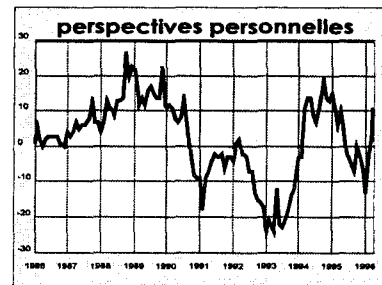
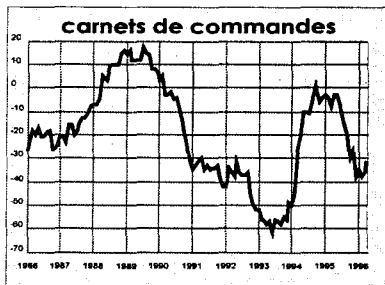


	1989	1990	1991	1992	1993	1994
TRAFIC FERROVIAIRE TOTAL (NATIONAL ET INTERNATIONAL)	373049	363856	357764	337532	298145	303147
TRAFIC TOTAL PAR CHEMIN DE FER EN TONNES KILOMÈTRES	78628	76044	74751	73626	66346	69133

5.3. Exemple 3: INSEE

Ensemble de l'industrie ENS

Selon les chefs d'entreprise interrogés à l'enquête mensuelle d'avril, l'activité industrielle a continué de diminuer légèrement, à un rythme plus rapide qu'au trimestre précédent. Les stocks de produits finis sont toujours jugés faiblement supérieurs à la normale mais ont semblé subir une baisse pour le deuxième mois consécutif. Les carnets de commandes sont restés faiblement dégarnis mais ont connu encore une légère amélioration. Les perspectives personnelles de production pour les prochains mois étaient plutôt optimistes; à l'opposé les perspectives générales étaient plutôt pessimistes. Les industriels interrogés prévoient des augmentations de prix à la production pour l'ensemble de l'industrie; à l'inverse pour leurs propres produits ils s'attendent à des baisses de prix à la production.



	novembre 1995	décembre 1995	janvier 1996	février 1996	mars 1996	avril 1996
ACTIVITÉ INDUSTRIELLE	-6	-9	-7	-11	-8	-9
STOCKS DE PRODUITS FINIS	18	22	19	24	23	19
CARNETS DE COMMANDES	-27	-38	-35	-38	-36	-31
PERSPECTIVES PERSONNELLES	-3	-6	-13	-2	2	11
PERSPECTIVES GÉNÉRALES	-21	-37	-33	-32	-23	-21

6. BIBLIOGRAPHIE

- [1] **Roos, J.L.** (1995). *ALIEN: Un outil pour modéliser la rédaction de diagnostics économiques*. Journées de Méthodologie Statistique - 15 et 16 décembre 1993, Publication Insee (1995).
- [2] **Roos, J.L.** (1995). *Comment assurer la crédibilité des discours économiques*. IA 95, Génie Linguistique, Montpellier, Juin 1995.
- [3] **Roos, J.L.** (1992). *Intelligence Artificielle en Langage C*. Editions Eyrolles, (1992).

Darwin et **ALIEN** ont bénéficié de l'aide et du financement de la Commission Européenne à travers les projets **DOSES** (contrat 2661006) et **SUPCOM** (contrat 56610005).

Ils ont toujours été par ailleurs soutenus, et ce depuis plusieurs années, par la Division Enquêtes de Conjoncture de l'INSEE.

La version en langue allemande a été faite par l'IFO (INSTITUT FÜR WIRTSCHAFTSFORSCHUNG, Munich).

ENGLISH SUMMARY

TEXT GENERATION FROM STATISTICAL DATA

Dr. JEAN LOUIS ROOS*

Institut National de la Statistique et des Etudes Economiques

ALIEN is a knowledge based system that modelizes the economist's behaviour. It is used at the French Statistical Institute and at Eurostat. It can be seen as an assistant that help economists to write economic diagnosis in natural language as well as a human economist could do.

Alien is very useful for analysis in very large data base and/or for regular diagnosis. It study the shape of all series in an application, the instantaneous relationship and determine what is important and what is not.

ALIEN is used under a user friendly software (DARWIN) developed for Windows 95. DARWIN build text generation with ALIEN, but also construct graphs and table, help the user to analyse the data base and the semantic of an application. Last, DARWIN build a nice layout of the mixed comment, graph and table with WINWORD or an equivalent software.

Keywords: Weighted Principal Components Analysis, Multiple Correspondence Analysis, Mixed Tables, Indicator Variables, Weighting

*Dr. Jean Louis Roos. Institut National de la Statistique et des Etudes Economiques. Avenue Albert Einstein BP 26000, 13791 Aix En Provence Cedex 3. Tel 42 16 29 59. Fax 42 16 29 00.
roos@cil3a.insee.atlas.fr

–Received october 1996.

–Accepted february 1997.

More and more frequently statistical databases are very large. If statisticians have to produce regular commentaries on the data contained in such databases, their work may be particularly protracted and tedious as they will have to examine substantial volumes of figures. Moreover, the data have a semantic aspect which never (or only very rarely) appears in the databases, so that the commentary produced by statisticians will to a great extent depend on their personal knowledge. For this reason, any change in the person responsible for producing regular statistical studies will sometimes dramatically affect the quality of the analysis texts. Strangely enough, the software manufacturers have put a lot of work into developing powerful consultation tools, but have done nothing about interpreting the contents of such databases.

As we see it, statisticians urgently need a tool to meet three requirements:

◆ **Statisticians need an assistant**

Statisticians should regularly be able to ask an assistant for a correctly drafted text analysing the information in a database. It should be possible to accompany texts of this kind with tables and graphs and it should be easy to modify them using a word-processing system. The assistant should be efficient enough not to overlook anything important and, obviously, the analysis must be «stable» over time and totally independent of external subjective factors.

◆ **Statisticians need to be able to deal with «Bulk Commentaires»**

Often it is not possible to use all the information available for producing a commentary, simply because the bulk of information to be processed is too great —as in the case of sectoral analyses or analysis by local-authority area or country etc. Sometimes for reasons of time alone, a human being is unable to produce dozens or hundreds of diagnoses. The difficulties will be even greater if the documents need to be accompanied by graphs and formatted.

◆ **Statisticians need guidelines**

The term «guidelines» means the definition of recommended good practice in various fields - in other words, what one should do, say or write in a given situation. Guidelines of this kind have been developed particularly in the field of medicine, by formalising diagnoses or treatments on which consensus had been reached. By a process of simplification, a guideline enables a situation and a treatment to be described in a formalised fashion which is accepted by a majority of practitioners. Admittedly very relative, this consensus nevertheless provides an efficient and reasoned response to a given disorder.

It is equally important for statisticians to have guidelines of this kind for each of the databases with which they work, as otherwise there is a great risk of interpretation errors.

As part of its specific work on formalising short-term analysis, INSEE has started to develop a programme to meet these requirements. EUROSTAT very quickly joined in with project and has provided a large part of the finance. The software is now in use at both INSEE and EUROSTAT. The text generation programme is «**Alien**» (Assistant Logique pour l'Interprétation Experte de données Numerique), and its operating environment under Windows is «**Darwin**» (Diagnostic, Analyse et Rapport sous Windows). The whole package is, as it were, a computerised assistant for statisticians.

1. TEXT GENERATION: ALIEN

Alien is a software which formalises—in fact models—the mechanisms leading a statistician to produce a given diagnosis from a set of quantitative data. This formalisation is based on the following observation: the majority of simple analysis texts have a stereotyped aspect—regardless of the language in which they are written. This means that very repetitive analysis schemes exist, which can be modelled.

1.1. Theoretical principles

The principles underlying **Alien** are fairly simple: statisticians no longer manipulate series but rather indicators, i.e. entities with a semantic aspect. Secondly, statisticians use stereotyped forms of discourse.

1.1.1. *The indicators*

Alien is based on the idea that economists do not manipulate statistical series, but rather a set of information about the series, including a knowledge of the vocabulary associated with each series.

1.1.2. *Stereotyped discourse*

Languages are so rich that it would be impossible, without using enormous resources, to define a method for writing all the possible discourses for a given diagnosis. However, experience shows that the types of discourse used for statistical commentaries are very stereotyped, because of laziness or at least out of habit.

1.1.3. *Dictionaries*

Dictionaries perform several functions in the drafting of a text. Obviously, their prime purpose is to provide the vocabulary, but they also show most of the most important grammatical rules, such as plural forms or conjugations. The same is true

of the **Alien** dictionaries. Just like a human being, **Alien** will also use dictionaries of synonyms or dictionaries of the dialects we have just mentioned.

1.2. The formal system: a double tree-structure

Writing is a particular human activity which involves frequent backtracking, and this is also true of statistical analysis. It is therefore difficult to produce a totally deterministic software. In practice, **Alien** is based on a double information tree-structure, text generation being handled by a series of small processing operations organised on the terminal nodes and using the dictionaries.

1.2.1. Indicator tree-structure

A main indicator defines the subject. The main indicators each have their own sub-sets of domain indicators which will constitute paragraphs of text. Each domain is made up of quantitative or qualitative indicators, some of which have tokens (for example, production may be analysed by branch or country, etc.). These indicators form the discourses on each series for which a commentary is required.

1.2.2. Content indicator tree-structure

All the indicators are structured in the same way. They have properties and each one is assigned a value (sometimes by default). An indicator has around 200 properties and it is possible to add original properties for each application.

1.2.3. The dictionaries

The dictionaries form a basic-tree structure based on «entries». They are specific entities which group the vocabulary common to the indicators and most of the associated grammatical rules and conjugations.

2. THE USER INTERFACE: DARWIN

Alien is still complicated to use, and therefore a user-friendly interface has been developed. This interface, known as **Darwin**, is a working environment under Windows. It provides easy use of the text generated via editors, a graphic tool for displaying series and final formatting of the text, graphs and tables using a word processing system, such as WINWORD.

Darwin and its associated programs are written in C with the 32-bit Windows Api. They are compatible with Windows 95 and Windows 3.1.

Apart from providing a text-generation environment, **Darwin** makes it possible to manipulate, generate and modify the elements in an application and to choose the language in which the text is to be produced (this section is still under development).

2.1. Capacities of the software

2.1.1. *The generation menu*

This uses the **Cia** artificial intelligence software. Generation may be displayed on the screen, or not as required.

2.1.2. *The result menu*

When generation has been completed, the Result menu enables the results to be displayed through, for example, text editors, and hence modifications are possible. It is also possible to display the associated graphs and tables. If the item generated is very large (sometimes a hundred or so files are constructed at the rate of one per paragraph), **Darwin** enables users to choose the paragraph they wish to view or in which they want to examine the graphs.

2.1.3. *The structural editors*

These specific editors permit structured viewing of the knowledge of an application. For example, a database is viewed in terms of the name of the series used by **Alien**; each series can be explored by clicking on the vertical scroll bar.

2.1.4. *Formatting*

A layout menu permits automatic formatting of **Alien** output using a different processing system, possibly with user modifications. A basic page structure and a list of commands must be defined in a given file. The graphs and tables can be inserted.

2.1.5. *Multilingual use*

Flags make it possible to display the dictionaries and choose generation in a given language. Only the English version is in fact being developed. Italian and Spanish are planned and a partial version in German has been developed by **Ifo**.

3. USES

INSEE and EUROSTAT are using this program on an experimental basis.

Biometria

Consell Directiu de la «Región Española de la International Biometric Society» (1996)

<i>President:</i>	C.M. Cuadras
<i>Vice-president:</i>	E. Carbonell
<i>Secretari:</i>	F. López-Santovería
<i>Treasurer:</i>	M. Dolores Sánchez
<i>Vocals:</i>	J.L. Chorro
	R. Estarelles
	G. Gómez (corresponsal del Biometric Bulletin)
	M. Ríos

Adreça electrònica: <http://www.iata.csic.es/IBSREsp/>

SOME UNRESOLVED ISSUES IN NON-LINEAR POPULATION DYNAMICS

JOE N. PERRY*

Rothamsted Experimental Station (U.K.)

The Lyapunov exponent is a statistic that measures the sensitive dependence of the dynamic behaviour of a system on its initial conditions. Estimates of Lyapunov exponents are often used to characterize the qualitative population dynamics of insect time series. The methodology for estimation of the exponent for an observed, noisy, short ecological time series is still under development. Some progress has been made recently in providing measures of error for these exponents. Studies that do not account for noise when reconstructing the dynamics of series must remain questionable.

Keywords: Chaos, population dynamics, non-linear dynamics, time series, Lyapunov exponent, randomization, aphids, moths.

*Joe N. Perry. Department of Entomology & Nematology. Rothamsted Experimental Station. Harpenden, Herts. AL5 2JQ. United Kingdom.

–Article rebut el febrer de 1997.

–Acceptat el juny de 1997.

Ecological time series concerned with the population fluctuations of organisms are usually short, consisting typically of no more than 30 generations, and often fewer. Such series are often based on counts, and estimate the population density of an insect population, once per generation. The data are therefore not continuous, but discrete. The populations they represent are often distributed exceedingly patchily, and frequently include many zero values. By contrast with physical variables, such populations are highly dynamic, and have usually evolved to shift ceaselessly in space and time for ecological reasons. They are often characterized by isolated clusters, which may be acting as metapopulations with varying degrees of inter-cluster dispersal (Perry & González-Andújar, 1993).

Turchin & Taylor (1992) proposed that the dynamics of ecological time series be represented by a response surface, with a dependent variable of logarithmically-transformed growth rates and with several independent variables, each a function of a lagged variable, i.e. a count from a previous generation. Their proposals were examined by Perry, Woïwod & Hanski (1993), who objected on various statistical grounds, pointed out difficulties over robustness, and presented alternative models. Both of these papers made predictions of the qualitative nature of the underlying, endogenous dynamics, by fitting the surface and inspecting the time series reconstructed from the deterministic skeleton represented by the estimated surface parameters. The behaviour of this time series, predicted for hundreds of generations into the future, was used to characterize the dynamics underlying the observed data.

However, this approach was severely criticized by Ellner & Turchin (1995), who emphasized that the previous methods were flawed because they did not allow for a random component in the dynamics, and this could lead to non-chaotic series being misidentified as chaotic. Ellner & Turchin (1995) identified three sources of variation that might influence the sensitivity of the system to initial conditions, namely: endogenous dynamics, exogenous dynamics and measurement error, and pointed out that fluctuations cannot be categorized as stochastic or dynamic by methodology that assumes the absence of noise. As an alternative, they suggested that dynamics should be characterized directly through the Lyapunov exponent, and presented Jacobean methods for calculation of the Lyapunov exponent that allow for dynamic noise. While this new methodology cannot disentangle the relative contributions of measurement error, which is usually assumed to be small, from exogenous dynamics, it does identify the effects of the exogenous dynamics, which is usually the aim of the exercise. Unfortunately, it is not easy to provide standard errors or measures of variability for such estimates. Although positive Lyapunov exponents indicate chaotic dynamics, positive values close to zero might occur due to measurement error, and the series might have a true value of zero, indicating quasiperiodic dynamics or a negative value, indicating stability. Furthermore, it is impossible to rely on theory or asymptotic results to supply measures of variability.

Zhou *et al.* (1997) provided some measures by using a randomization technique pioneered by Pollard, Lakhani & Rothery (1987), from the density dependence literature. Essentially, the series counts or growth rates are permuted under certain null hypotheses, and a randomization distribution of Lyapunov estimates built up from these permutations against which the observed value may be assessed. The advantage is that the approach is non-parametric and, much in the spirit of bootstrapping, uses the observed data and avoids the need to make too many, possibly dubious, assumptions. The disadvantage is that the variability of the randomization distribution is of interest only if the null hypothesis is plausible. Zhou *et al.* (1997) considered two null hypotheses. One, that the population undergoes complete compensation in the next generation for any exogenous density fluctuations suffered during the current one, is of interest as an extreme baseline, but is hardly likely to hold for any species in practice (Hanski, Woiwod & Perry, 1993). The other, that the population is behaving in a purely density-independent fashion, is more likely, and indeed is one possibility proposed in a debate that been waged in ecology for forty years. However, it is now thought much less likely than used to be the case, for the majority of species (Woiwod & Hanski, 1992).

The methods were applied to 46 time series comprising six aphid species from five sites and four moth species from six sites. There were few positive Lyapunov exponents and none was sufficiently large to characterize its time series as chaotic. Zhou *et al.* (1997) found that there were differences in the mean and the skewness of the randomization distributions of the Lyapunov exponent from the two hypotheses. Attempts have been made to derive confidence limits for Lyapunov exponent estimates by using replicated laboratory populations or populations from different geographic locations, and this approach might well prove a sensible way forward. However, this study showed that estimates of Lyapunov exponents may be very different at different localities, especially for species with a three-lag model. Of course, it must be stressed that these results may apply only at the spatial and temporal scales studied. Lyapunov exponent estimates were slightly greater for aphids, which have relatively complex life histories, than for moths, although it is still unclear whether the differences between life history strategies in aphids and moths are related to their strength of density-dependence (Hanski & Woiwod, 1993). Non-linearity is necessary, although not sufficient, for producing chaotic dynamics. Statistical methods have been developed to test for the non-linearity of a time series, although Zhou *et al.*'s methods to estimate Lyapunov exponents may be used independently of these.

A recent study (Costantino *et al.*, 1997) suggested that their laboratory data of the population dynamics of the flour beetle *Tribolium castaneum* showed convincing evidence of transitions to chaos. However, their methodology was similar to those earlier studies that assessed the population dynamics of a time series by fitting some mechanistic or empirical model and then inspecting realizations from the deterministic skeleton of the fitted model. The reported estimates of the Lyapunov exponents in

Costantino *et al.*'s data must be shown to be robust to the presence of the noise that they themselves estimate in their variance-covariance matrix Σ , for a valid characterization of the *Tribolium* dynamics. The next step will be for Costantino *et al.* to report such estimates for the stochastic version of their model and to compare their data with the output from this, rather than that from the deterministic skeleton. Until then, their characterization of the *Tribolium* dynamics must be viewed with caution.

REFERENCES

- [1] **Costantino, R.F., Desharnais, R.A., Cushing, J.M. & Dennis, B.** (1997). «Chaotic dynamics in an insect population». *Science*, **275**, 389–391.
- [2] **Ellner, S. & Turchin, P.** (1995). «Chaos in a 'noisy' world: new methods and evidence from time series analysis». *American Naturalist*, **145**, 343–375.
- [3] **Hanski, I. & Woiwod, I.P.** (1993). «Mean-related stochasticity and population variability». *Oikos*, **67**, 29–39.
- [4] **Hanski, I., Woiwod, I.P. & Perry, J.N.** (1993). «Density dependence, population persistence, and largely futile arguments». *Oecologia*, **95**, 595–598.
- [5] **Perry, J.N. & González-Andújar, J.L.** (1993). «Dispersal in a metapopulation neighbourhood model of an annual plant with a seedbank». *Journal of Ecology*, **81**, 453–463.
- [6] **Perry, J.N., Woiwod, I.P. & Hanski, I.** (1993). «Using response-surface methodology to detect chaos in ecological time series». *Oikos*, **68**, 329–339.
- [7] **Pollard, E., Lakhani, K.L. & Rothery, P.** (1987). «The detection of density dependence from a series of annual censuses». *Ecology*, **68**, 2046–2055.
- [8] **Turchin, P. & Taylor, A.D.** (1992). «Complex dynamics in ecological time-series». *Ecology*, **73**, 289–305.
- [9] **Woiwod, I.P. & Hanski, I.** (1992). «Patterns of density dependence in moths and aphids». *Journal of Animal Ecology*, **61**, 619–629.
- [10] **Zhou, X., Perry, J.N., Woiwod, I.P., Harrington, R., Bale, J.S. & Clarck, S.J.** (1997). «Detecting chaotic dynamics of insect populations from long-term survey data». *Ecological Entomology*. (in press).

STATISTICAL ANALYSIS OF THE GENERALIZED ADDITIVE SEMIPARAMETRIC SURVIVAL MODEL WITH RANDOM COVARIATES

V. BAGDONAVIČIUS*

University of Vilnius

M. NIKULIN**

University of Bordeaux-2

Generalizations of the additive hazards model are considered. Estimates of the regression parameters and baseline function are proposed, when covariates are random. The asymptotic properties of estimators are considered.

Keywords: Generalized additive model, semiparametric models, survival analysis, regression parameters, baseline function, asymptotic analysis, additive hazards.

Clasificación AMS: 60E15, 62F12, 62J99, 62P10

* V. Bagdonavičius. University of Vilnius, Lithuania. Steklov Mathematical Institute, St.Petersburg, Russia.

**M. Nikulin. University of Bordeaux-2, France.

– Article rebut el juliol de 1996.

– Acceptat el març de 1997.

1. INTRODUCTION

Models describing the dependence of lifetimes distributions on explanatory variables have been considered. A number of such models was proposed by Andersen, Borgan, Gill and Keiding (1993), Cox and Oakes (1984), Dabrowska and Doksum (1988), Dreesbeke, Fichet and Tassi (1989), Kalbfleisch and Prentice (1980), Lin and Ying (1994),(1996), etc. Bagdonavičius and Nikulin (1994)-(1996) proposed a general approach, which gives the possibility of formulating a number of new models and showing where known models fit into the proposed new classes.

Suppose that a time to failure $T_{x(\cdot)}$ is a nonnegative random variable with the survival function $S_{x(\cdot)}(t) = P\{T_{x(\cdot)} > t\}$ which depends on a vector of *stresses*

$$x: [0, +\infty) \rightarrow B \subset \mathbb{R}^m.$$

The time to failure $T_{x(\cdot)}$ could be called the resource of this item. But the notion of the resource should not depend on $x(\cdot)$. So we will call *the uniform resource* to the random variable $R^U = 1 - S_{x(\cdot)}(T_{x(\cdot)})$. It takes values in the interval $[0, 1)$ and does not depend on $x(\cdot)$. Note that $T_{x(\cdot)} = t$ if and only if $R^U = 1 - S_{x(\cdot)}(t)$. So the number $1 - S_{x(\cdot)}(t) \in [0, 1)$ is called *the uniform resource used until the moment t under the stress $x(\cdot)$* . The concrete item which failed at the moment t under the stress $x(\cdot)$ used $1 - S_{x(\cdot)}(t)$ of the resource. Instead of the uniform resource one can define a resource with any probability distribution, so we can consider a whole class of resources. Really, suppose that G is some fixed survival function, strictly decreasing and continuous on the support $[a, b]$, $-\infty \leq a < b \leq \infty$, $G(a) = 1$, $G(b) = 0$. $H = G^{-1}$. The functional

$$f_{x(\cdot)}^G(t) = (H \circ S_{x(\cdot)})(t),$$

is called the *G-transfer functional*. The survival function of the random variable $R^G = f_{x(\cdot)}^G(T_{x(\cdot)})$ is G and does not depend on $x(\cdot)$. The random variable R^G is called the *G-resource* and the number $f_{x(\cdot)}^G(t)$ is called the *G-resource used till the moment t* . Denote by $\mathcal{G}$ the family of survival functions, continuous and decreasing on their supports. Consider the class of transfer functionals $\mathcal{M} = \{f^G, G \in \mathcal{G}\}$. Models will be formulated in dependence on properties of the transfer functionals. Note that some assumptions may be satisfied by one transfer functional, but not satisfied by another. This is the cause of considering the *whole class of resources*.

In the case of $G(t) = e^{-t}1_{[0, \infty)}(t)$ the transfer functional has the form $f_{x(\cdot)}^G(t) = -\ln S_{x(\cdot)}(t)$ and the rate of resource use is

$$\frac{\partial f_{x(\cdot)}^G(t)}{\partial t} = \alpha_{x(\cdot)}(t),$$

where $\alpha_{x(\cdot)}(t) = -S'_{x(\cdot)}(t)/S_{x(\cdot)}(t)$ is the hazard rate of $T_{x(\cdot)}$.

The additive hazards model (AHM) (see Andersen et al. (1993)) holds on E if there exist functions λ_0 and a such that for all $x(\cdot) \in E$

$$\alpha_{x(\cdot)}(t) = \alpha_0(t) + a[x(t)].$$

Now we generalize this model.

Definition. *The G -generalized additive model holds on E if there exists a function $a(\cdot)$ on E and a survival function S_0 such that, for all $x(\cdot) \in E$, the transfer functional $f^G \in \mathcal{M}$ satisfies the differential equation*

$$(1) \quad \frac{\partial f_{x(\cdot)}^G(t)}{\partial t} = \frac{\partial f_0^G(t)}{\partial t} + a(x(t))$$

with the initial conditions $f_0^G(0) = f_{x(\cdot)}^G(0) = 0$.

So the stress influences additively the rate of resource use. Equation (1) implies that

$$S_{x(\cdot)}(t) = G\{f_0^G(t) + \int_0^t a[x(\tau)]d\tau\}.$$

Let us consider some particular models different from AHM.

1. Taking $G(t) = \exp\{-\exp\{t\}\}$, for $t \in R^1$, we obtain

$$\frac{\partial f_{x(\cdot)}^G(t)}{\partial t} = \frac{\alpha_{x(\cdot)}(t)}{A_{x(\cdot)}(t)}, \quad \frac{\partial f_0^G(t)}{\partial t} = \frac{\alpha_0(t)}{A_0(t)},$$

where

$$A_{x(\cdot)}(t) = \int_0^t \alpha_{x(\cdot)}(\tau)d\tau, \quad A_0(t) = \int_0^t \alpha_0(\tau)d\tau$$

are the cumulated hazards rates. So we have the model:

$$\frac{\alpha_{x(\cdot)}(t)}{A_{x(\cdot)}(t)} = \frac{\alpha_0(t)}{A_0(t)} + a(x(t)).$$

2. Taking $G(t) = 1/(1+t)$, for $t \geq 0$, we obtain

$$\frac{\alpha_{x(\cdot)}(t)}{S_{x(\cdot)}(t)} = \frac{\alpha_0(t)}{S_0(t)} + a(x(t)).$$

3. Taking $G(t) = 1/(1+e^t)$, for $t \in R^1$, we obtain

$$\frac{\alpha_{x(\cdot)}(t)}{1 - S_{x(\cdot)}(t)} = \frac{\alpha_0(t)}{1 - S_0(t)} + a(x(t)).$$

4. Take $G(t) = 1 - \Phi(\ln t)$, $t \geq 0$, where $\Phi(\cdot)$ is the normal $N(0,1)$ cumulative distribution function. In terms of survival functions we obtain the model :

$$\Phi^{-1}(1 - S_{x(\cdot)}(t)) = \ln \left\{ \int_0^t a[x(\tau)] d\tau + \exp[\Phi^{-1}(1 - S_0(t))] \right\}.$$

5. Taking $G = S_0$, we obtain

$$S_{x(\cdot)}(t) = S_0 \left\{ \int_0^t \sigma[x(\tau)] d\tau \right\},$$

where $\sigma[x(t)] = 1 + a[x(t)]$. It is the *accelerated life model*.

Other distributions of the resource can be taken.

So a number of alternatives to the AHM is proposed. For example, if at the beginning of life data follow well the AHM but later it is not so, the model of example 2 could be choosed. On the contrary, if at the beginning of life data do not follow the AHM but at the end of life this model suits well, the model of example 3 could be choosed.

2. ESTIMATION

2.1. Notations

Consider the model (1) with some specified G and an unknown baseline survival function S_0 .

The function $a[x(\cdot)]$ is parametrized as follows:

$$a[x(t)] = \gamma^T x(t),$$

where $\gamma = (\gamma_1, \dots, \gamma_m)^T$ is the vector of unknown regression parameters. So the following model is considered: for all $x(\cdot) \in E$

$$(2) \quad S_{x(\cdot)}(t) = G\{H(S_0(t)) + \int_0^t \gamma^T x(\tau) d\tau\}.$$

Suppose that n individuals are observed and assume that the vector of covariates for the i th individual is a random process $X_i(\cdot) = (X_{i1}(\cdot), \dots, X_{im}(\cdot))^T$.

Denote by $N_i(t)$ the univariate counting processes. This process counts the numbers of observed failures of each individual in the interval $[0, t]$, $t \geq 0$. Denote by

$Y_i(t)$ the numbers of individual “ at risk” (non-censored and non-failed) just prior to t for each i .

Suppose that $\{\mathcal{F}_t, t \geq 0\}$ is the filtration generated by

$$\{N_i(s), Y_i(s), X_i(s), i = 1, \dots, n; 0 \leq s \leq t\},$$

X_i are predictable and the intensities λ_i , given by

$$\lambda_i(t) = \lim_{\varepsilon \downarrow 0} \frac{1}{\varepsilon} P\{N_i(t + \varepsilon) - N_i(t) \geq 1 | \mathcal{F}_{t-}\},$$

exist. In this case the compensators of the counting processes $N_i(t)$ are

$$\Lambda_i(t) = \int_0^t \lambda_i(s) ds.$$

Assume that the random intensities λ_i satisfy the *multiplicative intensity model*

$$(3) \quad \lambda_i(t) = \alpha_{X_{i(\cdot)}}(t) Y_i(t).$$

Denote

$$N(t) = \sum_{i=1}^n N_i(t), \quad Y(t) = \sum_{i=1}^n Y_i(t), \quad \Lambda(t) = \sum_{i=1}^n \Lambda_i(t), \quad M_i(t) = N_i(t) - \Lambda_i(t),$$

$$M(t) = \sum_{i=1}^n M_i(t), \quad J(s) = I(Y(s) > 0), \quad \alpha = -\frac{G'}{G}, \quad \Psi = \alpha \circ H.$$

The hazard rates $\alpha_{X_{i(\cdot)}}(t)$ have the form

$$(4) \quad \alpha_i(t) = \alpha_{X_{i(\cdot)}}(t) = \Psi(S_{X_{i(\cdot)}}(t)) \{(H \circ S_0)'(t) + \gamma^T X_i(t)\}.$$

2.2. Estimating equation and estimators $\hat{H}_0, \hat{\gamma}$ and $\hat{S}_{X(\cdot)}$

From the Doob-Meyer decomposition $N = M + \Lambda$, equalities (3) and (4) imply

$$dN(t) = dM(t) + \sum_{i=1}^n \Psi(S_{X_{i(\cdot)}}(t)) Y_i(t) \{dH(S_0(t)) + \gamma^T X_i(t) dt\}$$

and

$$\int_0^t \frac{J(u)(dN(u) - S_*^{(0)}(\gamma, u) du)}{S^{(0)}(\gamma, u)} = \int_0^t J(u) dH(S_0(u)) + \int_0^t \frac{J(u) dM(u)}{S^{(0)}(\gamma, u)},$$

where

$$S^{(0)}(\gamma, u) = \sum_{i=1}^n Y_i(u) \psi(S_{X_i(\cdot)}(u)),$$

$$S_*^{(0)}(\gamma, u) = \sum_{i=1}^n Y_i(u) \psi(S_{X_i(\cdot)}(u)) \gamma^T X_i(u).$$

Under some mild assumptions M is a martingale, therefore

$$(5) \quad E \int_0^t \frac{J(u)(dN(u) - S_*^{(0)}(\gamma, u)du)}{S^{(0)}(\gamma, u)} = E \int_0^t J(u) dH(S_0(u)).$$

If $Y(t) > 0$, then

$$(6) \quad \int_0^t J(u) dH(S_0(u)) = H(S_0(t)).$$

Equalities (5) and (6) imply that a reasonable estimator $\hat{H}_0(t, \gamma)$ for $H_0(t) = H(S_0(t))$ (still depending on γ) is determined by the equation

$$\hat{H}_0(t, \gamma) = \int_0^t J(u) \frac{dN(u) - \tilde{S}_*^{(0)}(\gamma, u)du}{\tilde{S}^{(0)}(\gamma, u)},$$

where

$$\tilde{S}^{(0)}(\gamma, u) = \sum_{i=1}^n Y_i(u) \alpha(\hat{H}_i(u, \gamma)), \quad \tilde{S}_*^{(0)}(\gamma, u) = \sum_{i=1}^n Y_i(u) \alpha(\hat{H}_i(u, \gamma)) \gamma^T X_i(u),$$

$$H_i(u, \gamma) = H_0(u) + \int_0^u \gamma^T X_i(\tau) d\tau, \quad \hat{H}_i(u, \gamma) = \hat{H}_0(u-, \gamma) + \int_0^u \gamma^T X_i(\tau) d\tau.$$

Denote $\tau^* = \sup\{t : Y(t) > 0\}$. Suppose that at the non-random moment $\tau \in]0, \infty]$ all the individuals are censored, i.e., $\tau^* \leq \tau$. We propose to estimate the parameter γ by solving the estimating equations

$$U(\gamma, \tau) = 0,$$

where the estimating function is given by the formula below

$$U(\gamma, t) = \sum_{i=1}^n \int_0^t J(u) X_i(u) \{dN_i(u) - Y_i(u) \alpha(\hat{H}_i(u, \gamma)) d\hat{H}_i(u, \gamma)\}.$$

Those equations generalize the estimating equations of Lin and Ying (1994) for the additive hazards model (taking $\alpha(p) \equiv 1$).

If we denote by $\hat{\gamma}$ the estimator of γ , i.e. $U(\hat{\gamma}, \tau) = 0$, then the estimator of the function $H_0(t)$ is

$$(7) \quad \tilde{H}_0(t) = \hat{H}_0(t, \hat{\gamma})$$

and the estimator of the survival function $S_{x(\cdot)}(t)$ under any covariate $x(\cdot)$ is

$$(8) \quad \hat{S}_{x(\cdot)}(t) = G \left\{ \tilde{H}_0(t) + \int_0^t \hat{\gamma}^T x(u) du \right\}.$$

Let

$$\begin{aligned} \tilde{S}^{(1)}(\gamma, u) &= \sum_{i=1}^n X_i(u) Y_i(u) \alpha(\hat{H}_i(u, \gamma)), \\ \tilde{S}_*^{(1)}(\gamma, u) &= \sum_{i=1}^n X_i(u) Y_i(u) \alpha(\hat{H}_i(u, \gamma)) \gamma^T X_i(u), \\ \tilde{E}(\gamma, u) &= \frac{\tilde{S}^{(1)}(\gamma, u)}{\tilde{S}^{(0)}(\gamma, u)}. \end{aligned}$$

Then the estimating function $U(\gamma, t)$ can be written

$$(9) \quad U(\gamma, t) = \sum_{i=1}^n \int_0^t J(u) \{X_i(u) - \tilde{E}(\gamma, u)\} \{dN_i(u) - Y_i(u) \alpha(\hat{H}_i(u, \gamma)) \gamma^T X_i(u) du\}.$$

3. ASYMPTOTIC PROPERTIES OF THE ESTIMATORS

3.1. Asymptotic properties of the estimating function

Denote by $\|A\| = \sup_{i,j} |a_{ij}|$ the norm of any matrix A , by γ_0 the true value of the parameter γ ,

$$\begin{aligned} S^{(1)}(\gamma, u) &= \sum_{i=1}^n X_i(u) Y_i(u) \alpha(H_i(u, \gamma)), \\ S_*^{(1)}(\gamma, u) &= \sum_{i=1}^n X_i(u) Y_i(u) \alpha(H_i(u, \gamma)) \gamma^T X_i(u), \\ E(\gamma, u) &= \frac{S^{(1)}(\gamma, u)}{S^{(0)}(\gamma, u)}, \\ S^{(2)}(\gamma, u) &= \sum_{i=1}^n X_i(u) X_i^T(u) Y_i(u) \alpha(H_i(u)), \end{aligned}$$

$$\begin{aligned}\tilde{S}^{(2)}(\gamma, u) &= \sum_{i=1}^n X_i(u) X_i^T(u) Y_i(u) \alpha(\hat{H}_i(u, \gamma)), \\ S^{(3)}(\gamma, u) &= \sum_{i=1}^n Y_i(u) \alpha'(H_i(u, \gamma)), \\ S_*^{(3)}(\gamma, u) &= \sum_{i=1}^n Y_i(u) \alpha'(H_i(u, \gamma)) \gamma^T X_i(u).\end{aligned}$$

Assumptions A

1. Negligibility conditions.

There exist a neighborhood Γ of γ_0 and a scalar, vector or matrix functions $s^{(0)}$, $s_*^{(0)}$, $s^{(1)}$, $s_*^{(1)}$, $s^{(2)}$, $s_*^{(2)}$, $s^{(3)}$, $s_*^{(3)}$, such that for $m = 0, 1, 2$:

(a)

$$\begin{aligned}\sup_{\gamma \in \Gamma, t \in [0, \tau]} \left\| \frac{1}{n} S^{(m)}(\gamma, t) - s^{(m)}(\gamma, t) \right\| &\xrightarrow{P} 0, \\ \sup_{\gamma \in \Gamma, t \in [0, \tau]} \left\| \frac{1}{n} S_*^{(m)}(\gamma, t) - s_*^{(m)}(\gamma, t) \right\| &\xrightarrow{P} 0,\end{aligned}$$

(b) $s^{(0)}(\gamma_0, \cdot)$ is bounded away from zero on $[0, \tau]$,

(c) $s^{(m)}(\cdot, \cdot)$, $s_*^{(m)}(\cdot, \cdot)$ are continuous functions of $\gamma \in \Gamma$ uniformly in $t \in [0, \tau]$ and bounded on $\Gamma \times [0, \tau]$.

2.

$$H_0(\tau) = H(S_0(\tau)) < \infty,$$

3. α is a positive continuously differentiable on $]0, \infty[$ function,

4.

$$P\left\{ \sup_{t \in [0, \tau]} |X_{ij}(t)| < \infty \right\} = 1 \quad \text{for all } i, j.$$

5.

$$\sigma^2(t) = \int_0^t \frac{s^{(0)}(\gamma, u) dH_0(u) + s_*^{(0)}(\gamma, u) du}{s^{(0)2}(\gamma, u)} < +\infty.$$

Lemma. Suppose that assumptions A hold. Then

$$\sqrt{n} [\hat{H}_0(t, \gamma) - H_0(t)] \xrightarrow{D} h(\gamma, t) \int_0^t \frac{dV(u)}{h(\gamma, u)} \quad \text{as } n \rightarrow \infty,$$

where V is the Gaussian martingale with

$$EV(t) = 0 \quad \text{and} \quad \text{Cov}(V(t), V(s)) = \sigma^2(t \wedge s),$$

$$h(\gamma, t) = \exp \left\{ - \int_0^t \frac{1}{s^{(0)}(\gamma, u)} \left(s^{(3)}(\gamma, u) dH_0(u) + s_*^{(3)}(\gamma, u) du \right) \right\}.$$

Proof: Consider the difference:

$$\begin{aligned} \sqrt{n} [\hat{H}_0(t, \gamma) - H_0(t)] &= \sqrt{n} \left\{ \int_0^t J(u) \frac{dN(u) - \tilde{S}_*^{(0)}(\gamma, u) du}{\tilde{S}^{(0)}(\gamma, u)} - H_0(t) \right\} = \\ &= \sqrt{n} \left\{ \int_0^t J(u) \left[\frac{dM(u)}{\tilde{S}^{(0)}(\gamma, u)} + \frac{S^{(0)}(\gamma, u) - \tilde{S}^{(0)}(\gamma, u)}{\tilde{S}^{(0)}(\gamma, u)} dH_0(u) + \frac{S_*^{(0)}(\gamma, u) - \tilde{S}_*^{(0)}(\gamma, u)}{\tilde{S}^{(0)}(\gamma, u)} du \right] \right\}. \end{aligned}$$

Note that

$$\begin{aligned} \frac{1}{\sqrt{n}} [\tilde{S}^{(0)}(\gamma, u) - S^{(0)}(\gamma, u)] &= \\ \frac{1}{n} \sum_{i=1}^n Y_i(u) \alpha' \{H_i(u, \gamma)\} \sqrt{n} [\hat{H}_0(u-, \gamma) - H_0(u)] + o_p(1) &= \\ s^{(3)}(\gamma, u) \sqrt{n} [\hat{H}_0(u-, \gamma) - H_0(u)] + o_p(1). \end{aligned}$$

Similarly

$$\begin{aligned} \frac{1}{\sqrt{n}} [\tilde{S}_*^{(0)}(\gamma, u) - S_*^{(0)}(\gamma, u)] &= \\ s_*^{(3)}(\gamma, u) \sqrt{n} [\hat{H}_0(u-, \gamma) - H_0(u)] + o_p(1) \end{aligned}$$

and

$$\frac{n}{\tilde{S}^{(0)}(\gamma, u)} = \frac{1}{s^{(0)}(\gamma, u)} + o_p(1).$$

So

$$\begin{aligned} \sqrt{n} [\hat{H}_0(t, \gamma) - H_0(t)] &= \sqrt{n} \int_0^t \frac{J(u) dM(u)}{\tilde{S}^{(0)}(\gamma, u)} - \\ \int_0^t \sqrt{n} [\hat{H}_0(u-, \gamma) - H_0(u)] \frac{s^{(3)}(\gamma, u) dH_0(u) + s_*^{(3)}(\gamma, u) du}{s^{(0)}(\gamma, u)} + o_p(1). \end{aligned}$$

Note that the predictable variation

$$< \sqrt{n} \int_0^t \frac{J(u) dM(u)}{\tilde{S}^{(0)}(\gamma, u)} > = n \int_0^t J(u) \frac{S^{(0)}(\gamma, u) dH_0(u) + S_*^{(0)}(\gamma, u) du}{\tilde{S}^{(0)2}(\gamma, u)} \xrightarrow{P} \sigma^2(t)$$

and

$$\int_0^t J(u) \frac{n}{\tilde{S}^{(0)2}(\gamma, u)} I \left\{ \left| \frac{J(u)\sqrt{n}}{\tilde{S}^{(0)}(\gamma, u)} \right| > \varepsilon \right\} \left(S^{(0)}(\gamma, u) dH_0(u) + S_*^{(0)}(\gamma, u) du \right) \xrightarrow{P} 0.$$

By Rebolledo's theorem (see Andersen et al (1993)) we obtain the convergence :

$$\sqrt{n} \int_0^t \frac{J(u) dM(u)}{\tilde{S}^{(0)}(\gamma, u)} \xrightarrow{D} V(t) \quad \text{as } n \rightarrow \infty.$$

Then

$$\sqrt{n} [\hat{H}_0(t, \gamma) - H_0(t)] \xrightarrow{D} h(\gamma, t) \int_0^t \frac{dV(u)}{h(\gamma, u)} \quad \text{as } n \rightarrow \infty.$$

The proof is completed. ■

Corollary. Under assumptions A

$$(10) \quad \sqrt{n} [\hat{H}_0(t, \gamma) - H_0(t)] = h(\gamma, t) \sqrt{n} \int_0^t \frac{J(u) dM(u)}{h(\gamma, u) S^{(0)}(\gamma, u)} + o_p(1).$$

Consider the asymptotical distribution of the score function $U(\gamma_0, \tau)$. The Doob-Meyer decomposition and equality (9) imply

$$\frac{1}{\sqrt{n}} U(\gamma_0, t) = \frac{1}{\sqrt{n}} U^*(\gamma_0, t) + \frac{1}{\sqrt{n}} \Delta(\gamma_0, t) + o_p(1),$$

where

$$\frac{1}{\sqrt{n}} U^*(\gamma_0, t) = \frac{1}{\sqrt{n}} \sum_{i=1}^n \int_0^t J(u) \{X_i(u) - E(\gamma_0, u) + \frac{1}{h(\gamma_0, u) S^{(0)}(\gamma_0, u)} \times$$

$$\sum_{j=1}^n \left[\int_u^t J(s) [X_i(s) - E(\gamma_0, s)] \alpha'(H_j(s, \gamma)) h(\gamma_0, s) Y_j(s) \times$$

$$(11) \quad (dH_0(s) + \gamma_0^T X_j(s) ds) \} \} dM_i(u)$$

and

$$\frac{1}{\sqrt{n}} \Delta(\gamma_0, t) = - \int_0^t J(s) [\tilde{E}(\gamma_0, s) - E(\gamma_0, s)] dM(s) +$$

$$\int_0^t J(s)[\tilde{E}(\gamma_0, s) - E(\gamma_0, s)][\tilde{S}^{(0)}(\gamma_0, s) - S^{(0)}(\gamma_0, s)]dH_0(s) +$$

$$\int_0^t J(s)[\tilde{E}(\gamma_0, s) - E(\gamma_0, s)] \{ \alpha(\hat{H}_i(s, \gamma_0)) - \alpha(H_i(s)) \} \gamma_0^T X_i(s) Y_i(s) ds.$$

Thus the problem is to prove that $\frac{1}{\sqrt{n}}U^*(\gamma_0, t)$ converges to a Gaussian martingale and $\frac{1}{\sqrt{n}}\Delta(\gamma_0, t)$ converges to 0 in probability.

The integral (11) is a local martingale as an integral of a locally bounded predictable process with respect to a martingale, so the central limit theorem for martingales (Rebolledo's theorem, see Andersen et al. (1993), p.83) can be applied. The limit process of $\frac{1}{\sqrt{n}}U^*(\gamma_0, t)$ would be a Gaussian martingale if the predictable covariation process $\langle n^{-1/2}U^* \rangle (\gamma_0, t)$ would converge to some non-random matrix and the generalized Lindenberg condition would be satisfied. Note that the predictable covariation process

$$\langle n^{-1/2}U^* \rangle (\gamma_0, t) = \frac{1}{n} \sum_i \int_0^t J(u) \left\{ X_i(u) - E(\gamma_0, u) + \frac{1}{h(\gamma_0, u)S^{(0)}(\gamma_0, u)} \times \right.$$

$$\left. \sum_j \left[\int_u^t [X_j(s) - E(\gamma_0, s)] \alpha'(H_j(s, \gamma)) h(\gamma_0, s) Y_j(s) (dH_0(s) + \gamma_0^T X_j(s) ds) \right] \right\}^{\otimes 2} \times$$

$$(12) \quad Y_i(u) \alpha(H_i(u, \gamma)) [dH_0(u) + \gamma_0^T X_i(u) du].$$

Hence to obtain the limit distribution of $\frac{1}{\sqrt{n}}U^*(\gamma_0, t)$ the assumption that the process (12) converges in probability to some non-degenerate non-random matrix must be made. Note that this process has the form $\frac{1}{n} \sum_{i=1}^n \xi_i(t)$ and if the covariates $X_i(\cdot)$ are not behaving in some strange way (increasing very quickly and so on), this convergence is natural.

Therefore we formulate

Stability condition:

$$\langle n^{-1/2}U^* \rangle (\gamma_0, t) \xrightarrow{P} \Sigma(\gamma_0, t),$$

where $\Sigma(\gamma_0, t)$ is non-random and non-degenerated matrix.

Theorem 1. Suppose that assumptions A and the stability condition hold. Then, as $n \rightarrow \infty$

$$n^{-1/2}U(\gamma_0, \tau) \xrightarrow{D} N(0, \Sigma(\gamma_0, \tau)).$$

Proof: It is sufficient to prove that the Lindenberg condition is satisfied. Denote

$$H_{il}(u) = J(u) \{X_{il}(u) - E_l(\gamma_0, t) + \frac{1}{h(\gamma_0, u)S^{(0)}(\gamma_0, u)} \sum_{j=1}^n \int_u^t [X_{il}(s) - E_l(\gamma_0, s)] \alpha'(H_j(s, \gamma)) h(\gamma_0, s) Y_j(s) (dH(S_0(s)) + \gamma_0^T X_j(s) ds)\},$$

where $E_l(\gamma_0, s)$ is a l th component of the vector $E(\gamma_0, s)$. The Lindenberg condition in our case is

$$\frac{1}{n} \sum_i \int_0^t H_{il}^2(u) I\left\{\frac{1}{\sqrt{n}} |H_{il}(u)| > \varepsilon\right\} Y_i(u) \alpha_i(s) ds \xrightarrow{P} 0.$$

But it is obvious as by assumptions of the theorem $H_{ij}(u)$ are bounded on $[0, \tau]$ and

$$\int_0^\tau \alpha_i(s) ds < \infty.$$

By Rebolledo's theorem

$$n^{-1/2} U^*(\gamma_0, \tau) \xrightarrow{D} N(0, \Sigma(\gamma_0, \tau)).$$

The term $n^{-1/2} \Delta$ converges in probability to zero. It can be seen similarly as in Bagdonavičius & Nikulin (1996) by applying Lengart's inequality.

The proof of the theorem 1 is completed. ■

3.2. Asymptotic properties of the estimator $\hat{\gamma}$.

Let

$$e(\gamma, u) = \frac{s^{(1)}(\gamma, u)}{s^{(0)}(\gamma, u)},$$

$$S_*^{(4)}(\gamma, u) = \sum_{i=1}^n \{X_i(u) - E(\gamma, u)\} Y_i(u) \alpha'(H_i(u, \gamma)) \int_0^u X_i^T(\tau) d\tau \gamma^T X_i(u).$$

Assumptions B

1. Negligibility condition:

$$\sup_{\gamma \in \Gamma, t \in [0, T]} \left\| \frac{1}{n} S_*^{(4)}(\gamma, u) - s_*^{(4)}(\gamma, u) \right\| \xrightarrow{P} 0.$$

2. The matrix

$$\Sigma_1(\gamma_0, \tau) = \int_0^\tau \left\{ \frac{\partial}{\partial \gamma} e(\gamma_0, u) s^{(0)}(\gamma_0, u) dH_0(u) + \left[s_*^{(4)}(\gamma_0, u) + s^{(2)}(\gamma_0, u) - e(\gamma_0, u) s^{(1)T}(\gamma_0, u) \right] du \right\}$$

is non degenerated.

Theorem 2. Suppose assumptions B and those of the Theorem 1 hold. Then there exists a neighborhood of γ_0 within which, with probability tending to 1 as $n \rightarrow \infty$, the root $\hat{\gamma}$ of $U(\gamma, \tau) = 0$ is uniquely defined and

$$(13) \quad n^{1/2}(\hat{\gamma} - \gamma_0) \xrightarrow{D} N(0, \Sigma_2(\gamma_0, \tau)),$$

where

$$\Sigma_2(\gamma_0, t) = \Sigma_1^{-1}(\gamma_0, t) \Sigma(\gamma_0, t) \Sigma_1^{-1}(\gamma_0, t).$$

Proof: Using the Taylor expansion of $n^{-1/2}U(\gamma, \tau)$ around γ_0 in $\gamma = \hat{\gamma}$, we obtain

$$(14) \quad n^{1/2}(\hat{\gamma} - \gamma_0) = \left(-\frac{1}{n} \frac{\partial U(\gamma^*, \tau)}{\partial \gamma} \right)^{-1} n^{-1/2}U(\gamma_0, \tau),$$

where γ^* is on the line segment between $\hat{\gamma}$ and γ_0 . Theorem 1 implies, that the convergence (13) is proved if we prove that

$$-\frac{1}{n} \frac{\partial U(\gamma_0, \tau)}{\partial \gamma} \xrightarrow{P} \Sigma_1(\gamma_0, \tau) \quad \text{and} \quad \hat{\gamma} \xrightarrow{P} \gamma_0.$$

We have

$$\begin{aligned} & -\frac{1}{n} \frac{\partial U(\gamma, \tau)}{\partial \gamma} = \\ & \frac{1}{n} \sum_{i=1}^n \int_0^\tau J(u) \frac{\partial}{\partial \gamma} \tilde{E}(\gamma, u) \{ dM_i(u) + Y_i(u) \alpha(H_i(u, \gamma)) dH_0(u) + \\ & Y_i(u) \gamma^T X_i(u) [\alpha(H_i(u, \gamma)) - \alpha(\hat{H}_i(u, \gamma))] du \} + \frac{1}{n} \sum_{i=1}^n \int_0^\tau J(u) \{ X_i(u) - \tilde{E}(\gamma, u) \} \times \\ & Y_i(u) \left[\frac{\partial}{\partial \gamma} \alpha(\hat{H}_i(u, \gamma)) \gamma^T X_i(u) + \alpha(\hat{H}_i(u, \gamma)) X_i^T(u) \right] du. \end{aligned}$$

Note that

$$\frac{1}{n} \sum_{i=1}^n \int_0^\tau J(u) \frac{\partial}{\partial \gamma} \tilde{E}(\gamma, u) dM_i(u) \xrightarrow{P} 0,$$

$$\begin{aligned}
& \frac{1}{n} \sum_{i=1}^n \int_0^\tau J(u) \frac{\partial}{\partial \gamma} \tilde{E}(\gamma, u) Y_i(u) \alpha(H_i(u, \gamma)) dH_0(u) \xrightarrow{P} \int_0^\tau \frac{\partial}{\partial \gamma} e(\gamma, u) s^{(0)}(u, \gamma) dH_0(u), \\
& \frac{1}{n} \sum_{i=1}^n \int_0^\tau J(u) \frac{\partial}{\partial \gamma} \tilde{E}(\gamma, u) Y_i(u) \gamma^T X_i(u) [\alpha(\hat{H}_i(u, \gamma)) - \alpha(H_i(u, \gamma))] du \xrightarrow{P} 0, \\
& \frac{1}{n} \sum_{i=1}^n \int_0^\tau J(u) \{X_i(u) - \tilde{E}(\gamma, u)\} Y_i(u) \frac{\partial}{\partial \gamma} \alpha(\hat{H}_i(u, \gamma)) \gamma^T X_i(u) du \xrightarrow{P} \int_0^\tau s_*^{(4)}(\gamma, u) du, \\
& \frac{1}{n} \sum_{i=1}^n \int_0^\tau J(u) \{X_i(u) - \tilde{E}(\gamma, u)\} Y_i(u) \alpha[\hat{H}_i(u, \gamma)] \gamma^T X_i(u) du \xrightarrow{P} \\
& \int_0^\tau [s^{(2)}(\gamma, u) - e(\gamma, u) s^{(1)}(\gamma, u)] du.
\end{aligned}$$

Therefore

$$(15) \quad -\frac{1}{n} \frac{U(\gamma_0, \tau)}{\partial \gamma} \xrightarrow{P} \Sigma_1(\gamma_0, \tau).$$

In the rest of the proof we follow Lin and Ying (1996). By assumptions of the theorem, for any $\varepsilon > 0$, we can choose $\delta > 0$ such that for all n

$$\left\| \frac{1}{n} \frac{\partial}{\partial \gamma} U(\gamma, \tau) - \frac{1}{n} \frac{\partial}{\partial \gamma} U(\gamma_0, \tau) \right\| < \varepsilon$$

whenever $\|\gamma - \gamma_0\| < \delta$. On the other hand the convergence (15) implies that

$$(16) \quad P\left\{ \sup_{\|\gamma - \gamma_0\| \leq \delta} \left\| -\frac{1}{n} \frac{\partial}{\partial \gamma} U(\gamma, \tau) - \Sigma_1(\gamma_0, \tau) \right\| > 2\varepsilon \right\} \rightarrow 0$$

as $n \rightarrow \infty$.

Now we apply the theorem about the inverse mapping, which states that if $f(u), u \in R^p$ is continuously differentiable at u_0 and $\partial f(u_0)/\partial u$ is nonsingular, then there exist δ_0 and ε_0 that f is a one-to-one mapping on $B(u_0, \delta_0)$, the ball centered of the u_0 with radius δ_0 and $f(B(u_0, \delta_0)) \supset B(f(u_0), \varepsilon_0)$. As noted by Lin and Ying that result holds simultaneously for a family of such functions with common δ_0 and ε_0 as long as their derivatives at u_0 are sufficiently close. This result and (16) imply that there exist δ_1 and ε_1 such that $n^{-1}U(\cdot, \tau)$ is one-to-one mapping from the $B(\gamma_0; \delta_1)$ to $n^{-1}U(B(\gamma_0, \delta_1), \tau)$, which contains $B(n^{-1}U(\gamma_0, \tau), \varepsilon_1)$. Since

$$n^{-1}U(\gamma_0, \tau) \xrightarrow{P} \int_0^\tau [s^{(1)}(u, \gamma_0) - e(u, \gamma_0)] s^{(0)}(u, \gamma_0) dH_0(u) = 0,$$

we have

$$P\{0 \in B(n^{-1}U(\gamma_0; \tau); \varepsilon_1)\} \rightarrow 1 \quad \text{as } n \rightarrow \infty.$$

Therefore

$$P\{\hat{\gamma} \text{ exists and is unique in } B(\gamma_0, \delta_1)\} \rightarrow 1 \quad \text{and} \quad \hat{\gamma} \xrightarrow{P} \gamma_0.$$

Consider the equality (14). If $n \rightarrow \infty$, then

$$\gamma^* \xrightarrow{P} \gamma_0 \quad \text{and} \quad -\frac{1}{n} \frac{\partial U(\gamma^*, \tau)}{\partial \gamma} \xrightarrow{P} \Sigma_1(\gamma_0, \tau),$$

hence the convergence (13). The proof of the theorem 2 is complete. ■

3.3. Asymptotic properties of $\tilde{H}_0$ and $\hat{S}_{x(\cdot)}$

Now we will consider the asymptotic distribution of the estimator of the function H_0 . Denote

$$\begin{aligned} \sigma_1^2(\gamma_0, t) &= \int_0^t \frac{s^{(0)}(\gamma_0, u) dH_0(u) + s_*^{(0)}(\gamma_0, u) du}{(s^{(0)})^2(\gamma_0, u)}, \\ H_i(\gamma_0, u) &= J(u) \{X_i(u) - E(\gamma_0, u) + \frac{1}{h(\gamma_0, u) S^{(0)}(\gamma_0, u)} \times \\ &\sum_{j=1}^n \left[\int_u^t J(s) [X_j(s) - E(\gamma_0, s)] \alpha'(H_j(s)) h(\gamma_0, s) Y_j(s) (dH_0(s) + \gamma_0^T X_j(s) ds) \right] \}. \end{aligned}$$

Assumptions C (Stability conditions)

1.
$$\sum_{i=1}^n \int_0^t \frac{H_i(\gamma_0, u)}{\tilde{S}^{(0)}(\gamma_0, u)} \alpha(H_i(\gamma_0, u)) Y_i(u) (dH_0(u) + \gamma_0^T X_i(u) du) \xrightarrow{P} A(\gamma_0, t).$$
2.
$$\begin{aligned} \int_0^t J(u) \left\{ \frac{\partial}{\partial \gamma} \left(\frac{1}{\tilde{S}^{(0)}(\gamma^*, u)} \right) \left[S^{(0)}(\gamma_0, u) dH_0(u) + \tilde{S}_*^{(0)}(\gamma^*, u) du \right] - \right. \\ \left. \frac{\partial}{\partial \gamma} \left(\frac{\tilde{S}_*^{(0)}(\gamma^*, u)}{\tilde{S}^{(0)}(\gamma^*, u)} \right) du \right\} \xrightarrow{P} C(\gamma_0, t). \end{aligned}$$

Theorem 3. Under assumptions C and those of Theorem 2 for all $t \in [0, \tau]$

$$n^{1/2} \{\tilde{H}_0(t) - H_0(t)\} \xrightarrow{D} N(0, \sigma_2^2(\gamma_0, t)) \quad \text{as} \quad n \rightarrow \infty,$$

where

$$H_0(t) = H(S_0(t)), \quad \tilde{H}_0(t) = \hat{H}_0(t, \hat{\gamma}),$$

$$\sigma_2^2(\gamma_0, t) = \sigma_1^2(\gamma_0, t) + C(\gamma_0, t) \Sigma_2(\gamma_0, \tau) C^T(\gamma_0, t) - 2C(\gamma_0, t) \Sigma_1^{-1}(\gamma_0, \tau) A(\gamma_0, t).$$

Proof: By Taylor expansion around γ_0 the random process $\sqrt{n}\{\tilde{H}_0(t) - H_0(t)\}$ can be written in the following manner:

$$\begin{aligned} \sqrt{n}\{\tilde{H}_0(t) - H_0(t)\} &= n^{1/2} \left\{ \int_0^t \frac{J(u)}{\tilde{S}^{(0)}(\hat{\gamma}, u)} (dN(u) - \tilde{S}_*^{(0)}(\hat{\gamma}, u) du) - H_0(t) \right\} = \\ n^{1/2} \left\{ \int_0^t \frac{J(u)}{\tilde{S}^{(0)}(\gamma_0, u)} (dN(u) - \tilde{S}_*^{(0)}(\gamma_0, u) du) - \int_0^t J(u) \frac{\partial}{\partial \gamma} \left(\frac{1}{\tilde{S}^{(0)}(\gamma^*, u)} \right) dN(u) (\hat{\gamma} - \gamma_0) + \right. \\ &\quad \left. \int_0^t J(u) \frac{\partial}{\partial \gamma} \left(\frac{\tilde{S}_*^{(0)}(\gamma^*, u)}{\tilde{S}^{(0)}(\gamma^*, u)} \right) du (\hat{\gamma} - \gamma_0) - H_0(t) \right\} = \\ (17) \quad n^{1/2} \{ C(\gamma_0, t) (\hat{\gamma} - \gamma_0) + \int_0^t J(u) \frac{dM(u)}{\tilde{S}^{(0)}(\gamma_0, u)} \} + o_p(1), \end{aligned}$$

where γ^* is on the line segment between $\hat{\gamma}$ and γ_0 . Taking into account (15), we obtain

$$n^{1/2}(\hat{\gamma} - \gamma_0) = \Sigma_1^{-1}(\gamma_0, \tau) n^{-1/2} U(\gamma_0, \tau) + o_p(1).$$

From theorem 1:

$$n^{-1/2} U(\gamma_0, \tau) = n^{-1/2} \sum_{i=1}^n \int_0^\tau H_i(\gamma_0, u) dM_i(u) + o_p(1).$$

Therefore the predictable covariation

$$\begin{aligned} &< n^{-1/2} \sum_{i=1}^n \int_0^\tau H_i(\gamma_0, u) dM_i(u), n^{1/2} \int_0^t J(u) \frac{dM(u)}{\tilde{S}^{(0)}(\gamma_0, u)} > = \\ &\sum_{i=1}^n \int_0^t \frac{H_i(\gamma_0, u)}{\tilde{S}^{(0)}(\gamma_0, u)} \alpha(H_i(\gamma, u)) Y_i(u) (dH_0(u) + \gamma_0^T X_i(u) du) \xrightarrow{P} A(\gamma_0, t). \end{aligned}$$

The predictable variation

$$< n^{1/2} \int_0^t J(u) \frac{dM(u)}{\tilde{S}^{(0)}(\gamma_0, u)} > = n \int_0^t J(s) \frac{S^{(0)}(\gamma_0, s) dH(S_0(s)) + S_*^{(0)}(\gamma_0, s) ds}{\tilde{S}^{(0)2}(\gamma_0, s)} \xrightarrow{P} \sigma_1^2(\gamma_0, t).$$

Let

$$H_{1i}(\gamma_0, u) = C(\gamma_0, t) \Sigma_1^{-1}(\gamma_0, \tau) H_i(\gamma_0, u) + nJ(u)/\tilde{S}^{(0)}(\gamma_0, u).$$

Then

$$\sqrt{n}[\tilde{H}_0(t) - H_0(t)] = n^{-1/2} \sum_{i=1}^n \int_0^t H_{1i}(\gamma_0, u) dM_i(u) + o_p(1).$$

By assumptions of the theorem for all $\varepsilon > 0$

$$\frac{1}{n} \sum_{i=1}^n \int_0^t H_{1i}^2(\gamma_0, u) I\{|H_{1i}(\gamma_0, u)| > \varepsilon\} Y_i(u) \alpha_i(u) du \xrightarrow{P} 0.$$

The asymptotic normality of $\sqrt{n}[\tilde{H}_0(t) - H_0(t)]$ follows from the theorem of Rebolledo. The proof of theorem 3 is complete. ■

Consider the asymptotic distribution of the estimator $\hat{S}_{x(\cdot)}$ of the survival function $S_{x(\cdot)}(t)$ under any covariate $x(\cdot)$. Denote

$$C_x(\gamma_0, t) = C(\gamma_0, t) + \int_0^t x^T(u) du, \quad g = G'.$$

Theorem 4. *Under assumptions of theorem 1*

$$\sqrt{n}[\hat{S}_{x(\cdot)}(t) - S_{x(\cdot)}(t)] \xrightarrow{D} N(0, \sigma_x^2) \quad \text{as } n \rightarrow \infty,$$

where

$$\begin{aligned} \sigma_x^2(\gamma_0, t) = & \{\sigma_1^2(\gamma_0, t) + C_x(\gamma_0, t) \Sigma_2(\gamma_0, \tau) C_x^T(\gamma_0, t) - \\ & 2C_x(\gamma_0, t) \Sigma_1^{-1}(\gamma_0, \tau) A(\gamma_0, t)\} g^2(H(S_x(t))). \end{aligned}$$

Proof: By Taylor expansion around γ_0 we obtain (similarly as in the proof of theorem 3):

$$\begin{aligned} n^{1/2}[H(\hat{S}_{x(\cdot)}(t)) - H(S_{x(\cdot)}(t))] = \\ n^{1/2} \left\{ \tilde{H}_0(t, \gamma_0) - H_0(t) + \int_0^t x^T(u) du (\hat{\gamma} - \gamma_0) \right\} = \\ n^{1/2} \left\{ C_x(\gamma_0, t) \Sigma_1^{-1}(\gamma_0, \tau) U(\gamma_0, \tau) + \int_0^t J(u) \frac{dM(u)}{\tilde{S}^{(0)}(\gamma_0, u)} + o_p(1) \right\}. \end{aligned}$$

The proof of the asymptotical normality of $\sqrt{n}[H(\hat{S}_{x(\cdot)}(t)) - H(S_{x(\cdot)}(t))]$ is similar to the proof of the asymptotical normality of $\sqrt{n}[H(\hat{S}_0(t)) - H(S_0(t))]$ in the theorem 3. So we skip it. The asymptotical normality of $\sqrt{n}[\hat{S}_{x(\cdot)}(t) - S_{x(\cdot)}(t)]$ is obtained by the functional delta method. The proof is complete. ■

Remark 1. Replacing all the theoretical quantities by empirical ones in the theorems 1-4, we obtain consistent estimates for the limiting covariance matrix of $\sqrt{n}(\hat{\gamma} - \gamma_0)$ and the limiting variances of

$$\sqrt{n}(\tilde{H}_0(t) - H_0(t)) \quad \text{and} \quad \sqrt{n}(\hat{S}_{x(\cdot)}(t) - S_{x(\cdot)}(t)).$$

□

ACKNOWLEDGEMENT. The authors would like to thank Professor C.M. Cuadras and the referees for their constructive comments and remarks.

4. REFERENCES

- [1] **Andersen, P.K., Borgan, O., Gill, R.D. & Keiding, N.** (1993). *Statistical Models Based on Counting Processes*. New York: Springer-Verlag.
- [2] **Aranda-Ordaz, F.J.** (1983). «An extension of the proportional-hazards model for grouped data». *Biometrics*, **39**, 109–117.
- [3] **Bagdonavičius, V.** (1990). «Accelerated life models when the stress is not constant». *Kybernetika*, **26**, 289–295.
- [4] **Bagdonavičius, V. & Nikulin, M.** (1994). «Stochastic models of accelerated life». In: *Selected Topics on Stochastic Modelling* (ed. R. Gutiérrez, M. J. Valde-rrama), World Scient., Singapore, 73–87.
- [5] **Bagdonavičius, V. & Nikulin, M.** (1995). *Semiparametric models in accelerated life testing*. Queen's Papers in Pure and Applied Mathematics. Queen's University, Kingston, Ontario, Canada. **98**, 70 pp.
- [6] **Bagdonavičius, V. & Nikulin, M.** (1996). «Asymptotical analysis of semiparametric models in survival analysis and accelerated life testing». *Preprint 96009*, Mathématiques Appliquées de Bordeaux.
- [7] **Bagdonavičius, V. & Nikulin, M.** (1996). «Analyses of generalized additive semiparametric models». *Comptes Rendus, Academie des Sciences de Paris*. **323**, 9, Série I, 1079–1084.
- [8] **Bagdonavičius, V. & Nikulin, M.** (1996). «Asymptotical properties of estimators in the generalized additive-multiplicative semiparametric model». *Preprint 96012*, Mathématiques Appliquées de Bordeaux.
- [9] **Cox, D.R. & Oakes, D.** (1984). *Analysis of Survival Data*. London: Chapman and Hall.

- [10] **Dabrowska, D.M. & Doksum, K.A.** (1988). «Partial likelihood in transformations models with censored data». *Scand. Journal of Statist.*, **15**, 1–23.
- [11] **Kalbfleisch, J.D. & Prentice, R.L.** (1980). *The Statistical Analysis of Failure Time Data*. New York: Wiley.
- [12] **Lawless, J.F.** (1982). *Statistical Models and Methods for Lifetime Data*. New York: Wiley.
- [13] **Lin, D.Y. & Ying, Z.** (1994). «Semiparametrical analysis of the additive risk model». *Biometrika*, **81**, 61–71.
- [14] **Lin, D.Y. & Ying, Z.** (1996). «Semiparametric analysis of the general additive-multiplicative hazard models for counting processes». *The Annals of Statistics*, **23**, 5, 1712–1734.

EL PROBLEMA DE BEHRENS-FISHER EN LA INVESTIGACIÓN BIOMÉDICA. ANÁLISIS CRÍTICO DE UN ESTUDIO CLÍNICO MEDIANTE SIMULACIÓN

ESTEBAN VEGAS LOZANO*

Universitat de Barcelona

En el artículo se hace una revisión del problema de Behrens-Fisher, discutiendo los fundamentos inferenciales asociados a la dificultad de su resolución y exponiendo las soluciones prácticas más comunes, juntamente con una nueva solución basada en conceptos de geometría diferencial. A continuación, se realiza un estudio crítico de una investigación biomédica en donde las verdaderas probabilidades de error son distintas de las supuestas debido a que se ignoran probables diferencias entre las varianzas. En dicha investigación se rechazó la hipótesis nula de igualdad de medias ($p < 0.01$), si bien, la verdadera probabilidad de error de tipo I para valores próximos a los muestrales parece ser otra. Con este fin, se realiza un estudio de Monte Carlo para obtener estas estimaciones según se utilice el test t de Student de comparación de medias o diferentes soluciones más apropiadas para el problema de Behrens-Fisher. En este estudio de simulación se usa una técnica de reducción de la varianza específica para variables de respuesta dicotómica tales como contrastes de hipótesis (aceptar, no aceptar la hipótesis nula). Se presenta brevemente esta técnica y se ilustra como se diseña la simulación de acuerdo con la misma.

The Behrens-Fisher problem in biomedical research. Critical analysis of a clinical study by simulation

Keywords: Simulación de Monte Carlo, Reducción de la varianza, Curvatura estadística, Distancia de Rao, Bootstrap.

Clasificación AMS: 65C05

*Esteban Vegas Lozano. Departament d'Estadística. Facultat de Biologia. Universitat de Barcelona. Diagonal 645, 08028 Barcelona.

—Article rebut el gener de 1997.

—Acceptat el maig de 1997.

1. INTRODUCCIÓN

En la investigación experimental, en particular en la biomédica, es muy común que aparezca el problema de Behrens-Fisher: contrastar la igualdad de medias de dos poblaciones normales, sin hacer ninguna suposición acerca de las varianzas. De una manera más formal, el problema se puede exponer como sigue:

Sean X_1, X_2 dos variables aleatorias normales independientes con sus respectivas medias $\mu_i, i = 1, 2$, y con varianzas desconocidas y arbitrarias $\sigma_i^2, i = 1, 2$. Bajo estas premisas se considera el contraste de hipótesis:

$$\begin{aligned} H_0 : & \mu_1 = \mu_2 = \mu \quad \text{arbitrarios} \\ H_1 : & \mu_1 \neq \mu_2 \quad \sigma_1^2 > 0 \quad \sigma_2^2 > 0 \end{aligned}$$

Se supone que el contraste está basado en una muestra de $n_1 + n_2$ valores

$$\mathbf{x} = (x_{11}, \dots, x_{1n_1}; x_{21}, \dots, x_{2n_2}),$$

donde cada submuestra viene de n_i realizaciones independientes e idénticamente distribuidas (*iid*) de X_i , para $i = 1, 2$.

La diferencia con respecto al conocido test t de Student para la comparación de medias de dos poblaciones normales con varianzas desconocidas pero iguales es no hacer ninguna suposición respecto de las varianzas. Así pues, la diferencia fundamental está en que para aplicar el test t de Student se necesita *homocedasticidad* mientras que en el problema de Behrens-Fisher se incluye, además, el caso de *heterocedasticidad*. Como es bien conocido por los estadísticos, la aparición de heterocedasticidad suele complicar la inferencia estadística (algunos modelos típicos de inferencia clásica no se pueden aplicar como por ejemplo, el modelo lineal normal) y se entra en un campo donde existen numerosos problemas abiertos para los cuales, en el mejor de los casos, se han encontrado algunas soluciones aproximadas pero donde ninguna de ellas es óptima.

Este detalle produce un cambio cualitativo muy importante, que se refleja a un nivel más abstracto en la geometría de ciertas familias paramétricas de densidad. Así, se tiene que tanto el modelo paramétrico correspondiente al caso de varianzas iguales como el del caso de varianzas desiguales se pueden expresar como una familia exponencial de 3 y 4 parámetros, respectivamente. Pero es en el submodelo paramétrico asociado a la hipótesis nula donde se produce la diferencia a remarcar. El espacio de distribuciones de probabilidad asociado a la hipótesis nula en el caso de varianzas iguales es un «subespacio plano» del espacio de la familia exponencial de dimensión

3, mientras que en el caso de varianzas desiguales es un «subespacio curvado» de la familia exponencial de dimensión 4. Como consecuencia, se produce una ruptura de las buenas propiedades para los problemas de inferencia que tiene la familia exponencial, que se acrecienta cuanto mayor es la curvatura. Así, por ejemplo, la varianza del estimador máximo verosímil excede del valor de la cota de Cramer-Rao en proporción aproximada al valor de su curvatura al cuadrado. Tampoco existe un estadístico suficiente con la misma dimensión que la familia exponencial curvada (Efron, 1975).

Las familias exponenciales curvadas suelen aparecer bastante a menudo en la práctica y tiene gran importancia en la discusión de varios conceptos y métodos claves de inferencia estadística (ver Efron, 1975, 1978 y Efron and Hinkley, 1978).

2. SOLUCIONES PRÁCTICAS AL PROBLEMA DE BEHRENS-FISHER

El problema de Behrens-Fisher es una cuestión controvertida que no tiene solución óptima conocida y para el que las diferentes escuelas de pensamiento estadístico dan diferentes soluciones. Como breve resumen, se muestran algunas de las soluciones más habituales en la práctica. Para una revisión más exhaustiva véase Scheffé (1970) y Lee and Gurland (1975).

Para tamaños muestrales mayores que 10, las diferencias entre las diversas soluciones propuestas son generalmente mucho menores que sus diferencias con el test t de Student para la igualdad de medias. Por tanto, el uso de cualquiera de ellos es mejor que el uso del test t de Student, a menos que se garantice la validez de la suposición de igualdad de varianzas.

De todas estas soluciones, las más comúnmente utilizadas son:

- El test de Cochran y Cox.
- El test de Welch o el test de Welch-Aspin.

La primera se puede considerar, con el test de McCullough-Barnerjee, como aproximación a la solución de Behrens-Fisher, mientras que la mayoría de otros test discutidos por Lee and Gurland (1975) se pueden considerar como aproximaciones al test de Welch-Aspin.

En todos los casos, el estadístico de contraste es:

$$(1) \quad T' = \frac{\bar{x}_1 - \bar{x}_2}{\sqrt{(\hat{s}_1^2/n_1) + (\hat{s}_2^2/n_2)}}$$

donde, para cada submuestra $(x_{i1}, \dots, x_{in_i})$ $i = 1$ o 2 , $\bar{x}_1, \bar{x}_2$ representan las medias muestrales; s_1^2, s_2^2 corresponden a las varianzas muestrales con denominador n_i y $\hat{s}_1^2, \hat{s}_2^2$ designa a las varianzas muestrales con denominador $n_i - 1$.

El problema es que el estadístico T' no se distribuye, necesariamente, según el modelo de probabilidad t de Student con $n_1 + n_2 - 2$ grados de libertad. Por tanto, todo el esfuerzo se centra en conocer de forma aproximada la distribución muestral de T' .

Cochran y Cox (1950) diseñaron un método de aproximación a los puntos críticos de la distribución muestral T' . El método consiste en obtener estos puntos mediante:

$$(2) \quad t'_p = \frac{t(p, n_1 - 1)(\hat{s}_1^2/n_1) + t(p, n_2 - 1)(\hat{s}_2^2/n_2)}{(\hat{s}_1^2/n_1) + (\hat{s}_2^2/n_2)}$$

donde $t(p, v)$ es el p -percentil superior de una distribución t de Student con v grados de libertad. El proceso de decisión es *rechazar H_0 si $|T'| \geq t'_{\alpha/2}$* siendo α el nivel de significación nominal del test.

La principal ventaja de esta solución es su simplicidad. Tal vez es por esta razón que se ha convertido en el test más extensamente utilizado en la práctica, incluso sabiendo que la extensión del test puede diferir substancialmente del nivel de significación nominal, como fue mostrado por Cochran (1964). Afortunadamente, en la mayoría de ocasiones, la extensión del test está por debajo del nivel de significación nominal, es decir, se trata de un test conservador.

Welch (1938) propuso una aproximación alternativa. En esta aproximación T' se concibe como una variable aleatoria distribuida según la t de Student, pero con un número desconocido de grados de libertad. La solución pasa por determinar los grados de libertad (gl') que corresponden a la distribución de T' mediante la expresión:

$$(3) \quad gl' = \frac{\left(\frac{\hat{s}_1^2}{n_1} + \frac{\hat{s}_2^2}{n_2} \right)^2}{\frac{(\hat{s}_1^2/n_1)^2}{n_1 - 1} + \frac{(\hat{s}_2^2/n_2)^2}{n_2 - 1}}$$

El resultado obtenido para gl' se redondea al entero más próximo¹. Se obtienen así unos grados de libertad comprendidos entre un mínimo y un máximo conocidos: el mínimo es el valor más pequeño de $n_1 - 1$ y $n_2 - 1$; el máximo es $n_1 + n_2 - 2$. El criterio de decisión es *rechazar* H_0 si $|T'| \geq t(\alpha/2, gl')$ donde $t(\alpha/2, gl')$ es el $(\alpha/2)$ -percentil superior de una distribución t de Student con gl' grados de libertad.

El test de Welch-Aspin se basa en cálculos asintóticos sobre T' . Welch (1947) obtiene los puntos porcentuales de T' como una serie de potencias en $1/f_i = 1/(n_i - 1)$ para $i = 1, 2$ (véase el desarrollo de la fracción P en el apéndice A). Al año siguiente, Aspin (1948) extendió estos resultados hasta el orden 4. Por tanto, para poder utilizar este test se necesita tabular los puntos críticos. Se pueden encontrar en la tabla 11 de *Biometrika Tables* (1976) los valores tabulados para cuatro niveles de probabilidad ($\alpha/2 = 0.05, 0.025, 0.01$ y 0.005) para la cola superior. El criterio de decisión es *rechazar* H_0 si $|T'| \geq V(c; f_1, f_2, \alpha/2)$ donde $c = \frac{\hat{s}_1^2/n_1}{\hat{s}_1^2/n_1 + \hat{s}_2^2/n_2}$ y el valor $V(c; f_1, f_2, \alpha/2)$ se obtiene en las tablas anteriores.

Las diferencias entre los dos últimos test son mínimas siendo más común utilizar el test de Welch ya que no requiere de unas tablas específicas y en sólo un caso hace falta determinar los grados de libertad gl' . Dado que los valores de la distribución t de Student van disminuyendo a medida que van aumentando los grados de libertad, antes de calcular gl' se puede evaluar T' utilizando el gl' mínimo (es decir, el menor de $n_1 - 1$ y $n_2 - 1$); si se rechaza $H_0 : \mu_1 = \mu_2$, también se rechazará con el valor proporcionado por (3) para gl' ; si no se rechaza H_0 , podemos evaluar T' con el gl' máximo ($n_1 + n_2 - 2$); si se sigue sin rechazar H_0 , tampoco se rechazará calculando el valor exacto de gl' . De modo que el único caso en el que se necesita hacer uso de (3) para calcular el valor exacto de gl' será aquel en el que se mantenga H_0 con el gl' mínimo y se rechace con el gl' máximo.

A diferencia del test de Cochran y Cox, tanto el test de Welch como el test de Welch-Aspin, la probabilidad de error es muy cercana al valor nominal en todo el espacio paramétrico, aunque Fisher (1956) los critica por mostrar en un subconjunto relevante la existencia de un sesgo negativo en el sentido de Buehler (1959) (véase apéndice B).

¹El propio Welch (1947) sugirió posteriormente que hacer:

$$gl' = \left[\frac{\left(\frac{\hat{s}_1^2}{n_1} + \frac{\hat{s}_2^2}{n_2} \right)^2}{\frac{(\hat{s}_1^2/n_1)^2}{n_1 - 1} + \frac{(\hat{s}_2^2/n_2)^2}{n_2 - 1}} \right] - 2$$

puede ofrecer una solución más exacta para gl' . No obstante, la diferencia entre ambas soluciones es, en la mayor parte de los casos, insignificante.

Como una última alternativa, se cita el test denominado geodésico² descrito en Vegas y Ocaña (1996) y Vegas (1996) (véase un breve resumen en el apéndice C) que está basado en un enfoque totalmente distinto, concretamente en criterios de geometría diferencial ya que el estadístico propuesto, D^2 , expresa el nivel de disconformidad, en términos de distancia geodésica, entre los valores obtenidos en las estimaciones puntuales de la media y la desviación típica de la muestra $(\bar{x}_1, \bar{x}_2, s_1, s_2)$ y el conjunto de todos aquellos puntos paramétricos que cumplen la hipótesis nula de igualdad de medias con desviaciones típicas arbitrarias $(\mu, \mu, \sigma_1, \sigma_2)$.

3. ANÁLISIS CRÍTICO DE UNA INVESTIGACIÓN BIOMÉDICA

En esta sección se realiza un estudio crítico de las conclusiones de una investigación experimental en el entorno de las ciencias médicas y biológicas, en que las verdaderas probabilidades de error son distintas de las supuestas, precisamente a causa de ignorar probables diferencias de las varianzas. Es decir, es un ejemplo donde surge el problema de Behrens-Fisher en el que se obviaron las probables diferencias entre las varianzas.

3.1. Un estudio sobre la trombosis

Los datos que se presentan son una parte de un estudio más amplio que fue expuesto en un artículo de Oost *et al.* (1983). La presente revisión se centra en las medidas obtenidas sobre la excreción de β -tromboglobulina urinaria en 12 pacientes normales y 12 pacientes diabéticos.

Normal	Diabético	Normal	Diabético
4.1	11.5	11.5	33.9
6.3	12.1	12.0	40.7
7.8	16.1	13.8	51.3
8.5	17.8	17.6	56.2
8.9	24.0	24.3	61.7
10.4	28.8	37.2	69.2

En el artículo mencionado, se obtuvo como resultado que existía diferencia significativa ($p < 0.01$) en la excreción media de β -tromboglobulina urinaria entre los dos tipos de pacientes. A pesar de que no se indica claramente en el artículo, todo parece

²Para realizar el test geodésico se implementó un programa en Fortran v.3.2 que da como resultado la aceptación o rechazo de la igualdad de medias poblacionales. El programa esta a disposición de cualquiera que lo desee. Para solicitar una copia dirigirse al autor.

indicar que se realizó un test t de Student de comparación de medias para llegar a tal conclusión.

3.2. Comparación de medias entre los dos tipos de pacientes

El primer paso del estudio es comprobar que la comparación de medias entre los dos tipos de pacientes es, razonablemente, un caso del problema de Behrens-Fisher. Se verifica la normalidad utilizando la prueba de Lilliefors. Las tablas utilizadas no son las que obtuvo H. W. Lilliefors (1967) sino, las presentadas en A.L. Mason and C.B. Bell (1986) realizadas con un tamaño muestral de simulación mayor ($n = 20000$). Para ambas muestras, se acepta la hipótesis nula de normalidad para un nivel de significación de 0.05. Por otra parte, en el test F de comparación de varianzas se obtiene que existen diferencias significativas entre ellas ($p < 0.01$).

A continuación, se calcula algunas de las soluciones prácticas habituales para el problema de Behrens-Fisher: test de Cochran y Cox, Welch y Welch-Aspin. Todas ellas se basan en el estadístico T' (1) que da como resultado un valor de 3.3838 que, juntamente con el test geodésico (apéndice C), conducen a rechazar la igualdad de medias.

Incluso si no se tiene en cuenta la existencia de diferencias significativas entre las varianzas y se realiza el test t de Student de comparación de medias también se rechaza la hipótesis nula de igualdad de medias. Es decir, sea cual sea el test utilizado se rechaza la hipótesis nula, así pues, el estudio más interesante consiste en estimar cual es la verdadera probabilidad de error de tipo I.

3.3. Estimación de la verdadera probabilidad de error de tipo I por simulación

Como el resultado del contraste de medias ha sido rechazar la hipótesis nula el único error que podemos cometer es el de tipo I, es decir, rechazar la hipótesis nula cuando realmente la media poblacional de la primera población es igual que la segunda. En teoría esta probabilidad debería ser el nivel de significación nominal del test, que en este caso será de 0.05. Por tanto, lo que se desea verificar es si la extensión del test coincide con el nivel de significación nominal de 0.05.

3.3.1. Procedimiento

Se plantea un estudio de Monte Carlo que analiza el efecto de diferentes cocientes de varianzas (heterocedasticidad) alrededor de los valores muestrales en la estimación de la verdadera probabilidad de error de tipo I. Además, se incorpora una técnica de reducción de la varianza específica para variables de respuesta dicotómica (véase apéndice D para un breve resumen; más información en Ocaña y Vegas, 1995; Vegas 1996) para lograr unas estimaciones más precisas. En cada réplica de la simulación

Monte Carlo se obtiene, además del valor de la variable dicotómica de interés o respuesta Y (en este caso $Y = 1$ si se rechaza la hipótesis nula en el test estudiado, y $Y = 0$ en caso contrario), otra variable dicotómica de control C , correlacionada con Y , cuya esperanza, $E(C)$, es conocida. Esta variable C se basa en el resultado del test t de Student de comparación de medias ($C = 1$ si se rechaza la hipótesis nula y $C = 0$ en caso contrario) ya que este test está correlacionado con el test estudiado, para cualquiera de las soluciones habituales al problema de Behrens-Fisher, y además, se conoce la esperanza de C , es decir, su potencia. Es importante resaltar que interesa que la correlación entre la variable Y y la variable C sea elevada para que la precisión del estimador de la probabilidad de error de tipo I obtenido por esta técnica sea alta.

El algoritmo de simulación se divide en tres módulos:

1. Entrada de parámetros.
2. Obtención de los dos vectores de resultados (Y_k, C_k) $k = 1, \dots, n$ donde Y corresponde al resultado del test bajo estudio y C al resultado del test de control (test t de Student).
3. Estimación de la probabilidad de error de tipo I aplicando la técnica de reducción de la varianza indicada anteriormente.

En el apartado de la entrada de parámetros los valores que permanecen fijos se refieren a los tamaños muestrales (n_1, n_2) que son iguales a 12, al nivel de significación del test de respuesta y del test de control que valen 0.05, al número de réplicas de simulación ($n = 5000$) y réplicas del bootstrap paramétrico ($B = 1000$) necesario para estimar el punto crítico del test geodésico (véase apéndice C). Por otra parte, se utilizan tres pares de desviaciones típicas

$$(\sqrt{118.28}, \sqrt{84.54}), (\sqrt{1427.25}, \sqrt{84.54}) \text{ y } (\sqrt{410.87}, \sqrt{84.54})$$

que corresponden a posibles valores de varianzas tales que el cociente de cada pareja son los extremos (1.40, 16.88) y el punto medio (4.86) del intervalo de confianza del 95% para el cociente de varianzas poblacionales, respectivamente. Se escogieron aquellos valores de las varianzas en que al menos una de ellas era igual que la obtenida como varianza muestral a partir de los datos reales obtenidos por Oost *et al.* (1983). Al escoger estos tres pares de desviaciones típicas extremas se intenta observar cómo va variando la estimación de la probabilidad de error de tipo I con un número reducido de simulaciones.

El segundo apartado del algoritmo se divide en (véase fig. 1):

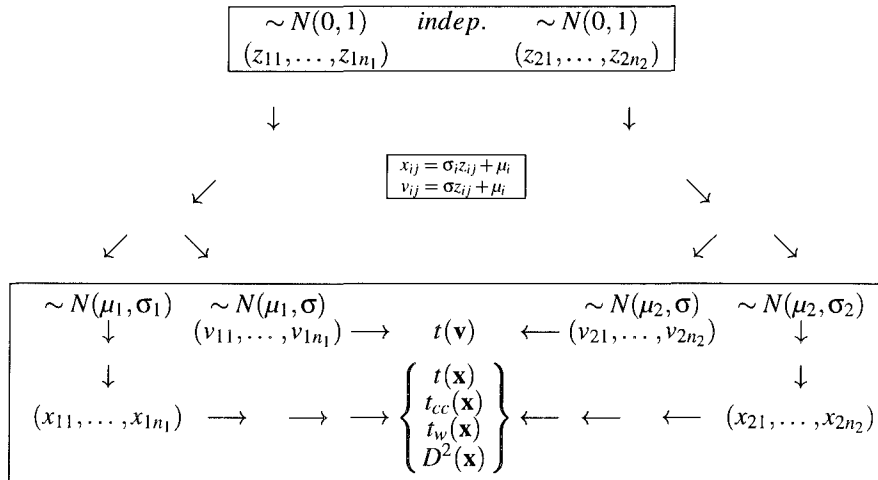
- Generación de las muestras.
- Calcular el test de estudio y el test t de Student de control ($t(\mathbf{v})$).
- Obtención de Y y C .

Como se desea estimar la verdadera probabilidad de error de tipo I para los siguientes tests:

- test t de Student de comparación de medias aunque las muestras provengan de normales con diferentes varianzas ($t(\mathbf{x})$).
- test de Cochran y Cox ($t_{cc}(\mathbf{x})$).
- test de Welch ($t_w(\mathbf{x})$).
- test geodésico ($D^2(\mathbf{x})$).

es necesario realizar tantos tipos de simulaciones como diferentes tipos de tests. Así, se substituye la frase «Calcular el test de estudio», que aparece en el parrafo anterior, por cada uno de ellos.

Generación de las muestras



Parejas diferentes:

$$\begin{array}{cccc} t(\mathbf{v}) \rightarrow C & t(\mathbf{v}) \rightarrow C & t(\mathbf{v}) \rightarrow C & t(\mathbf{v}) \rightarrow C \\ t(\mathbf{x}) \rightarrow Y^1 & t_{cc}(\mathbf{x}) \rightarrow Y^2 & t_w(\mathbf{x}) \rightarrow Y^3 & D^2(\mathbf{x}) \rightarrow Y^4 \end{array}$$

Figura 1. Obtención de los dos vectores de resultados (Y_k, C_k) $k = 1, \dots, n$ en la estimación de la verdadera probabilidad de error de tipo I para varios tests.

Para inducir la correlación requerida entre el resultado de cada uno de los anteriores tests ($t(\mathbf{x})$, $t_{cc}(\mathbf{x})$, $t_w(\mathbf{x})$ y $D^2(\mathbf{x})$) con el test t de Student de control ($t(\mathbf{v})$) se utiliza variables aleatorias comunes para evaluar cada uno de ellos.

A partir de n_1 valores normales estándar *iid*, $z_1 = (z_{11}, \dots, z_{1n_1})$, independientes de otros n_2 valores normales estándar *iid*, $z_2 = (z_{21}, \dots, z_{2n_2})$, se obtiene, gracias a la transformación $x_{ij} = \sigma_i z_{ij} + \mu_i$, $i = 1, 2$ y $j = 1, \dots, n_i$, los $n_1 + n_2$ valores $x = (x_{11}, \dots, x_{1n_1}; x_{21}, \dots, x_{2n_2})$ para calcular los cuatro tests de estudio en diferentes configuraciones de medias y varianzas, incluyendo el test t de Student evaluado bajo condiciones de desigualdad de varianzas. Y aplicando de nuevo la misma transformación, pero siendo la varianza utilizada $\sigma^2 = \frac{n_1\sigma_1^2 + n_2\sigma_2^2}{n_1 + n_2}$, un valor común e intermedio entre σ_1^2 y σ_2^2 , se obtiene $n_1 + n_2$ valores $v = (v_{11}, \dots, v_{1n_1}; v_{21}, \dots, v_{2n_2})$ que sirven para calcular el test t de Student ($t(v)$) utilizado como variable de control. De esta manera, los $n_1 + n_2$ valores de v cumplen las condiciones de aplicabilidad del test t de Student ($t(v)$): provenir de dos poblaciones normales con varianza común, y por tanto, se puede utilizar como variable de control al ser conocida su potencia.

Para cada tipo de test tendremos el conjunto de vectores de resultados (Y_k, C_k) $k = 1, \dots, n$ del cual se estima la probabilidad de error de tipo I aplicando la técnica de reducción de la varianza específica para variables dicotómicas.

3.3.2. Resultados

Los resultados presentados son los obtenidos a partir de $n = 5000$ réplicas de simulación y, en el caso del test geodésico además, se han hecho $B = 1000$ remuestras en cada réplica de simulación para estimar el valor crítico c_α (9, véase apéndice C). Se resumen en la tabla 1 donde se muestran, para cada uno de los test estudiados, las estimaciones de la probabilidad de error de tipo I en los tres niveles de heterocedastidad analizados:

$\bar{p}_1$: Representa la estimación puntual de la verdadera probabilidad de error de tipo I utilizando como estimador (11, véase apéndice D).

Tabla 1

Nivel de significación estimado en el test geodésico, test de Welch, test de Cochran y Cox y test t de Student, respectivamente, en relación al estudio sobre la trombosis con $n_1 = n_2 = 12$ bajo distintas combinaciones de parámetros con valores próximos a los muestrales correspondientes a la hipótesis nula.

	Resultado del test							
	Geodésico		Welch		Cochran y Cox		t de Student	
σ_1^2/σ_2^2	$\bar{p}_1$	$\pm 1.959964\hat{\sigma}_{GS}$	$\bar{p}_1$	$\pm 1.959964\hat{\sigma}_{GS}$	$\bar{p}_1$	$\pm 1.959964\hat{\sigma}_{GS}$	$\bar{p}_1$	$\pm 1.959964\hat{\sigma}_{GS}$
1.40	.0494	.002568	.0496	.002311	.0405†	.002539	.0504	.002191
4.86	.0535	.004796	.0535	.004775	.0506	.004488	.0507	.004838
16.88	.0546	.005631	.0530	.005540	.0466	.005408	.0620†	.005906

† = significativamente diferente del nivel de significación nominal.

$\pm 1.959964 \hat{\sigma}_{GS}$: Margen de error asociado a un nivel de confianza del 95% para la verdadera probabilidad de error de tipo I, $p_{1.}$, utilizando como estimador de la varianza de $\bar{p}_{1.}$ el estadístico $\hat{\sigma}_{GS}^2$ discutido en el apéndice D.

3.3.3. Conclusiones del estudio

Como se puede observar, tanto en las tablas como en el gráfico que se presenta posteriormente, sólo el test geodésico y el test de Welch incluyen el nivel de significación nominal de 0.05 en los respectivos intervalos de confianza aproximados del 95% para el nivel de significación a partir de una buena estimación del error estándar ($\hat{\sigma}_{GS}$, (15)) del estimador ($\bar{p}_{1.}$) para cualquiera de los tres cocientes de varianzas estudiados.

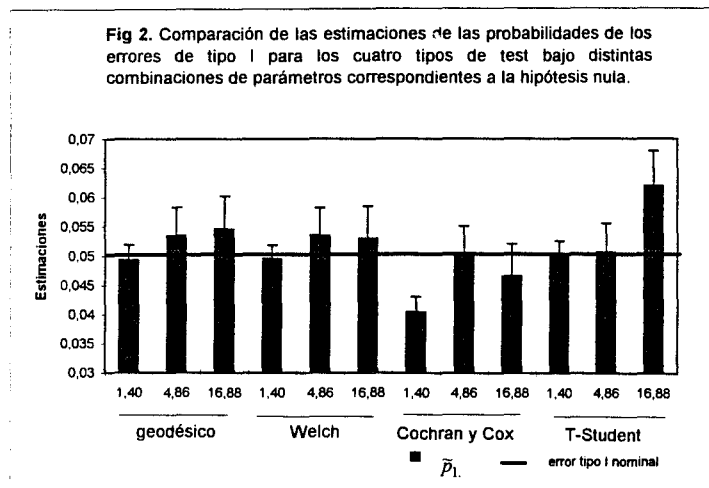


Figura 2. Comparación de las estimaciones de las probabilidades de los errores de tipo I para los cuatro tipos de test bajo distintas combinaciones de parámetros correspondientes a la hipótesis nula.

El test de Cochran y Cox es un test conservador. Su verdadera extensión tiende a ser menor que el nivel de significación nominal (Lee and Gurland, 1975) y esta tendencia es más fuerte en el caso de $n_1\sigma_1^2 = n_2\sigma_2^2$. Esta afirmación se refleja al estar excluido en el intervalo de confianza aproximado del 95% para el nivel de significación el valor de 0.05 cuando el cociente de varianzas poblacionales es de 1.40. Para valores superiores, sí se incluye el valor de 0.05 de nivel de significación nominal. Gráficamente se observa (fig. 2) como en el primer caso la línea horizontal

de valor 0.05 no corta el segmento que representa el intervalo de confianza aproximado del 95% para el nivel de significación construido a partir de $\tilde{p}_1$, correspondiente al cociente de varianzas poblacional igual a 1.40 para el test de Cochran y Cox mientras que para los valores de 4.86 y 16.88 sí que se corta.

Para el test t de Student ocurre al revés, como era de esperar. Cuando el cociente de varianzas es de 1.40 e incluso de 4.86 el intervalo de confianza aproximado del 95% para el nivel de significación obtenido a partir de $\tilde{p}_1$, incluye el valor de 0.05 de nivel de significación nominal. Sin embargo, cuando es de 16.88 la estimación de la probabilidad de error de tipo I es superior a 0.05.

De todo lo expuesto se puede concluir que tanto el test geodésico como el test de Welch mantienen el nivel de significación nominal de 0.05 para los diferentes cocientes de varianzas y que, además, no existen excesivas diferencias entre las diferentes estimaciones obtenidas entre ellos. En cambio, el test de Cochran y Cox falla con varianzas semejantes y el test t de Student va peor cuanto mayor es el valor del cociente de varianzas poblacionales.

4. CONCLUSIONES

Es bastante frecuente en trabajos experimentales, y en particular médico-biológicos, suponer homocedasticidad entre las poblaciones cuando en realidad un test de comparación de varianzas indicaría todo lo contrario. Así, en algunas ocasiones, cuando surge el problema de Behrens-Fisher se resuelve de manera inadecuada a partir del test t de Student de comparación de medias.

Esto implica un error, que repercute en que los p-valores no son los que corresponden, se rechaza la hipótesis nula de igualdad de medias falsamente más veces, incrementándose esta tendencia cuanto mayor sea el valor del cociente de varianzas (véase la fig. 2). Esto implica que el test t de Student será claramente no recomendable bajo condiciones de extrema diferencia de varianzas o cuando los valores del estadístico de test se sitúen en la frontera entre las dos posibles decisiones (aceptar o rechazar la hipótesis nula).

Desde otro punto de vista, si se comparan los diferentes tests según el grado de heterocedasticidad y utilizando los valores de los parámetros del estudio médico de Oost *et al.* (1983) como marco de referencia, se puede llegar a las siguientes conclusiones:

- Si el cociente de varianzas es elevado (16.88) el test que cometería más error, es decir, aquel que la verdadera probabilidad de error de tipo I se aleja más del

nivel de significación nominal de 0.05 es el test t de Student, ya que el intervalo de confianza del 95% siempre es superior a 0.05, entre 0.056 y 0.068. Por tanto, se rechaza la hipótesis nula cuando realmente es cierta un porcentaje mayor que el 5%. Si se observa los valores estimados en el test de geodésico (0.0546) y el test de Welch (0.0530) tiene cierta tendencia a ser superior a 0.05 aunque sus respectivos intervalos de confianza incluyan el valor de 0.05. Igual que sucede con el test de Cochran y Cox, aunque este caso la estimación puntual, 0.0466, es menor que 0.05.

- En el caso de un cociente de varianzas menor (4.86), todos los intervalos de confianza del 95% de la verdadera probabilidad de error de tipo I en los diferentes tests incluyen el nivel de significación nominal de 0.05. La diferencia entre los test aparece en sus estimaciones puntuales. El test t de Student (0.0507) y el test Cochran y Cox (0.0506) son los que tienen un valor más próximo a 0.05, mientras que, el test geodésico y el test de Welch con un valor de 0.0535 tienen una cierta predisposición a tener valores superiores a 0.05.
- Por último, para el cociente de varianzas (1.40) cercano a la homocedasticidad. El test que se aleja más del nivel de significación nominal es el de Cochran y Cox, el intervalo de confianza no incluye el valor de 0.05 y su estimación puntual vale 0.0405, es decir, que se rechaza la hipótesis de igualdad de medias falsamente un número menor de veces que lo que se esperaría (un 5%). En los otros tests los intervalos de confianza incluyen el nivel significación nominal. Las estimaciones del test geodésico (0.0494) esta más alejada del teórico 0.05 que el test de Welch (0.0496) y el test t de Student (0.0504).

En general, se debe recomendar utilizar tanto el test geodésico (véase apéndice C) como el test de Welch ya que son igualmente adecuados para cualquier valor del cociente de varianzas. El test de Cochran y Cox es muy conservador cuando el cociente de varianzas es próximo a uno, y el peor de todos es el test t de Student. En cualquier caso, las conclusiones finales del artículo que ha servido de base a la presente discusión (no así los «p-valores») son «correctas», en tanto que coinciden con las que se obtienen mediante una técnica estadística más apropiada.

En otros estudios paralelos (Vegas, 1996) en los cuales se hacía perturbar la normalidad no se llegaba a tener los p-valores tan alejados. Por tanto, parece indicar que la heterocedasticidad afecta más que la perturbación de la normalidad en la variación de los p-valores.

Por otro lado, la técnica de reducción de la varianza empleada para el estudio de la probabilidad de error de tipo I mediante simulación de Monte Carlo puede ser exportada hacia otros estudios similares con una eficacia adecuada y un mínimo esfuerzo en la complejidad del diseño de la simulación.

APÉNDICES

A. Desarrollo de la fracción P

La fracción P ($0 < P < 1$) de T' se desarrolla explícitamente hasta orden 2 como

$$(4) \quad \alpha \left[1 + \frac{(1+\alpha)^2}{4} \frac{\sum_{i=1}^2 (\hat{s}_i^4/n_i^2 f_i)}{(\sum_{i=1}^2 \hat{s}_i^2/n_i)^2} + \frac{(3+5\alpha^2+\alpha^4)}{3} \frac{\sum_{i=1}^2 (\hat{s}_i^6/n_i^3 f_i^2)}{(\sum_{i=1}^2 \hat{s}_i^2/n_i)^3} \right. \\ \left. - \frac{(15+32\alpha^2+9\alpha^4)}{32} \frac{\sum_{i=1}^2 (\hat{s}_i^4/n_i^2 f_i)^2}{(\sum_{i=1}^2 \hat{s}_i^2/n_i)^4} \right],$$

donde $\alpha = \Phi^{-1}(P)$ y Φ es la función de distribución de la normal estándar.

B. Existencia de sesgo negativo en el sentido de Buehler

Este hecho se muestra en que independientemente de $\hat{s}_1/\hat{s}_2$, $U = (n_1 - 1)(\hat{s}_1^2/\sigma_1^2) + (n_2 - 1)(\hat{s}_2^2/\sigma_2^2)$ sigue una distribución ji-cuadrado con $n_1 + n_2 - 2$ grados de libertad y $(\bar{x}_1 - \mu_1) + (\bar{x}_2 - \mu_2)$ está distribuido normalmente con media cero y varianza $\sigma_1^2/n_1 + \sigma_2^2/n_2$ siendo ambas variables aleatorias estocásticamente independientes. Por consiguiente,

$$(5) \quad T' \sqrt{\frac{\hat{s}_1^2/n_1 + \hat{s}_2^2/n_2}{\sigma_1^2/n_1 + \sigma_2^2/n_2} \frac{n_1 + n_2 - 2}{U}}$$

tiene una distribución t de Student con $n_1 + n_2 - 2$ grados de libertad dados $\hat{s}_1/\hat{s}_2, \sigma_1^2, \sigma_2^2$. Cuando $\hat{s}_1/\hat{s}_2 = 1, n_1 = n_2$ y $\sigma_1^2/\sigma_2^2 = w$, la expresión (5) se simplifica a

$$T' \sqrt{\frac{4}{2 + 1/w + w}}.$$

Por otra parte, puesto que $1/w + w \geq 2$, se tiene

$$\sqrt{\frac{4}{2 + 1/w + w}} \leq 1$$

y por tanto

$$Pr(|T'| > a | \hat{s}_1/\hat{s}_2 = 1) \geq Pr(|t_{(n_1+n_2-2)}| > a)$$

para todo $a > 0$, donde $t_{(n_1+n_2-2)}$ indica la distribución t de Student con $n_1 + n_2 - 2$ grados de libertad... En particular, si $n_1 = n_2 = 7$ y $a = 1.74$ (de la tabla 11 de *Biometrika Tables*),

$$Pr(|T'| > 1.74 | \hat{s}_1 / \hat{s}_2 = 1) \geq Pr(|t_{12}| > 1.74) > 0.1.$$

Entonces el conjunto con $\hat{s}_1 / \hat{s}_2 = 1$ es un subconjunto relevante donde el intervalo de confianza basado en el test de Welch-Aspin cubre el verdadero valor de $\mu_1 - \mu_2 = 0$ menos a menudo que el nivel de confianza sugerido.

C. Test geodésico

Se denomina así porque el proceso de decisión del nuevo test se basa en la distancia geodésica³ para resolver el problema de Behrens-Fisher.

Usando como base la distancia de Rao entre dos distribuciones normales univariantes (Atkinson y Mitchell, 1981; Burbea y Rao, 1982), una medida natural de discrepancia entre una muestra compuesta de dos submuestras *independientes*, caracterizadas por los puntos $(\bar{x}_1, s_1)$ y $(\bar{x}_2, s_2)$, y un punto paramétrico compuesto por (μ_1, σ_1) y (μ_2, σ_2) es la *distancia de Rao al cuadrado ponderada* por los tamaños muestrales n_1 y n_2 (es inmediato dada la expresión de d^2 para distribuciones normales independientes, véase Amari *et al.* (1987) página 237)

$$(6) \quad d^2[(\bar{x}_1, \bar{x}_2, s_1, s_2), (\mu_1, \mu_2, \sigma_1, \sigma_2)] = 2 \sum_{i=1}^2 n_i \left\{ \log \frac{1 + \Delta_i}{1 - \Delta_i} \right\}^2$$

donde

$$(7) \quad \Delta_i = \left\{ \frac{(\bar{x}_i - \mu_i)^2 + 2(s_i - \sigma_i)^2}{(\bar{x}_i - \mu_i)^2 + 2(s_i + \sigma_i)^2} \right\}^{1/2}.$$

Si se define:

$$(8) \quad D^2 = \min_{\mu, \sigma_1, \sigma_2} \{ d^2[(\bar{x}_1, \bar{x}_2, s_1, s_2), (\mu, \mu, \sigma_1, \sigma_2)] \},$$

intuitivamente, (8) representa el grado de discrepancia, en términos de distancia geodésica, entre los datos muestrales (resumidos por $\bar{x}_1, \bar{x}_2, s_1, s_2$) y el conjunto de todos los puntos paramétricos $(\mu, \mu, \sigma_1, \sigma_2)$ que satisfacen la hipótesis nula de igualdad de medias. Valores grandes de D^2 proporcionan evidencias contra esta hipótesis. Por

³Comúnmente llamada distancia de Rao en honor a su introductor en el ámbito estadístico (Rao, 1945)

tanto, el siguiente procedimiento de rechazo: «rechazar H_0 si $D^2 \geq c_\alpha$ », donde c_α debe satisfacer

$$(9) \quad P\{D^2 \geq c_\alpha | H_0\} = \alpha,$$

producirá un test «geodésico» (en el sentido de Burbea y Oller, (Burbea y Oller, 1989)) con un nivel de significación α para la hipótesis nula de igualdad de medias.

El test estadístico D^2 (para más información véase Vegas y Ocaña, 1996; Vegas 1996) no tiene una forma cerrada, quedando en función de un sistema de ecuaciones

$$(10) \quad \begin{cases} \sigma_1^2 = s_1^2 + \frac{(\bar{x}_1 - \mu)^2}{2} \\ \sigma_2^2 = s_2^2 + \frac{(\bar{x}_2 - \mu)^2}{2} \\ \frac{\sigma_1}{\sigma_2} = \frac{n_1 \cosh^{-1}(\sigma_1/s_1)}{n_2 \cosh^{-1}(\sigma_2/s_2)} \end{cases}$$

que debe ser resuelto, numéricamente, para obtener los valores $\hat{\mu}$, $\hat{\sigma}_1$ y $\hat{\sigma}_2$ que minimizan d^2 en (8), requerido para calcular $D^2 = d^2[(\bar{x}_1, \bar{x}_2, s_1, s_2), (\hat{\mu}, \hat{\sigma}_1, \hat{\sigma}_2)]$. Esto hace que sea muy difícil de resolver el problema distribucional de calcular c_α . Por consiguiente, se opta por aproximar su valor mediante bootstrap paramétrico. Es decir, se genera un número $B(= 1000)$ de remuestras del mismo tamaño, $n_1 + n_2$, del vector original de datos $x = (x_{11}, \dots, x_{1n_1}; x_{21}, \dots, x_{2n_2})$. Cada una de las remuestras $\mathbf{x}^*$ se produce por la generación de n_1 valores independientes e idénticamente distribuidos de una distribución normal de media

$$\bar{x} = \frac{n_1 \bar{x}_1 + n_2 \bar{x}_2}{n_1 + n_2}$$

y desviación estándar s_1 , junto con n_2 valores de una distribución normal con la misma media común, $\bar{x}$, y desviación estándar s_2 :

$$\mathbf{x}^* = (x_{11}^*, \dots, x_{1n_1}^*; x_{21}^*, \dots, x_{2n_2}^*).$$

Para cada remuestra $\mathbf{x}^*$, el valor correspondiente del test estadístico, D_*^2 , se calcula de la misma manera que D^2 y la probabilidad $\text{Prob}\{D^2 \geq D^2(\mathbf{x}) | H_0\}$ se estima por medio de:

$$P^* = \frac{\#[D_*^2 \geq D^2(\mathbf{x})] + 1}{B + 1}.$$

Para un nivel de significación fijo α la hipótesis nula se rechaza ($Y = 1$) si $P^* < \alpha$, y se acepta ($Y = 0$) en caso contrario.

La estimación P^* de $Prob\{D^2 \geq D^2(\mathbf{x}) \mid H_0\}$, en lugar de la frecuencia relativa $(\#[D_*^2 \geq D^2(\mathbf{x})]/B)$, garantiza que el nivel de significación bajo la distribución bootstrap sea realmente α en el sentido de que $Prob\{P^* \leq \alpha \mid N(\bar{x}, s_1, s_2)\} \leq \alpha$ (Dwass, 1957). No obstante este hecho no garantiza que el verdadero nivel de significación $Prob\{P^* \leq \alpha \mid N(\mu, \sigma_1, \sigma_2)\}$ sea realmente α .

D. Reducción de la varianza para variables de respuesta dicotómica en simulación

Supongamos que en cada réplica de una simulación de Monte Carlo se obtiene un valor de una variable de respuesta Y con distribución de Bernoulli. Por ejemplo, Y puede ser el resultado final de un test que acepta o rechaza una hipótesis, o un intervalo de confianza el cual incluye o no el verdadero valor del parámetro. Normalmente se realizará la simulación para estimar la esperanza $p_{1.} = E(Y)$, por ejemplo para estimar la verdadera probabilidad de rechazo de la hipótesis nula o la verdadera probabilidad de recubrimiento. El estimador insesgado más obvio es la frecuencia relativa $\hat{p}_{1.}$ del suceso $\{Y = 1\}$. Asumamos que, juntamente con un valor de Y , en cada réplica de simulación se produce un valor de otra variable dicotómica C , correlacionada con Y , cuya esperanza $p_{.1} = E(C) = P(C = 1)$ sea conocida. Por ejemplo C puede ser un test paramétrico cuya curva de potencia sea conocida, relacionado con el nuevo test Y bajo estudio de Monte Carlo o un intervalo de confianza relacionado cuyo verdadero recubrimiento sea conocido.

Al realizar el proceso de simulación se obtienen n pares de datos (Y_k, C_k) que se pueden resumir en forma de una tabla de contingencia 2×2 :

Tabla 1

	$C = 0$	$C = 1$	
$Y = 0$	n_{00}	n_{01}	$n_{0.}$
$Y = 1$	n_{10}	n_{11}	$n_{1.}$
	$n_{.0}$	$n_{.1}$	n

generada de acuerdo con las siguientes probabilidades p_{ij}

Tabla 2

	$C = 0$	$C = 1$	
$Y = 0$	p_{00}	p_{01}	$p_{0.}$
$Y = 1$	p_{10}	p_{11}	$p_{1.}$
	$p_{.0}$	$p_{.1}$	1

En estas condiciones, se puede obtener el estimador máximo verosímil de $p_{1.}$ que es

$$(11) \quad \tilde{p}_{1.} = p_{.0} \frac{n_{10}}{n_{00} + n_{10}} + p_{.1} \frac{n_{11}}{n_{01} + n_{11}} .$$

Se ha de notar que este estimador, $\tilde{p}_{1.}$, es válido sólo para valores distintos de cero de $n_{.j} = n_{0j} + n_{1j}$. Por tanto, condicionado por el suceso $\{0 < n_{.0} < n\}$.

Al aplicar los resultados clásicos de estimación máximo verosímil (Rothery, 1982) y del método delta (Ocaña y Vegas, 1995; Vegas 1996 donde se demuestra también que (11) coincide con el estimador que se obtendría mediante la técnica de reducción de la varianza conocida como «variables de control») se muestra que este estimador es asintóticamente normal y que su varianza puede ser estimada mediante

$$(12) \quad \hat{\sigma}_R^2 = \frac{1}{n} \left\{ \tilde{p}_{1.}(1 - \tilde{p}_{1.}) - \frac{(\tilde{p}_{1.}p_{.1} - \tilde{p}_{11})^2}{(1 - p_{.1})p_{.1}} \right\} .$$

El estimador (11) tiene como propiedades (Ocaña y Vegas, 1995; Vegas 1996): ser insesgado y con varianza nunca mayor a la varianza de la frecuencia relativa de $\{Y = 1\}$ e igual a

$$(13) \quad \text{var}(\tilde{p}_{1.}) = p_{00}p_{10}E\{n_{.0}^{-1}\} + p_{01}p_{11}E\{n_{.1}^{-1}\}$$

donde $E\{n_{.i}^{-1}\}$ es la esperanza de la inversa de la variable binomial $n_{.i}$, truncada a $\{0 < n_{.0} < n\}$. Además, $\tilde{p}_{1.}$ es asintóticamente eficiente.

En Ocaña y Vegas (1995) se demuestra que el estimador (12) tiene sesgo negativo mientras que el estimador

$$(14) \quad \hat{\sigma}_U^2 = \tilde{p}_{00}\tilde{p}_{10} \frac{E\{n_{.0}^{-1}\}}{1 - E\{n_{.0}^{-1}\}} + \tilde{p}_{01}\tilde{p}_{11} \frac{E\{n_{.1}^{-1}\}}{1 - E\{n_{.1}^{-1}\}}$$

es insesgado. Sin embargo, la esperanza $E\{n_{.i}^{-1}\}$ no tiene forma cerrada aunque se puede calcular numéricamente. Si se emplea la aproximación $E\{n_{.i}^{-1}\} \simeq (np_{.i} - (1 - p_{.i}))^{-1}$ (Grab and Savage, 1954) se obtiene un nuevo estimador

$$(15) \quad \hat{\sigma}_{GS}^2 = \frac{\tilde{p}_{00}\tilde{p}_{10}}{np_{.0} - (1 - p_{.0}) - 1} + \frac{\tilde{p}_{01}\tilde{p}_{11}}{np_{.1} - (1 - p_{.1}) - 1}$$

que es prácticamente insesgado y de cálculo más sencillo que (14). Por tanto, este último se escoge para ser utilizado como estimador de la varianza de $\hat{p}_1$ en los estudios de simulación.

Utilizando esta técnica se logran unas reducciones de la varianza importantes, de hasta un 95%, en comparación con las que se obtiene con el estimador habitual, frecuencia relativa ($\hat{p}_1$), con un incremento computacional y complejidad del diseño bajo. Un factor importante para conseguir una alta reducción de la varianza en la variable de estudio Y se basa en escoger adecuadamente la variable de control C de tal manera que exista una alta correlación entre ambas variables.

REFERENCIAS

- [1] **Amari, S-I, Barndorff-Nielsen, O.E., Kass, R.E., Lauritzen, S.L. and Rao, C.R.** (1987). *Differential geometry in statistical inference*. Volume 10 of *Lecture Notes-Monograph Series*, Institute of Mathematical Statistics, Hayward, California.
- [2] **Aspin, A.A.** (1948). «An examination and further development of a formula arising in the problem of comparing two mean values». *Biometrika*, **35**, 88–96.
- [3] **Atkinson, C. and Mitchell, A.F.S.** (1981). «Rao's distance measure». *Sankhyā*, **43**, A:345–365
- [4] **Buehler, R.J.** (1959). «Some validity criteria for statistical inferences». *Ann. Math. Statist.*, **30**, 845–867.
- [5] **Burbea, J. and Rao, C.R.** (1982). «Entropy differential metric, distance and divergence measures in probability spaces: a unified approach». *J. Multivariate Anal.*, **12**, 575–596.
- [6] **Burbea, J. and Oller, J.M.** (1989). *On Rao distance asymptotic distribution*. preprint series 67, Universitat de Barcelona, June.
- [7] **Cochran, W.G.** (1964). «Approximate significance levels of the Behrens-Fisher test». *Biometrics.*, **20**, 191–195.
- [8] **Cochran, W.G. and Cox, G.M.** (1950). *Experimental Designs*. John Wiley and Sons, New York.
- [9] **Dwass, M.** (1957). «Modified randomization tests for nonparametric hypotheses». *Ann. Math. Stat.*, **28**, 181–187.
- [10] **Efron, B.** (1975). «Defining the curvature of a statistical problem (with application to second order efficiency) (with discussion)». *Ann. Statist.*, **3**, 1189–1242.
- [11] **Efron, B.** (1978). «The geometry of exponential families». *Ann. Statist.*, **6**, 362–376.

- [12] **Efron, B. and Hinkley, D.V.** (1978). «Assessing the accuracy of the maximum likelihood estimator: Observed versus expected Fisher information». *Biometrika*, **65**, 457–487.
- [13] **Fisher, R. A.** (1956). «On a test of significance in Pearson's Biometrika Tables (No. 11)». *J. R. Statist. Soc. Serie B*, **18**, 56–60.
- [14] **Grab, E.L. and Savage, R.** (1954). «Tables of the expected value of $1/X$ for positive Bernoulli and Poisson variables». *J. Amer. Statist. Ass.*, **49**, 169–177.
- [15] **Lee, A.F.S. and Gurland, J.** (1975). «Size and power of tests for equality of means of two normal populations with unequal variances». *J. Amer. Statist. Ass.*, **70**, 933–941.
- [16] **Lilliefors, H.W.** (1967). «On the Kolmogorov-Smirnov test for normality with mean and variance unknown». *J. Amer. Statist. Ass.*, **63**, 339–402.
- [17] **Mason, A.L. and Bell, C.B.** (1986). «New Lilliefors and Srinivasan tables with applications». *Commun. Statist.- Simula.*, **15(2)** 451–477.
- [18] **Oost, B.A., Veldhuyzen, B., Timmermans, A.P.M. and Sixma, J.J.** (1983). «Increased urinary β -thromboglobulin excretion in diabetes assayed with a modified RIA kit-technique». *Thrombosis and Haemostasis*, **49**, 18–20.
- [19] **Pearson, E.S. and Hartley, H.O. (eds.)** (1976). *Biometrika Tables for Statisticians*. Volume 1, Cambridge University Press, Cambridge.
- [20] **Ocaña, J. and Vegas, E.** (1995). «Variance reduction for Bernoulli response variables in simulation». *Computational Statistics and Data Analysis*, **19**, 631–640.
- [21] **Rao, C.R.** (1945). «Information and accuracy attainable in the estimation of statistical parameters». *Bull. Calcutta Math. Soc.*, **37**, 81–91.
- [22] **Rothery, P.** (1982). «The use of control variates in Monte Carlo estimation of power». *Appl. Statist.*, **31**, 125–129.
- [23] **Scheffé, H.** (1970). «Practical solutions to the Behrens-Fisher problem». *J. Amer. Statist. Ass.*, **65**, 1501–1508.
- [24] **Vegas, E.** (1996). *Optimización en estudios de Monte Carlo en Estadística: aplicaciones al contraste de hipótesis*. Ph.D. Thesis, Universitat de Barcelona.
- [25] **Vegas, E. and Ocaña, J.** (1996). «Variance reduction in the study of a new test concerning the Behrens-Fisher problem». *International Journal of Computer Simulation*, in press.
- [26] **Welch, B.L.** (1938). «The significance of the difference between two means when the population variances are unequal». *Biometrika*, **29**, 350–362.
- [27] **Welch, B.L.** (1947). «The generalization of 'Students' problem when several different population variances are involved». *Biometrika*, **34**, 28–35.

ENGLISH SUMMARY

THE BEHRENS-FISHER PROBLEM IN BIOMEDICAL RESEARCH. CRITICAL ANALYSIS OF A CLINICAL STUDY BY SIMULATION

ESTEBAN VEGAS LOZANO*

Universitat de Barcelona

This paper reviews the Behrens-Fisher problem. Inferential foundations associated with the difficulty of its resolution are discussed and the most common practical solutions are exposed, together with a new solution based on concepts of differential geometry. Next, a critical study of biomedical research is presented in which the true probabilities of error are different from the expected values since probable differences between the variances are ignored. In this research the null hypothesis of equality of mean was rejected ($p < 0.01$), although, the true probability of type I error for values close to sample values may be different from nominal value (0.05). With this purpose, a Monte Carlo study is performed to obtain these estimations according to use of the t test for equality of means or other solutions more appropriate for Behrens-Fisher problem. In this simulation study a specific variance-reduction technique for dichotomous response variables such as statistical tests (to accept or reject the null hypothesis) is used. This technique is shown briefly and the design of the simulation is illustrated.

Keywords: Monte Carlo simulation, Variance reduction, Statistical curvature, Rao's distance, Bootstrap.

AMS Classification: 65C05

*Estepan Vegas Lozano. Departament d'Estadística. Facultat de Biologia. Universitat de Barcelona. Diagonal 645, 08028 Barcelona. Espanya.

–Received january 1997.

–Accepted may 1997.

1. INTRODUCTION

The Behrens-Fisher problem, i. e., testing for equality of means of two normal populations without making any assumption about variances, is a complex problem, without a known optimal solution.

The difficulty is reflected, on an abstract level, in the geometry of certain parametrical families of density. Thus, the parametrical submodel associated with the null hypothesis of equality of means is a curved exponential family of order 4. Therefore, this causes a breakdown of very nice properties for estimation, testing, and other inference problems associated with flat exponential families (see Efron (1975), Efron (1978) and Efron and Hinkley (1978)).

Some common practical solutions are: Cochran and Cox test, Welch test and Welch-Aspin test (see Scheffé (1970), Lee and Gurland (1975)). Another alternative is the geodesic test described in Vegas and Ocaña (1996) and Vegas (1996), based on differential geometry. The proposed statistic, D^2 , shows the discomformity level, in terms of geodesic distance, between the values obtained in the point estimations of the average and standard deviation of the sample $(\bar{x}_1, \bar{x}_2, s_1, s_2)$, and the set of all the parametric points in accordance with the null hypothesis of equality of means with arbitrary standard deviations $(\mu, \mu, \sigma_1, \sigma_2)$.

2. CRITICAL ANALYSIS OF A BIOMEDICAL STUDY

The data belong to a more extensive study about thrombosis (Oost *et al.* (1983)). This review is based on the measures obtained on urinary β -thromboglobulin excretion in 12 normal patients and 12 diabetic patients. The null hypothesis of no difference between means was rejected ($p < 0.01$) using the **t** test. The true probability of type I error for values close to sample values may be different since probable differences between the variances ($p < 0.01$ using the F test of variance comparison) are ignored.

With this purpose, a Monte Carlo study is performed to obtain the estimations of the true type I error probabilities when the **t** test for equality of means or other solutions (more appropriate for Behrens-Fisher problem) are used. Likewise, the effect of different variance quotients (1.40, 16.88 and 4.86, which correspond to the extreme values and middle point of the 95% confidence interval for variance quotients respectively) is analysed.

The precision of the estimations in the Monte Carlo study has been improved by the use of a variance-reduction technique, suitable for dichotomous response variables. This technique is described in Ocaña and Vegas (1995) and Vegas (1996)). In each replication of the simulation, besides the value of the dichotomous response variable

Y (in this case $Y = 1$ if the null hypothesis in the statistic test studied is rejected, and $Y = 0$ otherwise) an additional dichotomous control variable (C) is obtained. C should be a variable correlated with Y with known expectation. In the present study, C corresponds to the final outcome of the $\mathbf{t}$ test of comparison of means for equal variances.

Common random variates are used to induce the correlation required between the tests under study ($\mathbf{t}$ incorrectly used over samples from normal populations with different variances ($t(\mathbf{x})$), Cochran and Cox test ($t_{cc}(\mathbf{x})$), Welch test ($t_w(\mathbf{x})$) and geodesic test ($D^2(\mathbf{x})$)) with the «control» $\mathbf{t}$ test ($t(\mathbf{v})$), evaluated over samples generated according to equal variances (see fig. 1). More specifically, from n_1 standard normal *iid* values, $z_1 = (z_{11}, \dots, z_{1n_1})$, independent of other n_2 standard normal *iid* values, $z_2 = (z_{21}, \dots, z_{2n_2})$, are obtained the $n_1 + n_2$ values $x = (x_{11}, \dots, x_{1n_1}; x_{21}, \dots, x_{2n_2})$ using the transformation $x_{ij} = \sigma_i z_{ij} + \mu_i$, $i = 1, 2$ y $j = 1, \dots, n_i$. These samples are used to evaluate the four tests under different configurations of means and variances. The same transformation is applied again, but now with common $\sigma^2 = \frac{n_1\sigma_1^2 + n_2\sigma_2^2}{n_1 + n_2}$, an intermediate value between σ_1^2 and σ_2^2 . These new $n_1 + n_2$ values, $v = (v_{11}, \dots, v_{1n_1}; v_{21}, \dots, v_{2n_2})$, are used to evaluate the control $\mathbf{t}$ test ($t(v)$).

In this way, the $n_1 + n_2$ values of v fulfil the conditions of $\mathbf{t}$ test applicability ($t(v)$): they come from two normal populations with common variance, and therefore, they may be used as a control variable with known power (that is, expectation).

3. CONCLUSION

The geodesic and Welch tests maintain the nominal significance level of 0.05 under all variance ratios and, additionally, they are very similar. On the other hand, the Cochran and Cox test is very conservative under similar variances and the $\mathbf{t}$ test becomes worse as the variance quotients diverge from unity (see table 1 and figure 2).

Therefore, in the study of (Oost *et al.*(1983)), the true p-values do not correspond to the nominal ones, the null hypothesis of equality of means is rejected falsely more frequently than expected and this tendency increases with the variance quotient (see fig. 2). This means that the $\mathbf{t}$ test is not recommendable under conditions of extreme difference between variances or when the statistic values are on the border between the two possible decisions (to accept or reject the null hypothesis). In any case, the final conclusions of the paper (not the «p-values») are «correct», since they are consistent with the results obtained by all statistical techniques that are more appropriate than the $\mathbf{t}$ test.

On the other hand, the variance-reduction technique used to study the probability of type I error by Monte Carlo simulation may be used in other similar studies with a suitable effectiveness and a minimum effort in the complexity of the simulation design.

Secció docent i problemes

SECCIÓ DOCENT I PROBLEMES

La “SECCIÓ DOCENT I PROBLEMES” a la revista QÜESTIÓ té l'objectiu d'incloure una secció on es publiquen articles de caire docent, difícilment publicables en revistes de recerca. A cada número de QÜESTIÓ s'inclourà d'un a tres problemes i les solucions es donaran en el número següent.

Els lectors poden, si ho volen, proposar problemes amb les solucions pertinents i enviar-los a QÜESTIÓ, que farà una selecció i en publicarà els més adequats, fent la corresponent referència a l'autor.

També seran ben rebudes solucions alternatives a les propostes fetes per l'autor dels problemes; l'editorial es reservarà, però, el dret a publicar-les.

PROBLEMAS TIPO EN TEORÍA DE LA DECISIÓN

RAMÓN ALONSO SANZ*

Universidad Politécnica de Madrid

El objetivo de este artículo es ilustrar las técnicas y los conceptos básicos de la Teoría de la Decisión. Para lograr este objetivo se ha adoptado el problema tipo quizá más sencillo: el problema de la inversión. Éste queda caracterizado por dos decisiones alternativas (invertir o no invertir) y dos estados de la naturaleza (apreciación y depreciación). La información adicional adopta la forma de opinión de un experto. El problema se ha resuelto en forma extensa y normal. Las consecuencias se suponen expresadas en su forma primaria, como costes de oportunidad y como pérdidas. El árbol de decisión y la solución minimax también son considerados.

Type problems in Decision Theory

Keywords: Decision theory, problems.

Clasificación AMS: 62C05

*Ramón Alonso Sanz. ETSI Agrónomos. Unidad de Estadística. Universidad Politécnica de Madrid. C. Universitaria. 28040 Madrid. **e-mail:** ralonso@ccupm.upm.es

—Article rebut el febrer de 1995.

—Acceptat el desembre de 1996.

1. INTRODUCCIÓN

Los ejercicios de este artículo son casos particulares del planteado por la disyuntiva: invertir (d_1) o no invertir (d_2) una cierta cantidad x , en un proceso en el que se puede ganar (θ_1) o perder (θ_2) la cantidad y , con probabilidades $p(\theta_1) \equiv p$, $p(\theta_2) \equiv 1 - p$.¹

	θ_1	θ_2
d_1	$x + y$	$x - y$
d_2	x	x

La confianza del decisor en el experto al que se puede consultar se cuantifica mediante las verosimilitudes de acierto: $p(z_1/\theta_1) = p(z_2/\theta_2) = \pi$, siendo Z la variable aleatoria que designa la opinión de experto.

En los problemas 1.* se fija π a un valor que supone considerar al experto como tal: $\pi = 0,75$, siendo p un parámetro. Por el contrario, en los problemas 2.* , $p = 0,5$, valor que hace equivalentes ambas decisiones, siendo π el parámetro. En los problemas 3.* , tanto las probabilidades *a priori* como las verosimilitudes de acierto son parámetros. En el problema 4 se considera un contexto de incertidumbre total.

Índice de problemas

$\pi = 1/2$ $p = 3/4$

1.1.A	2.1.A	forma	extensa	con consecuencias primarias
1.1.B	2.1.B	«	«	« costes de oportunidad
1.2.A	2.2.A	forma	normal	« consecuencias primarias
1.2.B	2.2.B	«	«	« costes de oportunidad
1.2.C	2.2.C	«	«	« pérdidas

p y π parámetros

3.1	forma	normal	con costes de oportunidad
3.2	«	«	« pérdidas
3.3	árbol de decisión con consecuencias primarias		
3.4	«	«	« costes de oportunidad

π parámetro

4	incertidumbre total
---	---------------------

2. FUNDAMENTOS TEÓRICOS

Sea el problema de decisión (D, Θ, X) : un decisor ha de elegir una de entre m decisiones posibles (D); el resultado final (X) depende de su elección y de un conjunto de factores que él no controla, denominados genéricamente *Naturaleza* (Θ). Supondremos que ésta se puede presentar en n formas o *estados*.

		ESTADOS DE LA NATURALEZA				
		θ_1	$\dots$	θ_j	$\dots$	θ_n
DECISIONES	d_1					
	$\vdots$					
	d_i			x_{ij}		
	$\vdots$					
	d_m					

Si se decide d_i y la naturaleza se presenta en el estado θ_j , la consecuencia es x_{ij} ; en su *forma primaria* consideraremos que es una *ganancia*.

2.1. Decisión en ambiente de incertidumbre parcial o riesgo

Las probabilidades asociadas a los estados de la naturaleza, $p(\theta_j)$, miden el *grado de creencia* en su verificación. Supondremos que el *criterio* de decisión es el de la maximización del valor medio.

Análisis en forma extensa

$$d^* = d_i / \left\{ \max_i \left(V_i = \sum_{j=1}^n x_{ij} p(\theta_j) \right) \right\} = \frac{V_{\text{sin I}}}{\text{Valor medio sin Información}}$$

Decisión de Bayes

• Decisión con información adicional

Supongamos que, antes de tomar la decisión el decisor pide información (opinión) a un experto. Sea Z la variable aleatoria que describe la opinión del experto y sea $Z = \{z_1, \dots, z_k, \dots, z_n\}$ su dominio, donde z_k indica que el experto opina que se dará θ_k .

Ante la opinión del experto, el decisor modifica las probabilidades iniciales, *a priori*; los valores *a posteriori* se obtienen aplicando el teorema de Bayes:

$$p(\theta_j/z_k) = \frac{p(z_k/\theta_j) p(\theta_j)}{p(z_k)}$$

$\swarrow$ *a posteriori* $\swarrow$ *a priori*
 $\nwarrow$ *p(z_k/\theta_j)*

Verosimilitudes: miden el grado de confianza del decisor en el experto

El criterio de elección seguirá siendo el de maximización del valor medio, pero utilizando ahora las probabilidades *a posteriori*. Es decir, supuesto $Z = z_k$:

$$d_i^* = d_i / \max_i \left\{ V_i = \sum_{j=1}^n x_{ij} \underbrace{\frac{p(z_k/\theta_j) p(\theta_j)}{p(z_k)}}_{p(\theta_j/z_k)} \right\}$$

- *Valor medio del proceso de decisión con información*

$$V_{\text{conI}} = \sum_{k=1}^n \left(\max_i \left[\sum_{j=1}^n x_{ij} \frac{p(z_k/\theta_j) p(\theta_j)}{p(z_k)} \right] \right) p(z_k) = \sum_{k=1}^n \left(\max_i \left[\sum_{j=1}^n x_{ij} p(z_k/\theta_j) p(\theta_j) \right] \right)$$

Valor de la información: $V_{\text{deI}} = V_{\text{conI}} - V_{\text{sinI}}$

Propiedad: $V_{\text{deI}} \geq 0$.

- *Información Perfecta*: Sería la que aportaría un «experto perfecto». Aquél que verificase (siempre en opinión del decisor):

$$p(z_k/\theta_k) = 1 \quad \Leftrightarrow \quad p(\theta_k/z_k) = 1.$$

$$\text{Se cumple}^2 \quad V_{\text{conIP}} = \sum_{j=1}^n \left(\max_i x_{ij} \right) p(\theta_j). \quad \text{Ver Problema 1.1.A.}$$

En el extremo opuesto, restringiendo el problema a $n = 2$, el decisor puede creer que el «experto» siempre falla, y asignar las verosimilitudes: $p(z_2/\theta_1) = p(z_1/\theta_2) =$

1 $\Leftrightarrow p(\theta_2/z_1) = p(\theta_1/z_2) = 1$. En este caso, la regla de decisión es, naturalmente, hacer lo contrario de lo que opina el «experto».

- *Información irrelevante*

Si todas las verosimilitudes son iguales, $V_{\text{conI}} = V_{\text{sinI}}^3 \Leftrightarrow V_{\text{deI}} = 0$.

Costes de oportunidad

$$y_{ij} = \left(\max_i x_{ij} \right) - x_{ij}$$

Un mecanismo alternativo de resolver un problema de decisión, generalmente más cómodo, se basa en la *minimización del coste de oportunidad medio*⁴

$$d^* = d_i / \left(\min_i \left\{ C(d_i) = \sum_{j=1}^n y_{ij} p_j \right\} \right) = C_{\text{sinI}}$$

Se cumple (ver nota):

$$(1) \quad C_{\text{sinI}} = V_{\text{conIP}} - V_{\text{sinI}} = V_{\text{deIP}}$$

Con opinión del experto,

$$(2) \quad C_{\text{conI}} = V_{\text{conIP}} - V_{\text{conI}}^5$$

$$\text{De la diferencia [1] - [2]:} \quad V_{\text{deI}} = C_{\text{sinI}} - C_{\text{conI}}$$

- *Pérdidas (valores opuestos a las ganancias)*

$$\ell_{ij} = -x_{ij}$$

$$(3) \quad d^* = d_i / \underbrace{\min_i \left\{ L(d_i) = \sum_{j=1}^n \ell_{ij} p_j \right\}}_{L_{\text{sinI}}} = \underbrace{\min_i \left\{ -\sum_{j=1}^n x_{ij} p_j \right\}}_{-V_{\text{sinI}}} \equiv \max_i \{V_i\}$$

Análogamente:

$$(4) \quad L_{\text{conI}} = -V_{\text{conI}}.$$

$$\text{De la diferencia [3] - [4] resulta:} \quad V_{\text{deI}} = l_{\text{sinI}} - L_{\text{conI}}.$$

En los problemas resueltos con la matriz de consecuencias en forma de pérdidas, se ha supuesto que a todos los elementos de dicha matriz en su forma primaria se les ha restado la cantidad x . De forma general, si se efectúa una traslación en k unidades:

$$x'_{ij} = x_{ij} - k,$$

resulta:

$$d^* = d_i / \left(\max_i \left(\sum_{j=1}^n x'_{ij} p(\theta_j) \right) \right) = V \sin I' = V \sin I - k.$$

Análogamente: $V \cos I' = V \cos I - k$; de la diferencia de igualdades

$$V \text{de} I = V \cos I' - V \sin I'.$$

Análisis en forma normal

Se basa en la evaluación de las **estrategias** o **reglas de decisión**:

$$D = \{\sigma(z)/\delta : Z \longrightarrow D\}$$

$$\delta^* = \delta / \max_{\delta} \left(V(\delta) = \sum_{j=1}^n V(\delta/\theta_j) p(\theta_j) \right)$$

$$V(\delta/\theta_j) = \sum_{k=1}^n V(\delta/\theta_j, z_k) p(z_k/\theta_j)$$

En los problemas planteados con costes de oportunidad o pérdidas, la operación de *minimización* ha de sustituir a la de *maximización*.

2.2. Decisión en ambiente de incertidumbre total

El decisor no tiene ninguna apreciación de la probabilidad de los estados de la naturaleza.

Supondremos planteado el problema con las consecuencias como pérdidas ℓ_{ij} (problema (D, Θ, L)) y que el criterio de decisión es el *minimax*, en virtud del cual, la **decisión no aleatorizada** *minimax* es:

$$d^* = d_i / \min_i (\max_j \ell_{ij})$$

Decisiones aleatorizadas

Una distribución de probabilidad sobre D define una decisión aleatorizada. Así, $\Delta = \begin{pmatrix} d_1 & d_2 & \dots & d_m \\ \pi_1 & \pi_2 & \dots & \pi_m \end{pmatrix}$. Identificaremos Δ y Π .

La *pérdida media* asociada a la decisión aleatorizada Δ supuesto θ_j es

$$L_j(\Delta) = L(\Delta/\theta_j) = \sum_i \pi_i \ell_{ij}$$

Δ^* es una decisión *minimax* si:

$$\Delta^* = \Delta / \underbrace{\min_{\Delta} \left(\max_j L(\Delta/\theta_j) \right)}_v = \max_j L(\Delta^*/\theta_j) = V \quad (\text{valor minimax})$$

Métodos de obtención de Δ^*

① Programación Lineal

$V = \min V$	$\min v$	$\min \sum_{i=1}^m \pi'_i$
$L(\Delta/\theta_j) = \sum_{i=1}^m \pi_i \ell_{ij} \leq v$	$\sum_{i=1}^m \overbrace{\pi'_i}^{\pi_{ij}/v} \ell_{ij} \leq 1$	$\sum_{i=1}^m \pi'_i \ell_{ij} \leq 1$
$j = 1, 2, \dots, n$	$j = 1, 2, \dots, n$	$j = 1, 2, \dots, n$
$\pi_i \geq 0 \quad \forall i \quad 1' \pi = 1$	$\sum_{i=1}^m \pi'_i = \frac{1}{v}$	$\pi'_i \geq 0 \quad \forall i$
		(supuesto v positivo)

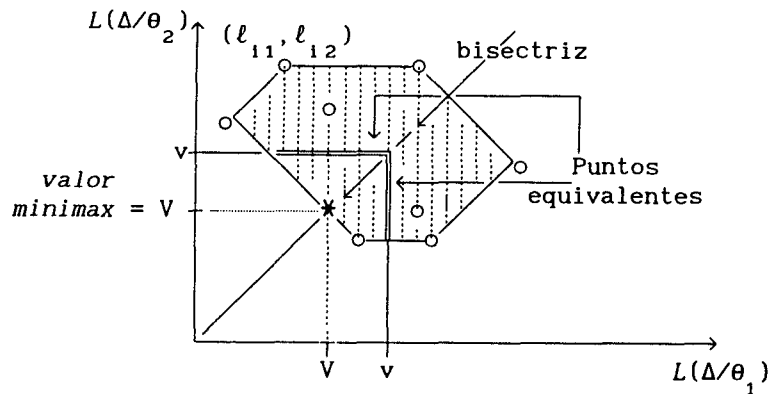
② Método gráfico (supuesto $n = 2$).

El **conjunto de pérdidas** asociado a Δ es $L(\Delta) = (L_1(\Delta), L_2(\Delta))$. Los conjuntos de pérdidas del conjunto de decisiones aleatorizadas forman la **región de pérdidas**. Si se aceptan únicamente decisiones puras, la región de pérdidas ($\mathcal{G}$) está formada por los puntos ($\circ$): $L(d_i) = (\ell_{i1}, \ell_{i2})$. Pero si se considera la posibilidad de aleatorizar las decisiones, problema extendido $(\mathcal{D}, \Theta, \Lambda)$, la *región de pérdidas* (G) está formada

por la *envolvente convexa* de G . Son **equivalentes** los puntos de la región de pérdidas que comparten el mismo valor máximo de la pérdida media: v ; así son equivalentes (al valor de su mayor coordenada, v), los puntos de la frontera de todo conjunto de puntos de la forma $\{(x,y)/x \leq v, y \leq v\}$. Designando por $N(v)$ a dicho conjunto frontera, el punto *minimax* se encuentra en la «última» intersección de $N(v)$ y G :

$$V = \min v / N(v) \cap G \neq \emptyset.$$

Gráficamente:



3. PROBLEMAS RESUELTOS

Enunciado común:

Sea el problema de decisión con matriz de consecuencias:

	θ_1	θ_2
d_1	$x+y$	$x-y$
d_2	x	x

$y > 0 \quad p \equiv p(\theta_1)$

El decisor pide opinión a un experto, que le merece una confianza cuantificada por las verosimilitudes de acierto:

$$p(z_1/\theta_1) = p(z_2/\theta_2) = \pi.$$

Problema 1.1.A

Sea $\pi = \frac{3}{4}$ en el enunciado común.

Obtener, aplicando el método de análisis en **forma extensa**:

- 1) a) la regla de decisión y el valor medio del proceso de decisión sin información [VsinI],
 b) el valor medio del proceso de decisión con información perfecta [VconIP],
 y
 c) el valor medio de la información perfecta [VdeIP]. Representar gráficamente $V_{sinI}(p)$, $V_{conIP}(p)$ y $V_{deIP}(p)$.
- 2) La regla de decisión y el valor medio del proceso de decisión contando con la información (opinión) del experto [VconI].
- 3) El valor medio de la información del experto [VdeI].
- 4) Representar gráficamente, en un cuadro conjunto, las expresiones de $V_{sinI}(p)$ y $V_{conI}(p)$. Destacar en dicho cuadro la expresión de $V_{deI}(p)$.

Solución:

- 1) a) **Sin información:**

$$d^* = d_i / \left\{ \max_i \left(V_i = \sum_{j=1}^2 x_{ij} p(\theta_j) \right) \right\} = V_{sinI}$$

$$\left. \begin{aligned} V_1 &= (x+y)p + (x-y)(1-p) = 2yp + x - y \\ V_2 &= xp + x(1-p) = x \end{aligned} \right\}$$

$$V_1 = V_2 \quad \Leftrightarrow \quad p_e = \frac{1}{2}, \quad V_e = x$$

$d^* = \begin{cases} d_2 & \left[0, \frac{1}{2} \right] \\ d_1 & \left[\frac{1}{2}, 1 \right] \end{cases}$	$\frac{p}{\left[0, \frac{1}{2} \right]}$	V_{sinI}
	$\frac{1}{2}, 1$	x $2yp + x - y$

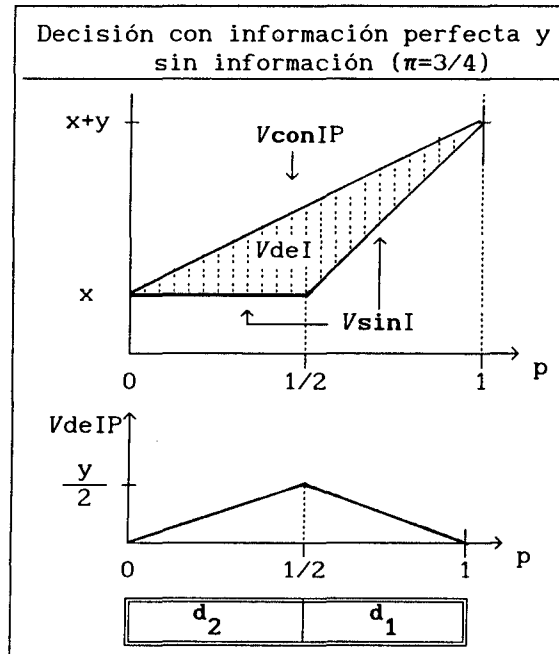
b) *Con información perfecta:*

$$V_{\text{conIP}} = \sum_{j=1}^2 \left(\max_i x_{ij} \right) p(\theta_j) = (x+y)p + x(1-p) = yp + x$$

c) *Valor medio de la información perfecta:*

$$V_{\text{deIP}} = V_{\text{conIP}} - V_{\text{sinI}}$$

$$V_{\text{deIP}} = \begin{cases} yp & \text{si } p \leq 1/2 \\ -yp + y & \text{si } p \geq 1/2 \end{cases}$$



2) *Con información parcial:*

$$\boxed{Z = z_k} \quad d^* = d_i / \max_i \left\{ V(d_i) = \sum_{j=1}^2 x_{ij} \underbrace{\frac{p(z_k/\theta_j)p(\theta_j)}{p(z_k)}}_{p(\theta_j/z_k)} \right\}$$

$$\boxed{Z = z_1}$$

$$\left. \begin{aligned} V_1 &= (x+y) \frac{p \frac{3}{4}}{p(z_1)} + (x-y) \frac{(1-p) \frac{1}{4}}{p(z_1)} = \frac{1}{p(z_1)} \frac{1}{4} ((2x+4y)p + x-y) \\ V_2 &= x \frac{p \frac{3}{4}}{p(z_1)} + x \frac{(1-p) \frac{1}{4}}{p(z_1)} = \frac{1}{p(z_1)} \frac{1}{4} (2xp + x) \end{aligned} \right\} d^* = \begin{cases} d_2 & \left[0, \frac{1}{4}\right] \\ d_1 & \left[\frac{1}{4}, 1\right] \end{cases}$$

$$V_1 = V_2 \quad \Leftrightarrow \quad p_e = \frac{1}{4}, \quad V_e = \frac{1}{p(z_1)} \frac{1}{4} \frac{3}{2} x$$

$$\boxed{Z = z_2}$$

$$\left. \begin{aligned} V_1 &= (x+y) \frac{p \frac{1}{4}}{p(z_2)} + (x-y) \frac{(1-p) \frac{3}{4}}{p(z_2)} = \frac{1}{p(z_2)} \frac{1}{4} (-(2x+4y)p + 3(x-y)) \\ V_2 &= x \frac{p \frac{1}{4}}{p(z_2)} + x \frac{(1-p) \frac{3}{4}}{p(z_2)} = \frac{1}{p(z_2)} \frac{1}{4} (-2xp + 3x) \end{aligned} \right\} d^* = \begin{cases} d_1 & \left[0, \frac{3}{4}\right] \\ d_2 & \left[\frac{3}{4}, 1\right] \end{cases}$$

$$V_1 = V_2 \quad \Leftrightarrow \quad p_e = \frac{3}{4}, \quad V_e = \frac{1}{p(z_2)} \frac{1}{4} \frac{7}{2} x$$

Las **estrategias** (o reglas de decisión) posibles son:

$$\delta_{12} = \begin{cases} d_1 & \text{si } Z = z_1 \\ d_2 & \text{si } Z = z_2 \end{cases} \quad \delta_{21} = \begin{cases} d_2 & \text{si } Z = z_1 \\ d_1 & \text{si } Z = z_2 \end{cases} \quad \delta_{11} = d_1 \quad \forall z \quad \delta_{22} = d_2 \quad \forall z$$

decisión inducida
por el experto

decisión opuesta a
la inducida por el
experto

caso omiso al experto

Las estrategias óptimas del problema son:

$$V_{\text{conI}} = \sum_{k=1}^2 \max_i \left\{ \sum_{j=1}^2 x_{ij} p(z_k/\theta_j) p(\theta_j) \right\}$$

$d^* = \begin{cases} \delta_{22} & \begin{bmatrix} p \\ 0, \frac{1}{4} \end{bmatrix} \\ \delta_{12} & \begin{bmatrix} 1 \\ \frac{1}{4}, \frac{3}{4} \end{bmatrix} \\ \delta_{11} & \begin{bmatrix} 3 \\ \frac{3}{4}, 1 \end{bmatrix} \end{cases}$	<table style="width: 100%; border-collapse: collapse;"> <tr> <th style="text-align: left; border-bottom: 1px solid black; padding: 5px;">V_{conI}</th> <th></th> </tr> <tr> <td style="padding: 5px;">x</td> <td style="padding: 5px;">$= \frac{1}{4} ((2x\cancel{p} + x) + (-2x\cancel{p} + 3x))$</td> </tr> <tr> <td style="padding: 5px;">$yp + x - \frac{y}{4}$</td> <td style="padding: 5px;">$= \frac{1}{4} ((2x + 4y)p + x - y) + (-2x\cancel{p} + 3x)$</td> </tr> <tr> <td style="padding: 5px;">$2yp + x - y$</td> <td style="padding: 5px;">$= \frac{1}{4} (((2x + 4y)p + x - y) + (-2x + 4y)p + 3(x - y))$</td> </tr> </table>	V_{conI}		x	$= \frac{1}{4} ((2x\cancel{p} + x) + (-2x\cancel{p} + 3x))$	$yp + x - \frac{y}{4}$	$= \frac{1}{4} ((2x + 4y)p + x - y) + (-2x\cancel{p} + 3x)$	$2yp + x - y$	$= \frac{1}{4} (((2x + 4y)p + x - y) + (-2x + 4y)p + 3(x - y))$
V_{conI}									
x	$= \frac{1}{4} ((2x\cancel{p} + x) + (-2x\cancel{p} + 3x))$								
$yp + x - \frac{y}{4}$	$= \frac{1}{4} ((2x + 4y)p + x - y) + (-2x\cancel{p} + 3x)$								
$2yp + x - y$	$= \frac{1}{4} (((2x + 4y)p + x - y) + (-2x + 4y)p + 3(x - y))$								

Obsérvese como para la resolución del problema, no ha sido necesario calcular las expresiones de $p(z_k)$.

En definitiva,

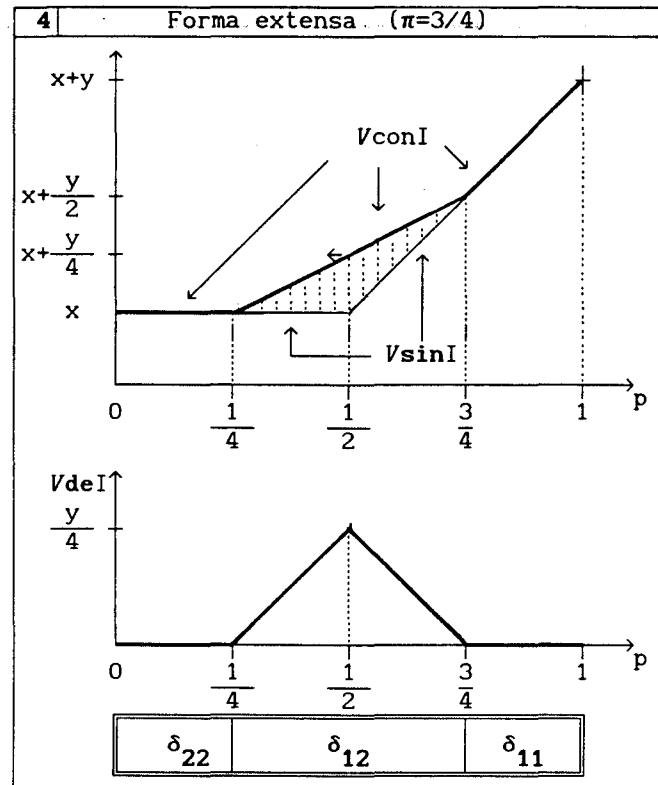
- a) si el decisor se encuentra *a priori* «muy seguro», adopta la decisión en la que confía,
- b) en caso contrario, adopta la decisión inducida por el experto.

En el primer supuesto, no aprecia en ninguna medida la opinión del experto ($V_{deI} = 0$), mientras que su valoración es creciente conforme se consideren valores de p próximos a 0,5. (ver apartado 3) y figura 4)).

3)

$V_{deI} = \begin{cases} 0 & \begin{bmatrix} p \\ 0, \frac{1}{4} \end{bmatrix} \\ yp - \frac{y}{4} & \begin{bmatrix} 1 \\ \frac{1}{4}, \frac{1}{2} \end{bmatrix} \\ -yp + \frac{3}{4}y & \begin{bmatrix} 1 \\ \frac{1}{2}, \frac{3}{4} \end{bmatrix} \\ 0 & \begin{bmatrix} 3 \\ \frac{3}{4}, 1 \end{bmatrix} \end{cases}$
--

4)



Problema 1.1.B

Resolver 1.1.A utilizando **costes de oportunidad**.

Solución:

Matriz de costes⁶:

$$y_{ij} = \left(\max_i x_{ij} \right) - x_{ij}$$

	θ_1	θ_2
d_1	0	y
d_2	y	0

1) Sin información:

$$d^* = d_i / \left(\min_i \left\{ C(d_i) = \sum_{j=1}^2 y_{ij} p(\theta_j) \right\} \right) = C \sin I$$

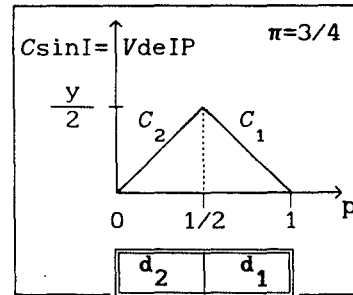
$$C(d_i) = V_{\text{conIP}} - C(d_i)$$

$$\left. \begin{aligned} C_1 &= 0p + y(1-p) = y - yp = (x+yp) - (x-y+2yp) \\ C_2 &= yp + 0(1-p) = yp = (x+yp) - (x) \end{aligned} \right\}$$

$$C_1 = C_2 \Leftrightarrow p_e = \frac{1}{2} \quad C_e = \frac{y}{2}$$

	p	$C \sin I$
d_2	$[0, \frac{1}{2}]$	yp
d_1	$[\frac{1}{2}, 1]$	$y-yp$

$$= V_{\text{deIP}} = V_{\text{conIP}} - V \sin I$$



2) Con información parcial:

$$\boxed{Z = z_k} \quad d^* = d_i / \min_i \left\{ C(d_i) = \sum_{j=1}^2 y_{ij} \frac{p(z_k / \theta_j) p(\theta_j)}{p(z_k)} \right\}$$

$$\boxed{Z = z_1}$$

$$\left. \begin{aligned} C_1 &= 0 \frac{p \frac{3}{4}}{p(z_1)} + y \frac{(1-p) \frac{1}{4}}{p(z_1)} = \frac{1}{p(z_1)} \frac{1}{4} (-yp + y) \\ C_2 &= y \frac{p \frac{3}{4}}{p(z_1)} + 0 \frac{(1-p) \frac{1}{4}}{p(z_1)} = \frac{1}{p(z_1)} \frac{1}{4} (3yp) \end{aligned} \right\} d^* = \begin{cases} d_2 \left[0, \frac{1}{4} \right] \\ d_1 \left[\frac{1}{4}, 1 \right] \end{cases}$$

$$C_1 = C_2 \Leftrightarrow p_e = \frac{1}{4}, \quad C_e = \frac{1}{p(z_1)} \frac{1}{4} \frac{3}{4} y$$

$$\boxed{Z = z_2}$$

$$\left. \begin{aligned} C_1 &= 0 \frac{p \frac{1}{4}}{p(z_2)} + y \frac{(1-p) \frac{3}{4}}{p(z_2)} = \frac{1}{p(z_2)} \frac{1}{4} (3y - 3yp) \\ C_2 &= y \frac{p \frac{1}{4}}{p(z_2)} + 0 \frac{(1-p) \frac{3}{4}}{p(z_2)} = \frac{1}{p(z_2)} \frac{1}{4} (yp) \end{aligned} \right\} d^* = \begin{cases} d_2 & \left[0, \frac{3}{4} \right] \\ d_1 & \left[\frac{3}{4}, 1 \right] \end{cases}$$

$$C_1 = C_2 \Leftrightarrow p_e = \frac{3}{4}, \quad C_e = \frac{1}{p(z_2)} \frac{1}{4} \frac{3}{4} y$$

Resulta:

$$C_{conI} = \sum_{k=1}^2 \min_i \left\{ \sum_{j=1}^2 y_{ij} p(z_k / \theta_j) p(\theta_j) \right\}$$

$$d^* = \begin{cases} \delta_{22} & \begin{bmatrix} p \\ 0, \frac{1}{4} \end{bmatrix} & \frac{C_{conI}}{yp} & = ((3yp) + (yp)) / 4 \\ \delta_{12} & \begin{bmatrix} 1 \\ \frac{1}{4}, \frac{3}{4} \end{bmatrix} & y/4 & = ((y - yp) + (yp)) / 4 \\ \delta_{11} & \begin{bmatrix} 3 \\ \frac{3}{4}, 1 \end{bmatrix} & y - yp & = (y - yp) + (3y - 3yp) / 4 \end{cases}$$

3) La expresión de VdeI calculada en 1.1.B se obtiene igualmente calculando la diferencia $C_{sinI} - C_{conI}$

4) Ver gráfico del problema 1.2.B.

Problema 1.2.A

Resolver 1.1.A aplicando el método de análisis en **forma normal**.

Solución:

$$\delta^* = \delta / \max_{\delta} \{V(\delta) = V(\delta/\theta_1)\theta_1 + V(\delta/\theta_2)(1-p)\}$$

$$V(\delta/\theta_1) = V(\delta/\theta_1, z_1)p(z_1/\theta_1) + V(\delta/\theta_1, z_2)p(z_2/\theta_1) = V(\delta/\theta_1, z_1)\frac{3}{4} + V(\delta/\theta_1, z_2)\frac{1}{4}$$

$$V(\delta/\theta_2) = V(\delta/\theta_2, z_1)p(z_1/\theta_2) + V(\delta/\theta_2, z_2)p(z_2/\theta_2) = V(\delta/\theta_2, z_1)\frac{1}{4} + V(\delta/\theta_2, z_2)\frac{3}{4}$$

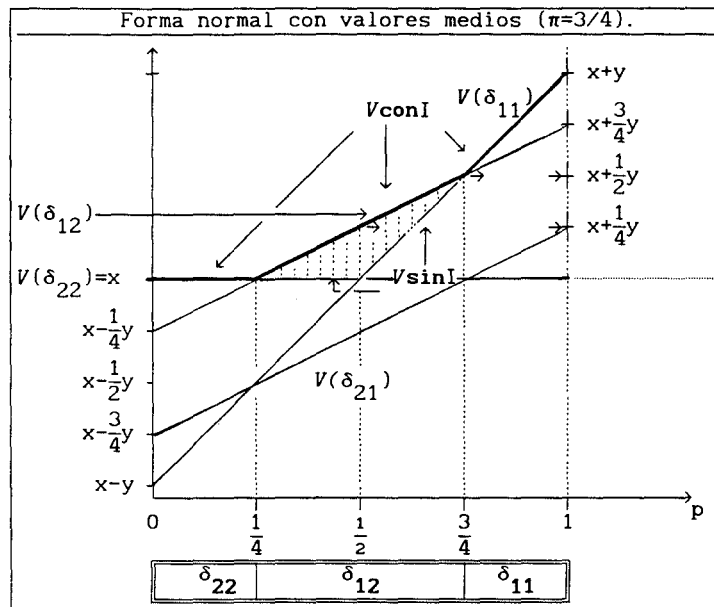
$$V(\delta_{12}/\theta_j, z_k) = x_{kj} \quad V(\delta_{21}/\theta_j, z_1) = x_{2j} \quad V(\delta_{21}/\theta_j, z_2) = x_{1j}$$

$$V(\delta_{11}/\theta_j, z_k) = x_{ij} \quad V(\delta_{22}/\theta_j, z_k) = x_{2j}$$

$$V(\delta_{12}) = \overline{yp + x - \frac{y}{4}} \quad V(\delta_{21}) = \overline{yp + x - \frac{3y}{4}} = V(\delta_{12}) - \frac{y}{2} \quad \Leftrightarrow \quad \delta_{21} \text{ dominada}$$

$$V(\delta_{11}) = \overline{2yp + x - y} \quad V(\delta_{22}) = \overline{x}$$

Resulta:

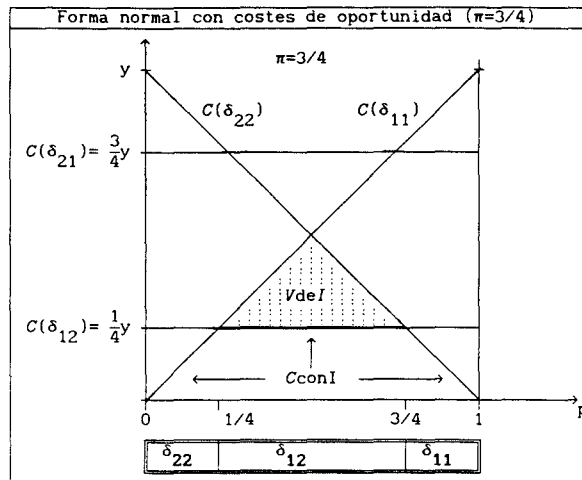


Problema 1.2.B

Resolver 1.2.A con **costes de oportunidad**.

Solución:

$$\delta^* = \delta / \min_{\delta} \left\{ C(\delta) = C(\delta/\theta_1)p + C(\delta/\theta_2)(1-p) = \underbrace{V_{\text{conIP}}}_{x+yp} - V(\delta) \right\}$$



$$C(\delta_{12}) = {}^{10} \overline{\frac{y}{4}p} \quad C(\delta_{21}) = {}^{11} \overline{\frac{3}{4}y = C(\delta_{12}) + \frac{y}{2}} \Leftrightarrow \delta_{21} \text{ dominada}$$

$$C(\delta_{11}) = {}^{12} \overline{y-yp} \quad C(\delta_{22}) = {}^{13} \overline{yp}$$

Problema 1.2.C

Resolver 1.2.A con **pérdidas**.

Solución:

La matriz de pérdidas, efectuada la traslación de magnitud x , resulta:

	θ_1	θ_2
d_1	$-y$	y
d_2	0	0

1) Sin información:

$$d^* = d_i / \left(\min_i \left\{ L(d_i) = \sum_{j=1}^2 \ell_{ij} p(\theta_j) \right\} \right) = L \sin I$$

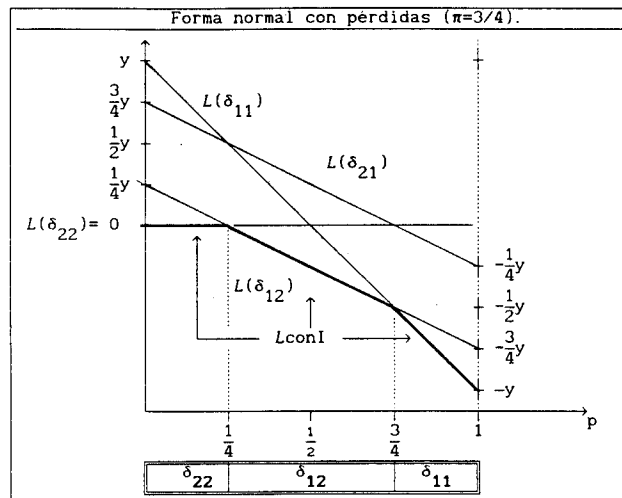
$$\left. \begin{array}{l} L_1 = -yp + y(1-p) = y - 2yp \\ L_2 = 0p + 0(1-p) = 0 \end{array} \right\} L_1 = L_2 \Leftrightarrow \begin{array}{l} p_e = \frac{1}{2} \\ L_e = 0 \end{array}$$

$$d^* = \begin{cases} d_2 & \begin{bmatrix} p \\ 0, \frac{1}{2} \end{bmatrix} & \frac{L \sin I}{0} \\ d_1 & \begin{bmatrix} \frac{1}{2}, 1 \end{bmatrix} & y - 2yp \end{cases}$$

2) $d^* = \delta / \min_{\delta} \{ L(\delta) = L(\delta/\theta_1)p + L(\delta/\theta_2)(1-p) \}$

$$L(\delta_{12}) = {}^{14} \boxed{-yp + y/4} \quad L(\delta_{21}) = {}^{15} \boxed{-yp + 3y/4} = L(\delta_{12}) + \frac{y}{2} \Leftrightarrow \delta_{21} \text{ dominada}$$

$$L(\delta_{11}) = {}^{16} \boxed{-2yp + y} \quad L(\delta_{22}) = \boxed{0} = -V(\delta_{22}) - x$$



Problema 2.1.A

Sea $p = \frac{1}{2}$ en el enunciado común.

Obtener, aplicando el método de análisis en **forma extensa**:

- 1) Obtener la regla de decisión y el valor medio del proceso de decisión sin contar con la información del experto $[V_{\sin I}]$, el valor medio del proceso contando con información perfecta $[V_{\text{conIP}}]$ y el valor medio de la información perfecta $[V_{\text{deIP}}]$.
- 2) La regla de decisión y el valor medio del proceso de decisión contando con la información (opinión) del experto $[V_{\text{conI}}(\pi)]$
- 3) El valor medio de la información del experto $[V_{\text{deI}}(\pi)]$.
- 4) Representar gráficamente, en un cuadro conjunto, las expresiones de $V_{\sin I}(\pi)$ y $V_{\text{conI}}(\pi)$. Destacar en dicho cuadro la expresión de $V_{\text{deI}}(\pi)$.

Solución:

1) Sin información:

$$\left. \begin{aligned} V_1 &= (x+y) \frac{1}{2} + (x-y) \frac{1}{2} = x \\ V_2 &= x \frac{1}{2} + x \frac{1}{2} = x \end{aligned} \right\} \Rightarrow \left\{ \begin{array}{l} V_{\sin I} = x \\ d_1 \text{ y } d_2 \text{ equivalentes} \end{array} \right.$$

Con información perfecta:

$$\begin{aligned} V_{\text{conIP}} &= (x+y) \frac{1}{2} + x \frac{1}{2} = x + \frac{y}{2}, \\ V_{\text{deIP}} &= V_{\text{conIP}} - V_{\sin I} = x + \frac{y}{2} - x = \frac{y}{2} \end{aligned}$$

2) Con información parcial:

$$\boxed{Z = z_1}$$

$$\left. \begin{aligned} V_1 &= (x+y) \frac{\frac{\pi}{2}}{p(z_1)} + (x-y) \frac{(1-\pi)\frac{1}{2}}{p(z_1)} = \frac{1}{p(z_1)} \frac{1}{2} (2y\pi + x - y) \\ V_2 &= x \frac{\frac{\pi}{2}}{p(z_1)} + x \frac{(1-\pi)\frac{1}{2}}{p(z_1)} = \frac{1}{p(z_1)} \frac{1}{2} x \end{aligned} \right\} d^* = \left\{ \begin{array}{l} d_2 \left[0, \frac{1}{2} \right] \\ d_1 \left[\frac{1}{2}, 1 \right] \end{array} \right.$$

$$V_1 = V_2 \quad \Leftrightarrow \quad \pi_e = \frac{1}{2} \quad V_e = \frac{1}{p(z_1)} \frac{1}{2} x$$

$$\boxed{Z = z_2}$$

$$\left. \begin{aligned} V_1 &= (x+y) \frac{(1-\pi)\frac{1}{2}}{p(z_2)} + (x-y) \frac{\pi\frac{1}{2}}{p(z_2)} = \frac{1}{p(z_2)} \frac{1}{2} (-2y\pi + x+y) \\ V_2 &= x \frac{(1-\pi)\frac{1}{2}}{p(z_2)} + x \frac{\pi\frac{1}{2}}{p(z_2)} = \frac{1}{p(z_2)} \frac{1}{2} x \end{aligned} \right\} d^* = \begin{cases} d_1 & \begin{bmatrix} \pi \\ 0, \frac{1}{2} \end{bmatrix} \\ d_2 & \begin{bmatrix} \frac{1}{2}, 1 \end{bmatrix} \end{cases}$$

$$V_1 = V_2 \quad \Leftrightarrow \quad \pi_e = \frac{1}{2} \quad V_e = \frac{1}{p(z_2)} \frac{1}{2} x$$

Resulta:

$$\boxed{d^* = \begin{cases} \delta_{21} & \begin{bmatrix} p \\ 0, \frac{1}{2} \end{bmatrix} \\ & \frac{1}{2} \\ \delta_{12} & \begin{bmatrix} \frac{1}{2}, 1 \end{bmatrix} \end{cases}} \quad \boxed{\begin{array}{c} \overline{V \text{ con I}} \\ -y\pi + x + \frac{y}{2} \\ x \\ y\pi + x - \frac{y}{2} \end{array}} = \begin{aligned} &= \frac{1}{2} (x + (-2y\pi + x + y)) \\ & \\ &= \frac{1}{2} ((2y\pi + x - y) + x) \end{aligned}$$

$$V \text{ con I}(\pi = 1) = V \text{ con IP} = y \times 1 + x - \frac{y}{2} = x + \frac{y}{2}.$$

$$3) \quad V \text{ de I} = \begin{pmatrix} -y\pi + x + \frac{y}{2} \\ x \\ y\pi + x + \frac{y}{2} \end{pmatrix} - \begin{pmatrix} x \\ x \\ x \end{pmatrix} = \begin{pmatrix} -y\pi + \frac{y}{2} \\ 0 \\ y\pi - \frac{y}{2} \end{pmatrix} \quad \begin{matrix} p \\ \begin{bmatrix} 0, \frac{1}{2} \end{bmatrix} \\ 1/2 \\ \begin{bmatrix} \frac{1}{2}, 1 \end{bmatrix} \end{matrix}$$

4) Ver figura del problema 2.2.A.

Problema 2.1.B

Resolver 2.1.A utilizando **costes de oportunidad**.

Solución:

1) Sin información:

$$\left. \begin{aligned} C_1 &= 0\frac{1}{2} + y\frac{1}{2} = \frac{y}{2} \\ C_2 &= y\frac{1}{2} + 0\frac{1}{2} = \frac{y}{2} \end{aligned} \right\} \Rightarrow \begin{cases} C \sin I = \frac{y}{2} \\ d_1 \quad y \quad d_2 \text{ equivalentes} \end{cases}$$

2) Con información:

$$Z = z_1$$

$$\left. \begin{aligned} C_1 &= 0\frac{\pi}{2} + y\frac{(1-\pi)}{2} = \frac{1}{p(z_1)}\frac{1}{2}(y-y\pi) \\ C_2 &= y\frac{\pi}{2} + 0\frac{(1-\pi)}{2} = \frac{1}{p(z_1)}\frac{1}{2}y\pi \end{aligned} \right\} d^* = \begin{cases} d_2 & \left[0, \frac{\pi}{2}\right] \\ d_1 & \left[\frac{1}{2}, 1\right] \end{cases}$$

$$C_1 = C_2 \Leftrightarrow \pi_e = \frac{1}{2} \quad C_e = \frac{1}{p(z_1)}\frac{1}{2}\frac{y}{2}$$

$$Z = z_2$$

$$\left. \begin{aligned} C_1 &= 0\frac{(1-\pi)}{2} + y\frac{\pi}{2} = \frac{1}{p(z_2)}\frac{1}{2}y\pi \\ C_2 &= y\frac{(1-\pi)}{2} + 0\frac{\pi}{2} = \frac{1}{p(z_2)}\frac{1}{2}(y-y\pi) \end{aligned} \right\} d^* = \begin{cases} d_1 & \left[0, \frac{\pi}{2}\right] \\ d_2 & \left[\frac{1}{2}, 1\right] \end{cases}$$

$$C_1 = C_2 \Leftrightarrow \pi_e = \frac{1}{2} \quad C_e = \frac{1}{p(z_2)}\frac{1}{2}\frac{y}{2}$$

Resulta:

$d^* = \begin{cases} \delta_{21} & \left[0, \frac{1}{2}\right] \\ & \frac{1}{2} \\ \delta_{12} & \left[\frac{1}{2}, 0\right] \end{cases}$	$\frac{p}{\quad}$	$\frac{C \sin I}{\quad}$	
	$y\pi$	$y\pi$	$= \frac{1}{2}(y\pi + y\pi)$
	$y - y\pi$	$y - y\pi$	$= \frac{1}{2}((y - y\pi) + (y - y\pi))$

4) Ver figura del problema 2.2.B.

Problema 2.2.A

Resolver 2.1.A aplicando el método de análisis en **forma normal**.

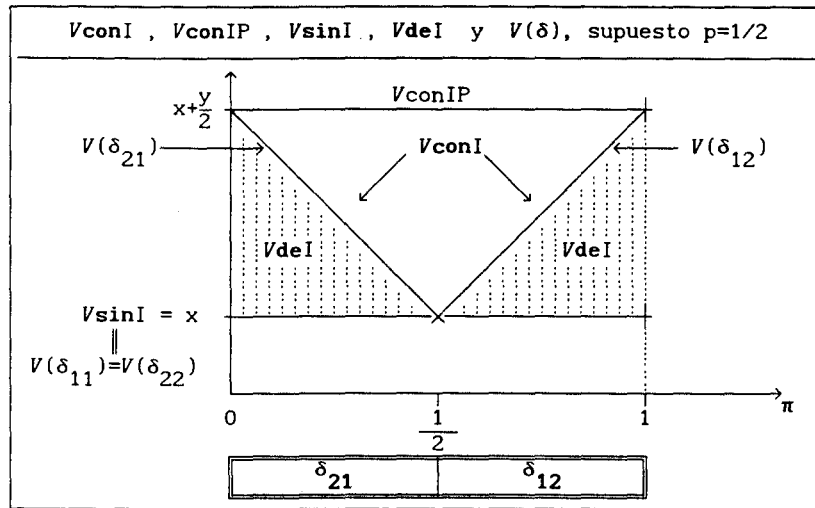
Solución:

$$\delta^* = \delta / \max_{\delta} \left\{ V(\delta) = V(\delta/\theta_1)p(\theta_1) + V(\delta/\theta_2)p(\theta_2) = \frac{1}{2} (V(\delta/\theta_1) + V(\delta/\theta_2)) \right\}$$

$$V(\delta/\theta_1) = V(\delta/\theta_1, z_1)p(z_1/\theta_1) + V(\delta/\theta_1, z_2)p(z_2/\theta_1) = V(\delta/\theta_1, z_1)\pi + V(\delta/\theta_1, z_2)(1 - \pi)$$

$$V(\delta/\theta_2) = V(\delta/\theta_2, z_1)p(z_1/\theta_2) + V(\delta/\theta_2, z_2)p(z_2/\theta_2) = V(\delta/\theta_2, z_1)(1 - \pi) + V(\delta/\theta_2, z_2)\pi$$

$$V(\delta_{12}) =^{17} \overline{y\pi + x - \frac{y}{2}} \quad V(\delta_{21}) =^{18} \overline{-y\pi + x + \frac{y}{2}} \quad V(\delta_{11}) =^{19} \overline{x} = V(\delta_{22})$$

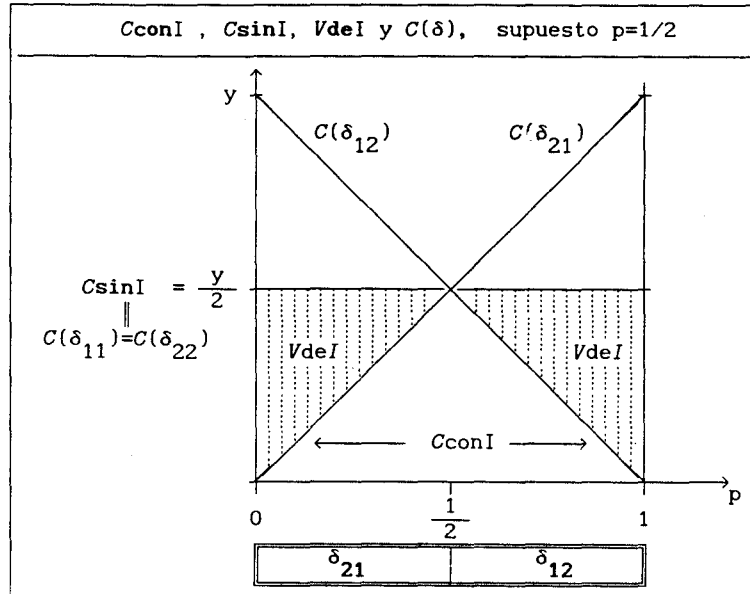


Problema 2.2.B Resolver 2.2.A con **costes de oportunidad**.

Solución:

$$C(\delta) = C(\delta/\theta_1)p(\theta_1) + C(\delta/\theta_2)p(\theta_2) = \frac{1}{2} (C(\delta/\theta_1) + C(\delta/\theta_2)) = VconIP - V(\delta)$$

$$C(\delta_{12}) =^{20} \overline{y - y\pi} \quad C(\delta_{21}) =^{21} \overline{y\pi} \quad C(\delta_{11}) =^{22} \overline{} \quad C(\delta_{22}) =^{23} \overline{y/2}$$



Las figuras precedentes resumen los resultados del tipo de problemas 2. Se observa en ellas una perfecta simetría en el valor de la información (V_{deI}). Para los casos extremos: $\pi = 1$ (experto) y $\pi = 0$ («antiexperto»), la información alcanza su máximo valor, ya que ambos son perfectamente útiles al decisor. El valor de la información es nulo para el «aexperto»: $\pi = 1/2$. Excepto en este caso, en el que se podría afirmar que «el experto está como el decisor»: $p = \pi = 1/2$, la opinión del experto, real o supuesto, tiene valor, bien para hacer lo que induce, bien para tomar la decisión contraria.

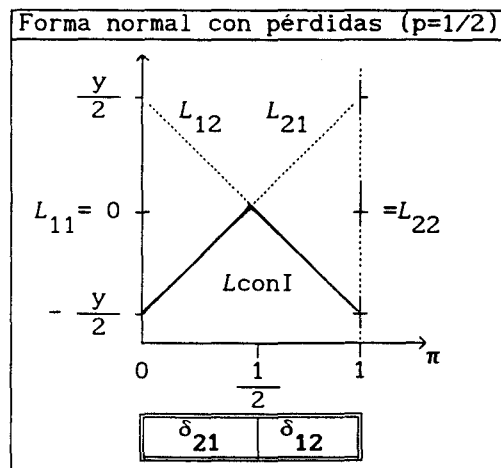
Problema 2.2.C

Resolver 2.2.B con **pérdidas**.

Solución:

$$L(\delta) = \frac{1}{2} (L(\delta/\theta_1) + L(\delta/\theta_2)) = -V(\delta) - x$$

$$L(\delta_{12}) = {}^{24} \left[\frac{y}{2} - y\pi \right] \quad L(\delta_{21}) = {}^{25} \left[y\pi - \frac{y}{2} \right] \quad L(\delta_{11}) = {}^{26} [0] = L(\delta_{22})$$



Problema 3.1

Considérese el problema planteado en su enunciado común.

Determinar las expresiones de los costes de oportunidad medios de las posibles estrategias y las funciones $CconI(p, \pi)$ y $Vdel(p, \pi)$.

Solución:

$$C(\delta) = C(\delta/\theta_1)p + C(\delta/\theta_2)(1-p)$$

$$C(\delta/\theta_1) = C(\delta/\theta_1, z_1)p(z_1/\theta_1) + C(\delta/\theta_1, z_2)p(z_2/\theta_1) = C(\delta/\theta_1, z_1)\pi + C(\delta/\theta_1, z_2)(1-\pi)$$

$$C(\delta/\theta_2) = C(\delta/\theta_2, z_1)p(z_1/\theta_2) + C(\delta/\theta_2, z_2)p(z_2/\theta_2) = C(\delta/\theta_2, z_1)(1-\pi) + C(\delta/\theta_2, z_2)\pi$$

Tomando del Problema 2.2.B las expresiones de $C(\delta/\theta, z)$, resulta:

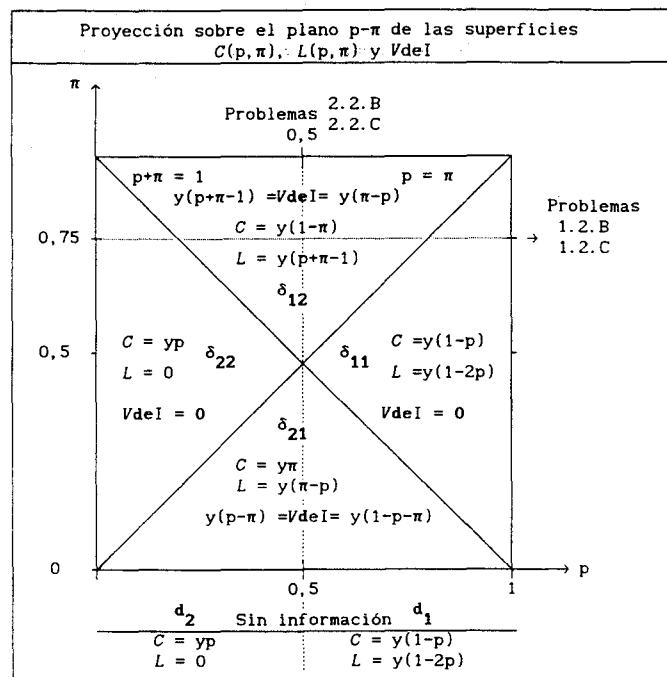
$$C(\delta_{12}) =^{27} \overline{y(1-\pi)} \quad C(\delta_{21}) =^{28} \overline{y\pi}$$

$$C(\delta_{11}) =^{29} \overline{y(1-p)} \quad C(\delta_{22}) =^{30} \overline{yp}$$

De 1.1.B:

$$d^* = \begin{cases} d_2 & \left[0, \frac{1}{2} \right] \\ d_1 & \left[\frac{1}{2}, 1 \right] \end{cases} \quad \begin{matrix} p \\ \frac{CsinI}{yp} \\ y-yp \end{matrix}$$

$$Vdel = CsinI - CconI$$



La superficie $C_{conI}(p,\pi)$ es una pirámide de base en el plano $C=0$ y de altura $y/2$ en $p=\pi=1/2$. Su proyección sobre el plano $p-\pi$ genera las cuatro regiones destacadas en la figura siguiente.

La forma de la función $VdeI$ queda caracterizada por las secciones visualizadas en los problemas 1.1.A y 2.2.A: es nula en las dos regiones en las que no influye el experto (δ_{11} y δ_{22}) y simétrica respecto a la vertical $p=0,5$ en el resto. Sobre esta línea se producen los máximos valores de $VdeI$ fijado π : $y(\pi-1/2)$ en δ_{12} e $y(-\pi+1/2)$ en δ_{21} . El máximo absoluto aparece cuando $\pi=1$ en el primer caso y $\pi=0$ en el segundo, es decir:

$$\max(VdeI(p,\pi)) = VdeI(1/2,1)) = VdeI(1/2,0)) = y/2.$$

Resultado ya obtenido en el problema 2.2.A; corresponde con una situación de máxima incertidumbre del decisor ($p=1/2$) y máxima seguridad de acierto/desacierto en el experto.

Problema 3.2

Resolver 3.1 con **pérdidas**.

Solución: Tomando del Problema 2.2.C las expresiones de $L(\delta/\theta, z)$, resulta:

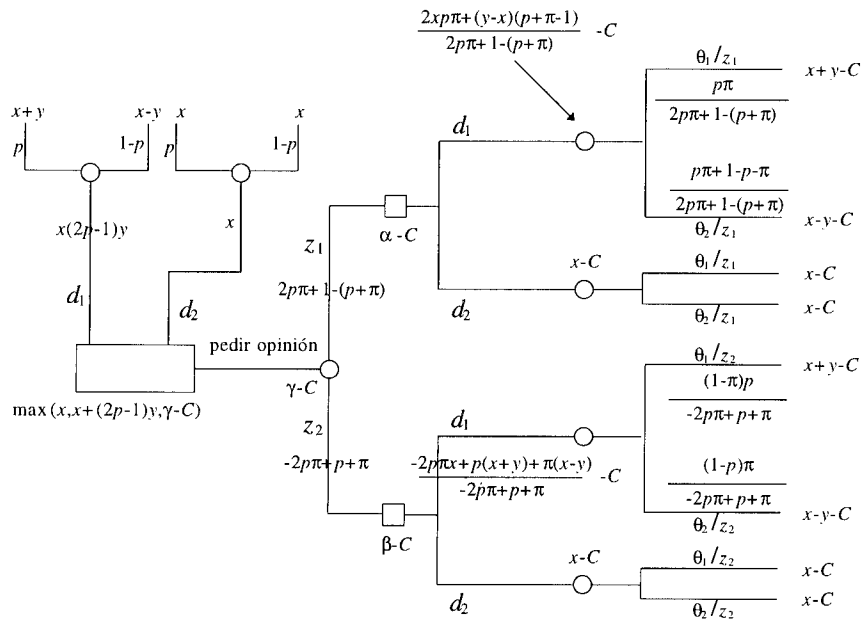
$$\begin{array}{ll} L(\delta_{12}) = {}^{31} \overline{y(1 - (p + \pi))} & L(\delta_{21}) = {}^{32} \overline{y(\pi - p)} \\ L(\delta_{11}) = {}^{33} \overline{y(1 - 2p)} & L(\delta_{22}) = \overline{0} \end{array}$$

La forma de la superficie $L\text{conI}(p, \pi)$ queda caracterizada por las secciones visualizadas en los problemas 1.2.C y 2.2.C. En la figura precedente aparecen también las pérdidas medias $L\text{conI} = L\sin I - V\text{deI}$.

Problema 3.3

Dibujar el árbol de decisión correspondiente al *Problema 3.1* planteado con consecuencias dadas en forma primaria. Discutir las decisiones a adoptar en el supuesto de que el experto pida C unidades monetarias.

Solución:



$$p(z_1) = 2p\pi + 1 - p - \pi \Rightarrow^{34} p(z_2) = -2p\pi + p + \pi$$

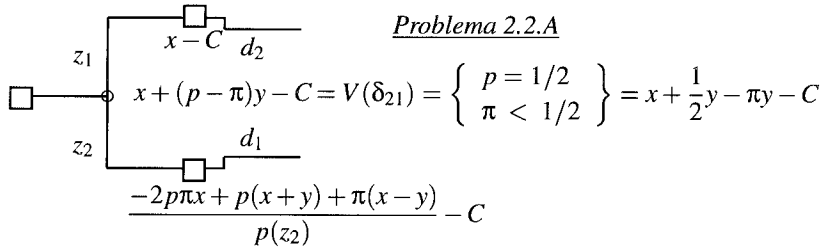
$$p(\theta_1/z_1) = \frac{p(z_1/\theta_1)p(\theta_1)}{p(z_1)} = \frac{\pi p}{2p\pi + 1 - (p + \pi)} \Rightarrow p(\theta_2/z_1) = \frac{p\pi + 1 - (p + \pi)}{2p\pi + 1 - (p + \pi)}$$

$$p(\theta_1/z_2) = \frac{p(z_2/\theta_1)p(\theta_1)}{p(z_2)} = \frac{(1 - \pi)p}{-2p\pi + p + \pi} \Rightarrow p(\theta_2/z_2) = \frac{-p\pi + \pi}{-2p\pi + p + \pi}$$

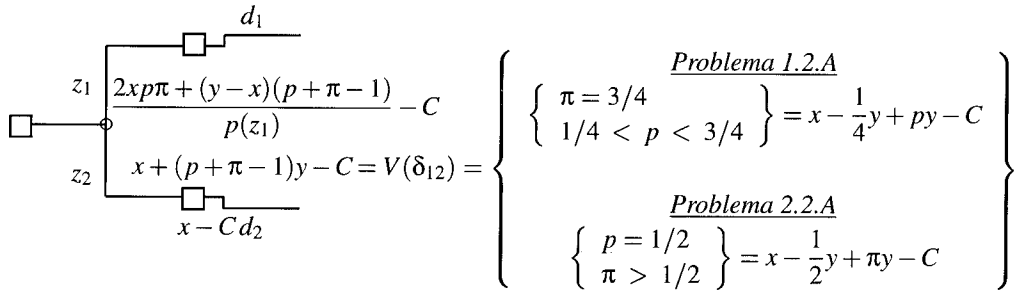
$$\frac{2xp\pi + (y - x)(p + \pi - 1)}{2p\pi + 1 - (p + \pi)} - C = x - C \Rightarrow \boxed{p + \pi = 1}$$

$$\frac{-2xp\pi + p(x + y) + \pi(x - y)}{-2p\pi + p + \pi} - C = x - C \Rightarrow \boxed{p = \pi}$$

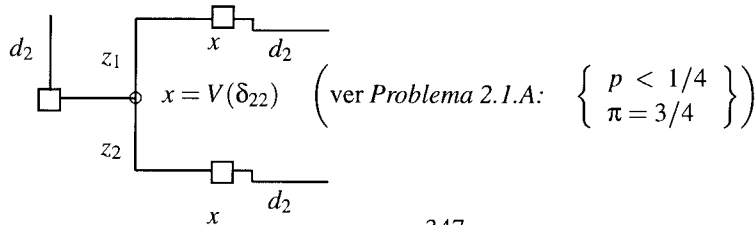
$$\frac{\pi < p, p + \pi < 1 \Rightarrow \delta_{21}}{C < \max\{(1 - (p + \pi))y, (p - \pi)y\}} \quad 35$$



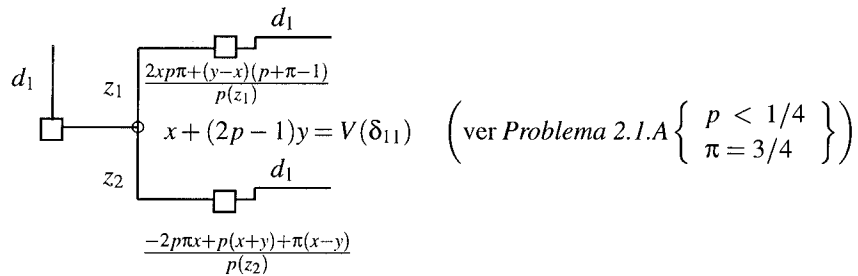
$$\frac{\pi > p, p + \pi > 1 \Rightarrow \delta_{12}}{C < \max\{(\pi - p)y, ((p + \pi) - 1)y\}} \quad 36$$



$$\frac{\pi > p, p + \pi > 1 \Rightarrow \delta_{22}}{C = 0} \equiv d_2 \quad \text{sin experto}^{37}$$



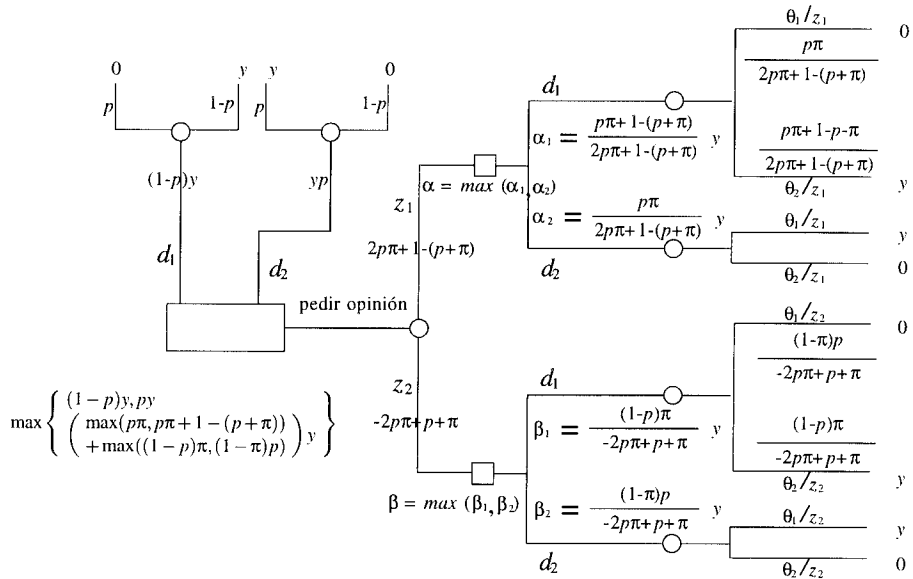
$$\frac{\pi < p, p + \pi > 1}{C = 0} \Rightarrow \delta_{11} \equiv d_1 \quad \text{sin experto}^{38}$$



Problema 3.4

Plantear el problema 3.3 utilizando **costes de oportunidad**.

Solución:



Problema 4

Resolver 3.2 en un contexto de **incertidumbre total**.

Solución:

Decisión no aleatorizada minimax

$$\min_i \left(\max_j \ell_{ij} \right) = \min(-y, 0) = 0 \Leftrightarrow d^* = d_2$$

Decisiones aleatorizadas

$$\Delta = \begin{pmatrix} \pi_1 & d_1 \\ \pi_2 & d_2 \end{pmatrix}$$

1) Programación lineal

Para asegurar que el parámetro v es positivo, se ha de tratar una matriz de pérdidas con todos sus valores no negativos. Sumando la cantidad y se consigue este objetivo. Al valor *minimax* obtenido en el problema transformado, $\mathcal{V}$, se le ha de restar y para obtener el del problema original.

$$\max z = \pi'_1 + \pi'_2$$

Forma estándar

Solución inicial:

$$\left. \begin{array}{l} 0\pi'_1 + y\pi'_2 \leq 1 \\ 2y\pi'_1 + y\pi'_2 \leq 1 \end{array} \right\} \Leftrightarrow \left. \begin{array}{l} 0\pi'_1 + y\pi'_2 + \pi'_3 = 1 \\ 2y\pi'_1 + y\pi'_2 + \pi'_4 = 1 \end{array} \right\} \quad \begin{array}{l} \pi'_1 = \pi'_2 = 0 \quad \pi'_3 = 1 \\ \pi'_4 = 1 \end{array}$$

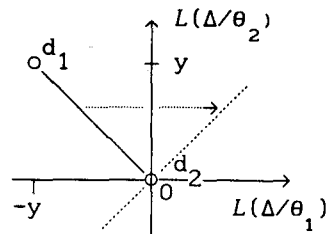
Método simplex:

π'_3	0	y	1	0	1	∞	θ
π'_4	$2y$	y	0	1	1	1	$\leftarrow$
	-1	-1	0	0	0	0	
	$\uparrow$						

π'_3	0	0	1	0	1	∞	θ
π'_1	1	$1/2$	0	$1/2y$	$1/2y$	1	$\leftarrow$
	0	$-1/2$	0	$1/2y$	$1/2y$	0	
	$\uparrow$						

π'_3	0	0	1	0	1	
π'_2	2	1	0	$1/y$	$1/y$	$\Rightarrow \pi_2 = 1$
	1	0	0	$1/y$	$1/y$	$\Rightarrow \mathcal{V} = y \Rightarrow V = 0$

2) Método gráfico



$$\Delta_{\minimax} = (0, 1) \equiv d_2$$

La decisión aleatorizada minimax es por tanto la pura d_2 : domina al resto. El valor *minimax* es nulo.

Estrategias aleatorizadas:

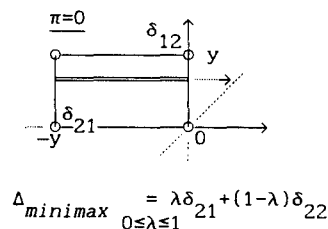
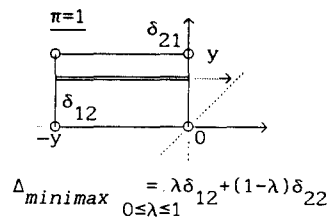
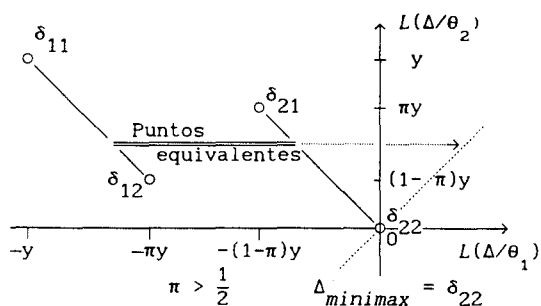
$$\Delta = \begin{pmatrix} \delta_{12} & \delta_{21} & \delta_{11} & \delta_{22} \\ \pi_1 & \pi_2 & \pi_3 & \pi_4 \end{pmatrix}$$

$$L(\Delta/\theta_j) = \sum_k \pi_k L(\delta_k/\theta_j) \quad \Delta^* = \Delta / \min_{\Delta} \left(\max_j L(\Delta/\theta_j) \right)$$

Del Problema 3.2:

$$L(\delta_{12}/\theta_1) = -y\pi \quad L(\delta_{12}/\theta_2) = y(1-\pi) \quad L(\delta_{11}/\theta_1) = -y = -L(\delta_{11}/\theta_2)$$

$$L(\delta_{21}/\theta_1) = -y(1-\pi) \quad L(\delta_{21}/\theta_2) = y\pi \quad L(\delta_{22}/\theta_1) = 0 = L(\delta_{22}/\theta_2)$$



En todos los supuestos, el valor *minimax* es nulo.

NOTAS

1. Este ejemplo es utilizado por Lindley, D.V. (1985). *Making Decisions*. Wiley.

$$2. V\text{conI} = \sum_{k=1}^n \left(\max_i \left[\sum_{j=1}^n x_{ij} p(z_k/\theta_j) p(\theta_j) \right] \right) = \sum_{k=1}^n \left(\max_i [x_{ik} p(\theta_k)] \right) = \underbrace{\sum_{k=1}^n \left(\left[\max_i x_{ik} \right] p(\theta_k) \right)}_{V\text{conIP}}$$

3. En efecto, si $p(z_k/\theta_j) = 1/n, \forall k, j,$:

$$V\text{conI} = \sum_{k=1}^n \left[\max_i \sum_{j=1}^n x_{ij} p(z_k/\theta_j) p(\theta_j) \right] = \sum_{k=1}^n \left[\max_i \sum_{j=1}^n x_{ij} \frac{1}{n} p(\theta_j) \right] =$$

$$= \frac{1}{n} \max_i \sum_j x_{ij} p_j = \max_i \sum_j x_{ij} p_j = V\text{sinI}$$

ya que ha desaparecido el subíndice k

$$4. C_i = \sum_{j=1}^n \left(\left(\max_i x_{ij} \right) - x_{ij} \right) p_j = \left(\sum_{j=1}^n \left(\max_i x_{ij} \right) p_j \right) - \left(\sum_{j=1}^n x_{ij} p_j \right) =$$

$$= V\text{conIP} - V(d_i) \Leftrightarrow \underbrace{\min_i (C(d_i))}_{C\text{sinI}} = V\text{conIP} - \underbrace{\max_i (V(d_i))}_{V\text{sinI}}.$$

$$5. Z = z_k \Leftrightarrow \min_i \left[\sum_{j=1}^n \left(\max_i x_{ij} \right) p(\theta_j/z_k) \right] = \left(\sum_{j=1}^n \left(\max_i x_{ij} \right) p(\theta_j/z_k) \right) - V\text{conI}_k$$

$$C\text{conI} = \sum_{k=1}^n \left(\left(\sum_{j=1}^n \left(\max_i x_{ij} \right) p(\theta_j/z_k) \right) - V\text{conI}_k \right) p(z_k) =$$

$$= \sum_{k=1}^n \left(\underbrace{\left(\sum_{j=1}^n \left(\max_i x_{ij} \right) p(\theta_j(z_k)) \right)}_{\text{}} - V\text{conI} \right) p(z_k)$$

$$\sum_{j=1}^n \left(\sum_{k=1}^n \left(\max_i x_{ij} \right) p(\theta_j/z_k) p(z_k) \right) = \sum_{j=1}^n \left(\left(\max_i x_{ij} \right) \underbrace{\sum_{k=1}^n p(\theta_j/z_k) p(z_k)}_{p(\theta_j)} \right) = V\text{conIP} p(z_k)$$

6. Al resolver el problema planteado con esta matriz de costes, queda resuelto el problema planteado con la misma matriz, en el supuesto de que las consecuencias fueran, originalmente, pérdidas. Este importante supuesto es analizado en el epígrafe 8.11 de De Groot (1970). *Optimal Statistical Decisions*. McGraw-Hill.

$$7. \quad V(\delta_{12}) = \left(x + \frac{3}{4}y\right)p + \left(x - \frac{1}{4}y\right)(1-p) = \boxed{yp + x - \frac{y}{4}}$$

$$V(\delta_{12}/\theta_1) = (x_{11} = x+y)\frac{3}{4} + (x_{21} = x)\frac{1}{4} = x + \frac{3}{4}y$$

$$V(\delta_{12}/\theta_2) = (x_{12} = x-y)\frac{1}{4} + (x_{22} = x)\frac{3}{4} = x + \frac{1}{4}y$$

$$8. \quad V(\delta_{21}) = \left(x + \frac{3}{4}y\right)p + \left(x + \frac{3}{4}y\right)(1-p) = \boxed{yp + x - \frac{3y}{4}}$$

$$V(\delta_{21}/\theta_1) = (x_{21} = x)\frac{3}{4} + (x_{11} = x+y)\frac{1}{4} = x + \frac{1}{4}y$$

$$V(\delta_{21}/\theta_2) = (x_{22} = x)\frac{1}{4} + (x_{12} = x-y)\frac{3}{4} = x - \frac{3}{4}y$$

$$9. \quad V(\delta_{11}) = (x+y)p + (x-y)(1-p) = \boxed{2yp + x - y}$$

$$V(\delta_{11}/\theta_1) = (x_{11} = x+y)\frac{3}{4} + (x_{11} = x+y)\frac{1}{4} = x+y$$

$$V(\delta_{11}/\theta_2) = (x_{12} = x-y)\frac{1}{4} + (x_{12} = x-y)\frac{3}{4} = x-y$$

$$10. \quad C(\delta_{12}) = \left(\frac{y}{4}\right)p + \left(\frac{y}{4}\right)(1-p) = \boxed{\frac{y}{4}p} = x + yp - \left(yp + x - \frac{y}{4}\right)$$

$$C(\delta_{12}/\theta_1) = (y_{11} = 0)\frac{3}{4} + (y_{21} = y)\frac{1}{4} = \frac{1}{4}y = C(\delta_{12}/\theta_2) = (y_{12} = y)\frac{1}{4} + (y_{22} = 0)\frac{3}{4}$$

$$11. \quad C(\delta_{21}) = \left(\frac{3}{4}y\right)p + \left(\frac{3}{4}y\right)(1-p) = \boxed{\frac{3}{4}y} = x + yp - \left(yp + x - \frac{y}{4}\right)$$

$$C(\delta_{21}/\theta_1) = (y_{21} = y)\frac{3}{4} + (y_{11} = 0)\frac{1}{4} = \frac{3}{4}y =$$

$$= C(\delta_{21}/\theta_2) = (y_{22} = 0)\frac{1}{4} + (y_{12} = y)\frac{3}{4}$$

$$12. \quad C(\delta_{11}) = 0(p) + (y)(1-p) = \boxed{y-yp} = x + yp - (2yp + x - y)$$

$$C(\delta_{11}/\theta_1) = (y_{11} = 0)\pi + (y_{11} = 0)(1-\pi) = 0$$

$$C(\delta_{11}/\theta_2) = (y_{21} = y)(1-\pi) + (y_{21} = y)\pi = y$$

$$13. \quad C(\delta_{22}) = y(p) + (0)(1-p) = \boxed{yp} = x + yp - x$$

$$C(\delta_{22}/\theta_1) = (y_{21} = y)\pi + (y_{21} = y)(1-\pi) = y$$

$$C(\delta_{22}/\theta_2) = (y_{22} = 0)(1-\pi) + (y_{21} = 0)\pi = 0$$

14. $L(\delta_{12}) = \left(-\frac{3}{4}y\right)p + \left(\frac{1}{4}y\right)(1-p) = \boxed{-yp + y/4} = -V(\delta_{12}) - x$
 $L(\delta_{12}/\theta_1) = (\ell_{11} = -y)\frac{3}{4} + (\ell_{21} = 0)\frac{1}{4} = -\frac{3}{4}y$
 $L(\delta_{12}/\theta_2) = (\ell_{12} = y)\frac{1}{4} + (\ell_{22} = 0)\frac{3}{4} = \frac{1}{4}y$
15. $L(\delta_{21}) = \left(-\frac{1}{4}y\right)p + \left(\frac{3}{4}y\right)(1-p) = \boxed{-yp + \frac{3}{4}y} = -VM(\delta_{21}) - x$
 $L(\delta_{21}/\theta_1) = (\ell_{21} = 0)\frac{3}{4} + (\ell_{11} = -y)\frac{1}{4} = -\frac{1}{4}y$
 $L(\delta_{21}/\theta_2) = (\ell_{22} = 0)\frac{1}{4} + (\ell_{12} = y)\frac{3}{4} = -\frac{3}{4}y$
16. $L(\delta_{11}) = (-y)p + (y)(1-p) = \boxed{-2yp + y} = -VM(\delta_{11}) - x$
 $L(\delta_{11}/\theta_1) = (\ell_{11} = -y)\pi + (\ell_{11} = -y)(1-\pi) = -y$ $L(\delta_{11}/\theta_2) = (\ell_{21} = y)(1-\pi) + (\ell_{21} = y)\pi = y$
17. $V(\delta_{12}) = \frac{1}{2} \left(2(y\pi + x - \frac{1}{2})\right) = \boxed{y\pi + x - \frac{1}{2}}$
 $V(\delta_{12}/\theta_1) = (x_{11} = x+y)\pi + (x_{21} = x)(1-\pi) = V(\delta_{12}/\theta_2) = (x_{12} = x-y)(1-\pi) + (x_{22} = x)\pi = yp + x - \frac{1}{2}$
18. $V(\delta_{21}) = \frac{1}{2} \left(2\left(-y\pi + x + \frac{1}{2}\right)\right) = \boxed{-y\pi + x + \frac{1}{2}}$
 $V(\delta_{21}/\theta_1) = (x_{21} = y)\pi + (x_{11} = x+y)(1-\pi) = V(\delta_{21}/\theta_2) = (x_{22} = x)(1-\pi) + (x_{12} = x-y)\pi = -yp + x - \frac{1}{2}$
19. $V(\delta_{11}) = \frac{1}{2} ((x+y) + (x-y)) = \boxed{x}$
 $V(\delta_{11}/\theta_1) = (x_{11} = x+y)\pi + (x_{11} = x+y)(1-\pi)$
 $V(\delta_{11}/\theta_2) = (x_{12} = x-y)(1-\pi) + (x_{12} = x-y)\pi$
20. $C(\delta_{12}) = \frac{1}{2} (2(y(1-\pi))) = \boxed{y-y\pi} = x + \frac{y}{2} - \left(y\pi + x - \frac{y}{2}\right) \quad V\text{conIP} \quad V(\delta_{12})$
 $C(\delta_{12}/\theta_1) = (y_{11} = 0)\pi + (y_{21} = y)(1-\pi) = C(\delta_{12}/\theta_2) = (y_{12} = y)(1-\pi) + (y_{22} = 0)\pi = y(1-\pi)$
21. $C(\delta_{21}) = \frac{1}{2} (2(y\pi)) = \boxed{y\pi} = x + \frac{y}{2} - \left(-y\pi + x + \frac{y}{2}\right)$
 $C(\delta_{21}/\theta_1) = (y_{21} = y)\pi + (y_{11} = 0)(1-\pi) = y\pi = C(\delta_{21}/\theta_2) = (y_{22} = 0)(1-\pi) + (y_{12} = y)\pi$
22. $C(\delta_{11}) = \frac{1}{2} (0+y) = \boxed{\frac{y}{2}} = x + \frac{y}{2} - x$
 $C(\delta_{11}/\theta_1) = (y_{11} = 0)\pi + (y_{11} = 0)(1-\pi) = 0$ $C(\delta_{11}/\theta_2) = (y_{21} = y)(1-\pi) + (y_{21} = y)\pi = y$

23. $C(\delta_{22}) = \frac{1}{2}(y+0) = \frac{y}{2} = x + \frac{y}{2} - x$
 $C(\delta_{22}/\theta_1) = (y_{11}=y)\pi + (y_{11}=0)(1-\pi) = y$ $C(\delta_{22}/\theta_2) = (y_{21}=0)(1-\pi) + (y_{21}=0)\pi = 0$
24. $L(\delta_{12}) = \frac{1}{2}((-y\pi) + (y(1-\pi))) = \boxed{\frac{y}{2} - y\pi}$
 $L(\delta_{12}/\theta_1) = (\ell_{11}=-y)\pi + (\ell_{21}=0)(1-\pi) = -y\pi$ $L(\delta_{12}/\theta_2) = (\ell_{12}=y)(1-\pi) + (\ell_{22}=0) = y(1-\pi)$
25. $L(\delta_{21}) = \frac{1}{2}((-y+y\pi) + (y\pi)) = \boxed{y\pi - \frac{y}{2}}$
 $L(\delta_{21}/\theta_1) = (\ell_{21}=0)\pi + (\ell_{11}=-y)(1-\pi) = -y(1-\pi)$ $L(\delta_{21}/\theta_2) = (\ell_{22}=0)(1-\pi) + (\ell_{12}=y) = y\pi$
26. $L(\delta_{11}) = \frac{1}{2}((-y) + (y)) = \boxed{0}$
 $L(\delta_{11}/\theta_1) = (\ell_{11}=-y)\pi + (\ell_{11}=-y)(1-\pi) = -y\pi$ $L(\delta_{11}/\theta_2) = (\ell_{21}=y)(1-\pi) + (\ell_{21}=y)\pi = y\pi$
27. $C(\delta_{12}) = C(\delta_{12}/\theta_1)p + C(\delta_{12}/\theta_2)(1-p) = y(1-\pi)p + y(1-\pi)(1-p) = \boxed{y(1-\pi)}$
28. $C(\delta_{21}) = C(\delta_{21}/\theta_1)p + C(\delta_{21}/\theta_2)(1-p) = y\pi p + y\pi(1-p) = \boxed{y\pi}$
29. $C(\delta_{11}) = C(\delta_{11}/\theta_1)p + C(\delta_{11}/\theta_2)(1-p) = 0p + y(1-p) = \boxed{y(1-p)}$
30. $C(\delta_{22}) = C(\delta_{22}/\theta_1)p + C(\delta_{22}/\theta_2)(1-p) = yp + 0(1-p) = \boxed{yp}$
31. $L(\delta_{12}) = L(\delta_{12}/\theta_1)p + L(\delta_{12}/\theta_2)(1-p) = -y\pi p + y(1-\pi) = y(1-(p+\pi))$
32. $L(\delta_{21}) = L(\delta_{21}/\theta_1)p + L(\delta_{21}/\theta_2)(1-p) = -y(1-\pi)p + y\pi(1-p) = y(\pi-p)$
33. $L(\delta_{11}) = L(\delta_{11}/\theta_1)p + L(\delta_{11}/\theta_2)(1-p) = -yp + y(1-p) = y(1-2p)$
34. $p(z_1) = p(z_1/\theta_1)p(\theta_1) + p(z_1/\theta_2)p(\theta_2) = p\pi + (1-\pi)(1-p) = p\pi + 1 - \pi - p + p\pi$
35. $V(\delta_{21}) - V(d_2) = (p-\pi)y > 0 \Rightarrow d_2$ sin experto descartada
 $V(\delta_{21}) - V(d_1) = (1-(p+\pi))y > 0 \Rightarrow d_1$ sin experto descartada
36. $V(\delta_{12}) - V(d_1) = (\pi-p)y > 0 \Rightarrow d_1$ sin experto descartada
 $V(\delta_{12}) - V(d_2) = (p+\pi-1)y > 0 \Rightarrow d_2$ sin experto descartada
37. $V(\delta_{22}) - V(d_1) = (1-2p)y > (1-(p+\pi))y > 0 \Rightarrow d_1$ sin experto descartada
38. $V(\delta_{11}) - V(d_2) = (2p-1)y > ((p+\pi)-1)y > 0 \Rightarrow d_2$ sin experto descartada

ENGLISH SUMMARY

TYPE PROBLEMS IN DECISION THEORY

RAMÓN ALONSO SANZ*

Universidad Politécnica de Madrid

This paper aims to illustrate the basic concepts and techniques underlying Decision Theory. To accomplish that, the somewhat simplest example is considered: the investment example. This problem is featured by

- a) two alternative decisions: to invest (d_1) or not to invest (d_2) the amount x , and*
- b) two uncertain events: appreciation (θ_1) measured by the gain y , and depreciation (θ_2) supposed also of extent y .*

Additional information has in the example the form of an expert opinion (Z). The problem is solved in the extensive and normal forms. Its decision tree is also specified. Consequences are supposed to be given in its primary form $(x + y, x - y, y)$, as regrets, and as losses. Decision tree and minimax solution are also obtained.

Keywords: Decision theory, problems.

AMS Classification: 62C05

*Ramón Alonso Sanz. ETSI Agrónomos. Unidad de Estadística. Universidad Politécnica de Madrid. C. Universitaria. 28040 Madrid. **e-mail:** ralonso@ccupm.upm.es

–Received february 1995.

–Accepted december 1996.

Extensive form

- Without additional Information

$$\begin{array}{l} d^* = d_i / \left\{ V_i = \max_i \left(\sum_{j=1}^n x_{ij} p(\theta_j) \right) \right\} = V_{\sin I} \\ \text{Bayes decision} \end{array} \quad \text{Expected Value without Inf.}$$

- With additional Information (expert opinion Z)

$$Z = z_k \quad \Leftrightarrow \quad d_i^* = d_i / \max_i \left\{ V_i = \sum_{j=1}^n x_{ij} \underbrace{\frac{p(z_k/\theta_j)p(\theta_j)}{p(z_k)}}_{p(\theta_j/z_k)} \right\}$$

Expected Value with Information

$$V_{\text{conI}} = \sum_{k=1}^n \left(\max_i \left[\sum_{j=1}^n x_{ij} p(z_k/\theta_j) p(\theta_j) \right] \right)$$

Expected Value of Information

$$V_{\text{deI}} = V_{\text{conI}} - V_{\sin I}$$

Perfect information ($p(z_k/\theta_k)$) :

$$\text{Expected Value with Perfect Information: } V_{\text{conIP}} = \sum_{j=1}^n \left(\max_i x_{ij} \right) p(\theta_j)$$

$$\text{Regrets: } y_{ij} = \left(\max_i x_{ij} \right) - x_{ij}$$

$$d^* = d_i / \left(\min_i \left\{ C(d_i) = \sum_{j=1}^n y_{ij} p_j \right\} \right) = C_{\sin I} = V_{\text{conIP}} - V_{\sin I} = V_{\text{deIP}}$$

Similarly:

$$C_{\text{conI}} = V_{\text{conIP}} - V_{\text{conI}} \quad \Leftrightarrow \quad V_{\text{deI}} = C_{\sin I} - C_{\text{conI}}$$

Losses: $\ell_{ij} = -x_{ij}$

$$d^* = d_i / \left(\min_i \left\{ L(d_i) = \sum_{j=1}^n \ell_{ij} p_j \right\} \right) = L_{\sin I} = -V_{\sin I}$$

Similarly, $L_{conI} = -V_{conI} \quad \Leftrightarrow \quad V_{deI} = L_{sinI} - L_{conI}$

Normal form

Strategies (decision rules):

$$D = \{\delta(z)/\delta: X \rightarrow D\}$$

$$\delta^* = \delta / \max_{\delta} \left(V(\delta) = \sum_{j=1}^n V(\delta/\theta_j) p(\theta_j) \right)$$

$$V(\delta/\theta_j) = \sum_{k=1}^n V(\delta/\theta_j, z_k) p(z_k/\theta_j)$$

Dealing with regrets or losses, *minimization* must substitute the above *maximization*.

Minimax solution

Pure decisions: $d^* = d_i / \min_i \left(\max_j \ell_{ij} \right)$

Random decisions: $\Delta = \begin{pmatrix} d_1 & d_2 & \dots & d_m \\ \pi_1 & \pi_2 & \dots & \pi_m \end{pmatrix}.$

minimax decision: $\Delta^* = \Delta / \min_{\Delta} \left(\max_j L(\Delta/\theta_j) \right)$

$$L(\Delta/\theta_j) = \sum_i \pi_i \ell_{ij}$$

Random strategies: $\Delta = \begin{pmatrix} \delta_1 & \delta_2 & \dots & \delta_k & \dots & \delta_{mn} \\ \pi_1 & \pi_2 & \dots & \pi_k & \dots & \pi_{mn} \end{pmatrix}$

minimax strategy: $\Delta^* = \Delta / \min_{\Delta} \left(\max_j L(\Delta/\theta_j) \right)$

$$L(\Delta/\theta_j) = \sum_k \pi_k L(\delta_k/\theta_j)$$

The graphic method based on the plotting of the $(L(\delta_k/\theta_1), L(\delta_k/\theta_2))$ region and the analytic based on linear programming, have been used to obtain Δ^* .

PROBLEMES PROPOSATS

PROBLEMA N° 64

1. Sigui X, Y dues variables aleatòries igualment distribuïdes, tals que

$$P(X > 0) = P(Y > 0) = 1$$

Sigui el valor esperat

$$\mu = E \left(\frac{X - Y}{X + Y} \right)$$

Demostra que $\mu = 0$.

C.M. Cuadras

Universitat de Barcelona

PROBLEMA N° 65

2. Sigui X una variable aleatòria normal $N(\mu, 1)$, on la mitjana μ és un paràmetre. Plantejant l'estimació puntual de la probabilitat $P(X > 0)$, demostra la següent desigualtat:

$$\frac{e^{-\mu^2}}{2\pi} < \Phi(\mu) (1 - \Phi(\mu)) \quad -\infty < \mu < +\infty$$

essent $\Phi(x)$ la funció de distribució $N(0, 1)$.

C.M. Cuadras

Universitat de Barcelona

SOLUCIONS ALS PROBLEMES PROPOSATS AL VOLUM 20. N° 3

PROBLEMA N° 62

- 1) Tenim que el jacobià del canvi a coordenades polars

$$x = r \cos \theta \quad y = r \sin \theta \quad 0 < \theta < 2\pi$$

és r , que verifica

$$r^2 = x^2 + y^2$$

La densitat conjunta de (R, θ) és

$$\begin{aligned} f_{R, \theta}(r, \theta) &= r f_{X, Y}(x, y) \\ &= r \frac{1}{2\pi} e^{-\frac{1}{2}(x^2 + y^2)} \\ &= \frac{1}{2\pi} r e^{-\frac{1}{2}r^2} \end{aligned}$$

Fàcilment veiem que les densitats marginals són

$$\begin{aligned} f_R(r) &= r e^{-\frac{1}{2}r^2} \quad r > 0 \quad (\text{gamma}) \\ f_\theta(\theta) &= \frac{1}{2\pi} \quad 0 < \theta < 2\pi \quad (\text{uniforme}) \end{aligned}$$

i per tant R i θ són estocàsticament independents, doncs:

$$f_{R, \theta}(r, \theta) = f_R(r) + f_\theta(\theta)$$

- 2) Un càlcul senzill demostra que

$$Z \sim \sin \theta \sim \cos \theta \quad 0 < \theta < 2\pi \quad \theta \text{ uniforme}$$

ón Z és una variable aleatòria amb densitat:

$$f_Z(Z) = \frac{1}{\pi} \frac{1}{\sqrt{1-Z^2}} \quad -1 < Z < +1$$

($\sim$ significa igualment distribuïdes). Atés que 2θ (mòdul 2π és també uniforme dins $(0, 2\pi)$), resulta que

$$Z \sim \sin 2\theta \sim \cos 2\theta \quad 0 < \theta < 2\pi \quad \theta \text{ uniforme}$$

Considerem la variable

$$W = \frac{\sin \theta + \cos \theta}{\sqrt{2}} \quad 0 < \theta < 2\pi$$

Aplicant fòrmules trigonomètriques conegudes:

$$\begin{aligned} 2W^2 &= \sin^2 \theta + \cos^2 \theta + 2 \sin \theta \cos \theta \\ &= 1 + \sin 2\theta \\ 2W^2 - 1 &= \sin 2\theta \sim Z \end{aligned}$$

Posant

$$\begin{aligned} W &= \cos \phi \\ 2\cos^2 \phi - 1 &= \cos 2\phi \sim Z \end{aligned}$$

Considerant ara $X = R \sin \theta$, $Y = R \cos \theta$, tenim que:

$$\begin{aligned} X + Y &= R(\sin \theta + \cos \theta) \\ &= \sqrt{2} R \cos \phi \\ &= \sqrt{2} V \quad 0 < \phi < 2\pi \quad \phi \text{ uniforme} \end{aligned}$$

ón $V = R \cos \phi$ és normal. Per tant, $X + Y$ és també normal.

C.M. Cuadras

Universitat de Barcelona

PROBLEMA N° 63

- 1) Si la distribució de U_1 és uniforme (0,1)

$$P(-2\log U_1 \leq u) = P\left(U_1 > e^{-u/2}\right) = 1 - e^{-u/2} \quad 0 < u < \infty,$$

i per tant, $-2\log U_1$ és gamma $G\left(\frac{1}{2}, \frac{1}{2}\right)$.

Sabem que la distribució khi-quadrat χ_2^2 és també $G\left(\frac{1}{2}, \frac{1}{2}\right)$, per tant

$$\chi_2^2 = \chi_1^2 + \chi_1'^2 \sim -2 \log U_1$$

on $\sim$ significa mateixa distribució i $\chi_1^2 \sim \chi_1'^2$ són independents. Aleshores, tenim que la distribució asimptòtica de $V = -2 \log \lambda$ és khi-quadrat amb $2k$ graus de llibertat

$$\chi_{2k}^2 = \sum_{j=1}^k \chi_1^2(j) + \chi_1'^2(j) = \sum_{j=1}^k -2 \log U_j$$

Per tant, la distribució asimptòtica de la raó de versemblança λ és

$$\lambda \sim \prod_{i=1}^k U_i$$

on les U_i són uniformes (0,1) independents.

- 2) En el cas $k = 1$ tenim que, asimptòticament, $\lambda \sim U$, on U és uniforme (0,1). A fi de que

$$P(\lambda < c | H_0) = \alpha$$

on $0 < \alpha < 1$ és el nivell de significació i $\{x_1, \dots, x_n / \lambda < c\}$ és la regió de rebuig de H_0 , veiem fàcilment que $c = \alpha$.

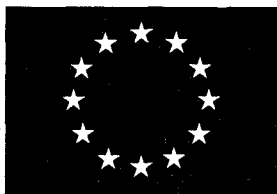
Per tant:

- a) acceptem H_0 si $\lambda \geq \alpha$,
- b) acceptem H_1 si $\lambda < \alpha$.

C.M. Cuadras

Universitat de Barcelona

Cursos i congressos d'estadística



EUROPEAN COMMISSION
EUROSTAT



Statistical Office of the European Communities

**Conference on
Statistical Data Protection '98**

***Call
for
Papers***

**25-27 March 1998
Lisbon, Portugal**

SDP'98 INTRODUCTION AND OBJECTIVES

The protection of sensitive data is a constant issue of concern for National Statistical Offices and for Eurostat. The 3rd International Seminar on Statistical Confidentiality held in Bled in 1997 and the Statistical Disclosure Control (SDC) project funded by the EC 4th Framework Program are two activities sponsored by Eurostat to promote research and encourage interaction between scientists and technicians in this field.

Two conferences are planned for 1998. One is the follow-up of the 3rd International Seminar on Statistical Confidentiality. The second is a new conference, christened "Statistical Data Protection'98" (SDP'98 for short).

The main goal of SDP'98 is to consolidate SDC as a statistical/computational research discipline and to promote exchange of experience among practitioners. Validation of the results of the ESPRIT SDC project is also an objective of SDP'98.

SDP'98 AREAS OF INTEREST

SDP'98 will concentrate on the technical aspects of Statistical Disclosure Control.

It will encompass:

1. Theory

The theoretical stream of the conference is expected to cover:

- statistical aspects of SDC (modelization, security assessment, etc.)
- operations research aspects of SDC (integer programming for data suppression, etc.)
- cryptographic aspects (encrypted data processing, shared control of confidential data, etc.)

2. Application/case studies

The application/case studies stream is intended for practitioners to present case studies dealing with:

- protection of data focusing on their degree of aggregation (micro/macrodata)
- protection of data focusing on their nature (business data, agricultural data, social data, etc.)
- protection of data focusing on their medium (off-line data, on-line data or database protection)

3. Computational problems

The computational problems stream is expected to cover:

- existing SDC software
- missing functionality in SDC software
- complexity of SDC algorithms
- rules used to encode national SDC regulations in software

SDP'98 SUBMISSIONS

Papers will be reviewed on the basis of an extended summary (five to ten pages, single-spaced) of sufficient detail to permit reasonable evaluation. Papers should be written in English and preferably typeset in LaTeX. Each submission should contain a contact E-Mail address and the cover letter should indicate which of the Conference streams fits best to the paper (theory, case studies, computational).

Send submissions by E-mail (in LaTeX format) to **sdp98@etse.urv.es**

Submissions may also be posted (on MS-DOS diskette in LaTeX format only) to the Program Chair at the address below:

Prof. Josep Domingo-Ferrer - (SDP'98 Chair)
Universitat Rovira i Virgili
E.T.S.E.- Autovia de Salou, s/n
E-43006 Tarragona, Catalonia, Spain.

All Submissions. For each submission send an E-Mail (in plain ASCII text) to **sdp98@etse.urv.es** containing:

- postal address and E-Mail address for communication
- complete Title, author and affiliation information
- the abstract of the paper (5 to 10 lines)
- a small selection of keywords and the preferred conference stream

Note: Submissions submitted on paper will not be considered. Diskettes used for submissions cannot be returned.

SDP'98 IMPORTANT DATES

Submission deadline:	30 September 1997 (postmark 27 September)
Acceptance notification:	20 November 1997
Camera ready papers due:	15 December 1997
Conference:	25-27 March 1998

SDP'98 INVITED SPEAKERS

Nabil Adam

(Director, Center for Information Management, Integration and Connectivity, Rutgers University, USA)

Subject: A survey on results on database protection.

William Winkler

(American Statistical Association Fellow and Principal researcher, U. S. Census Bureau)

Subject: Re-identification methods for evaluating the confidentiality of analytically valid microdata.

SDP'98 PROGRAM COMMITTEE

The Program Committee members represent worldwide statistical institutes and academia.

Program Chairman:

Joseph Domingo-Ferrer (U. Rovira i Virgili, Spain).

Program Committee members:

Dick Carter (Statistics Canada)
George E. Kokolakis (Nat. Tech. U. of Athens, Greece)
Daniel Defays (Eurostat)
Asterios Hatziparadissis (Eurostat)
Filipa Duarte de Carvalho (U. Tecnica Lisboa, Portugal)
Denise Lievesley (Director, The UK Data Archive)
Stephen Fienberg (Carnegie-Mellon U., USA)
Max Wigbout (Statistics New Zealand)
Robert Garfinkel (U. Connecticut, USA)
Leon Willenborg (Statistics Netherlands)
Sarah Giessing (Statistisches Bundesamt, Germany)
Laura Zayatz (U.S. Census Bureau)

SDP'98 CONFERENCE PROCEEDINGS

A pre-proceedings booklet will be available at the conference which will contain all presented papers

IOS Press will publish a selection of the presented papers in a proceedings book

The Program Committee will classify papers into three categories:

- REJECT
- ACCEPT FOR PRESENTATION
- ACCEPT FOR PRESENTATION AND PROCEEDINGS

A second reviewing round is likely for papers selected for the proceedings book

Note: We regret that diskettes used for Submissions cannot be returned

CONFERENCE FEE AND FURTHER INFORMATION

The conference fee will be 200,- ECU (Students free).

The preliminary program, exact conference venue, travelling indications, etc. will be posted on the conference Web Page in due time:

<http://www/businessoffice.lu/sdp98-home.html>

For other information contact

Fyfe Business Centre Luxembourg
29, rue Jean Pierre Brasseur
L-1258 Luxembourg
Fax: Luxembourg (+352) 45 59 05
E-mail: ron@telephonie.lu

08003 Barcelona

Exemplar per a l'entitat bancària

Autorització de domiciliació bancària per al pagament de les subscripcions anuals de la revista **Qüestió**

El sotasignat _____																				
autoritza el Banc/Caixa _____																				
Adreça _____																				
Codi postal _____ Ciutat _____																				
a abonar les subscripcions a la revista Qüestió amb càrrec al seu compte																				
número <table><tr><td> </td><td> </td><td> </td><td> </td><td> </td><td> </td><td> </td><td> </td><td> </td><td> </td><td> </td><td> </td><td> </td><td> </td><td> </td><td> </td><td> </td><td> </td><td> </td><td> </td></tr></table>																				
Data _____																				
Signatura																				

Qüestió:
Institut d'Estadística de Catalunya
Via Laietana, 58
08003 Barcelona

NORMES PER A L'ACCEPTACIÓ D'ARTICLES A QÜESTIÓ

GUIDELINES FOR THE SUBMISSION OF ARTICLES FOR QÜESTIÓ

GUIDELINES FOR THE SUBMISSION OF ARTICLES FOR QÜESTIÓ

The journal accepts for publication original articles that are not being considered for publication in any other journal in the fields of Statistics, Operational Research, Official Statistics or Biometrics. Articles may be theoretical or applied, including teaching aspects and applications and will be accepted in English, French, Catalan or any of the official languages in Spain. For all deliverings, the journal will issue a receipt certificate corresponding to the presentation date of article, which will appear as date of «received» in its final publication. *Qüestió* also accepts the publication of institutional advertisements on courses, seminars, conventions and other activities related to statistics or operational research that will be held in our country; they will be directly reproduced from the originals sended (that is to say, without any self-edition process or similar).

The originals assigned to the thematic sections of *Qüestió* will be sistematically reviewed by independent referees and members of the Editorial Board, whose result will be communicated to the main author of the article in order to correct, if necessary, any formal or content aspects. The «acceptance date» of the article, which will appear in its final publication, will be the date of sending the final version to the journal.

For the presentation of original articles, the author should send, to the Secretary of *Qüestió* (Institut d'Estadística de Catalunya), two nonreturnable copies of the paper typed on DIN-A4, one side of the paper only, double spaced and with wide margins. Each article should include a title, the name of the author or authors, their affiliation, full address and also an abstract of the paper (10-15 lines) at the beginning of the article and the translation of the title into English, followed by the main keywords (in the original language) and its assignation in the AMS classification as well. Bibliographical references should state the author's name followed by the year of publication in brackets [e.g.: Mahalanobis (1936), Rao (1928b)] and they should be listed at the end of the article in alphabetical order. Footnotes should be numbered in the article and appear at the foot of the corresponding page.

Once the evaluation has been passed, the author is required to provide the article on a diskette togheter with its printed version, processed by LATEX, preferably, or, failing that, Word Perfect (6.0A or earlier) or ASCII for text and tables. For grafics and figures, the appropriate formats of the editing programs PS, EPS, PCX o BMDP are strongly recommended. On the other hand, if the article is not written in English, the translation of its original title, abstract and keywords should be enclosed, as well as a full summary of the article in English (2-5 pages and with the same structure and titles that the original shows). The Secretary of *Qüestió* offers, to the authors who may request it, the precise rules of submission, the LATEX patterns for the edition and the AMS classification references.

QÜESTIÓ

Institut d'Estadística de Catalunya
Via Laietana, 58
08003 Barcelona