

Qüestió

Quaderns d'Estadística
i Investigació Operativa

Any 1999, volum 23, núm. 1
Segona època

Entitats patrocinadores:

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

Entitat col·laboradora:

International Biometric Society



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

Any 1999, volum 23, núm. 1

SUMARI

Editorial

Estadística

- Una revisión de diferentes aportaciones al diseño en poblaciones finitas 3
B. Font

Investigació Operativa

- Amplitude of weighted majority game strict representations 43
J. Freixas and A. Puente
On the compatibility of classical multiplier estimates with variable reduction techniques
when there are nonlinear inequality constraints 61
E. Mijangos and N. Nabona
Minimización global de un polinomio en la recta real 85
C. Beltrán Royo

Estadística Oficial: Tècniques de projeccions demogràfiques

- Le traitement de l'incertitude dans les perspectives démographiques 113
J. Duchêne
Mètodes i models per projectar els components del canvi demogràfic en les projeccions de
població de l'Institut d'Estadística de Catalunya 129
M. Farré Mallofré i J.A. Sánchez Cepeda

Biometria

- Papel ético del estadístico en la experimentación humana 155
E. Cobo

Secció docent i problemes

- Models gràfics d'independència 171
J.M. Duran Rúbies

Novetats de Software

- FORCE4/R, a new software product for forecasting and seasonal adjustment 199
A. Prat, I. Solé, J.M. Catot and J. Lorés

- Ressenyes d'activitats institucionals* 209

Informació per als autors i lectors



EDITORIAL

Aquest primer número del volum 23 recull un total de set articles, repartits entre les quatre seccions habituals, als quals cal afegir un vuitè original de caràcter docent i una àmplia ressenya a l'apartat dedicat a les novetats de programari estadístic. Tant el nombre d'articles com de pàgines publicades en el present número permet entreveure que l'edició d'aquest volum se situarà a l'entorn de la producció enregistrada en els darrers tres anys (23 articles i 580 pàgines) i, el que és més important, consolidar la tendència de publicar un nombre d'articles en el decurs d'un volum pràcticament idèntic al d'originals sotmesos anualment a *Qüestió*; amb el manteniment d'aquest esforç es garanteix als autors un temps d'espera de la seva publicació que gairebé depèn únicament de la durada del procés d'avaluació corresponent.

En relació als subscriptors de *Qüestió*, ens plau informar-los que, no obstant l'anunci d'un increment del preu de subscripció per al 1999 que s'esmentava en l'editorial del número 1 del volum 22 (1998), finalment es mantindrà per quart any consecutiu atès que l'Idescat i la resta de patrocinadors han pogut adequar les seves aportacions i, a la vegada, han augmentat les subscripcions i els altres ingressos. No obstant això, el creixement sostingut d'una part de les despeses materials fa preveure que l'import de la subscripció s'actualitzarà definitivament l'any 2000 (volum 24).

Des del punt de vista institucional cal suposar que, a partir d'ara, la producció i la difusió de la revista es reforçarà paulatinament amb la recent incorporació de la Universitat de Girona com a entitat patrocinadora, mitjançant l'acord de col·laboració signat amb l'Idescat el 8 de febrer de 1999. En aquest sentit, l'ampliació del Consell Editor de *Qüestió*, amb el Dr. Carles Barceló com a nou editor associat, i el suport de la Universitat de Girona s'encaminen decididament a millorar la qualitat i potenciar el desenvolupament de la revista.

C.M. Cuadras i E. Ripoll, editors executius

Comentari de les seccions

«Estadística», «Investigació Operativa», «Biometria» i «Secció Docent»

En aquesta ocasió, la secció «Estadística» conté només un original, *Una revisión de diferentes aportaciones al diseño en poblaciones finitas*, de B. Font, que presenta una revisió comparada de diferents enfocaments del mostreig òptim en poblacions finites, tant clàssic com bayesià.

A la secció «Investigació Operativa» s'hi publiquen tres articles; en el primer d'ells, *Amplitude of weighted majority game strict representations*, de J. Freixas i A. Puente,

es fa una notable aportació sobre jocs ponderats quan els pesos varien i es prova que algunes modificacions no canvien el joc. El segon article, *On the compatibility of classical multiplier estimates with variable reduction techniques when there are nonlinear inequality constraints*, d'E. Mijangos i N. Nabona, resol un problema de minimització d'una funció no lineal subjecta a restriccions, convertint desigualtats en igualtats. En el tercer, *Minimización global de un polinomio en la recta real*, de C. Beltrán Royo, es proposa un algorisme per minimitzar un polinomi real, per translació vertical del gràfic del polinomi a l'eix horitzontal i estudi de les arrels que resulten.

L'únic article de la secció «Biometria», *Papel ético del estadístico en la experimentación humana*, d'E. Cobo, és una dissertació sobre els principis ètics de l'experimentació científica amb éssers humans i el paper de l'estadístic, que s'hauria d'incorporar als comitès i organismes que controlen els assaigs clínics.

Finalment, la «Secció docent i problemes» inclou també un article, *Models gràfics d'independència*, de J.M. Duran Rúbies, que explica la independència i dependència condicional entre esdeveniments i variables, mitjançant símbols i grafs.

Carles M. Cuadras, editor executiu

Comentari de la secció «Estadística Oficial» i altres apartats

La secció «Estadística Oficial» inclou dos articles de caire metodològic sobre l'elaboració de les projeccions de població. El primer treball, *Le traitement de l'incertitude dans les perspectives démographiques*, de J. Duchêne, conté una relació de les mesures utilitzades per avaluar la precisió de les projeccions, així com dels seus avantatges i inconvenients; es proposen diferents tipus d'enfocaments per tractar el problema de la incertesa, com ara examinar els errors exposats de les projeccions amb una metodologia del tipus APC (edat, període cohort) o utilitzar el mètode de projeccions probabilistes basat en les opinions d'experts sobre la tendència dels components demogràfics, entre d'altres. El segon article, *Métodes i models per projectar els components del canvi demogràfic en les projeccions de població de l'Institut d'Estadística de Catalunya*, de M. Farré i J.A. Sánchez, presenta la metodologia utilitzada a l'Idescat en l'elaboració de les projeccions per al període 1996-2030, emprant mètodes de projecció específics per a cadascun dels components del creixement: així, es desenvolupa un model que integra la informació sobre la fecunditat de les generacions, l'ordre dels naixements i l'edat de les mares a fi de modelar la fecunditat; en el cas de la mortalitat, s'utilitza la tendència futura de l'esperança de vida a partir del mètode del traç cronològic proposat per Nacions Unides; per últim, la migració es tracta amb un model (Rogers) que permet aplicar uns perfils dinàmics en el temps.

A continuació, la «Secció docent i problemes» integra, a més de l'article ja referenciat, la presentació successiva de nous enunciats i la resolució dels problemes publicats en el número anterior. Seguidament, la secció «Novetats de software» acull una detallada presentació del sistema de modelització-predicció de sèries temporals anomenat *FORCE4/R*, que culmina un projecte del programa *ESPRIT IV* desenvolupat, entre d'altres, per l'equip investigador d'A. Prat de la UPC i la Universitat de Lleida. La recensió de l'arquitectura i prestacions, a càrrec dels seus autors, evidencia la diversitat d'eines i programari estadístic preexistent que integra el *FORCE4/R* i inclou una ressenya de la seva avaluació, que ha comptat amb la participació de l'Idescat.

Per últim, l'apartat dedicat a «Resenyes d'activitats institucionals» inclou, com ja és costum, una actualització de les activitats de la Sociedad Española de Biometría, amb l'anunci dels cursos monogràfics organitzats en col·laboració amb altres entitats. En segon lloc, s'ofereix una recensió del Training for European Statisticians Institute que *Qüestió* redifon des del volum 21 (1997); en aquesta ocasió, juntament amb els cursos del *Core Programme 1998-99* que s'impartiran fins el mes de juny d'enguany, s'anuncia el curs *Systems of Social Statistics (Core Programme 1999-2000)* que acollirà l'Idescat el proper mes d'octubre i que s'adreça, com la majoria d'activitats del TES Institute-Eurostat, als instituts d'estadística oficial en l'àmbit comunitari. Seguidament, es reproduïx el primer anunci del Tercer Congrés Europeu de Matemàtiques –*3ECM*– (Barcelona, 10-14 de juliol del 2000) que organitza la Societat Catalana de Matemàtiques-IEC i que, sota els auspicis de la European Mathematical Society, hi col·laboren dues universitats patrocinadores de *Qüestió* i l'Idescat, entre d'altres organismes.

Enric Ripoll, editor executiu

Estadística

UNA REVISIÓN DE DIFERENTES APORTACIONES AL DISEÑO EN POBLACIONES FINITAS

B. FONT

Universitat de València*

Este artículo es una revisión de los resultados más relevantes de diseño en poblaciones finitas. Los resultados presentados se clasifican en tres grupos: población fija, inferencia basada en modelos de superpoblación clásicos e inferencia basada en modelos de superpoblación Bayesianos, y se analizan de forma comparativa. Entre los resultados revisados señalemos las aportaciones sobre: procedimiento de muestreo, tamaño óptimo de la muestra y procedimientos de estratificación.

A comparative review of different contributions in survey sampling.

Palabras clave: Diseño óptimo, estrategias de muestreo, inferencia Bayesiana, modelos de superpoblación Bayesianos, modelos de superpoblación clásicos, población fija.

Clasificación AMS: 62D05.

*Begoña Font Belaire. Dept. Economia Financera i Matemàtica. Edifici Departamental Oriental. Av. dels Tarongers s/n. 46071 València (Espanya).

—Recibido en agosto de 1997.

—Aceptado en enero de 1999.

1. INTRODUCCIÓN

La obtención de información a través de encuestas de la población se ha convertido en una práctica frecuente en la investigación de entidades públicas y privadas que demandan respuestas a preguntas tan básicas como: ¿cuál es el método más apropiado para obtener la muestra?, ¿cuál es el tamaño apropiado de la muestra?, ¿cómo debo analizar los datos obtenidos? Los que estudiamos poblaciones finitas sabemos que las respuestas a estas preguntas no son únicas y universales, y además dependen de la perspectiva de análisis aplicada. Al escribir este artículo se ha pensado que era importante realizar un esfuerzo en recoger y analizar las respuestas que da la literatura a estas preguntas de diseño desde las distintas perspectivas. De este modo, el lector no especializado obtendrá con este trabajo una visión conjunta sobre las distintas respuestas de la literatura a sus preguntas sobre diseño, y el lector especializado tendrá una referencia en la que se analizan conjuntamente los resultados sobre diseño desde las diversas perspectivas, con la particularidad añadida de permitir el estudio por separado de cada grupo de aportaciones por el carácter autocontenido de las secciones sobre aportaciones desde las perspectivas de población fija, modelos de superpoblación clásicos y modelos de superpoblación Bayesianos.

Resumiendo y concretando, este artículo realiza una revisión de la literatura clásica y reciente sobre diseño en poblaciones finitas desde las perspectivas de población fija, y modelos de superpoblación desde las perspectivas clásica y Bayesiana. En concreto, se revisan las aportaciones en relación a cómo seleccionar la muestra, elegir el tamaño óptimo de unidades y elementos a muestrear en un muestreo en dos etapas, técnicas para estratificación de una población, etc., señalando las hipótesis asumidas en cada caso desde las distintas perspectivas. Otros artículos de revisión sobre las diferentes aportaciones al diseño y estimación son: Zacks(1970), Rao(1979), y Bellhouse(1984).

Debemos indicar que en este artículo no se pretende realizar una revisión de las aportaciones realizadas en el terreno de obtener el estimador óptimo para un muestreo dado o para un grupo de procedimientos de muestreo (aunque se indiquen como referencia algunos resultados), y por lo tanto que la exposición no se centrará en la forma de los estimadores sino en las características relativas al diseño. El lector interesado en un análisis más detallado sobre estimación puede dirigirse, por ejemplo, a los siguientes artículos y libros:

- Sobre muestreo en población fija, a los manuales de Hansen, Hurwitz and Madow(1953), Kish(1965), Cochran(1977), Hedayat and Sinha(1991) y Särndal, Swenson and Wretman(1992) y a los artículos de Chaudhuri(1988), Mukhopadhyay and Tracy(1993) y Godambe(1992).
- Sobre muestreo y estimación basada en un modelo de superpoblación, a los manuales de Cassel, Särndal and Wretman(1976), Bolfarine and Zacks(1992) y a los

artículos de Royall(1988,92a) desde las perspectivas clásica, y Ericson(1969a,88) y Font(1995b) desde la perspectiva Bayesiana.

Otras referencias de interés son el artículo de Särndal(1978) y las tesis de Murgui(1982) y Font(1995a) que recogen un resumen comparativo de las distintas aproximaciones al problema de estimación en poblaciones finitas.

El presente trabajo se organiza en 6 secciones incluyendo esta introducción. En la sección 2 se establecen los conceptos básicos y notaciones que se aplicarán en la descripción de los distintos resultados. En la sección 3 se presentan las aportaciones al diseño desde la perspectiva de población fija. En las secciones 4 y 5 las aportaciones al diseño que se apoyan en un modelo de superpoblación desde las metodologías clásica y Bayesiana respectivamente. Y en la sección 6 se presenta un resumen conclusión sobre la comparación realizada entre las distintas perspectivas y se referencian otras aportaciones al diseño no tratadas en las secciones 3, 4 y 5.

2. CONCEPTOS BÁSICOS Y NOTACIÓN

Consideremos una población finita de N (entero positivo conocido) elementos, $P = \{1, 2, \dots, N\}$, unas cantidades y_i , $i = 1, 2, \dots, N$ que representan el valor de una característica de interés asociada a cada uno de estos elementos, y supongamos que estamos interesados en inferir sobre la cantidad (desconocida) $\bar{y} = \frac{\sum_{i=1}^N y_i}{N}$. Para poder estimar $\bar{y}$ (denotaremos al estimador por $\hat{\bar{y}}$), obtendremos una muestra $s = \{i_1, i_2, \dots, i_n\}$ con probabilidad $p(s)$ mediante un mecanismo de muestreo no informativo de tamaño fijo n y mediremos para los índices muestreados el valor de la característica en estudio. En este artículo nos restringiremos a los resultados referidos a procedimientos de muestreo no informativos de tamaño fijo n . Dentro de este grupo indicaremos por srswr al muestreo aleatorio con reemplazamiento, ppswr al muestreo con reemplazamiento con $p_i = Pr(\text{seleccionar } i)$, $i = 1, 2, \dots, N$, srswor al muestreo aleatorio sin reemplazamiento, y ppswor al muestreo sin reemplazamiento con $\pi_i = Pr(i \in s)$, $i = 1, 2, \dots, N$. A partir de esta notación básica introduciremos a continuación las notaciones particulares correspondientes a las perspectivas de población fija y modelos de superpoblación.

La perspectiva de población fija considera como única fuente de aleatoriedad la debida al mecanismo de muestreo. Si representamos por $E_p(\cdot)$ y $V_p(\cdot)$ a los operadores media y varianza respecto al procedimiento de muestreo p , diremos que nuestro estimador $\hat{\bar{y}}$ es insesgado respecto al muestreo si:

$$(1) \quad E_p(\hat{\bar{y}}) = \bar{y},$$

y en este caso el error cuadrático medio cometido respecto al muestreo vendrá dado por:

$$(2) \quad E_p(\hat{y} - \bar{y})^2 = V_p(\hat{y}).$$

A partir de (1) y (2), Newman(1934) define como estrategia óptima (estrategia de uniforme mínima varianza) aquella estrategia $(\hat{y}, p)$ que minimice el error cuadrático medio respecto al muestreo p (ver (2)) entre todos los estimadores insesgados respecto a ese procedimiento de muestreo p (ver (1)). Los trabajos de Godambe(1955) y Godambe and Joshi(1965) demuestran la no existencia de estimadores óptimos según la definición anterior (más detalles en sección 3) y justifican la búsqueda de otros criterios de estrategia óptima (por ejemplo, la admisibilidad) o la introducción de modelos de superpoblación (ver otras razones en la sección 4) para delimitar en qué circunstancias relativas a la población un estimador era mejor que otro.

La perspectiva basada en modelos de superpoblación considera a los valores y_i como realizaciones de una variable aleatoria Y_i , $i = 1, 2, \dots, N$. Si asumimos esta segunda fuente de aleatoriedad y representamos por $E_m(\cdot)$ y $V_m(\cdot)$ a los operadores media y varianza respecto al modelo de superpoblación, diremos que el estimador $\hat{y}$ (manteniendo la notación inicial) es insesgado respecto al modelo si:

$$(3) \quad E_m(\hat{y} - \bar{Y}) = 0.$$

La literatura que considera un modelo de superpoblación propone dos criterios para seleccionar la estrategia óptima, según la importancia que el estadístico quiera dar a cada una de las dos fuentes de aleatoriedad (procedimiento de muestreo y modelo de superpoblación) que consisten en minimizar:

$$(4) \quad E_m E_p(\hat{y} - \bar{Y})^2,$$

o bien,

$$(5) \quad E_m(\hat{y} - \bar{Y})^2,$$

entre estimadores que satisfacen (1) ó (3).

A veces dispondremos también de una información complementaria sobre la variable en estudio y_i , en forma de una característica auxiliar x_i con $i = 1, 2, \dots, N$ conocida para todas las unidades de la población que se podrá aplicar en el mecanismo de muestreo y/o en la estimación para mejorar el proceso inferencial. Algunas notaciones adicionales de interés son:

$$\begin{aligned}\bar{y} &= \frac{\sum_{i=1}^N y_i}{N}, & \bar{y}_s &= \frac{\sum_{i \in s} y_i}{n}, & \bar{y}_u &= \frac{\sum_{i \notin s} y_i}{N-n}, \\ \bar{x} &= \frac{\sum_{i=1}^N x_i}{N}, & \bar{x}_s &= \frac{\sum_{i \in s} x_i}{n}, & \bar{x}_u &= \frac{\sum_{i \notin s} x_i}{N-n}, \\ S_y^2 &= \frac{\sum_{i=1}^N (y_i - \bar{y})^2}{N-1}, & S_x^2 &= \frac{\sum_{i=1}^N (x_i - \bar{x})^2}{N-1}.\end{aligned}$$

En algunas ocasiones los elementos de la población (de manera natural o bien tras un proceso de estratificación) aparecen agrupados en unidades más grandes (psu's). En este trabajo, denotaremos por K (entero positivo conocido) al número de unidades, M_i al número de elementos de la unidad i , $i = 1, 2, \dots, K$ (enteros positivos conocidos) con $\sum_{i=1}^K M_i = N$, y_{ij} al valor de la característica en estudio asociada al elemento j de la unidad i , para $j = 1, 2, \dots, M_i$, $i = 1, 2, \dots, K$. Las inferencias que se analizarán en este trabajo se refieren a la media de la población por elemento y se apoyarán en los valores observados y_{ij} para $(i, j) \in s$, con s una muestra obtenida a través de un muestreo en dos etapas, en el que primero se realiza un muestreo no informativo de tamaño fijo k (denotaremos a la muestra obtenida por s_0) y en segundo lugar un muestreo no informativo de tamaño fijo m_i para $i \in s_0$ con $\sum_{i \in s_0} m_i = n$ (denotaremos a la muestra en cada unidad por s_i , $i \in s_0$) ambos procedimientos de muestreo independientes, por tanto $s = \{(i, j) : j \in s_i, i \in s_0\}$. En este trabajo indicaremos por: stsrswor al muestreo estratificado con $k = K$ con elección de los elementos mediante srswor, clsrswor al muestreo cluster con $k < K$ y $m_i = M_i$, $i = 1, 2, \dots, K$ con elección de las unidades mediante srswor y 2ssrswor al muestreo en dos etapas con muestreo srswor de las unidades y de los elementos. Otras notaciones de interés son, para $i = 1, 2, \dots, K$:

$$\begin{aligned}\bar{y}_i &= \frac{\sum_{j=1}^{M_i} y_{ij}}{M_i}, & \bar{y}_{s_i} &= \frac{\sum_{j \in s_i} y_{ij}}{m_i}, & \bar{y}_{u_i} &= \frac{\sum_{j \notin s_i} y_{ij}}{M_i - m_i}, \\ \bar{x}_i &= \frac{\sum_{j=1}^{M_i} x_{ij}}{M_i}, & S_{yi}^2 &= \frac{\sum_{j=1}^{M_i} (y_{ij} - \bar{y}_i)^2}{M_i - 1}, & S_{xi}^2 &= \frac{\sum_{j=1}^{M_i} (x_{ij} - \bar{x}_i)^2}{M_i - 1},\end{aligned}$$

así como las siguientes definiciones de la varianza poblacional dentro de las unidades (S_{yw}^2), la varianza poblacional entre unidades (S_{yb}^2) y la correlación poblacional dentro de las unidades (R_y):

$$S_{yw}^2 = \frac{\sum_{i=1}^K (M_i - 1) S_{yi}^2}{N - K}, \quad S_{yb}^2 = \frac{\sum_{i=1}^K M_i (\bar{y}_i - \bar{y})^2}{K - 1}, \quad R_y = \frac{2 \sum_{i=1}^K \sum_{j < j^*}^{M_i} (y_{ij} - \bar{y})(y_{ij^*} - \bar{y})}{(\sum_{i=1}^K M_i^2 - 1)(N - 1) S_y^2},$$

esta última expresión (R_y), que puede interpretarse como un coeficiente de homogeneidad dentro de la unidad, ha sido aproximada por los diversos autores en términos de S_{yw}^2 , S_{yb}^2 y S_y^2 (por ejemplo, en el caso $M_i = M$, $i = 1, 2, \dots, K$, Hansen, Hurwitz and Madow(1953) usan la aproximación $\frac{M(K-1)}{N-1} \frac{S_{yb}^2 - S_y^2}{S_y^2}$ y Cochran(1977) la aproximación $\frac{S_{yb}^2 - S_y^2}{(M-1) S_y^2}$). Indicaremos también con R_y a todas estas aproximaciones cuando formen parte de la aproximación de un resultado.

3. APORTACIONES AL DISEÑO DESDE LA PERSPECTIVA DE POBLACIÓN FIJA

Las aportaciones al diseño de los autores que han trabajado desde esta perspectiva son muy numerosas, en concreto en esta sección se analizarán algunos resultados comparando procedimientos de muestreo distintos, expresiones sobre los tamaños de muestreo óptimos y algunos procedimientos para estratificar la población con el objetivo de obtener mejores inferencias.

Antes de pasar a la relación de las aportaciones en población fija conviene realizar una reflexión previa sobre el concepto de estimador «óptimo», para valorar mejor las consecuencias de la elección entre las distintas alternativas de diseño en base a los errores cuadráticos asociados a ese estimador «óptimo». La no existencia de una estrategia uniforme de mínima varianza demostrada por Godambe(1955) cuando se consideraba la minimización dentro de la clase de los estimadores insesgados, lineales y homogéneos, y por Godambe and Joshi(1965) cuando se consideraba la clase de todos los estimadores insesgados (más detalles en los artículos de Chaudhuri(1988) y Mukhopadhyay and Tracy(1993)), conduce a establecer el concepto de estimador admisible para un muestreo dado. En esta línea de investigación, Joshi(1965,67) demostró: (i) la admisibilidad del estimador de Horwitz-Thompson ($\hat{y}_{HT} = \sum_{i \in s} \frac{y_i}{\pi_i N}$) como estimador de $\bar{y}$ para muestreos ppswr con probabilidades de inclusión π_i estrictamente positivas, (ii) la admisibilidad de $\bar{y}_s$ para un diseño p entre todos los estimadores de $\bar{y}$, y (iii) la admisibilidad de $\bar{y}_s$, $\hat{y}_r = \frac{\bar{x}}{\bar{x}_s} \bar{y}_s$ (estimador de razón), e $\hat{y}_b = \bar{y}_s + \hat{\beta}(\bar{x} - \bar{x}_s)$ con $\hat{\beta} = \frac{\sum_{i \in s} (x_i - \bar{x}_s)(y_i - \bar{y}_s)}{\sum_{i \in s} (x_i - \bar{x}_s)^2}$ (estimador de regresión) entre los estimadores medibles (con algunas restricciones adicionales) de $\bar{y}$ para funciones de pérdida muy generales para procedimientos de muestreo sin reemplazamiento. En otra línea de investigación, Newman(1934) y Royall(1968) demostraron que para un muestreo srswr, el estimador $\bar{y}_s$ es el único estimador que minimiza el error cuadrático respecto al muestreo entre los estimadores insesgados lineales homogéneos y entre todos los estimadores insesgados respectivamente.

3.1. Comparación entre procedimientos de muestreo

3.1.1. Comparación entre muestreos en una etapa

Veremos dos grupos de resultados, los debidos a Lanke(1975) que comparan los muestreos srswor y srswr y los de Des Raj(1954), Reddy and Rao(1977) y Rao(1993) que permiten comparar muestreos srswr y ppswr.

Lanke(1975) define que un procedimiento de muestreo p' es mejor o al menos tan bueno como otro procedimiento p si dado un estimador insesgado de $\bar{y}$ respecto a p , $\hat{y}_p$, existe un estimador insesgado de $\bar{y}$ respecto a p' , $\hat{y}_{p'}$, tal que $V_{p'}(\hat{y}_{p'}) \leq V_p(\hat{y}_p)$, y obtiene una

condición necesaria para que un muestreo p' sea mejor o al menos tan bueno como p para la estimación de $N\bar{y}$. De la aplicación de esta condición a los muestreos aleatorios se llega a que en la estimación de $\bar{y}$ un muestreo srswr es mejor o al menos tan bueno como un muestreo srswr si $n \geq 2$.

Des Raj(1954) demuestra que el muestreo ppxwr (muestreo ppswr con $p_i = \frac{x_i}{N\bar{x}}$, $i = 1, 2, \dots, N$) puede producir peores estimaciones que el muestreo srswr cuando la relación entre y_i y x_i , $i = 1, 2, \dots, N$ se aleja de una línea recta a través del origen.

Rao(1993), considera dos muestreos proporcionales con reemplazamiento, ppswr y pps'wr con probabilidades de elegir i , p_i y p'_i , $i = 1, 2, \dots, N$ respectivamente y considera el muestreo pps''wr con probabilidades de elegir i , $p''_i = \alpha p_i + (1 - \alpha)p'_i$ con $0 \leq \alpha \leq 1$, $i = 1, 2, \dots, N$ obteniendo los siguientes resultados que representaremos por (1):

- (i) Si $V_{\text{ppswr}}(\hat{y}_{HH}) < V_{\text{pps'wr}}(\hat{y}_{HH})$, entonces $V_{\text{pps''wr}}(\hat{y}_{HH}) < V_{\text{pps'wr}}(\hat{y}_{HH})$.
- (ii) Si $V_{\text{pps'wr}}(\hat{y}_{HH}) < V_{\text{ppswr}}(\hat{y}_{HH})$, entonces $V_{\text{pps''wr}}(\hat{y}_{HH}) < V_{\text{ppswr}}(\hat{y}_{HH})$.
- (iii) $V_{\text{pps'wr}}(\hat{y}_{HH}) - V_{\text{ppswr}}(\hat{y}_{HH}) = \frac{1}{nN^2} \sum_{i=1}^N \frac{y_i^2(p_i - p'_i)}{p_i p'_i}$,

donde $\hat{y}_{HH} = \sum_{i \in s} \frac{y_i}{p_i N n}$ es el estimador de Hansen-Hurwitz asociado al muestreo ppswr (las primas representan los correspondientes estimadores Hansen-Hurwitz para los restantes muestreos).

De (1)(i) y (ii) se deduce que la estrategia $(\hat{y}_{HH}, \text{pps''wr})$ es mejor que la peor entre las estrategias $(\hat{y}_{HH}, \text{ppswr})$ y $(\hat{y}_{HH}, \text{pps'wr})$. De (1)(iii) tomando $p'_i = \frac{1}{N}$, $i = 1, 2, \dots, N$ (muestreo srswr) obtenemos una expresión de las diferencias en eficiencia entre un muestreo srswr y un muestreo ppswr.

3.1.2. Comparación entre muestreos srswr, stsrswor y clsrswor

La comparación entre los muestreos srswr y stsrswor es analizada en la mayoría de manuales sobre técnicas de muestreo. Por ejemplo, Cochran(1977) y Särndal, Swenson and Wretman(1992) obtienen la siguiente expresión:

$$(2) \quad V_{\text{srswr}}(\bar{y}_s) - V_{\text{stsrswor},(13)}(\bar{y}_s) = \frac{N-n}{nN(N-1)} [(K-1)S_{yb}^2 - \frac{1}{N} \sum_{i=1}^K (N - M_i)S_{yi}^2],$$

en la que el subíndice stsrswor (13) indica que se ha aplicado un muestreo aleatorio estratificado en el que el número de elementos muestreados en cada estrato es proporcional al tamaño del mismo (ver comentarios del apartado 3.2.2 y la expresión (13)).

A partir de (2) podemos observar que matemáticamente el muestreo stsrswor puede ser peor que el muestreo srswor. En particular, si suponemos que $S_{yi}^2 = S_{psu}^2$, $i = 1, 2, \dots, K$ tenemos que la expresión (2) se convierte en:

$$(3) \quad V_{\text{srswor}}(\bar{y}_s) - V_{\text{stsrswor},(13)}(\bar{y}_s) = \frac{(N-n)(K-1)}{nN(N-1)} [S_{yb}^2 - S_{psu}^2],$$

que será negativa (evidencia a favor del muestreo srswor) si la varianza poblacional entre estratos es más pequeña que la varianza dentro de los estratos. En la mayor parte de las situaciones prácticas la varianza entre estratos supera a la varianza dentro del estrato, (2) es positivo y se produce una ganancia en precisión al usar el muestreo stsrswor como alternativa al muestreo srswor.

Hansen, Hurwitz and Madow(1953) obtienen a partir de (2) la siguiente expresión aproximada de las ganancias derivadas del muestreo stsrswor frente al muestreo srswor:

$$(4) \quad V_{\text{srswor}}(\bar{y}_s) - V_{\text{stsrswor},(13)}(\bar{y}_s) \simeq \frac{N-n}{N} \frac{S_y^2}{n} R_y.$$

A partir de (2) y (4) podemos extraer las siguientes observaciones:

- Cuanto más grande sea la homogeneidad dentro de los estratos, las ganancias por estratificación serán mayores.
- Cuanto más grande sea la heterogeneidad entre las medias de los estratos, mayores serán las ganancias relativas por estratificación.
- La ganancia relativa no será muy alta salvo que las variaciones entre estratos sean bastante mayores que las variaciones dentro de los estratos.

En relación a la comparación entre los muestreos srswor y clsrswor tenemos, apoyándonos en los resultados de Cochran(1977), cuando $M_i = M$, $i = 1, 2, \dots, K$ que:

$$(5) \quad \begin{aligned} V_{\text{clsrswor}}(\bar{y}_s) - V_{\text{srswor}}(\bar{y}_s) &= \frac{N-n}{N} \frac{S_y^2}{n} \left\{ \frac{N-1}{M(K-1)} [1 + (M-1)R_y] - 1 \right\} \\ &\simeq \frac{N-n}{N} \frac{S_y^2}{n} (M-1)R_y. \end{aligned}$$

De (5) se deduce que si $R_y < 0$, el muestreo clsrswor es más eficiente que el srswor, y si $R_y > 0$ (la situación más habitual), el muestreo clsrswor es menos eficiente que

el srswor. Hansen, Hurwitz and Madow(1953) proporcionan una buena discusión sobre los valores de R_y para diferentes ítems y diferentes tamaños de cluster y Särndal, Swensson and Wretman(1992) estudian los muestreos clppswor (muestreos de dos etapas con $m_i = M_i$, $i \in s_0$ con muestreo para las unidades del tipo ppswor) y aplicando un estimador Horwitz-Thompson llegan a las siguientes conclusiones prácticas:

- Si podemos seleccionar los clusters con un muestreo ppswor con $\pi_i = \frac{k\bar{y}_i}{\sum_{i=1}^K \bar{y}_i}$, $i = 1, 2, \dots, K$, entonces este muestreo cluster será muy eficiente.
- Si seleccionamos los clusters con un muestreo ppswor con $\pi_i = \frac{kM_i}{N}$, $i = 1, 2, \dots, K$, el muestreo será bastante eficiente si las variaciones entre las medias poblacionales de cada cluster son pequeñas.
- Un muestreo clsrswor es a menudo poco eficiente cuando los clusters son de distinto tamaño, es más, si la correlación entre M_i y $M_i\bar{y}_i$, $i = 1, 2, \dots, K$ es positiva (que es el caso habitual), el muestreo clsrswor da peores resultados que en el caso de muestreo clsrswor con $M_i = M$, $i = 1, 2, \dots, K$.

3.1.3. Comparación entre muestreos srswor y 2ssrswor

Hansen, Hurwitz and Madow(1953) obtienen la siguiente expresión para comparar los dos tipos de muestreo en el caso: $M_i = M$, $i = 1, 2, \dots, K$ y $m_i = m$, $i \in s_0$

$$(6) \quad V_{2ssrswor}(\bar{y}_s) - V_{srswor}(\bar{y}_s) \simeq \frac{N-n}{N} \frac{S_y^2}{n} R_y (m-1).$$

A partir de (6) podemos concluir que un muestreo 2ssrswor será eficiente en la medida que la correlación entre los elementos dentro de cada unidad sea pequeña si es positiva, o en la medida que ésta sea negativa. En la práctica R_y es positiva y decrece al aumentar el tamaño de las unidades.

3.2. Tamaños de muestreo óptimos

La literatura clásica presenta tres criterios para determinar el tamaño de muestreo óptimo:

- Para un coste dado, determinar el tamaño que hace mínimo el error cuadrático medio en la estimación.
- Para una varianza respecto al muestreo dada del estimador de la media poblacional insesgado aplicado, determinar el tamaño que hace mínimo el coste de la estimación.
- El problema de decisión de minimizar una función de pérdida.

A continuación presentaremos los resultados obtenidos según estos criterios para muestreos srswor, stsrswor, clsrswor y 2ssrswor.

3.2.1. Tamaño de una muestra srswor

Konijn(1973) propone para una varianza V^2 dada la siguiente expresión para el tamaño n de una muestra srswor:

$$(7) \quad n = \frac{\frac{S_y^2}{V^2}}{\frac{N-1}{N} + \frac{S_y^2}{NV^2}},$$

esta expresión cuando N es bastante grande se simplifica tomando $\frac{N-1}{N} \simeq 1$ a la expresión presentada por Kish(1965).

Y Cochran(1977) considerando el criterio de minimizar la pérdida cuadrática multiplicada por la constante positiva γ más una función lineal de costes $C(n) = c_0 + c_1n$, obtiene la siguiente expresión para el tamaño óptimo n de una muestra srswor:

$$(8) \quad n = \frac{S_y^2 \sqrt{\gamma}}{\sqrt{c_1}}.$$

(Una forma más general de este resultado aparece en Yates(1960).) Debemos observar que las expresiones (7) y (8) proporcionan valores para n que son función de la varianza poblacional, e indicar que para poder fijar el tamaño óptimo la práctica habitual consiste en sustituir esta cantidad desconocida por un estimador.

3.2.2. Tamaño óptimo de una muestra stsrswor

Si consideramos la función de coste lineal: $C(m_1, \dots, m_K) = c_0 + \sum_{i=1}^K c_i m_i$ se tiene (ver, por ejemplo, Särndal, Swensson and Wretman(1992)) que:

- (i) Bajo la hipótesis de minimizar la varianza respecto al diseño para un muestreo stsrswor para un coste dado C , los tamaños óptimos m_i , $i = 1, 2, \dots, K$ vienen dados por:

$$(9) \quad m_i = \frac{(C - c_0) \frac{M_i S_{yi}}{\sqrt{c_i}}}{\sum_{i=1}^K M_i S_{yi} \sqrt{c_i}} = n \frac{\frac{M_i S_{yi}}{\sqrt{c_i}}}{\sum_{i=1}^K \frac{M_i S_{yi}}{\sqrt{c_i}}}.$$

- (ii) Bajo la hipótesis de minimizar el coste para una varianza dada V^2 para un muestreo stsrswor, los tamaños óptimos m_i , $i = 1, 2, \dots, K$ vienen dados por:

$$(10) \quad m_i = \frac{\frac{M_i S_{yi}}{\sqrt{c_i}} \sum_{i=1}^K M_i S_{yi} \sqrt{c_i}}{(NV)^2 + \sum_{i=1}^K M_i S_{yi}^2}.$$

Observemos que a partir de (9) se obtiene tomando $c_i = c$, $i = 1, \dots, K$, la asignación de Newman (el resultado sobre tamaño óptimo más popular), que consiste en la selección de muestras aleatorias estratificadas con tamaños de muestreo por estrato m_i , $i = 1, 2, \dots, K$ dados por:

$$(11) \quad m_i = n \frac{M_i S_{yi}}{\sum_{i=1}^K M_i S_{yi}}.$$

Esta expresión fue obtenida por Neyman(1934) y previamente por Tschuprow(1923).

A partir de (9) y (10) se obtiene una regla práctica para diseñar un muestreo aleatorio estratificado que consiste en seleccionar, dado un estrato, una muestra más grande cuanto más grande sea su tamaño, más variable sea y menor sea el coste de muestrearlo, y las expresiones de los tamaños de los estratos y tamaño total de la muestra. Sin embargo, la obtención de estos tamaños presenta algunas dificultades prácticas, al desconocimiento de S_{yi}^2 (habitual en este tipo de expresiones y que discutiremos en el próximo párrafo) hay que añadir ahora otros problemas que se producen por la no introducción en los problemas de optimización resueltos en (9) y (10) de restricciones sobre los tamaños de la muestra en los estratos y que son: la posibilidad de que el tamaño óptimo de muestreo en cada estrato no sea entero, sea negativo o sea superior al tamaño del estrato. Las recomendaciones habituales son las siguientes: si el tamaño óptimo de muestreo supera al tamaño del estrato, muestrear estos estratos totalmente y distribuir las restantes observaciones entre los otros estratos, y si el tamaño óptimo es negativo no seleccionar ningún elemento de estos estratos, modificar la ecuación de coste y obtener los nuevos óptimos.

En la práctica, se suelen emplear las siguientes 5 aproximaciones de la fórmula (11) para evitar el problema que supone desconocer S_{yi}^2 , $i = 1, 2, \dots, K$:

- (i) Asignación proporcional al total poblacional (que presenta problemas prácticos similares a (11) puesto que los totales de cada estrato son cantidades asimismo desconocidas)

$$(12) \quad m_i = n \frac{\sum_{j=1}^{M_i} y_{ij}}{\sum_{i=1}^K \sum_{j=1}^{M_i} y_{ij}}.$$

- (ii) Asignación proporcional al tamaño (que aplicábamos en (2) para establecer la comparación entre los muestreos srswor y stsrswor):

$$(13) \quad m_i = n \frac{M_i}{N}.$$

(iii) x -«óptima» asignación (con x una característica auxiliar que proporciona información acerca de y):

$$(14) \quad m_i = n \frac{M_i S_{xi}}{\sum_{i=1}^K M_i S_{xi}}.$$

(iv) Asignación proporcional a la raíz de las medias en cada estrato de una característica auxiliar:

$$(15) \quad m_i = n \frac{M_i \sqrt{\bar{x}_i}}{\sum_{i=1}^K M_i \sqrt{\bar{x}_i}}.$$

(v) Asignación proporcional al total de las x (con x una característica auxiliar que proporciona información acerca de y):

$$(16) \quad m_i = n \frac{\sum_{j=1}^{M_i} x_{ij}}{\sum_{i=1}^K \sum_{j=1}^{M_i} x_{ij}}.$$

La aplicación de las asignaciones (12) a (16) produce obviamente una pérdida en la precisión de las inferencias (en términos de la varianza) respecto a la aplicación de (11). Estas pérdidas de eficiencia dependerán en los casos (14) a (16) de la correlación entre la característica y en estudio y la característica auxiliar x (de valores conocidos para todos los elementos de la población) y en general de cuan buena sea la aproximación de S_{yi} por S_{xi} , $\sqrt{\bar{x}_i}$ o $\bar{x}_i$, $i = 1, 2, \dots, K$ respectivamente. En el caso de la asignación proporcional (13) (la aproximación más empleada, porque no requiere conocer una característica auxiliar), Hendayat and Sinha(1991) obtienen la siguiente expresión:

$$(17) \quad V_{\text{stsrswor},(13)}(\bar{y}_s) - V_{\text{stsrswor},(11)}\left(\frac{\sum_{i=1}^K M_i \bar{y}_{si}}{N}\right) = \frac{1}{nN} \sum_{i=1}^K M_i (S_{yi} - \frac{\sum_{i=1}^K M_i S_{yi}}{N})^2,$$

de la que se deduce que las diferencias de eficiencia entre las asignaciones de Newman y proporcionales pueden ser considerables, en especial si las varianzas dentro de cada estrato son muy diferentes entre sí.

3.2.3. Tamaño óptimo de una muestra clsrswor

La literatura no presenta muchos resultados en este campo y propone aplicar bajo la hipótesis de un coste cero para el muestreo dentro de las unidades los resultados que se presentaron en el apartado dedicado al tamaño para un muestreo srswor. Cochran(1977) para muestreos clsrswor con todos los clusters del mismo tamaño M propone una solución al problema de minimizar la varianza para un coste fijo con función de costes $C(k) = c_1 Mk + c_2 \sqrt{k}$.

3.2.4. Tamaño óptimo de una muestra 2ssrswor

Para presentar los resultados distinguiremos entre dos situaciones:

- Todas las unidades tienen el mismo número de elementos: $M_i = M, i = 1, 2, \dots, K$ y $m_i = m$ para $i \in s_0$.
- No todas las unidades tienen el mismo número de elementos.

Suponiendo que todas las unidades son del mismo tamaño, Cochran(1977) estudia el problema de minimizar la varianza para un coste dado, con función de costes lineal de la forma: $C(m, k) = c_1 k + c_2 km$ y propone como método para obtener los tamaños m y k óptimos el siguiente algoritmo sugerido por Eisenhart (ver Cameron(1951)).

Algoritmo (18). Dados $a = c_1 S_{yw}^2$ y $b = \frac{c_2}{M} (S_{yb}^2 - S_{yw}^2)$:

(i) Si $b > 0$ y $\sqrt{\frac{a}{b}} \leq M$, tomar el entero l que cumpla $l \leq \sqrt{\frac{a}{b}} \leq l + 1$ y entonces:

$$\begin{cases} \text{si } \frac{a}{b} \leq l(l+1) & \text{tomar } m = l \\ \text{en otro caso} & \text{tomar } m = l + 1 \end{cases}$$

(ii) Si $b \leq 0$ ó $\sqrt{\frac{a}{b}} > M$, tomar $m = M$.

(iii) Determinado m óptimo, se calcula k fijando el coste o la varianza.

Nuevamente, la determinación del tamaño óptimo depende de cantidades desconocidas (varianza entre los elementos dentro de las unidades y varianza entre las medias de las unidades) aunque en este caso los tamaños óptimos son enteros y positivos. Hedayat and Sinha(1991) se plantean nuevamente el problema de minimizar la varianza para un coste menor o igual a una cantidad fijada (la función de costes estudiada es prácticamente la misma, sólo añaden un coste fijo) y sujeta a la restricción $2 \leq m \leq M$ y $2 \leq k \leq K$ presentando un algoritmo alternativo para la obtención de los tamaños muestrales óptimos m y k .

En el caso más general, en el que el tamaño de todas las unidades no es el mismo, Cochran(1977) asumiendo la función de costes $C(\bar{m}, k) = c_1 k + c_2 k \bar{m}$ donde $\bar{m} = \frac{n}{k}$ y considerando una fracción de submuestreo constante para todas las unidades muestreadas ($\frac{m_i}{M_i} = \bar{m} \frac{K}{N}$, $i \in s_0$) propone aplicar el algoritmo (18) reemplazando m por $\bar{m}$ y tomando $a = c_1 \frac{\sum_{i=1}^K M_i S_{yi}^2}{N}$ y $b = c_2 \left(\frac{K^2 \sum_{i=1}^K M_i^2 (\bar{y}_i - \bar{y})^2}{N^2 (K-1)} - \frac{K \sum_{i=1}^K M_i S_{yi}^2}{N^2} \right)$. La solución obtenida $\bar{m}$ sirve para obtener los tamaños óptimos m_i , $i \in s_0$ aplicando que la fracción de submuestreo es constante para todas las unidades muestreadas.

Otra alternativa propuesta por Hansen, Hurwitz and Madow(1953) para la misma función de coste consiste en tomar, para $i \in s_0$:

$$(19) \quad m_i = M_i \frac{K}{N} \sqrt{\frac{c_1}{c_2} \frac{1 - R_y}{R_y}},$$

y calcular k fijando el costo o la varianza. (La expresión (19) es asimismo aplicable cuando todas las unidades tiene el mismo tamaño.)

Otras referencias adicionales de interés son: Cochran(1977), Konijn(1973) y Hansen, Hurwitz and Madow(1953) con funciones de coste alternativas y/o análisis de otros procedimientos de muestreo en dos etapas.

3.3. Técnicas de estratificación de una población

Desde un punto de vista de población fija, el muestreo stsrswor permite obtener mejores inferencias que los muestreos srswor y clsrswor, y sobre todo cuando además tomamos la muestra de acuerdo con un esquema de asignación óptima. Esta propiedad motiva la búsqueda de procedimientos para dividir los elementos de la población en unidades (cuando la población no está dividida de forma natural) o incluso estratificar a posteriori la población de modo que se puedan aprovechar las ganancias en eficiencia del muestreo stsrswor frente al srswor o mejorar los resultados de un muestreo clsrswor. Trataremos en este apartado de dos cuestiones: cómo construir los estratos y el número de estratos.

La idea más sencilla para construir K estratos con el objetivo de mejorar nuestras estimaciones consiste en establecer una partición en estratos, de modo que las diferencias $M_i S_{iy}^2$ entre los estratos que construyamos no sean muy grandes, de esta manera cuando realicemos el muestreo de los mismos con asignación proporcional al tamaño (por ejemplo), las pérdidas de eficiencia respecto a la asignación de Neyman (ver expresión (17)) serán menores y, en consecuencia, nuestra estimación será mejor. Som(1973) considera tres procedimientos para reducir las diferencias $M_i S_{iy}^2$ entre estratos:

- La regla de Dalenius and Gurney(1951) de construir estratos con valores de $M_i S_{iy}^2$ aproximadamente iguales.

- La regla de Ekman(1959) de construir estratos con valores de $M_i \text{rg}_i$ aproximadamente iguales, con rg_i el rango de la característica en estudio en el estrato i , $i = 1, 2, \dots, K$.
- Y la regla de Mahalanobis, Hansen, Hurwitz y Madow de construir estratos con valores de $M_i \bar{y}_i$ aproximadamente iguales.

Otra posibilidad consiste en hacer una estratificación de la población de modo que la varianza del estimador para la asignación de Pearson (ver (11)) sea directamente pequeña. Si representamos por $f(y)$ a una función de frecuencia sobre la característica en estudio y suponemos ordenada la población de acuerdo con esta característica, una manera de construir los estratos atendiendo a esta idea consistiría en determinar los límites intermedios $[y_{i-1}^0, y_i^0]$ de modo que se minimizara $V_{\text{stsrswor}, (11)}(\frac{\sum_{i=1}^{M_i} M_i \bar{y}_{s_i}}{N}) \simeq \sum_{i=1}^K \omega_i^2 \frac{\sigma_i^2}{m_i}$ con $\omega_i = \int_{y_{i-1}^0}^{y_i^0} f(y) dy$ y $\sigma_i^2 = \frac{1}{\omega_i} \int_{y_{i-1}^0}^{y_i^0} y^2 f(y) dy - [\frac{1}{\omega_i} \int_{y_{i-1}^0}^{y_i^0} y f(y) dy]^2$. Dalenius and Hodges(1959) plantean el problema en estos términos y proponen calcular la cumulativa de $\sqrt{f(y)}$ y tomar los límites intermedios de forma que los intervalos sean aproximadamente iguales en la escala $\text{cum} \sqrt{f(y)}$. (La idea de emplear la cumulativa $\sqrt{f(y)}$ y estos límites se apoya en la hipótesis de que si los estratos son numerosos y de rango pequeño, $f(y)$ dentro del estrato es aproximadamente constante.) Otros trabajos en esta línea (aunque en estos casos se propone muestrear un único elemento por estrato) consisten en la aplicación de los algoritmos de Kossack and Shiledar Baxi(1971) y sus modificaciones en Shiledar Baxi(1982,95).

Al enunciar los criterios de estratificación anteriores nos hemos apoyado en los valores de la característica en estudio, esta hipótesis no es realista ya que no conocemos la totalidad del vector poblacional. En la práctica estos métodos se aplican sobre una característica auxiliar x_i , $i = 1, 2, \dots, N$ y, obviamente, la eficiencia en la construcción depende de la correlación entre la característica en estudio y la auxiliar.

Los procedimientos anteriores se apoyan en construir un número dado, K de estratos. En principio, un aumento de K supondrá (basándonos en un mecanismo de construcción bueno) estratos internamente más homogéneos en sí y más heterogéneos entre sí y por tanto apoyándonos en (3) y (17) mayores ganancias de precisión. En esta dirección, Cochran(1977) apoyándose en una construcción de estratos basada en la regla de Dalenius and Hodges(1959) y suponiendo una distribución $f(y)$ dentro del estrato aproximadamente constante, obtiene la siguiente relación:

$$(20) \quad V_{\text{stsrswor}}(\frac{\sum_{i=1}^{M_i} M_i \bar{y}_{s_i}}{N}) \simeq \frac{V_{\text{srswor}}(\bar{y}_s)}{K^2},$$

de modo que la varianza estratificada es inversamente proporcional al cuadrado del número de estratos. Por otra parte, hay razones a favor de un K no muy alto, así cuan-

do para la construcción de los estratos se emplea una característica auxiliar (que es lo habitual en la práctica) existe un tamaño K' a partir del cual no es posible reducir la varianza por estratificación (este tamaño tope se debe a la varianza de la característica en estudio no explicable por la característica auxiliar). Otro motivo es el coste, mayor número de estratos supone mayores costes por trabajos extra en planificación y cálculo, y estos costes adicionales pueden no justificar las reducciones consiguientes (si se producen) de la varianza.

4. APORTACIONES AL DISEÑO DESDE LA PERSPECTIVA DE MODELOS DE SUPERPOBLACIÓN CLÁSICOS

Las aportaciones desde la perspectiva de modelos de superpoblación se apoyan en considerar que la característica en estudio y_i es realización de una variable aleatoria Y_i , $i = 1, 2, \dots, N$ que se distribuye según un determinado modelo probabilístico. Desde el punto de vista del diseño podemos señalar dos grupos de motivaciones para introducir esta nueva fuente de aleatoriedad. Desde la perspectiva de población fija, la introducción del modelo de superpoblación permite la obtención de estrategias óptimas de diseño definidas, como aquellas que se obtienen minimizando el valor esperado respecto al modelo del error cuadrático medio respecto al diseño (minimizar (sección 2,4)), dentro de la clase de estimadores insesgados respecto al mecanismo de muestreo (ver expresión (sección 2,1)). Y permite valorar en presencia de información auxiliar en forma de una covariable x_i , $i = 1, 2, \dots, N$ conocida, la eficiencia de los estimadores de razón o regresión frente al estimador media, y la eficiencia de las aproximaciones a los tamaños óptimos de la muestra y de construcción de estratos basados en la información proporcionada por esta covariable. Por otra parte, desde una perspectiva predictiva, otro grupo importante de autores defienden que el conocimiento del estadístico sobre las características de la población en estudio le permite argumentar que la población finita es la realización hoy de un determinado modelo de superpoblación o la realización de un modelo que describe un mecanismo o proceso en el mundo real; afirman que el modelo asumido constituye la fuente de aleatoriedad más importante ya que describe la población y definen como estimador óptimo (estimador óptimo desde la perspectiva predictiva) a aquel que minimiza el error cuadrático respecto al modelo (minimiza (sección 2,5)) dentro de la clase de estimadores insesgados respecto al modelo (ver sección 2,3). (Notemos que, para un muestreo no informativo es equivalente minimizar (sección 2,4) o (sección 2,5) para una clase dada de estimadores.)

Las diferencias conceptuales entre los defensores de cada una de estas perspectivas son importantes. De una manera resumida, los argumentos esgrimidos por la perspectiva población fija-modelo de superpoblación son que al realizar el análisis, la población está fijada, no es resultado de un proceso aleatorio y en consecuencia la fuente de aleatoriedad está en el mecanismo de selección de los elementos de las muestras sobre los

que vamos a basar nuestra estimación, y que la estimación predictiva está sujeta a mayores sesgos en la estimación y mayores errores por especificación incorrecta del modelo. Los argumentos respuesta de los autores predictivos son que el modelo de superpoblación tiene las mismas características de objetividad y de interpretación en términos de frecuencias que los modelos considerados de forma habitual en los trabajos estadísticos en poblaciones infinitas, que la circunstancia de que las observaciones de la población finita ya estén realizadas aunque siguen siendo desconocidas para nosotros, no convierte en menos apropiada la aplicación del modelo de superpoblación que si hubieramos considerado el mismo modelo pero antes de que la población finita fuera realizada, y la realización de estudios sobre estimadores robustos frente a errores en la especificación del modelo.

En esta sección distinguiremos entre las aportaciones respecto al diseño realizadas desde las dos perspectivas anteriores: aportaciones desde la metodología población fija-superpoblación y desde la perspectiva predictiva, y nos apoyaremos en los siguientes dos modelos:

— **Modelo Clásico Base** ($C_1(\beta_0, \beta_1; v(x_i), \rho)$):

$$Y_i = \beta_0 + \beta_1 x_i + E_i, \quad i = 1, 2, \dots, N,$$

donde $E_i, i = 1, 2, \dots, N$ es un error aleatorio tal que:

$$E_m(E_i) = 0 \quad C_m(E_i, E_j) = \begin{cases} \sigma^2 v(x_i) & \text{si } i = j \\ \rho \sigma^2 \sqrt{v(x_i)} \sqrt{v(x_j)} & \text{si } i \neq j \end{cases}$$

con $\beta_0, \beta_1, \sigma^2 > 0$ y ρ constantes desconocidas y $v(\cdot)$ una función de x_i conocida con imagen en los reales positivos; también asumimos que $x_i > 0, i = 1, 2, \dots, N$.

— **Modelo Clásico Estratificado** ($C_2(\beta_{0i}, \beta_{1i}; v(x_{ij}), \rho_i)$):

$$Y_{ij} = \beta_{0i} + \beta_{1i} x_{ij} + E_{ij}, \quad j = 1, 2, \dots, M_i, \quad i = 1, 2, \dots, K,$$

donde $E_{ij}, j = 1, 2, \dots, M_i, i = 1, 2, \dots, K$ es un error aleatorio tal que:

$$E_m(E_{ij}) = 0 \quad C_m(E_{ij}, E_{i^*j^*}) = \begin{cases} \sigma_i^2 v(x_{ij}) & \text{si } i = i^*, j = j^* \\ \rho_i \sigma_i^2 \sqrt{v(x_{ij})} \sqrt{v(x_{i^*j^*})} & \text{si } i = i^*, j \neq j^* \\ 0 & \text{si } i \neq i^* \end{cases}$$

con $\beta_{0i}, \beta_{1i}, \sigma_i^2 > 0$ y $\rho_i, i = 1, 2, \dots, K$, constantes desconocidas y $v(\cdot)$ una función de x_{ij} conocida con imagen en los reales positivos; también asumimos que $x_{ij} > 0, j = 1, 2, \dots, M_i, i = 1, 2, \dots, K$.

4.1. Aportaciones población fija-modelo de superpoblación

4.1.1. Estrategias óptimas que se apoyan en el modelo clásico base

Cassel, Särndal and Wretman(1976,77) demuestran la optimalidad de la estrategia $(\bar{y}_s, \text{pswor})$ para el modelo $C_1(0, \mu; v(x_i) = 1, \rho)$ sin características auxiliares $(x_i = 1, i = 1, 2, \dots, N)$ entre los muestreos ppswor con probabilidades de inclusión estrictamente positivas y estimadores lineales e insesgados respecto al muestreo.

Para poblaciones finitas en las que disponemos de información auxiliar en forma de una covariable, Godambe and Joshi(1965) obtienen la optimalidad de la estrategia $(\hat{y}_{HT}, \text{pswor})$ con $\pi_i = \frac{n\sqrt{v_i}}{\sum_{i=1}^N \sqrt{v_i}}$ para el modelo $C_1(0, \beta_1; v(x_i) = v_i, 0)$ con variables error independientes entre los muestreos ppswor con probabilidades de inclusión estrictamente positivas y estimadores de la forma $\hat{y} = \frac{n}{N}\bar{y}_s + \frac{N-n}{N}\hat{y}_u$ con $\hat{y}_u$ un estimador de $\bar{Y}_u$ insesgado respecto al procedimiento de muestreo. Debemos citar también los trabajos anteriores de Godambe(1955) y Hájek(1959) en los que se obtenía la optimalidad de esta misma estrategia para el modelo $C_1(0, \beta_1; v(x_i) = x_i^2, 0)$ minimizando $E_p V_m(\hat{y})$ entre los estimadores lineales y homogéneos insesgados respecto al procedimiento de muestreo (también sin reemplazamiento y con probabilidades de inclusión estrictamente positivas) y que fueran además insesgados respecto al modelo (se minimizaba así (sección 2,4) entre los estimadores que cumplían (sección 2,1) y (sección 2,3)). (Hájek(1959) obtiene este resultado como caso particular de un teorema en el que considerando un modelo más general que el modelo clásico base prueba que la estrategia $(\hat{y}_{HT}, \text{pswor})$ con $\pi_i \propto \sqrt{V_m(Y_i)}/\sqrt{c_i}$, $i = 1, 2, \dots, N$ minimiza $E_p V_m(\hat{y})$ entre los estimadores lineales y homogéneos insesgados respecto al procedimiento de muestreo y los procedimientos de muestreo sin reemplazamiento que cumplen que $C = \sum_{i=1}^N c_i \pi_i$.)

Si revisamos estos tres resultados observamos que la estrategia óptima de muestreo consiste en un muestreo ppswor con $\pi_i \propto \sqrt{V_m(Y_i)}$, $i = 1, 2, \dots, N$. Este último resultado se obtiene también para los modelos de transformación e intercambiabiles estudiados por Cassel, Särndal and Wretman(1976,77), el modelo de regresión no independiente con coeficientes de regresión conocidos de Tam(1984), el modelo de Mukherjee and Sengupta(1989) y el modelo de permutación aleatoria de Rao(1975).

4.1.2. Estrategias óptimas que se apoyan en el modelo clásico estratificado

En poblaciones estratificadas, Cassel, Wretman and Särndal(1977) obtienen, para el modelo $C_2(0, \mu_i; v(x_{ij}) = 1, 0)$ sin características auxiliares $(x_{ij} = 1, j = 1, 2, \dots, M_i, i = 1, 2, \dots, K)$, la optimalidad de la estrategia $(\hat{y}_{HT}, \text{pswor})$ con $\pi_{ij} = \frac{n\sigma_i}{\sum_{i=1}^K M_i \sigma_i}$ entre los muestreos ppswor con probabilidades de inclusión estrictamente positivas y estimadores lineales e insesgados respecto al muestreo.

Observemos que el muestreo strswor con asignación de Neyman (ver (sección 3,11)) satisface las condiciones anteriores, y que la estrategia de muestreo óptima vuelve a ser un muestreo ppswor con $\pi_{ij} \propto \sqrt{V_m(Y_{ij})}$, $j = 1, 2, \dots, M_i$, $i = 1, 2, \dots, K$. Se pueden obtener resultados similares respecto a la estrategia óptima de muestreo para modelos intercambiables (ver Thompson(1978)), modelos de transformación (ver Cassel, Särndal and Wretman(1976,77)) y para modelos de permutación aleatoria (ver Rao and Bellhouse(1978)).

Los modelos de superpoblación se han aplicado también cuando tenemos una característica auxiliar asociada a la característica en estudio para:

- Valorar la utilidad que puede obtenerse al aplicar esta información auxiliar en la construcción del estimador o en la estratificación de una población. Destacar en esta línea de investigación el trabajo de Särndal, Swensson and Wretman(1992).
- Valorar comparativamente entre dos grupos de estrategias en presencia de información auxiliar: las que introducen dicha información en la construcción del estimador pero no en el mecanismo de muestreo y las que la introducen en el estimador y también en el mecanismo de muestreo aplicado. En relación a esta línea de investigación podemos citar el análisis presentado en el manual de Cassel, Särndal and Wretman(1977) y en el artículo de Mukhopadhyay and Tracy(1993), y los artículos de Särndal(1980) y Padmawar(1981), entre otros.

4.2. Aportaciones predictivas

4.2.1. Estrategias predictivas basadas en el modelo clásico básico

Cassel, Särndal and Wretman(1977), comprueban la optimalidad de $\bar{y}_s$ para el modelo $C_1(0, \mu; v(x_i) = 1, \rho)$ sin características auxiliares ($x_i = 1, i = 1, 2, \dots, N$) entre los estimadores lineales e insesgados.

Para poblaciones finitas en las que se dispone de información adicional en forma de una covariable, Brewer(1963) y Royall(1970) considerando el modelo $C_1(0, \beta_1; v(x_i), 0)$ con variables error independientes proponen el estimador:

$$(1) \quad \hat{\bar{y}}_{b1} = \frac{n}{N} \bar{y}_s + \frac{N-n}{N} \hat{\beta}_1 \bar{x}_u,$$

donde $\hat{\beta}_1 = \frac{\sum_{i \in s} \frac{w_i y_i}{v(x_i)}}{\sum_{i \in s} \frac{w_i^2}{v(x_i)}}$, demostrando su optimalidad para el modelo anterior entre los estimadores lineales e insesgados. En este mismo trabajo, Royall(1970) propone la

estrategia óptima de muestreo demostrando que si $v(x_i)$ y $\frac{v(x_i)}{x_i^2}$ son funciones no decrecientes la estrategia de muestreo óptima consiste en un muestreo intencionado en el que se seleccionan con probabilidad 1 los n índices distintos de la población con valor de la característica auxiliar más alto. En otro trabajo, Royall(1971) asumiendo el modelo $C_1(\beta_0, \beta_1; v(x_i) = 1, 0)$ con variables error independientes probó la optimalidad del estimador de regresión:

$$(2) \quad \hat{y}_{b2} = \bar{y}_s + \hat{\beta}_2(\bar{x} - \bar{x}_s),$$

con $\hat{\beta}_2 = \frac{\sum_{i \in s} (x_i - \bar{x}_s) y_i}{\sum_{i \in s} (x_i - \bar{x}_s)^2}$, en la clase de los estimadores lineales e insesgados. Demostrando además que la estrategia óptima de muestreo consiste en seleccionar una muestra balanceada de orden 1 (ver (3)).

El procedimiento de muestreo intencionado propuesto por Royall(1970) basa su optimalidad en el modelo $C_1(0, \beta_1; v(x_i), 0)$ con variables error independientes y en principio, si el vector poblacional no fuera realización de este modelo de superpoblación exactamente, los resultados podrían presentar sesgos considerables (críticas de Cox(1971) y Neyman(1971) a los trabajos de Royall(1970,71)). Como respuesta a esta crítica el artículo de Royall and Herson(1973a) estudia la robustez de los estimadores $\hat{y}_{b1}$ (expresión (1)) e $\hat{y}_{b2}$ (expresión (2)). Definiendo el concepto de muestra balanceada de orden L , como aquella muestra que cumple la ecuación:

$$(3) \quad \frac{\sum_{i \in s} x_i^l}{n} = \frac{\sum_{i=1}^N x_i^l}{N}, \quad l = 1, 2, \dots, L,$$

estos autores obtienen los siguientes resultados sobre la robustez de los estimadores de regresión propuestos:

- (i) Si la muestra cumple (3) para $L = 1$, el estimador $\hat{y}_{b1}$ para $v(x_i) = x_i$ permanece insesgado y óptimo bajo los modelos: $C_1(\mu, 0; v(x_i) = 1, 0)$, $C_1(\beta_0, \beta_1; v(x_i) = 1, 0)$ y $C_1(\beta_0, \beta_1; v(x_i) = x_i, 0)$ todos ellos con variables error independientes.
- (ii) Si la muestra cumple (3), el estimador $\hat{y}_{b1}$ para $v(x_i) = x_i$ permanece insesgado y óptimo para modelos de regresión polinómica de orden L (o inferior) con residuos E_i independientes tales que $V_m(E_i) = \sigma^2 p(x_i)$, $i = 1, 2, \dots, N$ con $p(x)$ un polinomio de x de grado L (o inferior) que no puede contener potencias de x que no aparezcan en la regresión polinómica.
- (iii) Si el procedimiento de muestreo (no necesariamente aleatorio) produce muestras balanceadas, el estimador $\hat{y}_{b1}$ es óptimo para los modelos indicados en (i), (ii) y para el modelo $C_1(0, \beta_1; v(x_i) = x_i, 0)$ con variables error independientes. Pero esta robustez tiene un coste de eficiencia (el mismo para cualquier $L \geq 1$) cuando

el modelo correcto es $C_1(0, \beta_1; v(x_i) = x_i, 0)$, debido a que el procedimiento de muestreo aplicado no es el óptimo (ver párrafos anteriores para ver el muestreo intencionado óptimo).

- (iv) Si la muestra cumple (3), el estimador $\hat{y}_{b2}$ permanece insesgado y óptimo para todos los modelos de regresión polinomiales descritos en el epígrafe (ii), con la característica de que esta robustez no supone en este caso costes en eficiencia cuando el modelo correcto para la población es el modelo $C_1(\beta_0, \beta_1; v(x_i) = 1, 0)$ con variables error independientes.

Debemos señalar que la obtención práctica de muestras balanceadas no es sencilla y que el cumplimiento de la relación (3) es en la práctica muy difícil de conseguir. Royall(1992) propone el uso de procedimientos aleatorios restringidos como protección ante muestras poco balanceadas (no como protección de una mala elección de modelo) que garanticen que la muestra sea balanceada. Un ejemplo de estos procedimientos de muestreo es el «basket method» de Wallenius(1980).

4.2.2. Una técnica predictiva de estratificación

Royall and Herson (1973b), considerando que la población está modelizada correctamente de acuerdo con el modelo $C_1(0, \beta_1; v(x_i) = x_i, 0)$ con variables error independientes, proponen un procedimiento de estratificación y una estrategia conjunta de estimador y mecanismo de muestreo estratificado que son más eficientes que la estrategia de estimador $\hat{y}_{b1}$ con $v(x_i) = x_i$ y muestra balanceada (ver (3)) que comentábamos en la subsección 4.2.1. La propuesta de Royall and Herson también tiene la propiedad de robustez ante desviaciones polinómicas del modelo. Sin entrar en detalles sobre la forma del estimador, el procedimiento de estratificación y muestreo óptimo se obtienen a partir de las siguientes reglas:

- (i) Los estratos se construyen en función de los valores de la característica auxiliar x_i , $i = 1, 2, \dots, N$ ordenando los elementos de la población de menor a mayor valor de la covariable y formando los estratos con los M_1 primeros elementos, los M_2 elementos siguientes y así sucesivamente hasta formar el K -ésimo estrato con los M_K últimos elementos.
- (ii) La muestra se obtiene mediante un muestreo estratificado (se muestrean todos los estratos) seleccionando (por un procedimiento de muestreo aleatorio o intencionado) en cada estrato una muestra balanceada de m_i elementos distintos.
- (iii) Bajo el criterio de minimizar el error cuadrático del estimador respecto al modelo para un coste $C = c_0 + \sum_{i=1}^K c_i m_i$ dado, Royall and Herson(1973b) prueban que el número óptimo de elementos a seleccionar en cada estrato i , $i = 1, 2, \dots, K$, viene dado por la siguiente expresión:

$$(4) \quad m_i = n \frac{M_i \sqrt{\frac{\bar{x}_i}{c_i}}}{\sum_{i=1}^K M_i \sqrt{\frac{\bar{x}_i}{c_i}}}.$$

Observemos que esta expresión coincide cuando $c_i = c$, $i = 1, 2, \dots, K$ con la asignación proporcional a la raíz de las medias en cada estrato (ver (sección 3,15)) que se propone en los trabajos desde la perspectiva de población fija para aproximar la asignación óptima de Newman.

4.2.3. Estrategias predictivas basadas en el modelo clásico estratificado

Royall(1976) obtiene para el modelo $C_2(0, \beta_{1i} = \mu; v(x_{ij}) = 1, \rho_i)$ sin características auxiliares ($x_{ij} = 1$, $j = 1, 2, \dots, M_i$, $i = 1, 2, \dots, K$) la expresión del estimador óptimo para $\bar{Y}$ entre los estimadores lineales e insesgados respecto al modelo para un muestreo en dos etapas sin reemplazamiento. En este mismo trabajo se realizan los siguientes comentarios sobre los tamaños de muestreo óptimos en un muestreo en dos etapas:

- (i) Si $\sigma_i^2 = \sigma^2$ y $\rho_i = \rho > 0$ fijados k y n propone escoger en la primera etapa del muestreo las k unidades distintas de mayor tamaño.
- (ii) Si $\sigma_i^2 = \sigma^2$ y $\rho_i = \rho = 0$ fijados k y n , el tamaño óptimo de muestreo de cada unidad seleccionada cumple que:

$$(5) \quad m_i = \omega \frac{M_i n}{\sum_{i \in s_0} M_i} + (1 - \omega) \frac{n}{k},$$

$$\text{con } \omega = \frac{\sum_{i \in s_0} M_i}{N}.$$

- (iii) Si $\sigma_i^2 = \sigma^2$ y $\rho_i = \rho = 1$ fijados k y n , el tamaño óptimo de muestreo de cada unidad seleccionada es aquél que cumple (5) con $\omega = \frac{\sum_{i \in s_0} M_i}{\sum_{i \in s_0} M_i + k(N - \sum_{i \in s_0} M_i)}$.
- (iv) Si $\sigma_i^2 = \sigma^2$, $\rho_i = \rho = 0$ ó 1 y $M_i = M$, $i = 1, 2, \dots, K$, si tomamos $m_i = m$, $i \in s_0$, el tamaño óptimo de muestreo de cada unidad seleccionada es aquel que cumple para una función de costes $C(k, m) = c_1 k + c_2 n$ la siguiente expresión:

$$(6) \quad m = \frac{n}{k} = \sqrt{\frac{c_1}{c_2} \frac{1 - \rho}{\rho}}.$$

Observemos que esta expresión coincide con la obtenida por Hansen, Hurwitz and Madow(1953) para población fija (ver (sección 3,19)) sustituyendo R_y por ρ .

Observemos que los puntos (i), (ii) y (iii) solamente dan respuestas parciales al problema de tamaño óptimo porque no presentan propuestas para determinar el número óptimo de unidades a muestrear k , el punto (iv), en cambio, proporciona una respuesta a este problema más satisfactoria, ya que a partir de (6) y asumiendo un coste fijo C tendríamos que el número óptimo de unidades a muestrear sería $k = C/(c_1 + c_2\sqrt{[c_1(1-\rho)]/(c_2\rho)})$.

5. APORTACIONES AL DISEÑO DESDE LA PERSPECTIVA DE MODELOS DE SUPERPOBLACIÓN BAYESIANOS

Las aportaciones Bayesianas que presentaremos a continuación se basan en asumir un modelo de superpoblación Bayesiano, desde esta perspectiva la solución inferencial al problema de estimar $\bar{Y} = \frac{n}{N}\bar{y}_s + \frac{N-n}{N}\bar{Y}_u$ consiste en minimizar, dada una muestra s formada con n índices distintos, el error cuadrático del estimador $\hat{\bar{y}} = \frac{n}{N}\bar{y}_s + \frac{N-n}{N}\hat{\bar{y}}_u$ respecto a la distribución predictiva del vector $\mathbf{Y}_u|\mathbf{y}_s$. (En las expresiones anteriores representamos por $\hat{\bar{y}}_u$ a un estimador de $\bar{Y}_u$ y por $\mathbf{Y}_u = (Y_i, i \notin s)$ y $\mathbf{y}_s = (y_i, i \in s)$ a los vectores poblacionales no observados y observados dada la muestra respectivamente.) De acuerdo con esta definición, el estimador Bayes se obtiene tomando $\hat{\bar{y}}_u = E(\bar{Y}_u|\mathbf{y}_s)$ y el error cuadrático medio cometido viene dado por la expresión $V(\bar{Y}|\mathbf{y}_s) = \frac{(N-n)^2}{N^2}V(\bar{Y}_u|\mathbf{y}_s)$.

Los resultados que presentaremos a continuación se apoyan en los siguientes tres modelos Bayesianos:

– **Modelo Bayes Base** ($B_1(\mu_0; \sigma^2, \sigma_0^2)$):

$$Y_i|\mu \sim N(Y_i|\mu, \sigma^2), \text{ indep., } i = 1, 2, \dots, N,$$

$$\mu \sim N(\mu_0, \sigma_0^2),$$

con $\sigma^2 > 0$, $\sigma_0^2 > 0$ y μ_0 escalares conocidos.

– **Modelo Bayes Estratificado** ($B_2(\mu_0; \sigma_i^2, \sigma_0^2)$):

$$Y_{ij}|\mu_i \sim N(Y_{ij}|\mu_i, \sigma_i^2), \text{ indep., } j = 1, 2, \dots, M_i, i = 1, 2, \dots, K,$$

$$\mu_i \sim N(\mu_0, \sigma_0^2), \text{ iid, } i = 1, 2, \dots, K,$$

con $\sigma_i^2 > 0, i = 1, 2, \dots, K, \sigma_0^2 > 0$ y μ_0 escalares conocidos.

– **Modelo Bayes en Dos Etapas** ($B_3(\mu_0; \sigma_i^2, \sigma_\mu^2, \sigma_0^2)$):

$$\begin{aligned} Y_{ij}|\mu_i &\sim N(Y_{ij}|\mu_i, \sigma_i^2), \text{ indep.}, j = 1, 2, \dots, M_i, i = 1, 2, \dots, K, \\ \mu_i|\mu &\sim N(\mu_i|\mu, \sigma_\mu^2), \text{ iid}, i = 1, 2, \dots, K, \\ \mu &\sim N(\mu_0, \sigma_0^2), \end{aligned}$$

con $\sigma_i^2 > 0, i = 1, 2, \dots, K, \sigma_\mu^2 > 0, \sigma_0^2 > 0$ y μ_0 escalares conocidos.

Las notaciones indep. e iid significan respectivamente independientes e igualmente distribuidos, y las notaciones $B_1(\mu_0; \sigma^2, \infty)$ y $B_3(\mu_0; \sigma_i^2, \sigma_\mu^2, \infty)$ que aparecerán más adelante en el texto, representan la distribución inicial de referencia $\pi(\mu) \propto 1$.

A partir de las definiciones anteriores, observemos que:

- En el modelo $B_1(\mu_0; \sigma^2, \sigma_0^2)$, las variables $Y_i, i = 1, 2, \dots, N$ son intercambiables.
- En el modelo $B_2(\mu_0; \sigma_i^2, \sigma_0^2)$, las variables $Y_{ij}, j = 1, 2, \dots, M_i$ son intercambiables para $i = 1, 2, \dots, K$.
- En el modelo $B_3(\mu_0; \sigma_i^2, \sigma_\mu^2, \sigma_0^2)$, las variables $Y_{ij}|\mu_i, j = 1, 2, \dots, M_i$ son intercambiables con medias por unidad, $\mu_i, i = 1, 2, \dots, K$ asimismo intercambiables.

La intercambiabilidad expresa, desde un punto de vista subjetivo, que no hay información sobre la característica en estudio y_i en el índice $i, i = 1, 2, \dots, N$ y por lo tanto no hay razón para considerar que las inferencias basadas en una muestra de tamaño fijo n (con índices distintos) son mejores que las basadas en otra muestra distinta también del mismo tamaño (y también con índices distintos). Esta noción de intercambiabilidad ajusta con la idea objetiva de un muestreo srswor, y en consecuencia podemos asociar al modelo $B_1(\mu_0; \sigma^2, \sigma_0^2)$ un muestreo srswor, al modelo $B_2(\mu_0; \sigma_i^2, \sigma_0^2)$ un muestreo stsrswor y al modelo $B_3(\mu_0; \sigma_i^2, \sigma_\mu^2, \sigma_0^2)$ un muestreo 2ssrswor.

Distinguiremos a continuación entre dos grupos de aportaciones: las relativas a la comparación desde una perspectiva Bayes de los muestreos srswor, stsrswor y clsrswor, y las relativas a los tamaños óptimos de muestreo.

5.1. Comparación Bayesiana entre los muestreos srswor, stsrswor y clsrswor

Bayarri and Font(1994,98) y Font(1995a) apoyándose en los modelos $B_1(0; \sigma^2, \infty)$ para el modelo srswor y $B_3(0; \lambda\sigma^2, (1-\lambda)\sigma^2, \infty)$ con $0 < \lambda < 1$ conocido para un muestreo 2ssrswor (que tiene como casos particulares para $k = K$ el muestreo stsrswor y para $m_i = M_i, i \in s_0$ el muestreo clsrswor) obtienen los siguientes resultados:

- (i) Si $K = k$, $M_i = M$, $m_i = m$, $i = 1, 2, \dots, K$, los estimadores Bayes para los modelos $B_1(0; \sigma^2, \infty)$ y $B_3(0; \lambda\sigma^2, (1 - \lambda)\sigma^2, \infty)$ coinciden con la media muestral $\bar{y}_s$, y además:

$$(1) \quad V_{B_1(\text{srswor})}(\bar{Y}|\mathbf{y}_s) - V_{B_3(\text{stsrswor})}(\bar{Y}|\mathbf{y}_s) = \frac{N - n}{N} \frac{\sigma^2}{n} \rho.$$

- (ii) Si $K < k$, $M_i = M$, $i = 1, 2, \dots, K$ y $m_i = M$, $i \in s_0$, los estimadores Bayes para los modelos $B_1(0; \sigma^2, \infty)$ y $B_3(0; \lambda\sigma^2, (1 - \lambda)\sigma^2, \infty)$ coinciden con la media muestral $\bar{y}_s$, y además:

$$(2) \quad V_{B_3(\text{clsrswor})}(\bar{Y}|\mathbf{y}_s) - V_{B_1(\text{srswor})}(\bar{Y}|\mathbf{y}_s) = \frac{N - n}{N} \frac{\sigma^2}{n} (M - 1) \rho.$$

Donde $\rho = C(Y_{ij}, Y_{ij^*} | \mu)$, $j \neq j^*$, $i = 1, 2, \dots, K$ para el modelo $B_3(0; \lambda\sigma^2, (1 - \lambda)\sigma^2, \infty)$.

Observemos que (1) y (2) coinciden con los resultados aproximados desde la perspectiva de población fija (sección 3,4) y (sección 3,5) respectivamente, sustituyendo ρ por R_y y σ^2 por S_y^2 (la correlación y varianza del modelo por las poblacionales).

El artículo de Bayarri and Font(1998) estudia también la estimación de $\lambda = 1 - \rho$ para muestreos stsrswor y clsrswor, considerando el modelo $B_3(\mu_0; \lambda\sigma^2, (1 - \lambda)\sigma^2, \sigma_0^2)$ y una distribución sobre λ uniforme $(0,1)$.

5.2. Tamaños Bayes de muestreo óptimos

Murgui(1982) presenta para los modelos $B_1(\mu_0; \sigma^2, \sigma_0^2)$ y $B_2(\mu_0; \sigma_i^2, \sigma_0^2)$ los siguientes resultados sobre tamaño óptimo:

- (i) Considerando el modelo $B_1(\mu_0; \sigma^2, \sigma_0^2)$, se obtiene bajo el criterio de minimizar el valor esperado respecto a un muestreo srswor de la función $P(n) = \gamma V(\bar{Y}|\mathbf{y}_s) + c_0 + c_1 n$ la siguiente expresión para el tamaño óptimo:

$$(3) \quad n = \frac{\sigma \sqrt{\gamma}}{\sqrt{c_1}} \left(1 + \frac{\sigma^2}{N \sigma_0^2} \right) - \frac{\sigma^2}{\sigma_0^2},$$

que para $\sigma_0^2 \rightarrow \infty$ (distribución inicial no informativa) coincide con el tamaño óptimo desde la perspectiva de población fija para una muestra srswor (ver expresión (sección 3,8)) sustituyendo σ por S_y .

- (ii) Considerando el modelo $B_2(\mu_0; \sigma_i^2, \sigma_0^2)$, bajo la hipótesis de minimizar $E_{\text{stsrswor}} V(\bar{Y} | y_s)$ para un coste $C = c_0 + \sum_{i=1}^K c_i m_i$ obtiene la siguiente expresión para los tamaños de muestreo óptimos para cada estrato:

$$(4) \quad m_i = \frac{(\sigma_i^2 + M_i \sigma_0^2) \sigma_i}{\sqrt{c_i} \sigma_0^2} \frac{C - c_0 + \sum_{i=1}^K \frac{\sigma_i^2 c_i}{\sigma_0^2}}{\sum_{i=1}^K \frac{\sigma_i^3 \sqrt{c_i}}{\sigma_0^2} + \sum_{i=1}^K \sigma_i M_i \sqrt{c_i}} - \frac{\sigma_i^2}{\sigma_0^2},$$

que para $\sigma_0^2 \rightarrow \infty$ (distribución inicial no informativa) coincide con los tamaños óptimos desde la perspectiva de población fija para un muestreo stsrswor (ver expresión (sección 3,9)) sustituyendo σ_i por S_{yi} , $i = 1, 2, \dots, K$ y con la asignación de Neyman (ver expresión (sección 3,11)) cuando $c_i = c$, $i = 1, 2, \dots, K$ con la misma sustitución.

Ericson(1969b) obtiene para un modelo estratificado general intercambiable una solución alternativa a (4) (exacta para algunos valores de C y un algoritmo iterativo para los restantes valores) para obtener el tamaño óptimo de los estratos que minimiza $E_p V(\bar{Y} | y_s)$ bajo las restricciones: $\sum_{i=1}^K c_i m_i \leq C$ y $0 \leq m_i \leq M_i$, $i = 1, 2, \dots, K$. Además, en un trabajo posterior, Ericson(1988), considerando un modelo en 2 etapas intercambiable más general que $B_3(\mu_0; \sigma_i^2, \sigma_\mu^2, \sigma_0^2)$ y bajo la hipótesis de unidades del mismo tamaño ($M_i = M$, $i = 1, 2, \dots, K$), resuelve el problema de minimizar $E_{2\text{ssrswor}} V(\bar{Y} | y_s)$ bajo las restricciones: $c_1 k + c_2 n \leq C$, $1 \leq k \leq K$ y $0 \leq m_i \leq M_i$, $i \in s_0$. Traduciendo estos últimos resultados a nuestro modelo con $M_i = M$ y $\sigma_i^2 = \sigma^2$, $i = 1, 2, \dots, K$, podemos escribir el siguiente algoritmo para determinar el número de unidades y tamaño de las unidades óptimos según este criterio:

Algoritmo (5)

- (i) Si $\sigma_\mu^2 \leq \frac{\sigma^2}{M-1}$, tomar el entero l que cumpla $l \leq \frac{C}{c_1 + M c_2} \leq l + 1$ y entonces:

$$\begin{cases} \text{si } C < (l+1)c_1 + (lM+1)c_2 & \text{tomar } k = l \\ \text{en otro caso} & \text{tomar } k = l + 1 \end{cases}$$

Además $m_i = M$, $i \in \{i_1, \dots, i_l\}$ y m_{l+1} (tamaño de la unidad $l+1$ a muestrear si $k = l+1$) es un entero tal que $m_{l+1} \leq \frac{C - (k+1)c_1 - kMc_2}{c_2} \leq m_{l+1} + 1$.

- (ii) Si $\sigma_\mu^2 > \frac{\sigma^2}{M-1}$, encontrar k y m , con $m_i = m$, $i \in s_0$ que minimicen la siguiente expresión:

$$V^*(k, m) = \frac{(K-k)\sigma_\mu^2}{Kk} + \frac{(M-m)\sigma^2}{(M-1)mk}.$$

La expresión $V^*(k, m)$ tiene la misma forma que la expresión de la varianza para una muestra 2ssrswor desde la perspectiva de población fija, sustituyendo S_{yw}^2 por $\frac{M}{M-1}\sigma^2 = E(\frac{\sum_{i=1}^K \sum_{j=1}^M (y_{ij} - \mu_i)^2}{N-K})$ y S_{yb}^2 por $M\sigma_\mu^2 = E(\frac{\sum_{i=1}^K M(\mu_i - \bar{\mu})^2}{K-1})$, (con $\bar{\mu} = \frac{\sum_{i=1}^K \mu_i}{K}$). Esta identificación entre las dos expresiones permite aplicar para minimizar $V^*(k, m)$ el algoritmo de Eisenhart (algoritmo (sección 3,18)) tomando $a = c_1 \frac{M}{M-1} \sigma^2$ y $b = c_2 (\sigma_\mu^2 - \frac{\sigma^2}{M-1})$. Como observación general a todos estos resultados debidos a Ericson, conviene recordar que aunque el criterio de estimador óptimo aplicado por Ericson(1969b,88) consiste en obtener el estimador lineal Bayesiano (el estimador lineal que minimiza el error cuadrático medio entre los estimadores lineales), éste coincide bajo la hipótesis de normalidad con el estimador Bayes (el aplicado por los restantes autores citados).

Murgui(1982) aporta también resultados de tamaño óptimo para un modelo Bayesiano normal-gamma sin características auxiliares, un modelo Bayesiano estratificado normal-gamma sin características auxiliares y una aproximación al tamaño óptimo para el modelo $B_3(\mu_0; \sigma_i^2, \sigma_\mu^2, \sigma_0^2)$.

Observemos también que para modelos de regresión lineal (cuando conocemos una característica auxiliar de la población), el procedimiento de muestreo óptimo que minimiza $E_p V(\bar{Y} | \mathbf{y}_s)$ y los objetivos de robustez respecto al modelo dan lugar a procedimientos de muestreo intencionados y a muestras balanceadas Bayesianas respectivamente (ver Murgui(1982)). En general, a partir de los resultados basados en modelos de superpoblación Bayesianos y considerando distribuciones iniciales de referencia, se recuperan los resultados clásicos predictivos sobre estimación y diseño.

6. CONCLUSIONES Y REFERENCIAS ADICIONALES SOBRE OTRAS APORTACIONES AL DISEÑO EN POBLACIONES FINITAS

Las metodologías basadas en población fija y en modelos de superpoblación clásicos y Bayesianos obtienen, en ausencia de características auxiliares, resultados de diseño muy «similares». Si revisamos esta exposición observamos que las únicas diferencias entre los resultados de población fija y de superpoblación clásicos surgen de tomar la varianzas respecto al modelo en vez de las varianzas poblacionales y entre los resultados de superpoblación clásicos y Bayesianos en la presencia en estos últimos de parámetros de la distribución inicial que modifican levemente los resultados. En cambio, la coincidencia entre los resultados basados en las dos aproximaciones al problema se «rompe» al introducir en la inferencia la información derivada de una característica auxiliar, defendiendo la primera un muestreo aleatorio y la segunda un muestreo intencionado.

Se ofrecen a continuación, precedidas por un título orientativo de su contenido, algunas referencias adicionales sobre aportaciones no analizadas en este trabajo en referencia al diseño en poblaciones finitas:

- **La solución minimax a los problemas de tamaño óptimo.** Consultar Solomon and Zacks(1970), Bellhouse(84), Chaudhuri(1988) y Bolfarine y Zacks(1991).
- **La aplicación del diseño experimental a los problemas de muestreo en poblaciones finitas.** Ver, por ejemplo, Rao(1979) que presenta un resumen con referencias sobre las aplicaciones más extendidas y dentro de las aportaciones a muestreo controlado, los artículos de Chakrabarti(1963), Srivastava and Collins(1985) y Srivastava and Ouyang(1992).
- **Aportaciones al diseño multiobjetivo.** Ver Malec(1995).
- **Obtención de estrategias óptimas asintóticas.** Ver Bellhouse(1984) y Bolfarine and Zacks(1992).
- **Aportaciones al diseño con más de una característica auxiliar.** Ver, por ejemplo, Rao(1993) en población fija, Royall and Pfeffermann(1982), Royall(1992b), Bolfarine and Zacks(1992) y Tam(1995) desde la perspectiva predictiva, y Draper and Guttman(1968b) desde la perspectiva Bayesiana para un muestreo en dos fases.
- **Procedimientos de muestreo en dos fases y secuenciales óptimos.** Consultar los manuales clásicos de población fija (por ejemplo Cochran(1977)) y desde un punto de vista Bayesiano, Zacks(1969) para la aproximación teórica al problema y De-Groot and Starr(1969) y Draper and Guttman(1968a,b) en un trabajo aplicado a la estimación de las medias poblacionales de una población finita.
- **Inferencias sobre la media poblacional a partir de una muestra obtenida a partir de una ruta aleatoria.** Ver: Kish(1965) desde la aproximación de población fija, y Font(1995a) y Bayarri and Font(1996) desde la aproximación Bayes.

7. AGRADECIMIENTOS

Quiero expresar aquí mi agradecimiento a la Dra. M.J. Bayarri por la lectura previa de este trabajo y sus muy útiles y adecuados comentarios. También deseo expresar mi gratitud al Dr. Carles M. Cuadras y a dos referees anónimos de la revista *Qüestió* por sus comentarios y sugerencias. Este trabajo ha sido parcialmente financiado por el proyecto de la DGES PB96-0776.

8. REFERENCIAS

- [1] **Bayarri, M.J. and Font, B.** (1994). «A (Bayesian) note on non-random samples from finite populations». In *1994 Proceedings of the section on Bayesian statistical science. ASA.*, 246-251.
- [2] **Bayarri, M.J. and Font, B.** (1996). «Bayesian hierarchical models for random routes in finite populations». In *Data analysis and information system.* Bock and Polasek, eds. Springer. 301-312.
- [3] **Bayarri, M.J. and Font, B.** (1998). «A Bayesian comparison of cluster, strata and random samples». *Journal of Statistical Planning and Inference.* Forthcoming.
- [4] **Bellhouse, D.R.** (1984). «A review of optimal designs in survey sampling». *The Canadian Journal of Statistics*, **12**, 53-65.
- [5] **Bolfarine, H. and Zacks, S.** (1991). «Bayes and minimax prediction in finite population». *J. Statist. planning and inference*, **28**, 139-151.
- [6] **Bolfarine, H. and Zacks, S.** (1992). *Prediction Theory for finite populations.* Springer-Verlag. New York.
- [7] **Brewer, K.R.W.** (1963). «Ratio estimation and finite population. Some results deducible from the assumption of an underlying stochastic process». *Aust. J. Statist.*, **5**, 93-105.
- [8] **Cameron, J.M.** (1951). «Use of variance components in preparing schedules for the sampling of baled wool». *Biometrics*, **7**, 83-96.
- [9] **Cassel, C.M, Särndal, C.E. and Wretman, J.H.** (1976). «Some results on generalized difference estimation and generalized regression estimation for finite populations». *Biometrika*, **63**, 615-620.
- [10] **Cassel, C.M, Särndal, C.E. and Wretman, J.H.** (1977). *Foundations of inference in survey sampling.* John Wiley and Sons. New York.
- [11] **Chakrabarti, M.C.** (1963). «On the use of incidence matrices in sampling from a finite population». *Journal of the Indian Statist. Assoc.*, **1**, 78-85.
- [12] **Chaudhuri, A.** (1988). «Optimality of sampling strategies». *Handbook of Statistics*, **6**, 47-96.
- [13] **Cochran, W.G.** (1977). *Sampling Techniques. Third Edition.* John Wiley and Sons, New York.
- [14] **Cox, D.R.** (1971). «Discussion of Royall(1971)». *Foundations of statistical inference.* V.P. Godambe and D.A. Sprott, eds. Holt, Rinehart and Winston. Toronto. 275.
- [15] **Dalenius, T. and Gurney, M.** (1951). «The problem of optimum stratification. II». *Skand. Akt.*, **34**, 133-148.

- [16] **Dalenius, T. and Hodges, J.L.Jr.** (1959). «Minimum variance stratification». *Jour. Amer. Stat. Assoc.*, **54**, 88-101.
- [17] **DeGroot, M.H. and Starr, N.** (1969). «Optimal two-stage stratified sampling». *The Annals of Math. Statist.*, **40**, 575-582.
- [18] **Des Raj** (1954). «On sampling with probabilities proportionate to size». *Ganita*, **5**, 175-182.
- [19] **Draper, N.R. and Guttman, I.** (1968a). «Some Bayesian stratified two-phase sampling results». *Biometrika*, **55**, 131-139.
- [20] **Draper, N.R. and Guttman, I.** (1968b). «Bayesian stratified two-phase sampling results: k characteristics». *Biometrika*, **55**, 587-589.
- [21] **Ekman, G.** (1959). «An approximation useful in univariate stratification». *Ann. Math. Stat.*, **30**, 219-229.
- [22] **Ericson, W.A.** (1969a). «Subjective Bayesian models in sampling finite populations». *J. Roy. Statist. Soc. (Ser. B)*, **31**, 195-233.
- [23] **Ericson, W.A.** (1969b). «Subjective Bayesian models in sampling finite populations. Stratification». In *New developments in survey sampling*. N.L. Johnson and Harry Smith Jr., eds. John Wiley and sons.
- [24] **Ericson, W.A.** (1988). «Bayesian inference in finite populations». *Handbook of Statistics*, **6**, 213-246.
- [25] **Font, B.** (1995a). *Análisis Bayesiano de muestras no aleatorias en poblaciones finitas*. Tesis doctoral. Universitat de València.
- [26] **Font, B.** (1995b). «Inferencia Bayesiana en poblaciones finitas: un análisis comparativo». *Estadística Española*, **37**, **120**, 309-320.
- [27] **Godambe, V.P.** (1955). «A unified theory of sampling from finite population». *J. Roy. Statist. Soc. (Ser. B)*, **17**, 269-278.
- [28] **Godambe, V.P.** (1992). «Survey Sampling - as I understand it». *Current issues in statistical inference. Essays in honor of D. Basu*. Ed. M. Ghosh and P.K. Pathak, **17**, 187-195.
- [29] **Godambe, V.P. and Joshi, V.M.** (1965). «Admissibility and Bayes estimation in sampling finite population I». *Ann. Math. Statist.*, **36**, 1707-1722.
- [30] **Hájek, J.** (1959). «Optimum strategy and other problems in probability sampling». *Časopis Pěst. Mat.*, **84**, 387-423.
- [31] **Hansen, M.H., Hurwitz, W.N. and Madow, W.G.** (1953). *Sample survey methods and theory*. Vol. I and II. John Wiley and Sons. New York.
- [32] **Hedayat, A.S. and Sinha, B.K.** (1991). *Design and inference in finite population sampling*. John Wiley and Sons. New York.

- [33] **Joshi, V.M.** (1965). «Admissibility and Bayes estimation in sampling finite population II, III». *Ann. Math. Statist.*, **36**, 1723-1742.
- [34] **Joshi, V.M.** (1967). «Confidence intervals for the mean of a finite population». *Ann. Math. Statist.*, **38**, 1180-1207.
- [35] **Kish, L.** (1965). *Survey Sampling*. John Wiley and Sons. New York.
- [36] **Konijn, H.S.** (1973). *Statistical theory of sample survey design and analysis*. North Holland. Amsterdam.
- [37] **Kossack, C.F.** and **Shiledar Baxi, H.R.** (1971). «On designing of a unit-stratified survey design for discrete set of observations». *Internat. Statist. Rev.*, **39**, 46-56.
- [38] **Lanke, J.** (1975). *Some contributions to the theory of survey sampling*. Ph. D. thesis. University of Lund.
- [39] **Malec, D.** (1995). «Selecting multiple-objective fixed-cost sample designs using an admissibility criterion». *Journal of statistical planning and inference*, **48**, 229-240.
- [40] **Mukherjee, R.** and **Sengupta, S.** (1989). «Optimal estimation of a finite population total under a general correlated model». *Biometrika*, **76**, 789-794.
- [41] **Mukhopadhyay, P.** and **Tracy, D.S.** (1993). «Some inferential aspects in sampling finite populations under a fixed population set up: a review». *Journal of Applied Statistical Science*, **1**, 95-123.
- [42] **Murgui, J.S.** (1982). *Diseño e inferencia en poblaciones finitas: modelos de superpoblación*. Tesis doctoral. Universitat de València.
- [43] **Newman, J.** (1934). «On the two different aspects of the representative method. The method of stratified sampling and the method of purposive selection». *J. Roy. Stat. Soc.*, **97**, 558-625.
- [44] **Neyman, J.** (1971). «Discussion of Royall(1971)». *Foundations of statistical inference*. V.P. Godambe and D.A. Sprott, eds. Holt, Rinehart and Winston. Toronto, 276-278.
- [45] **Padmawar, V.R.** (1981). «A note on the comparison of certain sampling strategies». *J. Roy. Statist. Soc. (Ser. B)*, **43**, 321-326.
- [46] **Rao, J.N.K.** (1975). «On the foundations of survey sampling». *A Survey of statistical design and linear models*. J.N. Srivastava, ed. North Holland, Amsterdam, 489-506.
- [47] **Rao, J.N.K.** (1979). «Optimization in the desing of sample surveys». In J.S. Rustagi (ed.), *Optimizing Methods in Statistics: Proceedings of an International Conference*. New York. Academic Press. 419-434.
- [48] **Rao, J.N.K.** and **Bellhouse, D.R.** (1978). «Optimal estimation of a finite population mean under generalized random permutation models». *Journal of statistical planning and inference*, **2**, 125-141.

- [49] **Rao, T.J.** (1993). «On certain problems of sampling designs and estimation for multiple characteristics». *Sankhyā*, Ser. B, **55**, 372-384.
- [50] **Reddy, V.N. and Rao, T.J.** (1977). «Modified PPS method of estimation». *Sankhyā*, Ser. C, **39**, 185-197.
- [51] **Royall, R.M.** (1968). «An old approach to finite population sampling theory». *J. Amer. Stat. Ass.*, **63**, 1269-1279.
- [52] **Royall, R.M.** (1970). «On finite population sampling theory under certain linear regression models». *Biometrika*, **57**, 377-387.
- [53] **Royall, R.M.** (1971). «Linear regression models in finite populations sampling theory». *Foundations of statistical inference*. V.P. Godambe and D.A. Sprott, eds. Holt, Rinehart and Winston. Toronto.
- [54] **Royall, R.M.** (1976). «The linear least-squares prediction approach to two-stage sampling». *J. Amer. Stat. Ass.*, **71**, 657-664.
- [55] **Royall, R.M.** (1982). «Balanced samples and robust Bayesian inference in finite population sampling». *Biometrika*, **69**, 401-409.
- [56] **Royall, R.M.** (1988). «The prediction approach to sampling theory». *Handbook of Statistics*, **6**, 399-413.
- [57] **Royall, R.M.** (1992a). «The model based (prediction) approach to finite population sampling theory». *Current issues in statistical inference. Essays in honor of D. Basu*. Ed. M. Ghosh and P.K. Pathak, **17**, 225-240.
- [58] **Royall, R.M.** (1992b). «Robustness and optimal design under prediction models for finite populations». *Survey methodology*, **18**, 179-185.
- [59] **Royall, R.M. and Herson, J.** (1973a). «Robust estimation in finite population I». *J. Amer. Stat. Ass.*, **68**, 880-889.
- [60] **Royall, R.M. and Herson, J.** (1973b). «Robust estimation in finite population II: stratification on a size variable». *J. Amer. Stat. Ass.*, **68**, 890-893.
- [61] **Särndal, C.E.** (1978). «Design-based and model-based inference in survey sampling». *Scand. J. Statist.*, **5**, 27-52.
- [62] **Särndal, C.E.** (1980). «A two-way classification of regression estimation strategies in probability sampling». *The Canadian Journal of Statistics*, **8**, 165-177.
- [63] **Särndal, C.E., Swensson, B. and Wretman, J.** (1992). *Model assisted survey sampling*. Springer-Verlag. New York.
- [64] **Shilekar Baxi, H.R.** (1982). «Approximately optimum stratified design for a finite population». *Sankhyā*, Ser. B, **44**, 102-113.
- [65] **Shilekar Baxi, H.R.** (1995). «Approximately optimum stratified design for a finite population II». *Sankhyā*, Ser. B, **57**, 391-404.

- [66] **Solomon, H. and Zacks, S.** (1970). «Optimal design of sampling from finite populations: a critical review and indication of new research areas». *J. Amer. Stat. Ass.*, **65**, 653-677.
- [67] **Som, R.K.** (1973). *A manual of sampling techniques*. Heinemann. London.
- [68] **Srivastava, J. and Collins, F.** (1985). «On a general theory of sampling, using experimental design concepts I: estimation». *Bulletin de l'Institut International of Statistique*, **51**, 10.3, 1-16.
- [69] **Srivastava, J. and Ouyang, Z.** (1992). «Sampling theory using experimental design concepts». In *Current issues in statistical inference: essays in honor of D. Basu*. M. Ghosh and P.K. Pathak, eds. 241-264.
- [70] **Tam, S.M.** (1984). «Optimal estimation in survey sampling under a regression superpopulation model». *Biometrika*, **71**, 645-647.
- [71] **Tam, S.M.** (1995). «Optimal and robust strategies for cluster sampling». *J. Amer. Stat. Ass.*, **90**, 379-382.
- [72] **Thompson, M.E.** (1978). «Stratified sampling with exchangeable prior distribution». *The Annals of Statistics*, **6**, 1168-1169.
- [73] **Tschuprow, A.A.** (1923). «On the mathematical expectation of the moments of frequency distributions in the case of correlated observation». *Metron*, **2**, 461-493, 646-683.
- [74] **Wallenius, K.T.** (1980). «Statistical methods in sole source contract negotiation». *Journal of undergraduate mathematics and applications*, **0**, 35-47.
- [75] **Yates, F.** (1960). *Sampling methods for censuses and surveys*. Charles Griffin and Co. London. Third Edition.
- [76] **Zacks, S.** (1969). «Bayes sequential designs of fixed size samples from finite populations». *J. Amer. Stat. Ass.*, **64**, 1342-1349.

ENGLISH SUMMARY

A COMPARATIVE REVIEW OF DIFFERENT CONTRIBUTIONS IN SURVEY SAMPLING

B. FONT

Universitat de València*

Design is an important and applied area in finite populations. This article is a comparative review of some results in designing based on three different approaches: fixed population approach, inference based on classical superpopulation models and inference based on Bayesian superpopulation models. The article is organized in six sections, including an introduction. The contents of each section are described in this summary.

Keywords: Bayesian superpopulation models, classical superpopulation models, optimal design, sampling strategies, survey sampling.

AMS Classification: 62D05.

*Begonia Font Belaire. Dept. Economia Financera i Matemàtica. Edifici Departamental Oriental. Av. dels Tarongers s/n. 46071 València (Espanya).

–Received August 1997.

–Accepted January 1999.

Section 2: Basic concepts and notation. We review the more basic concepts and definitions, and establish the usual notation.

Section 3: Contributions based on fixed population approach. A review of the more known results in this approximation is presented including sampling designs, optimum size and stratification procedures. We point out the following results for later comparative reference in this summary:

(i) Earnings from stratification:

$$(1) \quad V_{\text{srswor}}(\bar{y}_s) - V_{\text{stratification}}(\bar{y}_s) \simeq \frac{N-n}{N} \frac{S_y^2}{n} R_y.$$

[Hansen et al(1953)]

(ii) Loses from cluster:

$$(2) \quad V_{\text{cluster}}(\bar{y}_s) - V_{\text{srswor}}(\bar{y}_s) \simeq \frac{N-n}{N} \frac{S_y^2}{n} (M-1) R_y.$$

[Cochran(1977)]

(iii) Random sampling optimum size:

$$(3) \quad n = \frac{S_y^2 \sqrt{\gamma}}{\sqrt{c_1}}.$$

[Cochran(1977)]

(iv) Stratification optimum size:

$$(4) \quad m_i = n \frac{M_i S_{yi}}{\sum_{i=1}^K M_i S_{yi}}.$$

[Neyman(1934)]

(v) Two-stage sampling optimum size:

$$(5) \quad m_i = M_i \frac{K}{N} \sqrt{\frac{c_1}{c_2} \frac{1 - R_y}{R_y}},$$

[Hansen et al(1953)]

Section 4: Contributions based on classical superpopulation models. In this section we distinguish between contributions based on design-unbiased estimation approach and on predictive approach. Point out:

- (i) The agreement assuming a superpopulation model without auxiliary variable information with results in the fixed population approach. By example, Cassel et al(1977) with a design-unbiased estimation show the optimality of stratification with Neyman's assignment (see (4)). And Royall(1976) for the model:

$$Y_{ij} = \mu + E_{ij}, \quad E(E_{ij}) = 0, \quad C(E_{ij}, E_{i^*j^*}) = \begin{cases} \sigma^2 & i = i^*, j = j^* \\ 0 & \text{otherwise} \end{cases}$$

gets the following two-stage sampling optimum size with $M_i = M$ and $m_i = m$:

$$(6) \quad m = \frac{n}{k} = \sqrt{\frac{c_1}{c_2} \frac{1-\rho}{\rho}}.$$

Compare (6) with (5).

- (ii) The disagreement assuming a superpopulation model with auxiliary variable information with the results based on the fixed population approach about optimum sampling design. So, Royall(1970, 71) and Royall and Heston(1973a,b) showed that the optimal sampling is a purposive design.

Section 5: Contributions based on Bayesian superpopulation models. We review two groups of contributions:

- (i) The results of Bayarri and Font(1994) and Font(1995a) based on a Bayesian model without auxiliary variable information showing with $M_i = M$ and $m_i = m$ that:

$$(7) \quad V_{\text{srswor}}(\bar{Y}|\mathbf{y}_s) - V_{\text{stratification}}(\bar{Y}|\mathbf{y}_s) = \frac{N-n}{N} \frac{\sigma^2}{n} \rho,$$

and

$$(8) \quad V_{\text{cluster}}(\bar{Y}|\mathbf{y}_s) - V_{\text{srswor}}(\bar{Y}|\mathbf{y}_s) = \frac{N-n}{N} \frac{\sigma^2}{n} (M-1)\rho.$$

Compare (7) with (1), and (8) with (2).

- (ii) The results about optimum size. Emphasize the results of Murgui(1982) and Ericson(1969b, 88). Point out the agreement assuming models without auxiliary variable information with optimum size results in fixed populations too, by example: Murgui(1982) shows for a random sampling plan that:

$$(9) \quad n = \frac{\sigma\sqrt{\gamma}}{\sqrt{c_1}} \left(1 + \frac{\sigma^2}{N\sigma_0^2}\right) - \frac{\sigma^2}{\sigma_0^2},$$

and for a stratified sampling plan that:

$$(10) \quad m_i = \frac{(\sigma_i^2 + M_i \sigma_0^2) \sigma_i}{\sqrt{c_i} \sigma_0^2} \frac{C - c_0 + \sum_{i=1}^K \frac{\sigma_i^2 c_i}{\sigma_0^2}}{\sum_{i=1}^K \frac{\sigma_i^3 \sqrt{c_i}}{\sigma_0^2} + \sum_{i=1}^K \sigma_i M_i \sqrt{c_i}} - \frac{\sigma_i^2}{\sigma_0^2}.$$

Compare (9) with (3), and (10) with (4) for $\sigma_0^2 \rightarrow \infty$ (noninformative distribution) and $c_i = c$. And the disagreement assuming models with auxiliary variable information with the results based on fixed population approach.

Section 6: Conclusions and additional references. This section gives a main conclusion: «agreement between the three approaches when there aren't auxiliary variable information and disagreement in the otherwise», and additional references about other approximations to the design problems: minimax sampling designs, experimental designs, random routes, etc.

Investigació Operativa

AMPLITUDE OF WEIGHTED MAJORITY GAME STRICT REPRESENTATIONS

J. FREIXAS

A. PUENTE

Universitat Politècnica de Catalunya*

Some real situations which may be described as weighted majority games can be modified when some players increase or decrease their weights and/or the quota is modified. Nevertheless, some of these modifications do not change the game. In the present work we shall estimate the maximal percentage variations in the weights and the quota which may be allowed without changing the game (amplitude). For this purpose, we have to use strict representations of weighted majority games.

Keywords: Simple games, weighted majority games, strict representations, linearly separable function, tolerance, amplitude.

AMS Classification: 90D12, 90B25, 94C10

*Department of Applied Mathematics III. Polytechnic University of Catalonia. Av. Bases de Manresa 61-73. 08240 Manresa (Espanya).

Research partially supported by projects PB96-0493 of DGES and PR98-06 of U.P.C.

—Received May 1996.

—Accepted November 1998.

1. INTRODUCTION

Shareholder societies, political models and even electronical applications can be described by means of weighted majority games, which can be modified if the weights and/or the quota which define them are modified.

For instance, the increase of capital in a shareholder society makes investors to increase or decrease their shares, at the same time that it provokes a change of the quota to adapt it to the new situation, previous consensus. To assure that the new distribution does not interfere in the fight for the control of the company it is necessary to estimate the maximum percentage in the variations of weights and quota which leave invariant the game associated to the initial situation.

Analogously, in the realization of a linearly separable switching function by means of an electronic device, the components used to fix the weights and the threshold cannot be completely accurate. Hence, in determining the required accuracy of these components, it is necessary to estimate the maximum percentage errors in the weights and the threshold which may be allowed without disturbing the function to be realized.

In the reliability of systems it is interesting to know which subsets of components make an additive system to work when these subsets work, and which of them make the system to fail when all of them fail. The additive system is characterized by the weight of each component and the threshold of fail. This problem can be naturally transferred to the game theory. To do this it is only necessary to consider the components as players, and the subsets of components as coalitions. So, the additive reliability systems become weighted majority games.

In this paper we start from the tolerance, solution obtained by Hu in the resolution of this problem assigned to the field of electronics, and we improve it from the amplitude when we transfer such a problem to the field of game theory.

The paper is organized as follows: In Section 2 we explain the basical definitions that permit the pursuit of the work. In Section 3 we summarize the results on tolerance obtained by Hu and, at the same time, we improve them. In Section 4 we define the amplitude for strict representations of weighted majority games, which will be the maximum percentage in the variation of the weights and the quota which leave the game unchanged. As an immediate consequence we deduce that such a value improves the tolerance. In Section 5 we obtain the simplified expression for the amplitude for strict representations of monotonic weighted majority games. Section 6 is devoted to find the quota which allows us to find the maximum value for the amplitude when the weights are given. Finally, Section 7 includes two examples to illustrate the preceding results.

2. BASIC NOTATIONS AND DEFINITIONS

Let Q be the set $\{0, 1\}$. For any given positive integer n , consider the cartesian power

$$Q^n = Q \times \cdots \times Q.$$

Thus, the elements of Q^n are the 2^n ordered n -tuples

$$(x_1, \dots, x_n).$$

By a *switching function* of n variables, we mean a function

$$f : Q^n \longrightarrow Q$$

from the n -cube Q^n into Q .

A switching function $f : Q^n \longrightarrow Q$ is *linearly separable* if it admits a *system*

$$[T; w_1, \dots, w_n]$$

such that for an arbitrary point

$$x = (x_1, \dots, x_n)$$

of the n -cube Q^n we have

$$\begin{aligned} w_1 x_1 + \cdots + w_n x_n &\geq T, \text{ if } f(x) = 1 \\ w_1 x_1 + \cdots + w_n x_n &< T, \text{ if } f(x) = 0. \end{aligned}$$

The n real numbers $w_1, \dots, w_n$ in this system are called the *weights*, and the first real number T is referred to as the *threshold* or *quota*.

It is always possible to modify the quota in such a way that the previous definition could be rewritten using strict inequalities. In this case, the system is called a *strict separating system* for the linearly separable function f .

We will see the way in which we can transfer this concept to the field of game theory.

A *simple n -person game* is a pair (N, v) where N is described as $N = \{1, 2, \dots, n\}$ and is called the set of *players*. Every $S \subseteq N$ is a *coalition*, $C(N)$ is the set of all coalitions and $v : C(N) \longrightarrow \{0, 1\}$ such that $v(\emptyset) = 0$ is the *characteristic function*. We will suppose that v is not identically equal to zero. A coalition S is *winning* if

$v(S) = 1$ and *losing* otherwise. The set of winning coalitions is denoted by $\mathcal{W}$ and the set of losing coalitions by $\mathcal{L}$.

A simple game (N, v) is a *weighted majority game* iff there are real numbers $T, w_1, \dots, w_n$ such that $v(S) = 1$ if $w(S) \geq T$ and $v(S) = 0$ otherwise, where $w(S) = \sum_{i \in S} w_i$. Then $[T; w_1, \dots, w_n]$ is called a *representation* of the weighted majority game (N, v) .

If T is such that $v(S) = 1$ if $w(S) > T$ and $v(S) = 0$ if $w(S) < T$, then $[T; w_1, \dots, w_n]$ is called a *strict representation* of (N, v) .

As is obvious, every weighted majority game admits a strict representation.

From $v(\emptyset) = 0$ it is clear that $T > 0$.

A simple game is *monotonic* if all subcoalitions of the losing coalitions are losing. If each proper subcoalition of a winning coalition is losing, this winning coalition is called *minimal*. It should be noted that a monotonic simple game is completely determined by its minimal winning coalitions. The set of minimal winning coalitions is denoted by $\mathcal{W}^m$. For monotonic simple games a player $i \in N$ is *null* if $i \notin S$ for all $S \in \mathcal{W}^m$. In a weighted majority game, we will denote by D the set of null players with non-positive weight (if any).

Throughout this paper, let $[T; w_1, w_2, \dots, w_n]$ be a strict representation of a weighted majority game.

Formally, a switching function f with $f(0, \dots, 0) = 0$ is equivalent to a simple game and a strict separating system is a strict representation of a weighted majority game without condition $T > 0$.

Gambarelli (1983) studied the effects on the game when a player increases his weight in prejudice of others, or decreases in favour. This situation can be generalized in case that there exist variations in each one of the weights and the quota, which is what we want to study. Carreras (1993) studied the effects on the Shapley value of a weighted majority game in which weights are given and the quota is modified. He particularly studied those effects on the European Parliament. These two articles have a relation with our paper in the sense that in both of them we can see variations either in weights or in the quota.

3. TOLERANCE

Throughout the section, let $f : Q^n \rightarrow Q$ be an arbitrarily given linearly separable switching function of n variables and let $[T; w_1, \dots, w_n]$ be a given strict separating system for f .

For each point $x = (x_1, \dots, x_n)$ in Q^n , let

$$w(x) = w_1x_1 + \dots + w_nx_n.$$

Then it follows from the definition of a strict separating system that

$$\begin{aligned} w(x) &> T, \text{ if } f(x) = 1, \\ w(x) &< T, \text{ if } f(x) = 0. \end{aligned}$$

Let A denote the maximum of the function $w(x)$ for all $x \in f^{-1}(0)$ and let B denote the minimum of the function $w(x)$ for all $x \in f^{-1}(1)$. If $f^{-1}(0)$ is empty, we set $A = -\infty$; if $f^{-1}(1)$ is empty, we set $B = \infty$. Then we have $A < T < B$.

Adapting the definitions for strict representations of weighted majority games we have the following results.

Let A denote the maximum of $w(S)$ for all $S \in \mathcal{L}$ and let B denote the minimum of $w(S)$ for all $S \in \mathcal{W}$. Then, we have $A < T < B$ and $A \geq 0$.

Now let m denote the smallest of the two positive numbers $T - A$ and $B - T$. On the other hand, let

$$M = T + |w_1| + \dots + |w_n|.$$

Let $\lambda_1, \dots, \lambda_n$ and Λ be $n + 1$ arbitrary real numbers and let

$$\begin{aligned} w'_i &= (1 + \lambda_i)w_i \quad i = 1, \dots, n \\ T' &= (1 + \Lambda)T. \end{aligned}$$

Then, the real numbers $\lambda_1, \dots, \lambda_n$ and Λ represent the relative variations if we use the numbers $w'_1, \dots, w'_n$ and T' instead of the original numbers $w_1, \dots, w_n$ and T as weights and quota. In this paper we are going to find the maximum of those positive real numbers δ such that if

$$|\Lambda| < \delta, \quad |\lambda_i| < \delta \quad i = 1, \dots, n$$

then $[T'; w'_1, w'_2, \dots, w'_n]$ is still a representation to the given game. Such a positive real number δ was given by Hu (1965) for strict separating systems. He defined the number $\frac{m}{M}$ (taking $|T|$ in M instead of T), which is completely determined by the set of real numbers $[T; w_1, \dots, w_n]$.

Theorem 3.1. (Hu, 1965) Let $f : Q^n \rightarrow Q$ be an arbitrarily given linearly separable switching function of n variables and let $[T; w_1, \dots, w_n]$ be a given strict separating

system for f . If $|\lambda_i| < \frac{m}{M}$ for each $i=1, \dots, n$ and if $|\Lambda| < \frac{m}{M}$, then $[T'; w'_1, \dots, w'_n]$ is a strict separating system for the given linearly separable switching function.

He called this positive number the *tolerance* of the system and denoted

$$\tau[T; w_1, \dots, w_n] = \frac{m}{M}.$$

Theorem 3.2. (Hu, 1965) Let $f : Q^n \rightarrow Q$ be an arbitrarily given linearly separable switching function of n variables and let $[T; w_1, \dots, w_n]$ be a given strict separating system for f . Then:

- a) $\tau[T; w_1, \dots, w_n] \leq 1$.
- b) $\tau[T; w_1, \dots, w_n] \leq \tau[C; w_1, \dots, w_n]$, where C stands for $\frac{A+B}{2}$. If $T \neq C$ the inequality is strict.

Adapting Hu's results for strict representations of weighted majority games we have the following theorem.

Theorem 3.3. Let $[T; w_1, \dots, w_n]$ be a strict representation of a weighted majority game. Then

- a) $\tau[T; w_1, \dots, w_n] \leq 1$.
- b) $\tau[T; w_1, \dots, w_n] \leq \tau[C; w_1, \dots, w_n]$, where C stands for $\frac{A+B}{2}$. If $T \neq C$ the inequality is strict.

Our main objective is to find the greatest value for δ such that if

$$|\Lambda| < \delta, \quad |\lambda_i| < \delta \quad i = 1, \dots, n$$

then $[T'; w'_1, \dots, w'_n]$ is still a representation to the given game. We will distinguish the monotonic case from the non-monotonic case.

First we will see how the bound given by Hu for the tolerance can be improved when we are restricted to strict representations of weighted majority games.

Theorem 3.4. Let $[T; w_1, \dots, w_n]$ be a strict representation of a weighted majority game. Then,

$$\tau[T; w_1, \dots, w_n] \leq \frac{1}{3}.$$

Proof: From Theorem 3.3 the tolerance reaches its maximum when T is the arithmetic mean

$$T = \frac{A + B}{2}$$

of the real numbers A and B . Then it follows that $m = \frac{B-A}{2}$ and $M = \frac{A+B}{2} + |w_1| + \dots + |w_n|$. Due to the fact that v is not identically equal to zero it exists a coalition S such that $w(S) \geq B$ and, consequently, $|w_1| + \dots + |w_n| \geq B$.

Theorefore, we obtain

$$\begin{aligned} \tau[T; w_1, \dots, w_n] &= \frac{m}{M} \leq \frac{\frac{B-A}{2}}{\frac{A+B}{2} + |w_1| + \dots + |w_n|} \leq \\ &\leq \frac{B-A}{A+B+2B} \leq \frac{B}{A+3B} \leq \frac{1}{3}. \end{aligned}$$

■

The following result proves that $\frac{1}{3}$ is reached and it characterizes the strict representations of monotonic weighted majority games which reach it.

Proposition 3.5. *The set of strict representations of monotonic weighted majority games with tolerance $\frac{1}{3}$ is*

$$[T; 2T, 0, \dots, 0].$$

Proof: Let $[T; w_1, \dots, w_n]$ be a strict representation of a weighted majority game. Its tolerance is:

$$\tau[T; w_1, \dots, w_n] = \frac{m}{M}$$

where $m = \min\{T - A, B - T\}$ and $M = T + |w_1| + \dots + |w_n|$.

We want to determine which are the strict representations of monotonic weighted majority with tolerance $\frac{1}{3}$.

From Theorem 3.3 we must consider $T = \frac{A+B}{2}$.

$$\tau[T; w_1, \dots, w_n] = \frac{B - A}{A + B + 2(|w_1| + \dots + |w_n|)} = \frac{1}{3} \Leftrightarrow$$

$$B - 2A = |w_1| + \dots + |w_n|.$$

Because $|w_1| + \dots + |w_n| \geq B$ and $A \geq 0$, we obtain that $A = 0$ and $|w_1| + \dots + |w_n| = B$.

Therefore from $A = 0$, $B = 2T$ and $|w_1| + \dots + |w_n| = 2T$, we can deduce that the game is:

$$[T; w_1, \dots, w_n] \equiv [T; 2T, 0, \dots, 0].$$

■

For the set of non-monotonic games this maximum is smaller.

Theorem 3.6. *Let $[T; w_1, \dots, w_n]$ be a strict representation of a non-monotonic weighted majority game. Then,*

$$\tau[T; w_1, \dots, w_n] \leq \frac{1}{5}.$$

Proof: The tolerance reaches its maximum when

$$T = \frac{A + B}{2}.$$

Then it follows that $m = \frac{B-A}{2}$ and $M = \frac{A+B}{2} + |w_1| + \dots + |w_n|$. Due to the fact that the game is non-monotonic there exist coalitions $R \subset S$ such that $v(R) = 1$, $v(S) = 0$ and $w_i < 0 \forall i \in S - R$. Hence,

$$\sum_{i \in R} w_i \geq B \text{ and } \sum_{i \in S} w_i \leq A,$$

and therefore

$$\sum_{i \in S} |w_i| = \sum_{i \in R} |w_i| + \sum_{i \in S-R} |w_i| \geq B + (B - A) = 2B - A.$$

We obtain

$$\tau[T; w_1, \dots, w_n] = \frac{m}{M} \leq \frac{\frac{B-A}{2}}{\frac{A+B}{2} + |w_1| + \dots + |w_n|} \leq \frac{B - A}{A + B + 2 \sum_{i \in S} |w_i|} \leq$$

$$\leq \frac{B - A}{A + B + 2(2B - A)} = \frac{B - A}{5B - A} \leq \frac{1}{5}.$$

■

Analogously to Proposition 3.5, we characterize the strict representations of non-monotonic weighted majority game with tolerance $\frac{1}{5}$.

Proposition 3.7. *The set of strict representations of non-monotonic weighted majority games with tolerance $\frac{1}{5}$ is*

$$[T; 2T, 0, \dots, 0, -2T].$$

Proof: Let $[T; w_1, \dots, w_n]$ be a strict representation of a non-monotonic weighted majority game. We want to determine which representations of this type have tolerance $\frac{1}{5}$.

From Theorem 3.3 we must consider $T = \frac{A+B}{2}$.

$$\begin{aligned} \tau[T; w_1, \dots, w_n] &= \frac{B - A}{B + A + 2(|w_1| + \dots + |w_n|)} = \frac{1}{5} \Leftrightarrow \\ 2B - 3A &= |w_1| + \dots + |w_n|. \end{aligned}$$

In the proof of Theorem 3.6 we showed that $|w_1| + \dots + |w_n| \geq 2B - A$, and since $A \geq 0$ we obtain that $A = 0$ and $|w_1| + \dots + |w_n| = 2B$.

Therefore from $A = 0$, $B = 2T$ and $|w_1| + \dots + |w_n| = 4T$ we can deduce that:

$$[T; w_1, \dots, w_n] \equiv [T; 2T, 0, \dots, 0, -2T].$$

■

4. AMPLITUDE FOR STRICT REPRESENTATIONS OF WEIGHTED MAJORITY GAMES

Our main objective in the present section is to find, for strict representations of weighted majority games, the greatest positive real number δ such that if

$$|\Lambda| < \delta, \quad |\lambda_i| < \delta \quad i = 1, 2, \dots, n$$

then $[T'; w'_1, \dots, w'_n]$ is equivalent to $[T; w_1, \dots, w_n]$.

We will call this constant amplitude of the representation and we will see that it is the maximum of rate one in the variation of weights and in the quota, so as the game remains invariant. As the tolerance provides a bound which guarantees that the game remains invariant, the tolerance has to be smaller than, or equal to, the amplitude.

Given a strict representation of a weighted majority game $[T; w_1, \dots, w_n]$, for each coalition $S \subseteq N$ let

$$\begin{aligned} a(S) &= |w(S) - T| \\ b(S) &= T + \sum_{i \in S} |w_i|. \end{aligned}$$

Note that these are positive numbers. Take

$$P = \min_{S \subseteq N} \frac{a(S)}{b(S)}.$$

We call this number the *amplitude* of the representation $[T; w_1, \dots, w_n]$ and denote

$$\mu[T; w_1, \dots, w_n] = P.$$

The minimum P is attained for, at least, some coalition, namely from now on S_0 .

Theorem 4.1. *If $|\lambda_i| < P$ for each $i = 1, 2, \dots, n$ and $|\Lambda| < P$, then $[T'; w'_1, \dots, w'_n]$ is equivalent to $[T; w_1, \dots, w_n]$ and P is the greatest upper bound for the constants $\lambda_1, \dots, \lambda_n, \Lambda$.*

Proof: First of all, we observe that, from the definition, $P < 1$ and because $T' = (1 + \Lambda)T$, if $|\Lambda| < P$ we obtain that $T' > 0$.

For each coalition $S \subseteq N$, let

$$w'(S) = \sum_{i \in S} w'_i.$$

For the first part it suffices to prove that $w'(S) > T'$ for every $S \in \mathcal{W}$ and $w'(S) < T'$ for every $S \in \mathcal{L}$.

First, let us assume that $S \in \mathcal{W}$. Then we have

$$a(S) = w(S) - T.$$

By definition of $w'(S)$, we have

$$\begin{aligned} w'(S) - T' &= \sum_{i \in S} w'_i - T' = \sum_{i \in S} (1 + \lambda_i)w_i - (1 + \Lambda)T = [w(S) - T] + \\ &+ [\sum_{i \in S} \lambda_i w_i - \Lambda T]. \end{aligned}$$

Since $w(S) - T = a(S)$ and

$$\left| \sum_{i \in S} \lambda_i w_i - \Lambda T \right| \leq \sum_{i \in S} |\lambda_i| |w_i| + |\Lambda| T < P[\sum_{i \in S} |w_i| + T] = \frac{a(S_0)}{b(S_0)} b(S) \leq a(S),$$

it follows that $w'(S) - T' > 0$ and hence $w'(S) > T'$.

Next, let us assume that $S \in \mathcal{L}$. Then we have

$$a(S) = T - w(S).$$

As above, we have

$$T' - w'(S) = (1 + \Lambda)T - \sum_{i \in S} (1 + \lambda_i)w_i = [T - w(S)] + [\Lambda T - \sum_{i \in S} \lambda_i w_i].$$

Since $T - w(S) = a(S)$ and

$$\left| \Lambda T - \sum_{i \in S} \lambda_i w_i \right| \leq |\Lambda| T + \sum_{i \in S} |\lambda_i| |w_i| < P[T + \sum_{i \in S} |w_i|] = \frac{a(S_0)}{b(S_0)} b(S) \leq a(S)$$

it follows that $T' - w'(S) > 0$ and hence $w'(S) < T'$.

For the second part we will suppose $Q > P$, and then we will demonstrate that the game given by

$$[T(1 + \Lambda); (1 + \lambda_1)w_1, \dots, (1 + \lambda_n)w_n]$$

is not equivalent to $[T; w_1, \dots, w_n]$ for all Λ and λ_i with $|\Lambda| < Q$ and $|\lambda_i| < Q$ for each $i = 1, \dots, n$.

Let $S_0 \subseteq N$ such that $\frac{a(S_0)}{b(S_0)} = P$. If $S_0 \in \mathcal{W}$, taking

$$\Lambda = \epsilon \quad \text{and} \quad \lambda_i = \begin{cases} -\epsilon & \text{if } w_i \geq 0 \\ \epsilon & \text{if } w_i < 0 \end{cases}$$

with $P < \epsilon < Q$, we will obtain a contradiction concerning S_0 .

$$\begin{aligned} w'(S_0) - T' &= [w(S_0) - T] - \epsilon[T + \sum_{i \in S_0} |w_i|] = a(S_0) - \epsilon b(S_0) < a(S_0) - \\ \frac{a(S_0)}{b(S_0)} b(S_0) &= 0 \end{aligned}$$

and hence $S_0 \notin \mathcal{W}$.

Analogously, is $S_0 \in \mathcal{L}$, taking

$$\Lambda = -\epsilon \quad \text{and} \quad \lambda_i = \begin{cases} \epsilon & \text{if } w_i \geq 0 \\ -\epsilon & \text{if } w_i < 0 \end{cases}$$

with $P < \epsilon < Q$, we will obtain a contradiction concerning S_0 .

$$\begin{aligned} T' - w'(S_0) &= [T - w(S_0)] - \epsilon[T + \sum_{i \in S_0} |w_i|] = a(S_0) - \epsilon b(S_0) < a(S_0) - \\ \frac{a(S_0)}{b(S_0)} b(S_0) &= 0 \end{aligned}$$

and hence $S_0 \notin \mathcal{L}$.

As a consequence of the maximality of the amplitude, we can deduce that the tolerance is smaller than, or equal to, the amplitude.

5. AMPLITUDE OF MONOTONIC WEIGHTED MAJORITY GAME STRICT REPRESENTATIONS

When we are restricted to strict representations of monotonic weighted majority games, the weight of each no null player is positive (we only have to compare a minimal winning coalition, which possesses this no null player, with the same coalition without him; the resulting inequality tells us the weight must be positive). Thus, a weight may only be non-positive if it belongs to a null player. Taking into account these facts, for monotonic games we are going to find a simpler expression for the amplitude.

Theorem 5.1. *If $[T; w_1, \dots, w_n]$ is a strict representation of a monotonic weighted majority game with amplitude P , then*

$$P = \min \left\{ \frac{B - T}{B + T - 2w(D)}, \frac{T - A}{T + A} \right\}$$

where D is the set of null players with negative weight (if any).

Proof: First we are going to demonstrate that

$$P \geq \min \left\{ \frac{B - T}{B + T - 2w(D)}, \frac{T - A}{T + A} \right\}.$$

Taking into account that $P = \min_{S \subseteq N} \frac{a(S)}{b(S)}$, it suffices to prove for $S \in \mathcal{L}$ that

$$\frac{a(S)}{b(S)} \geq \frac{T - A}{T + A}$$

and for $S \in \mathcal{W}$ that

$$\frac{a(S)}{b(S)} \geq \frac{B - T}{B + T - 2w(D)}.$$

For this purpose we have to check for $S \in \mathcal{L}$ that

$$2AT + A \left(\sum_{i \in S} |w_i| - w(S) \right) - T \left(\sum_{i \in S} |w_i| + w(S) \right) \geq 0.$$

Since $\sum_{i \in S} |w_i| - w(S) = -2w(S \cap D)$, $\sum_{i \in S} |w_i| + w(S) = 2w(S - D)$ and by definition of null player $w(S - D) \leq A$, it follows that

$$2[T(A - w(S - D)) - Aw(S \cap D)] \geq 0.$$

Next, let us assume that $S \in \mathcal{W}$. Then, taking into account $\sum_{i \in S} |w_i| - w(S) = -2w(S \cap D)$ and $\sum_{i \in S} |w_i| + w(S) = 2w(S - D)$, we have to check

$$2[-BT + Tw(D) - w(D)w(S) + Bw(S \cap D) + Tw(S - D)] \geq 0.$$

Regrouping terms it is clear that we have to check

$$2[T(w(D) + w(S - D) - B) - w(D)w(S) + Bw(S \cap D)] \geq 0.$$

Taking into account $w(S) \geq B$, it is enough to check

$$2[T(w(D) + w(S - D) - B) + B(w(S \cap D) - w(D))] \geq 0.$$

The first member of the sum is positive because for any winning coalition S we have that $(S - D) \cup D$ is also winning. The second member of the sum is non-negative because $w(S \cap D) \geq w(D)$. So, it is clear that:

$$2[T(w(D) + w(S - D) - B) + B(w(S \cap D) - w(D))] \geq 0.$$

Therefore,

$$P \geq \min \left\{ \frac{B - T}{B + T - 2w(D)}, \frac{T - A}{T + A} \right\}.$$

Now, we are going to demonstrate that $P \leq \min \left\{ \frac{B - T}{B + T - 2w(D)}, \frac{T - A}{T + A} \right\}$.

Let $S_0 \in \mathcal{W}$ be such that $w(S_0) = B$. Then $D \cap S_0 = D$ and

$$P = \min_{S \subseteq N} \frac{a(S)}{b(S)} \leq \frac{a(S_0)}{b(S_0)} = \frac{B - T}{T + B - 2w(D)}.$$

Let $S_0 \in \mathcal{L}$ such that $w(S_0) = A$. Then $D \cap S_0 = \emptyset$ and

$$P = \min_{S \subseteq N} \frac{a(S)}{b(S)} \leq \frac{a(S_0)}{b(S_0)} = \frac{T - A}{T + A}.$$

Therefore,

$$P \leq \min \left\{ \frac{B - T}{B + T - 2w(D)}, \frac{T - A}{T + A} \right\}.$$

■

Shareholder societies and most models in political science can be described by specifying non-negative weights for the voters and a positive quota. These situations give rise to $w(D) = 0$ and therefore the amplitude is $P = \min \left\{ \frac{B - T}{B + T}, \frac{T - A}{T + A} \right\}$.

Notice that from the definition of P the amplitude satisfies $0 < \mu < 1$, and for each number $x \in (0, 1)$ it exists a 2-game

$$\left[\frac{1+x}{1-x}; \left(\frac{1+x}{1-x} \right)^2, 1 \right]$$

whose amplitude is x .

6. MAXIMUM AMPLITUDE

In this section we want to determine the maximum value that the amplitude for a strict representation of a weighted majority game can reach, when the weights of players are invariable.

For a given strict representation of a weighted majority game $[T; w_1, \dots, w_n]$, the quota T may be any real number between A and B . Let us suppose $A > 0$, this is,

$$0 < \max_{S \in \mathcal{L}} w(S) = A < T < B = \min_{S \in \mathcal{W}} w(S).$$

We define the function

$$f(S, T) = \frac{a(S, T)}{b(S, T)} \quad \text{if } S \subseteq N \text{ and } T \in (A, B),$$

where $a(S, T) = |w(S) - T|$, $b(S, T) = T + \sum_{i \in S} |w_i|$ and let

$$\begin{aligned} F(T) &= \min_{S \in \mathcal{L}} f(S, T), \\ G(T) &= \min_{S \in \mathcal{W}} f(S, T). \end{aligned}$$

Then $f(S, T) \leq 1$, and using $A > 0$ it follows that $F(T) < 1$.

It turns out that the amplitude $\mu[T; w_1, \dots, w_n]$ reaches its maximum when T is the unique number such that

$$F(T) = G(T).$$

Precisely, we have the following theorem.

Theorem 6.1. *For every strict representation of a weighted majority game*

$[T; w_1, \dots, w_n]$ with $0 < A < T < B$ we have

$$\mu[T; w_1, \dots, w_n] \leq \mu[T^*; w_1, \dots, w_n]$$

where T^ stands for the unique number such that $F(T) = G(T)$. If $T \neq T^*$ the inequality is strict.*

Proof: For every coalition S , the function $f(S, T)$ is continuous and derivable with respect to T in (A, B) , and consequently $F(T)$ and $G(T)$ are continuous too. Fixing S , the derivative of $f(S, T)$ with respect to T is

$$\frac{\sum_{i \in S} |w_i| + w_i}{\left(T + \sum_{i \in S} |w_i|\right)^2} \geq 0 \quad \text{if } S \in \mathcal{L},$$

$$-\frac{\sum_{i \in S} |w_i| + w_i}{\left(T + \sum_{i \in S} |w_i|\right)^2} < 0 \quad \text{if } S \in \mathcal{W}.$$

Then, $F(T)$ is a nondecreasing function and $G(T)$ is a strictly decreasing function. To obtain $T \in (A, B)$ which defines the maximum amplitude for the given representation, we consider the function

$$P(T) = \min\{F(T), G(T)\} \text{ for } A < T < B$$

and we demonstrate there is just one number T^* which reaches the $\max_{T \in (A, B)} P(T)$. The uniqueness is due to the above considerations about nondecreasing behaviour of F and decreasing behaviour of G . The existence can be proved using Bolzano's Theorem:

$$\lim_{T \rightarrow B^-} F(T) - G(T) = \lim_{T \rightarrow B^-} F(T) > 0$$

and due to the fact that $A > 0$, it follows that $\lim_{T \rightarrow A^+} F(T) = 0$. Then

$$\lim_{T \rightarrow A^+} F(T) - G(T) = \lim_{T \rightarrow A^+} -G(T) < 0.$$

■

In particular, for monotonic games the amplitude $P = \min \left\{ \frac{B-T}{B+T-2w(D)}, \frac{T-A}{T+A} \right\}$ reaches its maximum when $\frac{B-T}{B+T-2w(D)} = \frac{T-A}{T+A}$ and therefore

$$T = \frac{w(D) + \sqrt{(w(D))^2 + 4AB - 4Aw(D)}}{2}.$$

If $w(D) = 0$, T is the geometric mean of the real numbers A and B .

7. APPLICATIONS

To illustrate the ideas of this paper, two examples of amplitude of strict representations of weighted majority games are presented.

Example 7.1. A town signed a biannual agreement with three gas companies, X, Y and Z for which the supply of gas to the town is guaranteed by the collaboration of at least two of them. The first year the needs of the town were of 75 Km^3 , and each one of the firms offered a fixed quantity of 60 Km^3 , 30 Km^3 and 60 Km^3 , respectively.

This situation can be described by the strict representation of the weighted majority game

$$[75; 60, 30, 60].$$

As it can be seen any coalition made by two or more of the firms is sufficient for the town needs, that's to say: $\mathcal{W}^m = \{\{1, 2\}, \{1, 3\}, \{2, 3\}\}$ and $A = 60$, $B = 90$.

Apart from the coalition formed, the amplitude of the representation is

$$\mu = \min \left\{ \frac{T - A}{T + A}, \frac{B - T}{B + T} \right\} = \frac{1}{11}.$$

From this, we can assure that for the second year, bearing in mind that the town necessities and the disponibilities of the companies will slightly vary, we can describe this situation as follows:

$$[75(1 + \Lambda); 60(1 + \lambda_1), 30(1 + \lambda_2), 60(1 + \lambda_3)].$$

From the amplitude which we obtain, we can assure that the maximum percentage of such variations is a 9.09% so as to guarantee the fulfilment of the agreement.

But, if we had used the tolerance, $\tau = \frac{1}{15}$, the maximum estimated percentage in the possible modifications would have been just of a 6.66%.

□

Example 7.2. We suppose that a shareholder society is formed by three majority shareholders (each one of them has respectively 50.000, 25.000 and 25.000 shares) and an ocean of small shareholders which possess a total of 5.000 shares. A bill of the company is passed if the sum of the shares belonging to the holders which vote in favour is more than 60.000 shares. It can be foreseen that at the end of the year there will be a variation of the capital which will affect the distribution of the actions as well as the

quota. If we exclude the possibility of the entry of new investors, the situation can be described with the following game:

$$[60000(1 + \Lambda); 50000(1 + \lambda_1), 25000(1 + \lambda_2), 25000(1 + \lambda_3), w_4(1 + \lambda_4), \dots, w_n(1 + \lambda_n)],$$

where the subindices $4, \dots, n$ represent the smaller players, $\sum_{i=4}^n w_i = 5.000$ and $w_i > 0$ for $i = 4, \dots, n$.

As the amplitude is $\frac{1}{23}$, any variation such as $|\Lambda| < \frac{1}{23}$, $|\lambda_i| < \frac{1}{23}$ for each

$i = 1, 2, \dots, n$, assures that the process of taking decisions in the company will not vary.

For example, the distribution

$$[62.500; 47.916, 26.041, 23.958, w_4(1 + \lambda_4), \dots, w_n(1 + \lambda_n)]$$

represents the same situation as the first game.

□

8. ACKNOWLEDGMENTS

The authors wish to thank two anonymous referees for their helpful comments.

9. REFERENCES

- [1] **Carreras, F.** (1993). «Cooperation and National Defense». *Qüestió*, **17,1**, 103-134.
- [2] **Gambarelli, G.** (1983). «Common Behaviour of Power Indices». *Int. J. of Game Theory*, **12**, 237-244.
- [3] **Hu, S.** (1965). *Threshold logic*. U. of California Press.

ON THE COMPATIBILITY OF CLASSICAL MULTIPLIER ESTIMATES WITH VARIABLE REDUCTION TECHNIQUES WHEN THERE ARE NONLINEAR INEQUALITY CONSTRAINTS

E. MIJANGOS

Zientzi Fakultatea (EHU)*

N. NABONA

Universitat Politècnica de Catalunya**

The minimization of a nonlinear function subject to linear and nonlinear equality constraints and simple bounds can be performed through minimizing a partial augmented Lagrangian function subject only to linear constraints and simple bounds by variable reduction techniques. The first-order procedure for estimating the multiplier of the nonlinear equality constraints through the Kuhn-Tucker conditions is analyzed and compared to that of Hestenes-Powell. There is a method which identifies those major iterations where the procedure based on the Kuhn-Tucker conditions can be safely used and also computes these estimates. This work justifies the extension of the former results to the case of general inequality constraints. To this end two procedures that convert inequalities into equalities are considered.

Keywords: Nonlinear programming, general inequality constraints, augmented lagrangian, variable reduction, Lagrange multiplier estimates.

* Department of Applied Mathematic, Statistics and Operations Research, Zientzi Fakultatea (EHU), P.O. Box 644, 48080 Bilbao (Espanya). Corresponding author e-mail: mepmifee@lg.ehu.es.

** Department of Statistics and Operations Research, Universitat Politècnica de Catalunya, Pau Gargallo, 5. 08028 Barcelona (Espanya).

– Received February 1998.

– Accepted October 1998.

1. INTRODUCTION AND MOTIVATION

Consider the linearly and nonlinearly constrained problem

$$\begin{aligned}
 (1) \quad & \text{minimize } f(x) \\
 (2) \quad & \text{subject to: } Ax = b \\
 (3) \quad & c(x) = 0 \\
 (4) \quad & l \leq x \leq u,
 \end{aligned} \tag{EP}$$

where

- (1) $f : \mathbb{R}^n \rightarrow \mathbb{R}$, $f(x)$ is nonlinear and twice continuously differentiable on the feasible set defined by constraints (2–4).
- (2) A is an $m \times n$ matrix and b an m -vector.
- (3) $c : \mathbb{R}^n \rightarrow \mathbb{R}^r$, is such that $c = [c_1, \dots, c_r]^t$, $c_i(x)$ being linear or nonlinear and twice continuously differentiable on the feasible set defined by constraints (2) and (4) $\forall i = 1, \dots, r$.
- (4) $n \gg m + r$.

To solve this problem one could use, among others, partial augmented Lagrangian techniques [1, 2, 3] as in [9, 11, 6], where only the general constraints (3) are included in the Lagrangian. In the application of these techniques there are two fundamental steps. The first solving

$$\begin{aligned}
 (5) \quad & \text{minimize}_x L_\rho(x, \mu) \\
 (6) \quad & \text{subject to: } Ax = b \\
 (7) \quad & l \leq x \leq u,
 \end{aligned} \tag{ES}$$

where $\rho > 0$ and μ are fixed,

$$L_\rho(x, \mu) = f(x) + \mu^t c(x) + \frac{1}{2} \rho c(x)^t c(x).$$

Should the solution $\tilde{x}$ obtained be infeasible with respect to (3), the second step, which is the updating of the estimate μ of the Lagrange multipliers of constraints (3), is carried out, also updating, if necessary, the penalty coefficient ρ , then going back to the first step. Should $\tilde{x}$ be feasible (or the violation of constraints (3) be sufficiently small)

the procedure ends. It is of paramount importance that the multipliers estimate μ be as accurate as possible, otherwise the convergence of the algorithm can be severely handicapped, as shown in [1, 2, 3].

In practice there are two first-order procedures to estimate μ . On the one hand the method put forward by Hestenes [8] and Powell [15]

$$\tilde{\mu} = \mu + \rho c(\tilde{x}),$$

and on the other hand μ_L obtained through the classical solution to the system of Kuhn-Tucker necessary conditions

$$(8) \quad \nabla f(\tilde{x}) + \nabla c(\tilde{x})\mu + A^t\pi + \lambda = 0$$

by least squares, as suggested in [7]. However, whereas the first procedure can be always used for any $\tilde{x}$, without hampering the convergence [2, 3], this is not the case with the second procedure, as system (8) is only known to be compatible at the optimizer x^* , not being necessarily so at $\tilde{x}$, thus possibly giving rise to bad estimates μ_L , shown up by large residuals for system (8).

Section 2 of this work presents a study of the viability of using this multiplier estimation technique within the minimization of a Partial Augmented Lagrangian subject to linear constraints and bounds by the Murtagh and Saunders procedure [13] for problem **EP**. (An alternative development of the contents of Section 2 can be found in [10].)

Sections 3 and 4 consider two ways of extending these results to problem

$$\begin{aligned} (9) \quad & \text{minimize } f(x) \\ (10) \quad & \text{subject to: } Ax = b \\ (11) \quad & \underline{c} \leq c(x) \leq \bar{c} \\ (12) \quad & l \leq x \leq u, \end{aligned} \tag{IP}$$

where $\underline{c}_j < \bar{c}_j$, for $j = 1, \dots, r$. In Section 3, a vector of slacks « y » is used to convert constraints (11) into equalities, and, in Section 4, slacks « y » and artificial variables « w » are used with the same aim. Section 5 contains the conclusions.

2. ANALYSIS OF COMPATIBILITY

The compatibility of the multiplier estimate obtained through the classical solution to the system of Kuhn-Tucker necessary conditions with the variable reduction techniques and its relationship with that of Hestenes and Powell are analyzed along this section.

Let us consider the first-order conditions associated with a local optimizer x^* of the problem **EP**, which are

$$\begin{aligned}\nabla f(x^*) + \nabla c(x^*)\mu^* + A^t\pi^* + \lambda^* &= 0 \\ Ax^* &= b \\ c(x^*) &= 0 \\ \lambda_i^* &\leq 0, \quad \text{if } x_i^* = l_i \\ \lambda_i^* &\geq 0, \quad \text{if } x_i^* = u_i \\ \lambda_i^* &= 0, \quad \text{otherwise}\end{aligned}$$

so that unique vectors μ^* , π^* and λ^* exist, such that the first equation holds. These vectors are denoted *Lagrange multipliers*; being $\nabla c(x) = [\nabla c_1(x), \dots, \nabla c_r(x)]$. (The gradient is considered to be a column vector).

Throughout this work we assume

AS1. x^* is a *regular point* — i.e., the Jacobian of the active constraints at x^* has full rank.

Solving problem **EP** through a partial augmented Lagrangian techniques consists basically of the following algorithm, where for given $\rho > 0$ and μ the subproblem **ES** (5–7) is successively solved.

Algorithm 2.1.

1. For an initial point x_0 (not necessarily feasible with respect to $c(x) = 0$), a given scalar $\rho > 0$ and vector μ , solve subproblem **ES** and obtain its optimizer $\tilde{x} = x(\mu, \rho)$.
2. Should this $\tilde{x}$ make $c(\tilde{x})$ to be zero or nearly so for a prespecified tolerance, $x^* = \tilde{x}$ and the problem is solved, otherwise,
3. μ is updated by estimating μ^* , for which two possibilities are considered:
 - U_{KT} . Solving system

$$(13) \quad \nabla f(\tilde{x}) + \nabla c(\tilde{x})\mu + A^t\pi + \lambda = 0,$$

(which could be solved through least squares using QR factorization as justified in [7], obtaining $\tilde{\mu}$, and also $\tilde{\pi}$).

– U_{HP} . Setting

$$(14) \quad \tilde{\mu} = \mu + \rho c(\tilde{x}),$$

as established in [2, 3].

4. Should the solution of **ES** not reduce $\|c(x)\|$ sufficiently, ρ would be updated as $\rho = \nu\rho$, where $\nu > 1$.
5. Make $\mu = \tilde{\mu}$ and $x_0 = \tilde{x}$, and return to 1.

The issue is now the compatibility of system (13) at $\tilde{x}$, and in this event the relationship between the procedures U_{KT} and U_{HP} to estimate vector μ^* at that point.

Each time subproblem **ES** is solved exactly by Murtagh and Saunders's active set method [13], we obtain an optimizer $\tilde{x}$ and an associated partition of matrix $A = [B_A \mid S_A \mid N_A]$, and, thus, the variable reduction matrix Z_A shown below:

$$(15) \quad Z_A = \begin{bmatrix} -B_A^{-1}S_A \\ \mathbf{1} \\ \mathbf{0} \end{bmatrix}, \quad \text{which satisfies} \quad AZ_A = 0.$$

Since $\tilde{x}$ is an optimizer of problem **ES** the necessary first-order optimality conditions must hold; i.e., there exist unique vectors $\tilde{\pi}$ and $\tilde{\lambda}$ such that:

$$(16) \quad \nabla_x L_\rho(\tilde{x}, \mu) + A^t \tilde{\pi} + \tilde{\lambda} = 0$$

$$(17) \quad A\tilde{x} = b,$$

$$(18) \quad \tilde{x}_i = l_i, \quad i = t+1, \dots, \bar{t}$$

$$(19) \quad \tilde{x}_i = u_i, \quad i = \bar{t}+1, \dots, n$$

where t is the number of basic and superbasic variables, and

$$\begin{aligned} \tilde{\lambda}_i &\leq 0, & \text{if } \tilde{x}_i &= l_i \\ \tilde{\lambda}_i &\geq 0, & \text{if } \tilde{x}_i &= u_i \\ \tilde{\lambda}_i &= 0, & \text{otherwise.} \end{aligned}$$

Expression (16) is equivalent to

$$(20) \quad A^t \tilde{\pi} + \nabla c(\tilde{x})[\mu + \rho c(\tilde{x})] + \tilde{\lambda} = -\nabla f(\tilde{x}),$$

which in matrix form yields

$$\begin{bmatrix} B_A^t & \nabla_{B_A} c(\tilde{x}) & \mathbf{0} \\ S_A^t & \nabla_{S_A} c(\tilde{x}) & \mathbf{0} \\ N_A^t & \nabla_{N_A} c(\tilde{x}) & \mathbf{1} \end{bmatrix} \begin{bmatrix} \tilde{\pi} \\ \mu + \rho c(\tilde{x}) \\ \tilde{\lambda} \end{bmatrix} = - \begin{bmatrix} \nabla_{B_A} f(\tilde{x}) \\ \nabla_{S_A} f(\tilde{x}) \\ \nabla_{N_A} f(\tilde{x}) \end{bmatrix},$$

where $\nabla_{B_A} c(\tilde{x})$ stands for the rows of $\nabla c(\tilde{x})$ associated with the rows of B_A^t . Similarly $\nabla_{S_A} c(\tilde{x})$ and $\nabla_{N_A} c(\tilde{x})$, and also the partition of $\nabla f(\tilde{x})$, are defined.

Let us assume that matrix

$$A(\tilde{x}) = \underbrace{\begin{array}{|c|c|c|} \hline \overbrace{B_A}^m & \overbrace{S_A}^s & N_A \\ \hline \nabla_{B_A} c(\tilde{x})^t & \nabla_{S_A} c(\tilde{x})^t & \nabla_{N_A} c(\tilde{x})^t \\ \hline \end{array}}_n \begin{array}{l} \}m \\ \}r \end{array}$$

has full row rank. It is now possible to get from $A(\tilde{x})$ a full rank basic matrix B such that it contains matrix B_A , which was obtained at the end of the first step of Algorithm 2.1. Once the basic matrix B has been defined, matrix S can be established such that $[B \mid S]$ contains B_A and S_A as submatrices, thus:

(21)

$$B = \underbrace{\begin{array}{|c|c|c|} \hline \overbrace{B_A}^{m+r} & \overbrace{S'_A}^{s'} & \overbrace{\tilde{N}_A}^{m_1} \\ \hline \nabla_{B_A} c(\tilde{x})^t & \nabla_{S'_A} c(\tilde{x})^t & \nabla_{\tilde{N}_A} c(\tilde{x})^t \\ \hline \end{array}}_{\substack{m \\ s' \\ m_1}} \begin{array}{l} \}m \\ \}r \end{array} \quad \text{and} \quad S = \underbrace{\begin{array}{|c|} \hline \overbrace{S''_A}^{s-s'} \\ \hline \nabla_{S''_A} c(\tilde{x})^t \\ \hline \end{array}}_{s''},$$

with $S_A = [S'_A \mid S''_A]$. The rest of columns from $A(\tilde{x})$ makes up submatrix N . Let $\nabla_B v(x)$ denote the gradient with respect to the variables associated with B of any differentiable function $v(x)$. Similarly $\nabla_S v(x)$ and $\nabla_N v(x)$. See [10] for an efficient procedure to build up B from data available at $\tilde{x}$ when subproblem **ES** has been solved.

From this partition $[B \mid S \mid N]$ of $A(\tilde{x})$ we have the variable reduction matrix

$$Z(\tilde{x}) = \begin{bmatrix} -B^{-1}S \\ \mathbf{1} \\ \mathbf{0} \end{bmatrix}, \quad \text{which satisfies} \quad A(\tilde{x})Z(\tilde{x}) = \mathbf{0}.$$

Premultiplying by $Z(\tilde{x})$ both sides of (20) we get the equivalent expression

$$(22) \quad Z(\tilde{x})^t A^t \tilde{\pi} + Z(\tilde{x})^t \nabla c(\tilde{x}) [\mu + \rho c(\tilde{x})] + Z(\tilde{x})^t \tilde{\lambda} = -Z(\tilde{x})^t \nabla f(\tilde{x}).$$

According to the definition of $Z(\tilde{x})$ the following must hold:

$$(23) \quad Z(\tilde{x})^t \nabla c(\tilde{x}) = 0 \quad \text{and} \quad Z(\tilde{x})^t A^t = 0.$$

Furthermore,

$$(24) \quad Z(\tilde{x})^t \tilde{\lambda} = \begin{bmatrix} -(B^{-1}S)^t & \mathbf{1} & \mathbf{0} \end{bmatrix} \begin{bmatrix} \tilde{\lambda}^B \\ \tilde{\lambda}^S \\ \tilde{\lambda}^N \end{bmatrix} = -(B^{-1}S)^t \tilde{\lambda}^B + \tilde{\lambda}^S.$$

As we can see, $Z(\tilde{x})^t \nabla f(\tilde{x})$ will vanish only if the following condition holds

$$(25) \quad \tilde{\lambda}^S = (B^{-1}S)^t \tilde{\lambda}^B,$$

which in general is not satisfied, as can be easily proved; see [11].

Let $\sigma_k = [\sigma_{1k}, \dots, \sigma_{\bar{m}k}]^t$ (with $\bar{m} = m + r$) be the column of $B^{-1}S$ associated with the superbasic variable x_k ; then (25) can be recast as

$$\tilde{\lambda}_k^S = \sum_{j=1}^{\bar{m}} \sigma_{jk} \tilde{\lambda}_j^B, \quad k \in S,$$

S being the set of indices associated with the columns of S .

The *basic equivalent path* β_k of superbasic variable x_k is defined as the set of basic variables x_i that have a nonzero entry σ_{ik} in the column of $B^{-1}S$ corresponding to variable x_k .

Taking into account the expressions (22)-(24) we get

$$(26) \quad Z(\tilde{x})^t \nabla f(\tilde{x}) = \alpha,$$

where α is a vector (whose dimension is the number of columns of S) such that

$$(27) \quad \alpha_k = \sum_{i \in \beta_k \cap \tilde{N}} \sigma_{ik} \tilde{\lambda}_i^B - \tilde{\lambda}_k^S,$$

$\tilde{N}$ being the index set of the columns of matrix

$$(28) \quad \begin{bmatrix} \tilde{N}_A \\ \nabla_{\tilde{N}} c(\tilde{x})^t \end{bmatrix}$$

selected to make up the basis matrix of $A(\tilde{x})$ (see expression (21)) and σ_{ik} , as previously defined, entry (i, k) of matrix $B^{-1}S$. Furthermore, since by construction the

set of indices associated with the columns of S is a subset of the set of indices associated with the columns of S_A , $\tilde{\lambda}_k = 0$ holds for all k corresponding to a column of S , hence (27) becomes

$$(29) \quad \alpha_k = \sum_{i \in \beta_k \cap \tilde{N}} \sigma_{ik} \tilde{\lambda}_i^B.$$

It must be pointed out that vector α turns out to be the nonzero part of the residual vector corresponding to system (13), when, once the partition of matrix $A(\tilde{x})$ is fixed, it is solved calculating first (π, μ) through the solution of the compatible system

$$B^t \begin{bmatrix} \pi \\ \mu \end{bmatrix} = -\nabla_B f(\tilde{x}),$$

and then computing:

$$\lambda^N = -N^t \begin{bmatrix} \pi \\ \mu \end{bmatrix} - \nabla_N f(\tilde{x}),$$

as, by definition of $Z(\tilde{x})$ and in view of (13) we are led to

$$(30) \quad \begin{aligned} Z(\tilde{x})^t \nabla f(\tilde{x}) &= -(B^{-1}S)^t \nabla_B f(\tilde{x}) + \nabla_S f(\tilde{x}) \\ &= S^t [-(B^t)^{-1} \nabla_B f(\tilde{x})] + \nabla_S f(\tilde{x}) \\ &= S^t \begin{bmatrix} \pi \\ \mu \end{bmatrix} + \nabla_S f(\tilde{x}). \end{aligned}$$

Here a series of propositions are presented to be used later.

Let us consider now, for any x , a full-row-rank matrix

$$A(x) = \begin{bmatrix} A \\ \nabla c(x)^t \end{bmatrix}$$

partitioned as $A(x) = [B \ S \ N]$, where B is a nonsingular matrix, and S and N have at least one column. Let

$$(31) \quad \hat{A}(x) = \begin{bmatrix} B & S & N \\ \mathbf{0} & \mathbf{0} & \mathbf{1} \end{bmatrix},$$

and

$$Z(x) = \begin{bmatrix} -B^{-1}S \\ \mathbf{1} \\ \mathbf{0} \end{bmatrix}, \quad \text{which satisfies} \quad \hat{A}(x)Z(x) = 0.$$

Let $x = \bar{x}$ be a vector such that system

$$(32) \quad A^t \pi + \nabla c(\bar{x}) \mu + \lambda + \nabla f(\bar{x}) = 0$$

is compatible. Suppose that $\lambda_i = 0$ for all i associated with a column of either B or S . Then, premultiplying this system by $Z(\bar{x})^t$ we have

$$Z(\bar{x})^t A^t \pi + Z(\bar{x})^t \nabla c(\bar{x}) \mu + Z(\bar{x})^t \lambda + Z(\bar{x})^t \nabla f(\bar{x}) = 0,$$

which, since $\hat{A}(\bar{x})Z(\bar{x}) = 0$, implies $Z(\bar{x})^t \nabla f(\bar{x}) = 0$.

Proposition 2.1. *System (32) is compatible if and only if*

$$(33) \quad Z(\bar{x})^t \nabla f(\bar{x}) = 0$$

is verified.

Proof: The previous result to this proposition proves its first part.

To prove the second part we consider the nonsingular $(n \times n)$ -matrix

$$W = \begin{bmatrix} \mathbf{1} & \mathbf{0} & \mathbf{0} \\ -(B^{-1}S)^t & \mathbf{1} & \mathbf{0} \\ \mathbf{0} & \mathbf{0} & \mathbf{1} \end{bmatrix},$$

where $\mathbf{0}$ stands for a zero matrix of suitable dimensions, and such that the unit matrix at the bottom is $(n - t) \times (n - t)$, $(n - t)$ being the number of columns of N .

Let us consider now that system (32) is not compatible. Hence, the matrix $[\hat{A}(\bar{x})^t \quad \nabla f(\bar{x})]$ of this system has full column rank and we have

$$W[\hat{A}(\bar{x})^t \quad \nabla f(\bar{x})] = W \begin{bmatrix} B^t & \mathbf{0} & \nabla_B f(\bar{x}) \\ S^t & \mathbf{0} & \nabla_S f(\bar{x}) \\ N^t & \mathbf{1} & \nabla_N f(\bar{x}) \end{bmatrix} = \begin{bmatrix} B^t & \mathbf{0} & \nabla_B f(\bar{x}) \\ \mathbf{0} & \mathbf{0} & Z(\bar{x})^t \nabla f(\bar{x}) \\ N^t & \mathbf{1} & \nabla_N f(\bar{x}) \end{bmatrix},$$

where $\nabla_M f(\bar{x})$ is the gradient of $f(x)$ at $x = \bar{x}$ with respect to the subset of variables associated with M , for all submatrix M formed by columns in $A(\bar{x})$.

Since the product of a nonsingular matrix of order $n \times n$ multiplied by a matrix with full column rank of order $n \times (m + r + (n - t) + 1)$ gives rise to a matrix with these same characteristics, it is shown that the product $Z(\bar{x})^t \nabla f(\bar{x})$ cannot be the null vector.

Therefore, if (33) holds, system (32) is compatible. ■

Let B_A be a nonsingular square submatrix of B made up by columns of A and let S_A be the submatrix formed by the columns common to $[B \ S]$ and A once removed the columns of B_A . Moreover, N_A is a submatrix constituted by the rest of columns of A , thus their associated variables are the same as those associated with the columns of N .

Let

$$(34) \quad BS = \begin{bmatrix} B_A & S_A \\ \nabla_{B_A} c(\bar{x})^t & \nabla_{S_A} c(\bar{x})^t \end{bmatrix}$$

and Z_A a matrix defined by expression (15) (although using the current B_A and S_A).

Proposition 2.2. *Matrix BS has full row rank if and only if $\nabla c(\bar{x})^t Z_A$ has also full row rank.*

Proof: It is sufficient to take into account the matrix product

$$\begin{bmatrix} B_A & S_A \\ \nabla_{B_A} c(\bar{x})^t & \nabla_{S_A} c(\bar{x})^t \end{bmatrix} \begin{bmatrix} \mathbf{1} & -B_A^{-1} S_A \\ \mathbf{0} & \mathbf{1} \end{bmatrix} = \begin{bmatrix} B_A & \mathbf{0} \\ \nabla_{B_A} c(\bar{x})^t & \nabla c(\bar{x})^t Z_A \end{bmatrix},$$

where the first matrix is BS , see (34), and Z_A is defined by (15). ■

Premultiplying equation (32) by matrix Z_A^t and moving $Z_A^t \nabla f(\bar{x})$ to the right hand side we obtain

$$(35) \quad Z_A^t \nabla c(\bar{x}) \mu = -Z_A^t \nabla f(\bar{x}).$$

Proposition 2.3. *Let BS be a full-row-rank matrix. System (32) is compatible if and only if system (35) is compatible*

Proof: To prove the «only if» part it is enough to premultiply the system (32) by Z_A .

Now, to show that the «if» part holds, let us consider the matrix

$$W_A = \begin{bmatrix} \mathbf{1} & \mathbf{0} & \mathbf{0} \\ -(B_A^{-1} S_A)^t & \mathbf{1} & \mathbf{0} \\ \mathbf{0} & \mathbf{0} & \mathbf{1} \end{bmatrix}.$$

Suppose that (32) is not compatible, then, since BS has full row rank, matrix $[\hat{A}(\bar{x})^t \quad \nabla f(\bar{x})]$ has full column rank. Therefore, to prove this part it is sufficient to operate as in proposition 2.1, but now replacing the matrix W with the matrix W_A , thus

$$\begin{aligned} W_A[\hat{A}(\bar{x})^t \quad \nabla f(\bar{x})] &= W_A \begin{bmatrix} B_A^t & \nabla_{B_A} c(\bar{x}) & \mathbf{0} & \nabla_{B_A} f(\bar{x}) \\ S_A^t & \nabla_{S_A} c(\bar{x}) & \mathbf{0} & \nabla_{S_A} f(\bar{x}) \\ N_A^t & \nabla_{N_A} c(\bar{x}) & \mathbf{1} & \nabla_{N_A} f(\bar{x}) \end{bmatrix} \\ &= \begin{bmatrix} B_A^t & \nabla_{B_A} c(\bar{x}) & \mathbf{0} & \nabla_{B_A} f(\bar{x}) \\ \mathbf{0} & Z_A^t \nabla c(\bar{x}) & \mathbf{0} & Z_A^t \nabla f(\bar{x}) \\ N_A^t & \nabla_{N_A} c(\bar{x}) & \mathbf{1} & \nabla_{N_A} f(\bar{x}) \end{bmatrix}. \end{aligned}$$

Note that if matrix $[\hat{A}(\bar{x})^t \quad \nabla f(\bar{x})]$ has full column rank, $[Z_A^t \nabla c(\bar{x}) \quad Z_A^t \nabla f(\bar{x})]$ has also full column rank. ■

Corollary 2.1. *Under the same conditions of the previous propositions, system (35) is compatible if and only if (33) holds.*

Proof: It is a direct result of the propositions 2.1 and 2.3. ■

Now we consider $\bar{x} = \tilde{x} = x(\mu, \rho)$ (i.e. the optimizer of **ES** considered at the beginning of this section).

As a result of these propositions and corollary, we have the following consequences:

- If $\alpha \neq 0$ (see expression (26)), system (13) is not compatible, and neither is system

$$(36) \quad Z_A^t \nabla c(\tilde{x})\mu = -Z_A^t \nabla f(\tilde{x}).$$

Therefore, in this case the estimate of μ^* by the U_{KT} procedure is not reliable, because system (13) is not compatible.

- If when building up the basic matrix B of $A(\tilde{x})$ it is not necessary to take columns of $A(\tilde{x})$ that contain columns of N_A , we have

$$(37) \quad Z(\tilde{x})^t \nabla f(\tilde{x}) = 0,$$

and the associated systems (13) and (36) are compatible. Moreover, the μ obtained by solving both (13) and (36) will be the same, see proposition 2.2 and corollary 2.1. Therefore, it is enough to solve (36) to calculate μ in procedure U_{KT} .

- If problem **EP** has not simple bounds, i.e., if it is only defined by (1-3), then it automatically holds $\alpha = 0$ for all $\tilde{x} = x(\mu, \rho)$ optimizing

$$\begin{aligned} & \underset{x}{\text{minimize}} && L_\rho(x, \mu) \\ & \text{subject to:} && Ax = b, \end{aligned}$$

where $L_\rho(x, \mu) = f(x) + \mu^t c(x) + \frac{\rho}{2} \|c(x)\|_2^2$ is the augmented Lagrangian function. Therefore, for all vector $\tilde{x}$ obtained in this way, it is enough to solve (36) — with $\bar{x} = \tilde{x}$ — to calculate μ in procedure U_{KT} .

- It is clear from the above results that in order to build up matrix B it is preferable to use only columns of BS . Should this submatrix of $A(\tilde{x})$ not have full row rank, one must initially search for suitable columns among the l columns associated with N_A such that $\tilde{\lambda}_l = 0$, for $l \in \tilde{N}$, if any, see (29).

Next, by means of a proposition we analyze the general case considered at the beginning of this section.

Proposition 2.4. *If $\alpha = 0$ and **ES** (5–7) is solved with exact minimization, the U_{KT} procedure is valid for estimating μ^* . Furthermore, when B is built up following the rule given at the last item, the results obtained coincide with those found by means of the U_{HP} procedure. Otherwise, even if $\alpha = 0$, the estimates of μ^* obtained through U_{KT} and U_{HP} are not the same.*

Proof: As shown above, $\alpha = 0$ implies the validity of procedure U_{KT} for estimating μ^* , because of (26) and Proposition 2.1.

The second part of this proposition is proved next.

The solution of **ES** by means of exact minimization implies that $Z_A^t \nabla_x L_\rho(\tilde{x}, \mu) = 0$. This is equivalent to

$$(38) \quad Z_A^t \nabla c(\tilde{x})(\mu + \rho c(\tilde{x})) = -Z_A^t \nabla f(\tilde{x}).$$

If the set $\overline{BS}$ of indices associated with $[B \ S]$ coincides with the set $\overline{BS}_A$ of the indices associated with $[B_A \ S_A]$, $\alpha = 0$ and the conditions of Propositions 2.1, 2.2 and 2.3 are fulfilled directly, and therefore the system

$$(39) \quad Z_A^t \nabla c(\tilde{x})\mu = -Z_A^t \nabla f(\tilde{x})$$

is compatible and has a single solution $\bar{\mu}$, which compared with (38) implies that

$$\bar{\mu} = \mu + \rho c(\tilde{x}),$$

whose second member corresponds to the estimate of μ^* by means of the U_{HP} procedure.

Due to the length of the part of the proof corresponding to the case in which $\overline{BS}$ contains strictly $\overline{BS}_A$, it is divided into several sections.

(i) *Definition of matrix $\tilde{N}_A$ to fill up the row rank of $[B \ S]$.*

Let us consider again the submatrix given in (28)

$$\begin{bmatrix} \tilde{N}_A \\ \nabla_{\tilde{N}} c(\tilde{x})^t \end{bmatrix},$$

whose columns are selected among the columns of $A(\tilde{x})$ associated with nonbasic variables (with regard to subproblem **ES**) so that the matrix B of expression (21) is nonsingular. Let $\tilde{N}$ be, as above, the set of indices associated with the columns of $\tilde{N}_A$. Hence $\overline{BS}_A \cup \tilde{N} = \overline{BS}$. From now on, we consider that the columns of this submatrix appear arranged in A immediately after those of S_A , without loss of generality.

(ii) *Definition of matrix $\overline{Z}_A$.*

A new reduction matrix $\overline{Z}_A$ is also defined for matrix A such that

$$(40) \quad \overline{Z}_A = [Z_A \ Z_{\tilde{N}}] = \begin{bmatrix} -B_A^{-1} S_A & -B_A^{-1} \tilde{N}_A \\ \mathbf{1} & \mathbf{0} \\ \mathbf{0} & \mathbf{1} \\ \mathbf{0} & \mathbf{0} \end{bmatrix}.$$

The matrices $\mathbf{1}$, from left to right, have their dimensions fixed, respectively, by the number of columns in S_A and the number of columns in $\tilde{N}_A$. The matrix $\overline{Z}_A$ has full column rank.

(iii) *Compatibility of system $\overline{Z}_A^t \nabla c(\tilde{x}) \mu = -\overline{Z}_A^t \nabla f(\tilde{x})$.*

Since $\alpha = 0$, system

$$A^t \pi + \nabla c(\tilde{x}) \mu + \lambda = -\nabla f(\tilde{x}),$$

has one only solution $(\bar{\pi}, \bar{\mu}, \bar{\lambda})$ (see Proposition 2.1). Moreover, by Proposition 2.3, the following system is compatible:

$$(41) \quad \bar{Z}_A^t \nabla c(\tilde{x}) \mu = -\bar{Z}_A^t \nabla f(\tilde{x}).$$

The solution of this is $\bar{\mu}$, which is unique due to matrix $\bar{Z}_A^t \nabla c(\tilde{x})$ having full column rank. To prove this last, it is sufficient to take into account that

$$(42) \quad [B \ S] = \begin{bmatrix} B_A & S_A & \tilde{N}_A \\ \nabla_{Bc}(\tilde{x})^t & \nabla_{Sc}(\tilde{x})^t & \nabla_{\tilde{N}c}(\tilde{x})^t \end{bmatrix},$$

is of full rank $(m + r)$ and to construct a suitable matrix W . Let W be

$$W = \begin{bmatrix} \mathbf{1} & -B_A^{-1} S_A & -B_A^{-1} \tilde{N}_A \\ \mathbf{0} & \mathbf{1} & \mathbf{0} \\ \mathbf{0} & \mathbf{0} & \mathbf{1} \end{bmatrix}.$$

This matrix is nonsingular and of order $(m + s + m_1) \times (m + s + m_1)$, being m_1 the number of columns in $\tilde{N}_A$ and such that $s + m_1 \geq r$, by construction of $[B \ S]$ (42).

Multiplying and taking into account (40)

$$[B \ S]W = \begin{bmatrix} B_A & \mathbf{0} & \mathbf{0} \\ \nabla_{Bc}(\tilde{x})^t & \nabla c(\tilde{x})^t Z_A & \nabla c(\tilde{x})^t Z_{\tilde{N}} \end{bmatrix} = \begin{bmatrix} B_A & \mathbf{0} \\ \nabla_{Bc}(\tilde{x})^t & \nabla c(\tilde{x})^t \bar{Z}_A \end{bmatrix},$$

then because of the features of the factor matrices with respect to the rank, the product is a $(m + r) \times (m + s + m_1)$ -matrix that has full row rank, and since B_A is a nonsingular matrix with rank m , the submatrix $\nabla c(\tilde{x})^t \bar{Z}_A$ has rank r , or in other words, it has full row rank.

(iv) *Calculation of $\bar{Z}_A^t \nabla_x L_\rho(\tilde{x}, \mu)$.*

Let us return to the subproblem **ES** and perform the product $\bar{Z}_A^t \nabla_x L_\rho(\tilde{x}, \mu)$, the result is

$$\bar{Z}_A^t \nabla_x L_\rho(\tilde{x}, \mu) = \begin{bmatrix} Z_A^t \\ Z_{\tilde{N}}^t \end{bmatrix} \nabla_x L_\rho(\tilde{x}, \mu) = \begin{bmatrix} Z_A^t \nabla_x L_\rho(\tilde{x}, \mu) \\ Z_{\tilde{N}}^t \nabla_x L_\rho(\tilde{x}, \mu) \end{bmatrix}.$$

$Z_A^t \nabla_x L_\rho(\tilde{x}, \mu)$ is null because **ES** is solved by means of exact minimization. Besides,

$$Z_{\tilde{N}}^t \nabla_x L_\rho(\tilde{x}, \mu) + Z_{\tilde{N}}^t A^t \tilde{\pi} + Z_{\tilde{N}}^t \tilde{\lambda} = 0.$$

Let us also observe that the second term on the left is zero, as $Z_{\tilde{N}}$ expands a null subspace of A and

$$Z_{\tilde{N}}^t \tilde{\lambda} = (-B_A^{-1} \tilde{N}_A)^t \tilde{\lambda}^B + \tilde{\lambda}^{\tilde{N}} = \tilde{\lambda}^{\tilde{N}},$$

since $\tilde{\lambda}^B = 0$ — which is associated with B_A . Thus,

$$\bar{Z}_A^t \nabla_x L_\rho(\tilde{x}, \mu) = \begin{bmatrix} \mathbf{0} \\ -\tilde{\lambda}^{\tilde{N}} \end{bmatrix},$$

which, by developing $\nabla_x L_\rho(\tilde{x}, \mu)$, is equivalent to

$$(43) \quad \bar{Z}_A^t \nabla c(\tilde{x})(\mu + \rho c(\tilde{x})) = -\bar{Z}_A^t \nabla f(\tilde{x}) + \begin{bmatrix} \mathbf{0} \\ -\tilde{\lambda}^{\tilde{N}} \end{bmatrix}.$$

If $\tilde{\lambda}^{\tilde{N}}$ is null — i.e., if following the rule put forward in the item previous to this proposition we are able to make up a basis matrix B —, then we have

$$(44) \quad \bar{Z}_A^t \nabla c(\tilde{x})(\mu + \rho c(\tilde{x})) = -\bar{Z}_A^t \nabla f(\tilde{x}).$$

(v) *Comparison of expressions (41) and (44).*

Finally, if we compare expressions (41) and (44), taking into account the compatibility of system (41) and the uniqueness of its solution, the conclusion reached is that $\bar{\mu} = \mu + \rho c(\tilde{x})$ is fulfilled if $\tilde{\lambda}^{\tilde{N}}$ is null. In this case the procedures U_{KT} and U_{HP} produce the same estimate of μ^* .

Note that if $\tilde{\lambda}^{\tilde{N}}$ is not null and $\alpha = 0$ with exact minimization, the U_{KT} procedure is reliable (in the sense of that the residuals of both systems (13) and (41) are null), although the estimate of $\bar{\mu}$ that it provides is different from that given by U_{HP} , see (43). ■

In reference [10] there is also an efficient procedure for computing α , and practicalities related to the implementation of Algorithm 2.1 with the problem considered here (code PFNRN [12]) and computational results.

3. EXTENSION BY USING VECTOR y

In this section the above results are extended to the case of problems with general inequality constraints (problem **IP** (9–12)) by using a vector y , taking into account the

technique put forward by Conn *et al* in [5], which is also employed by Murtagh and Saunders in [14].

Through this section and the following we add to assumption **AS1**, given in §2, the new one:

AS2. μ^* satisfies the strict complementarity condition

$$\begin{aligned} \text{if } c_j(x^*) = \underline{c}_j &\Rightarrow \mu_j^* < 0, \\ \text{if } c_j(x^*) = \bar{c}_j &\Rightarrow \mu_j^* > 0, \\ &\text{otherwise } \mu_j^* = 0. \end{aligned}$$

Solving problem **IP** is equivalent to solving:

$$\begin{aligned} (45) \quad & \text{minimize } f(x) \\ (46) \quad & \text{subject to: } Ax = b \\ (47) \quad & c(x) - y = 0 \quad (\mathbf{EPy}) \\ (48) \quad & l \leq x \leq u \\ (49) \quad & \underline{c} \leq y \leq \bar{c}, \end{aligned}$$

where y represents the slacks vector that turns the inequalities (11) into equalities.

As regards problem **EP** (1–4), the difference between this and problem **EPy** is that in the latter there are constraints (47) and (49) instead of constraints (3). Therefore we must analyze the effect of replacing (3) with (47) and (49) in the results of the former Section to see whether they can still be applied to the current problem.

First, it can be observed that the Kuhn-Tucker conditions of an optimal solution (x^*, y^*) to problem **EPy** are

$$\begin{aligned} (50) \quad & \nabla f(x^*) + \nabla c(x^*)\mu + A^t\pi + \lambda = 0 \\ (51) \quad & -\mu + \gamma = 0 \\ (52) \quad & Ax^* = b \\ (53) \quad & c(x^*) - y^* = 0, \end{aligned}$$

such that unique vectors μ^* , π^* , λ^* and γ^* exist that satisfy equations (50) and (51), and with γ^* also satisfying

$$(54) \quad \begin{aligned} \gamma_i^* &\leq 0 \text{ if } y_i^* = \underline{c}_i \\ \gamma_i^* &\geq 0 \text{ if } y_i^* = \bar{c}_i \\ \gamma_i^* &= 0 \quad \text{otherwise,} \end{aligned}$$

and λ^*

$$(55) \quad \begin{aligned} \lambda_i^* &\leq 0 \text{ if } x_i^* = l_i \\ \lambda_i^* &\geq 0 \text{ if } x_i^* = u_i \\ \lambda_i^* &= 0 \quad \text{otherwise.} \end{aligned}$$

Note that the only consequence on the multipliers of the introduction of the slacks y are expressions (51) and (54), the first being used to determine γ^* once the rest of the variables have been found through (50). Therefore, in order to obtain a first-order estimate of all multipliers at point (x, y) it suffices, as pointed out in [7], to solve

$$(56) \quad \nabla f(x) + \nabla c(x)\mu + A^t\pi + \lambda = 0,$$

just as in the case of problem **EP**, obtaining γ afterwards through (51). Furthermore, expression (13) does not change — let us compare it with (56).

Here the associated subproblem is defined as

$$\begin{aligned} (57) \quad & \underset{x, y}{\text{minimize}} \quad L_\rho(x, y, \mu) \\ (58) \quad & \text{subject to: } Ax = b \\ (59) \quad & l \leq x \leq u \\ (60) \quad & \underline{c} \leq y \leq \bar{c} \end{aligned} \tag{ESy}$$

where

$$(61) \quad L_\rho(x, y, \mu) = f(x) + \mu^t[c(x) - y] + \frac{1}{2}\rho\|c(x) - y\|_2^2$$

is the augmented Lagrangian function. Furthermore, in Algorithm 2.1, subproblem **ES** (5–7) is replaced by subproblem **ESy** and $c(x)$ with $c(x) - y$, hence the U_{HP} type estimate at the optimizer $(\tilde{x}, \tilde{y})$, obtained from the solution of **ESy**, can be written as:

$$(62) \quad \tilde{\mu} = \mu + \rho[c(\tilde{x}) - \tilde{y}],$$

and the U_{KT} type estimate is obtained solving

$$\nabla f(\tilde{x}) + \nabla c(\tilde{x})\tilde{\mu} + A^t\pi + \lambda = 0,$$

as in §2 with (13). Thus, as before in §2, we can now define matrices $A(\tilde{x})$, $Z(\tilde{x})$ and Z_A , arriving at the results analogous to those obtained in §2, by propositions 2.1–2.3 and corollary 2.1.

Other modifications of the algorithm are associated with the substitution of x with x, y . From the Kuhn-Tucker conditions for this subproblem we obtain

$$(63) \quad \nabla f(\tilde{x}) + \nabla c(\tilde{x})\{\mu + \rho[c(\tilde{x}) - \tilde{y}]\} + A^t \pi + \lambda = 0$$

$$(64) \quad -\{\mu + \rho[c(\tilde{x}) - \tilde{y}]\} + \gamma = 0$$

such that vectors $\pi = \tilde{\pi}$ and $\lambda = \tilde{\lambda}$ exist satisfying the first equation and a vector $\tilde{\gamma}$ can be obtained through the second one. As in §2 for expression (20), here (63) is the key expression for studying the compatibility of estimate U_{KT} when using the active set techniques of Murtagh and Saunders [13] to solve **ESy**, obtaining the results equivalent to those of the proposition 2.4.

In conclusion, all the analysis after expression (20) in §2 is also applicable to this case, with the only exception that expression $\mu + \rho c(\tilde{x})$ must be replaced by $\mu + \rho[c(\tilde{x}) - \tilde{y}]$.

4. EXTENSION BY USING VECTORS y AND w

Here the results of §2 are extended to the case of problems with general inequality constraints by using vectors y and w .

This section describes an alternative way of dealing with problem **IP** (9–12). Solving problem **IP**, as put forward in [9] (using slack variables raised to the square, as done by Rockafellar in [16]), is equivalent to solving problem:

$$\begin{aligned} (65) \quad & \text{minimize } f(x) \\ (66) \quad & \text{subject to: } Ax = b \\ (67) \quad & (c(x) - \underline{c}) - y^2 = 0 \\ (68) \quad & l \leq x \leq u \\ (69) \quad & y^2 + w^2 = \bar{c} - \underline{c}, \end{aligned} \quad (\text{EPyw})$$

where $y, w \in \mathbb{R}^r$ are auxiliary vectors of free variables — without bounds — that through vectors

$$y^2 = \begin{bmatrix} y_1^2 \\ \vdots \\ y_r^2 \end{bmatrix} \quad \text{and} \quad w^2 = \begin{bmatrix} w_1^2 \\ \vdots \\ w_r^2 \end{bmatrix}$$

allow the transformation of inequality (11) into an equality.

As regards the reference problem **EP** (1–4), the difference with problem **EPyw** (65–69) is that in the latter instead of constraint (3) there are constraints (67) and (69), but contrary to what happens in Section 2, the number of bounds does not increase. Therefore, we must examine the effect of replacing (3) with (67) and (69) in the main calculation steps involved in the results of Section 2 and extend them to the type of problem now considered.

It must be noticed first that the Kuhn-Tucker conditions to be satisfied by an optimizer (x^*, y^*, w^*) of problem **EPyw** are

$$(70) \quad \nabla f(x^*) + \nabla c(x^*)\mu + A^t\pi + \lambda = 0$$

$$(71) \quad (-\mu + \gamma)^t y^* = 0$$

$$(72) \quad \gamma^t w^* = 0$$

$$(73) \quad Ax^* = b$$

$$(74) \quad (c(x^*) - \underline{c}) - (y^*)^2 = 0,$$

$$(75) \quad (y^*)^2 + (w^*)^2 = \bar{c} - \underline{c},$$

such that unique vectors μ^* , π^* , λ and γ^* exist that satisfy equations (70), (71) and (72), and λ^* such that

$$\begin{aligned} \lambda_i^* &\leq 0 \text{ if } x_i^* = l_i \\ \lambda_i^* &\geq 0 \text{ if } x_i^* = u_i \\ \lambda_i^* &= 0 \quad \text{otherwise.} \end{aligned}$$

Note that the only consequence of the introduction of slacks y (apart from the constraints where they appear) are expressions (71) and (72), and γ^* can be obtained from these once the remaining multipliers have been computed from (70), given that in case $y^* \neq 0$, $\gamma^* = \mu^*$ (due to (71)), otherwise, $w^* \neq 0$ yields $\gamma^* = 0$ (due to (72)). Therefore, to obtain a first-order estimate of the U_{KT} type at point (x, y, w) of all multipliers it suffices to solve

$$(76) \quad \nabla f(x) + \nabla c(x)\mu + A^t\pi + \lambda = 0,$$

as in problem **EP** (1–4) and then compute γ through (71) and (72). Furthermore, expression (13), which coincides with (76), does not change.

Here the associated subproblem would be

$$\begin{aligned} &\underset{x, y}{\text{minimize}} \quad \bar{L}_\rho(x, y, \mu) \\ &\text{subject to: } Ax = b \\ &\quad l \leq x \leq u \\ &\quad y^2 + w^2 = \bar{c} - \underline{c}, \end{aligned} \tag{ESyw}$$

where

$$(77) \quad \bar{L}_\rho(x, y, \mu) = f(x) + \mu^t[(c(x) - \underline{c}) - y^2] + \frac{\rho}{2} \| (c(x) - \underline{c}) - y^2 \|_2^2$$

is the augmented Lagrangian function. Furthermore, see [9], the constraints where variables (y, w) appear can be eliminated by replacing the above augmented Lagrangian by

$$(78) \quad L_\rho(x, \mu) = f(x) + \sum_{j=1}^r \left\{ \mu_j \varphi_j[c_j(x), \mu_j, \rho] + \frac{1}{2} \rho |\varphi_j[c_j(x), \mu, \rho]|^2 \right\},$$

where

$$\varphi_j[c_j(x), \mu_j, \rho] = \begin{cases} c_j(x) - \bar{c}_j & \text{if } \mu_j + \rho[c_j(x) - \bar{c}_j] > 0 \\ c_j(x) - \underline{c}_j & \text{if } \mu_j + \rho[c_j(x) - \underline{c}_j] < 0 \\ -\mu_j/\rho & \text{otherwise,} \end{cases}$$

the expression in braces that appears in (78) being continuously differentiable with respect to x , for f and c defined as in §1, and (bearing in mind the strict complementarity assumption **AS2** in §2) twice continuously differentiable with respect to x if $x \in \mathcal{X}$, where

$$\mathcal{X} = \{x \mid \mu_j + \rho[c_j(x) - \bar{c}_j] \neq 0, \mu_j + \rho[c_j(x) - \underline{c}_j] \neq 0, \forall j = 1, \dots, r\}.$$

Therefore, according to [9], solving problem **ESy_w** is equivalent to solving problem

$$\begin{aligned} & \underset{x}{\text{minimize}} \quad L_\rho(x, \mu) \\ & \text{subject to:} \quad Ax = b \\ & \quad \quad \quad l \leq x \leq u. \end{aligned} \tag{ESx}$$

Moreover, in Algorithm 2.1, replacing subproblem **ES** (5–7) with **ESx** and $c_j(x)$ with $\varphi_j[c_j(x), \mu_j, \rho]$, for $j = 1, \dots, r$, we have that the U_{HP} type estimate at the optimizer $\tilde{x}$, obtained through the solution of **ESx**, can be expressed by (see [2, 9])

$$(79) \quad \tilde{\mu} = \mu + \rho \varphi[c(\tilde{x}), \mu, \rho],$$

such that $\varphi \equiv [\varphi_1, \dots, \varphi_r]^t$, and the U_{KT} type estimate is obtained solving

$$\nabla f(\tilde{x}) + \nabla c(\tilde{x})\mu + A^t\pi + \lambda = 0,$$

as in §2 with (13). Thus, as before in §2, we can now define matrices $A(\tilde{x})$, $Z(\tilde{x})$ and Z_A , arriving at results analogous to those obtained in §2, by propositions 2.1-2.3 and corollary 2.1.

In addition, from the Kuhn-Tucker conditions associated with subproblem **ESx** we obtain:

$$(80) \quad \nabla f(\tilde{x}) + \nabla c(\tilde{x}) \{ \mu + \rho \varphi[c(\tilde{x}), \mu, \rho] \} + A^t \pi + \lambda = 0,$$

where vectors $\pi = \tilde{\pi}$ and $\lambda = \tilde{\lambda}$ exist satisfying this equation.

Hence, from (80), by an analogous process to that followed from expression (20), we are led to the same results as are obtained in §2 when using the active set techniques of Murtagh and Saunders [13] to solve **ESx**, obtaining the equivalent results to those of the proposition 2.4.

In conclusion, all the analysis after expression (20) in §2 is also applicable to this case, with the only exception that expression $\mu + \rho c(\tilde{x})$ must be replaced by $\mu + \rho \varphi[c(\tilde{x}), \mu, \rho]$.

5. CONCLUSIONS

This work shows when one may compute under certain guarantees the first-order multiplier estimate based on the Kuhn-Tucker conditions. In addition, it proves the equivalence between this type of first-order estimate and that obtained through the original multiplier method (Hestenes and Powell's method) when the exact minimization is used to solve the subproblem and it is not necessary to use columns of the constraint matrix that correspond to strongly active variables (with respect to the subproblem) to obtain a submatrix of the Jacobian (not including the simple bounds) having full row rank. Nevertheless, in practice not even in the latter case both procedures give the same estimate, as usually an inexact minimization of the subproblem is carried out. This work puts forward also a procedure to compute the multiplier estimates by solving a reduced system (41), instead of having to solve the large system (13), if the conditions so permit (see first lines of this paragraph).

In the previous two sections specific procedures for transforming problems of type **IP** (9–12) into problems of type **EP** (1–4) have been considered.

The procedure described in Section 4 as compared with that described in section 3 has the advantage that it does not increase subproblem size with respect to the original problem **IP**. Results of numerical tests comparing both procedures have not been included yet as, in our view, the construction of appropriate software to exploit the technique put forward by Conn *et al* in [5] —once fitted to the structure of problem **EPy** (45–49)— is by no means trivial and its coding is still underway.

Using these procedures, inequalities in general constraints are eliminated. The validity for these problems of results and procedures put forward in [10] for type **EP** problems only has also been proved.

6. REFERENCES

- [1] **Bertsekas D.P.** (1977). «Approximation procedures based on the method of multipliers». *J. Opt. Th. and Appl.*, **23**, 487–510.
- [2] **Bertsekas D.P.** (1982). *Constrained Optimization and Lagrange Multiplier Methods*. Academic Press, New York.
- [3] **Bertsekas D.P.** (1995). *Nonlinear Programming*. Athena Scientific, Belmont, Massachusetts.
- [4] **Bertsekas D.P.** (1998). *Network Optimization: Continuous and Discrete Models*. Athena Scientific, Belmont, Massachusetts.
- [5] **Conn, A.R., Gould, N. and Toint, Ph.L.** (1994). «A note on exploiting structure when using slack variables». *Mathematical Programming*, **67**, 89–97.
- [6] **Conn, A.R., Gould, N., Saternaer, A. and Toint, Ph.L.** (1996). «Convergence properties of an augmented Lagrangian algorithm for optimization with a combination of general equality and linear constraints». *SIAM J. on Optimization*, **6**, **3**, 674–703.
- [7] **Gill, P.E., Murray, W. and Wright, M.H.** (1981). *Practical Optimization*. Academic Press, New York.
- [8] **Hestenes, M.R.** (1969). «Multiplier and gradient methods». *J. Opt. Th. and Appl.*, **4**, 303–320.
- [9] **Mijangos, E. and Nabona, N.** (1996). «The application of the multipliers method in nonlinear network flows with side constraints». *Technical Report 96/10*, Dept. of Statistics and Operations Research. Universitat Politècnica de Catalunya, 08028 Barcelona, Spain.
- [10] **Mijangos, E. and Nabona, N.** (1997). «On the compatibility of classical first order multiplier estimates with variable reduction techniques.» *Technical Report 97/06*, Dept. of Statistics and Operations Research. Universitat Politècnica de Catalunya, 08028 Barcelona, Spain.
- [11] **Mijangos, E.** (1996). «Optimización de flujos no lineales en redes con restricciones laterales mediante técnicas de multiplicadores». *Ph.D. Thesis, Statistics and Operations Research Dept.*. Universitat Politècnica de Catalunya, 08028 Barcelona, Spain.
- [12] **Mijangos, E.** (1997). «PFNRN03 user's guide». *Technical Report 97/05*, Dept. of Statistics and Operations Research. Universitat Politècnica de Catalunya, 08028 Barcelona, Spain.
- [13] **Murtagh, B.A. and Saunders, M.A.** (1978). «Large-scale linearly constrained optimization». *Mathematical Programming*, **14**, 41–72.

- [14] **Murtagh, B.A.** and **Saunders, M.A.** (1983). *MINOS 5.0. User's guide*. Dept. of Operations Research, Stanford University, CA 9430, USA.
- [15] **Powell, M.J.D.** (1969). «A method for nonlinear constraints in minimization problems». *Optimization* (R. Fletcher, ed.), 283–298. Academic Press New, York.
- [16] **Rockafellar, R.T.** (1971). «New applications of duality in convex programming». *Proc. Confer. Probab.*, 4th, Brasov, Romania, 73–81.

MINIMIZACIÓN GLOBAL DE UN POLINOMIO EN LA RECTA REAL

C. BELTRÁN ROYO

Universitat Politècnica de Catalunya*

En este artículo presentamos y probamos numéricamente un nuevo algoritmo para la minimización global de un polinomio de grado par. El algoritmo está basado en la simple idea de trasladar verticalmente el grafo del polinomio hasta que el eje OX sea tangente al grafo del polinomio trasladado. En esta privilegiada posición, cualquier raíz real del polinomio trasladado es un mínimo global del polinomio original.

Globally minimizing a polynomial on the real line.

Palabras clave: Optimización global, máximo común divisor de polinomios, algoritmo de Euclides, secuencia de Sturm, división sintética de dos polinomios.

Clasificación AMS: 49J05

*Dept. d'Estadística i Investigació Operativa. Secció d'Informàtica. Universitat Politècnica de Catalunya. Pau Gargallo, 5. 08028 Barcelona (Espanya).

—Recibido en octubre de 1997.

—Aceptado en julio de 1998.

1. RESUMEN

Dado un polinomio $P(x)$ de grado par, queremos resolver el problema de minimizar *globalmente* $P(x)$ en el conjunto de los reales (problema P-1). Para resolver el problema hemos diseñado un algoritmo o método nuevo denominado método del máximo común divisor (MCD).

El algoritmo del máximo común divisor se basa en:

- a) Conocer el número de raíces reales diferentes del polinomio sin tener que calcularlas (método de Sturm).
- b) En el cálculo del máximo común divisor de dos polinomios que denotamos por MCD.

La idea del método es trasladar verticalmente el polinomio P hasta conseguir que el eje OX sea tangente al grafo de P trasladado (denotamos por P^* este P trasladado). Una vez ya tenemos P^* calculamos el máximo común divisor de P^* y su derivada $P^{*'} (Q = \text{MCD}(P^*, P^{*'}))$ que en la mayoría de los casos será un polinomio de grado uno. Finalmente, cualquier raíz del polinomio Q es un mínimo global del polinomio P .

Para poner en práctica el anterior esquema usaremos conocidos métodos y algoritmos:

- a) Vía método de Sturm podemos conocer el número de raíces reales distintas de P y un uso iterativo de esta información nos conducirá a P^* .
- b) Vía el algoritmo de Euclides calcularemos Q .
- c) La división sintética de dos polinomios será requerida para aplicar el algoritmo de Euclides.
- d) Finalmente usaremos el método de Birge-Vieta para obtener una raíz de Q .

2. PRECEDENTES

En el artículo Goldstein (1971) resuelven el problema P-1 de forma directa i.e. se obtiene un mínimo global extrayendo de P cada mínimo local encontrado vía división de polinomios. Esta estrategia puede llevar a tener que resolver un gran número de minimizaciones locales, además de que si no se hace minimización local exacta, los errores acumulados pueden degradar muy seriamente la solución.

Otros autores, resuelven el problema P-1 localmente i.e. minimizan globalmente $P(x)$ en un intervalo pero no en todo $\mathbb{R}$ (denominamos este problema P-2). Repasemos tres trabajos:

- 1) En el artículo Wingo (1985) se propone un algoritmo sencillo para el problema P-2 que no requiere la evaluación de derivadas.
- 2) En el artículo Floudas (1992) se resuelve el problema P-2 vía programación lineal con restricciones no lineales. La idea del método es cambiar en el polinomio cada potencia de x en una nueva variable. Así pasan de la minimización unidimensional de $P(x)$ en un intervalo (a, b) , a la minimización de una función lineal $c'x$ en un dominio de R^{n+1} .
- 3) En el artículo Bromberg (1992) se propone un algoritmo para la minimización global de una función $f(x)$ en un intervalo (a, b) . La idea de su método consiste en reducir el intervalo inicial (a, b) aplicando técnicas de minimización local.

Por lo tanto, a excepción del trabajo de Goldstein (1971), ninguno de los trabajos citados resuelve el problema de minimizar $P(x)$ en toda la recta real (problema P-1). El método que ahora se propone resuelve el problema P-1 completamente.

3. NOTACIÓN

- Consideramos polinomios $P(x) = a_1 \cdot x^n + a_2 \cdot x^{n-1} + \dots + a_n \cdot x + a_{n+1}$ de coeficientes reales y variable real.
- Con $P(x, c)$ denotemos el polinomio que tiene $a_{n+1} = c$ i.e.

$$P(x, c) = a_1 \cdot x^n + a_2 \cdot x^{n-1} + a_3 \cdot x^{n-2} + \dots + a_n \cdot x + c$$

- Conjunto de las raíces reales del polinomio P

$$\text{Raíces}(P) = \{x \text{ de } \mathbb{R} / P(x) = 0\}$$

- Cardinal de raíces(P) $\# \text{Raíces}(P)$
- Grafo de un polinomio $\text{grafo}(P) = \{(x, P(x)) \text{ en } \mathbb{R}^2 / x \text{ en } \mathbb{R}\}$
- Semiespacio positivo $S^+ = \{(x, y) \text{ de } \mathbb{R}^2 / y \text{ positivo}\}$
- Semiespacio negativo $S^- = \{(x, y) \text{ en } \mathbb{R}^2 / y \text{ negativo}\}$

4. MODELO PROPIO - RESULTADOS TEÓRICOS

En todo el trabajo se está presuponiendo:

- La minimización de un polinomio $P(x)$ de grado par, dado que un polinomio de grado impar no está acotado inferiormente en $\mathbb{R}$. Por la misma razón, supondremos que el coeficiente a_1 es positivo.

- La minimización de $P(x)$ se hace en todo $\mathbb{R}$.
- Todas las raíces de $P(x)$ que se consideran son reales.
- Un polinomio puede alcanzar el mínimo global en más de un punto. El método desarrollado da como resultado solamente uno de estos puntos, aunque el método fácilmente se puede adaptar para dar todos los puntos donde el polinomio alcanza el mínimo absoluto.

A continuación veremos unas definiciones, unos procedimientos y unos resultados que están demostrados en el reporte de investigación Beltran (1997) y que usaremos en el algoritmo del MCD.

La idea del algoritmo es en síntesis: dado un polinomio $P(x)$ de grado par, siempre podemos trasladarlo verticalmente de forma que $\text{Grafo}(P)$ sea tangente al eje OX obteniendo así el polinomio $P^*(x)$. Este último polinomio tiene la característica de que cualquiera de sus raíces reales es un punto de mínimo global de $P^*(x)$ y de $P(x)$ en todo $\mathbb{R}$. La forma de buscar una raíz cualquiera de $P^*(x)$ se hace teniendo en cuenta que en cada punto donde se anula $P^*(x)$ también se anula su derivada $P^{*'}(x)$. Dado que los ceros comunes de $P^*(x)$ y $P^{*'}(x)$ coinciden con los ceros del polinomio $\text{MCD}(P^*, P^{*'})$, cualquier cero de este último polinomio será un mínimo global de $P^*(x)$ y por tanto de $P(x)$.

4.1. Polinomios pobres, ricos y buenos

Definición DES-1 (polinomio pobre, rico y bueno):

Consideremos un polinomio de grado par y con $a_1 > 0$:

$$P(x) = a_1 \cdot x^n + a_2 \cdot x^{n-1} + a_3 \cdot x^{n-2} + \cdots + a_n \cdot x + a_{n+1}$$

- $P(x)$ es pobre si el eje OX no corta el grafo(P) (escribiremos P^-).
- $P(x)$ es rico si el eje OX corta el grafo(P) (escribiremos P^+).
- $P(x)$ es bueno si el eje OX es tangente al grafo(P) (escribiremos P^*).

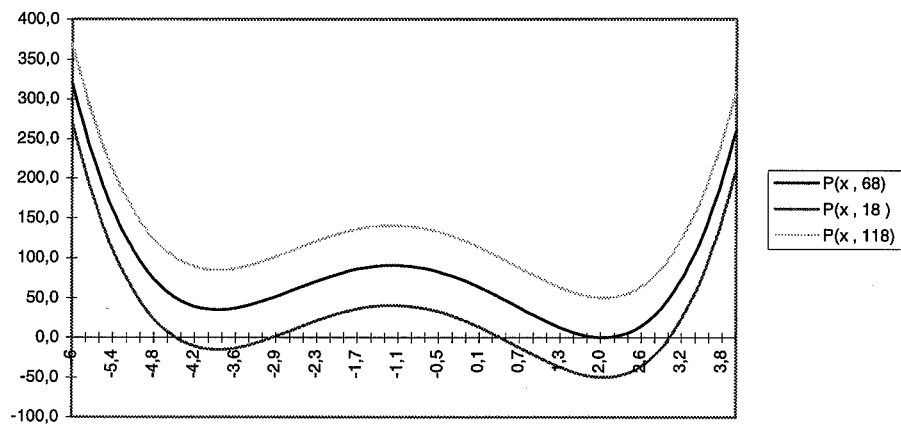
Corolario DES- 1. Dado un polinomio $P(x)$ de grado n :

- Siempre existe una traslación vertical de $P(x)$ que lo convierte en polinomio bueno $P^*(x)$. Además $P(x)$ difiere de $P^*(x)$ solamente en el coeficiente a_{n+1} .
- x^* es un óptimo global de $P^*(x)$ si y solo si x^* es un óptimo global de $P(x)$. Por tanto podemos limitar la búsqueda de óptimos globales en el caso de polinomios buenos.

EJEMPLO REP-1.

$$P(x, c) = x^4 + 4x^3 - 11x^2 - 36x + c$$

$P(x, 68)$	polinomio bueno	(1 raíz real)
$P(x, 18)$	polinomio rico	(4 raíces reales)
$P(x, 118)$	polinomio pobre	(0 raíces reales)



4.2. Obtención de un óptimo global de un polinomio bueno $P^*(x)$

Teorema DES-1. (caracterización de óptimo global vía MCD) Sean:

$P^*(x)$ polinomio bueno y $P^{*'}(x)$ su derivada,
 $M(x)$ polinomio definido como el MCD(P^* , $P^{*'}(x)$),
 x^* punto de $\mathbb{R}$.

Entonces:

$$M(x^*) = 0 \quad \text{sii} \quad x^* \text{ mínimo global de } P^*(x).$$

EJEMPLO REP-2. Continuamos con el polinomio del ejemplo REP-1.

$$P^*(x) = x^4 + 4x^3 - 11x^2 - 36x + 68$$

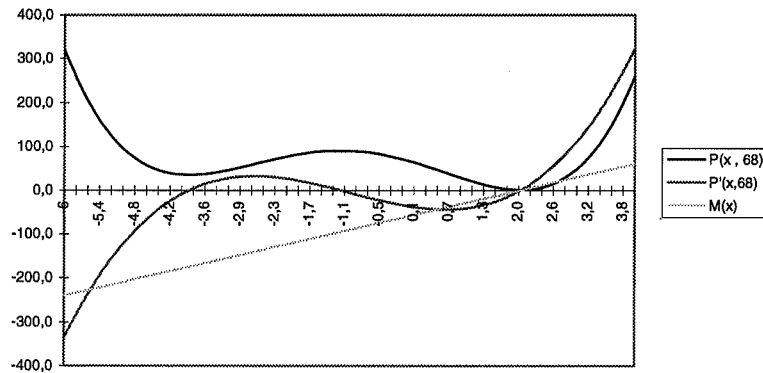
$$P^{*'}(x) = 4x^3 + 12x^2 - 22x - 36$$

$$M(x) = \text{MCD}(P^*, P^{*'}) = (x - 2)$$

En el siguiente gráfico se ve como:

- $x = 2$ es el mínimo global de $P(x, 68)$.
- $x = 2$ es el único real donde se anula simultáneamente P^* y $P^{*'}$.
- por lo tanto $(x - 2)$ es un factor de M (en realidad el único) y así $M(2) = 0$.
- Podemos concluir que $x^* = 2$ óptimo global de $P^*(x)$ y por lo tanto de $P(x)$.

NOTA: En lugar de representar $y = M(x)$ hemos representado $y = 30 \cdot M(x)$, pues la representación de $M(x)$ queda casi pegada al eje OX .



4.3. Obtención del polinomio bueno P^* asociado a un polinomio P

Tenemos dos tipos de información que usaremos en el algoritmo OGP.

4.3.1. Primera información - #Raíces(P)

El siguiente resultado puede encontrarse en Henrici (1974) pág. 448.

Teorema de Sturm DES-2. (Número de ceros reales y distintos de $P(x)$ en (a, b)). Definimos $P_0(x) := P(x)$ y $P_1(x) := P'(x)$. Sea $\{P_0, P_1, P_2, \dots, P_m\}$ la secuencia de polinomios obtenidos mediante el algoritmo de Euclides para calcular el MCD(P_0, P_1) (ver apéndice). Entonces:

El número de ceros reales y distintos de P_0 en el intervalo (a, b) es igual a la diferencia entre el número de cambios de signo en $\{P_0(a), P_1(a), P_2(a), \dots, P_m(a)\}$ y el número de cambios de signo en $\{P_0(b), P_1(b), P_2(b), \dots, P_m(b)\}$ suponiendo que $P_0(a) \cdot P_0(b)$ sea diferente de cero.

Definición DES-2. (Secuencia de Sturm asociada a un polinomio). El conjunto de polinomios $\{P_0, P_1, P_2, \dots, P_m\}$ que aparece en el enunciado del teorema anterior, se denomina secuencia de Sturm asociada a P_0 . Escribiremos $\text{Sturm}(P)$ para hacer referencia a la secuencia de Sturm asociada al polinomio $P(x)$.

EJEMPLO REP-3. Continúa del ejemplo REP-2.

$$P^*(x) = x^4 + 4x^3 - 11x^2 - 36x + 68$$

Veamos $\text{Sturm}(P^*)$:

$$\begin{aligned} P_0(x) &:= P^*(x) = x^4 + 4x^3 - 11x^2 - 36x + 68 \\ P_1(x) &:= P^{*'}(x) = 4x^3 + 12x^2 - 22x - 36 \\ P_2(x) &= 8.5x^2 - 21.5x - 77 \\ P_3(x) &= -9.4x + 18.8 = -9.4 \cdot (x - 2) \\ P_4(x) &= 0 \end{aligned}$$

$$\begin{aligned} N(-\infty) &:= \text{Cambios de signo en } \{P_0(-\infty), P_1(-\infty), P_2(-\infty), P_3(-\infty), P_4(-\infty)\} = \\ &= \text{Cambios de signo en } \{+, -, +, +, 0\} = \\ &= 2 \end{aligned}$$

$$\begin{aligned} N(\infty) &:= \text{Cambios de signo en } \{P_0(\infty), P_1(\infty), P_2(\infty), P_3(\infty), P_4(\infty)\} = \\ &= \text{Cambios de signo en } \{+, +, +, -, 0\} = \\ &= 1 \end{aligned}$$

Entonces por el teorema de Sturm tenemos que $\# \text{Raíces}(P^*) = N(-\infty) - N(\infty) = 1$ y podemos asegurar que P^* es polinomio bueno. Podemos ver la representación de $P^*(x) = P(x, 68)$ en Ejemplo REP-1 donde efectivamente se ve que P^* tiene una sola raíz real.

Nótese que al ser $P_4(x) = 0$ entonces $P_3(x) = -9.4 \cdot (x - 2)$ nos da el MCD($P^*, P^{*'} = (x - 2)$).

4.3.2. Segunda información - la función residuo

Definición DES-4. (función residuo asociada a un polinomio). Dada la familia de polinomios $F = \{P(x, c) : c \text{ de } \mathbb{R}\}$ definimos la función residuo $R_P(c)$ asociada a F , como el resto de la división $P_{m-1}(x, c) : P_m(x, c)$ donde P_{m-1} y P_m son los dos últimos polinomios no constantes de Sturm ($P(x, c)$).

Corolario DES-4. (Carácter de un polinomio vía función residuo). Sea c^* la mayor raíz real de $R_P(c)$. Entonces:

- a) $P(x, c)$ pobre sii $c > c^*$.
- b) $P(x, c)$ bueno sii $c = c^*$.
- c) $P(x, c)$ rico sii $c < c^*$.

EJEMPLO REP-4. Continúa del ejemplo REP-3.

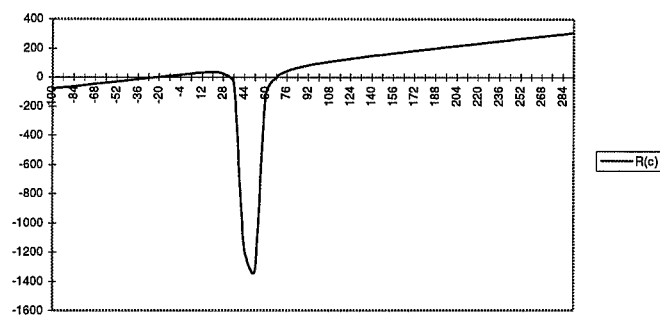
$$P(x, c) = x^4 + 4x^3 - 11x^2 - 36x + c$$

Con la ayuda del paquete «Mathematica» (Wolfram (1992)) hemos calculado $R_P(c)$ obteniendo la fracción algebraica siguiente:

$$R_P(c) = A(c)/B(c)$$

donde
$$A(c) = 289 \cdot (16c^3 - 1256c^2 - 483c + 809676)$$
$$B(c) = (68c - 3255)^2$$

En la siguiente representación de $R_P(c)$ se ve como su mayor raíz real es $c^* = 68$, que confirma el hecho de que $P(x, 68)$ sea un polinomio bueno (como vimos en el ejemplo REP-1).



5. MÉTODO DEL MCD PARA LA MINIMIZACIÓN GLOBAL DE UN POLINOMIO

Dado un polinomio de grado par $P(x)$, buscamos un x^* mínimo global de $P(x)$ en $\mathbb{R}$. Las etapas del método del MCD se explican en 5.1, 5.2 y 5.3.

Entonces, en virtud del teorema DES-1, x^* es un óptimo global de $P(x)$ en $\mathbb{R}$.

5.1. Obtención del polinomio bueno $P^*(x)$ asociado a un polinomio $P(x)$

Aplicando el corolario DES-4 es suficiente encontrar c^* raíz máxima de la función residuo $R_P(c)$ y entonces $P(x, c^*)$ polinomio bueno asociado a P^* . Estamos interesados, pues, en un intervalo (a_0, b_0) que contenga c^* . Dado $P(x)$ hacemos:

5.1.1. Cálculo del intervalo inicial (a_0, b_0) .

5.1.2. Cálculo de c^* .

5.1.1. Cálculo del intervalo inicial tal que contenga c^*

Vía el teorema de Sturm, podemos conocer $\#Raíces(P(x, c_k))$ $k = 0, 1, 2, 3, \dots$ y así determinar el carácter rico, pobre o bueno de cada $P(x, c_k)$. Podemos incrementar c_k hasta que lleguemos a un $P(x, c_p)$ polinomio pobre y así tendremos garantizado $c^* < c_p$.

Tomaremos $b_0 = c_p$ como extremo superior del intervalo inicial.

En resumen, hemos empobrecido el polinomio inicial $P(x, c_0)$ para buscar b_0 .

Por otra parte, dado que $P(x, 0)$ siempre tiene una raíz real en $x = 0$, podemos tomar como extremo inferior del intervalo inicial, cualquier valor a_0 positivo, lo más cerca de b_0 posible y de forma que $P(x, a_0)$ sea un polinomio rico.

EJEMPLO REP-5. Continúa del ejemplo REP-4.

Dado que $P(x, 18)$ es un polinomio rico, después de 2 iteraciones de la rutina EMPOBRINT, obtenemos $P(x, 72)$ que es un polinomio pobre. De esta forma tenemos que $18 < c^* < 72$ con lo cual podemos tomar $(18, 72)$ como intervalo inicial.

A continuación tenemos Sturm ($P(x, \underline{72})$).

***** EMPOBRINT 0

Polinomio a optimizar $P(x, 18)$

1.00 4.00 -11.00 -36.00 18.00

OGP/out	sturm					
	1.00	4.00	-11.00	-36.00	<u>72.00</u>	P_0
	.00	4.00	12.00	-22.00	-36.00	P_1
	.00	.00	8.50	1.50	-81.00	P_2
	.00	.00	.00	11.36	18.06	P_3
	.00	.00	.00	.00	25.30	P_4

***** EMPOBRINT 2

5.1.2. Cálculo de c^*

5.1.2.1. Primera información, $\#Raíces(P)$. Vía teorema de Sturm podemos conocer el número de raíces de cada polinomio $P(x, c_k)$. Al mismo tiempo por el corolario DES-4 sabemos que c^* es un valor límite en el sentido de que si $c < c^*$ entonces $P(x, c)$ es un polinomio rico, y en cambio, si $c > c^*$ entonces $P(x, c)$ es un polinomio pobre. Un primer algoritmo básico para calcular c^* , sería seguir un esquema tipo bisección con la información $\#Raíces(P(x, c_k))$ $k = 0, 1, 2, \dots$, de manera que si $P(x, c_k)$ es pobre, disminuiríamos c_k , y en cambio, si $P(x, c_k)$ es rico, aumentaríamos c_k i.e. vamos trasladando verticalmente el grafo de P hasta que sea tangente al eje OX . Siguiendo este proceso de forma iterativa, acabaríamos obteniendo $P(x, c_n)$ polinomio bueno.

5.1.2.2. Segunda información, la función residuo. Recordemos la definición de función residuo asociada a un polinomio (definición DES-4):

Dada la familia de polinomios $F = \{P(x, c) : c \text{ de } \mathbb{R}\}$ definimos la función residuo $R_P(c)$ asociada a F , como el residuo de la división $P_{m-1}(x, c) : P_m(x, c)$ donde P_{m-1} y P_m son los dos últimos polinomios no constantes de Sturm ($P(x, c)$).

La forma de utilizar la información de la función residuo $R_P(c)$ nos la da el siguiente teorema:

Teorema DES-3. (Caracterización de polinomio bueno vía la función de residuo $R_P(c)$). Sea $F = \{P(x, c) : c \text{ de } \mathbb{R}\}$ y definimos $P^*(x) = P(x, c^*)$, entonces:

P^* polinomio bueno sii c^* es la mayor raíz real de $R_P(c)$.

Se trata pues de calcular la mayor raíz de $R_P(c)$.

Llegados hasta este punto, de las alternativas para calcular esta raíz hemos escogido el método de la secante por:

- No disponer de $R_P(c)$ de forma explícita ni de su derivada. El valor de $R_P(c_0)$ se obtiene como un subproducto al calcular Sturm ($P(x, c_0)$).
- El método se ha mostrado robusto y rápido en la práctica.

El método de la secante está descrito en el apéndice y programado en la rutina SECANT.

5.1.2.3. Combinación de las dos informaciones

ALGORITMO - 3. (Rutina BISSECCIO). Combinaremos las dos informaciones anteriores para obtener un algoritmo rápido y robusto.

0. Partimos del intervalo inicial (a_0, b_0) tal que $a_0 < c^* < b_0$. Hacemos $k = 0$.
1. Aplicamos el método de la secante y podemos tener los tres casos siguientes:
 - 1.1. Si encontramos la mayor raíz c de $R_P(c)$ STOP.
 - 1.2. Si encontramos una raíz c_0 de $R_P(c)$ pero no es máxima. Entonces, hacemos $k = k + 1$, tomando como nuevo intervalo $(a_k, b_k) = (c_0 + \epsilon, b_k)$ y volvemos al punto a 1.
 - 1.3. Si el algoritmo no converge o lo hace a un punto externo al intervalo (a_k, b_k) , entonces vamos a 2.
2. Ha fracasado el método de la secante, y aplicamos una etapa del método de la bisección al intervalo (a_k, b_k) .
3. Volvemos a 1.

EJEMPLO REP-5. El intervalo inicial es $(36, 72)$. La rutina BISSECCIO encuentra $c^* = 68$ con 8 iteraciones del método de la secante. Al final tenemos Sturm ($P(x, 68)$).

***** BISSECCIÓ

c	$R_P(c)$
<u>36.0000</u>	-39.4913
<u>72.0000</u>	25.3022
57.9418	-198.6757
70.4119	16.4875
69.4563	10.4699
67.7938	-1.6308
68.0178	.1392
68.0002	.0017
<u>68.0000</u>	<u>.0000</u>

BIS/SEC	iter	8					
OGP/out	sturm						
	1.00	4.00	-11.00	-36.00	<u>68.00</u>	P_0	
	.00	4.00	12.00	-22.00	-36.00	P_1	
	.00	.00	8.50	21.50	-77.00	P_2	
	.00	.00	.00	-9.47	18.95	P_3	
	.00	.00	.00	.00	<u>.00</u>	P_4	

***** BISSECCIÓ

5.2. Calcular $M(x) := \text{MCD}(P^*, P^{*'})$

Una vez conseguido $P^* = P(x, c^*)$ polinomio bueno, necesitamos $M(x) := \text{MCD}(P^*, P^{*'})$.

Usaremos el método de Euclides. Construimos $\text{Sturm}(P^*) = \{P_0^*, P_1^*, \dots, P_m^*\}$ y tal como se describe en el apéndice, P_m^* coincide con el $\text{MCD}(P^*, P^{*'})$ salvo una constante y entonces podemos tomar $M(x) = P_m^*(x)$ ya que la constante no afecta a las raíces de $M(x)$. Nótese que para calcular $M(x)$ no necesitamos hacer ningún cálculo adicional una vez hemos construido $\text{Sturm}(P^*)$.

En el ejemplo anterior veíamos como al ser $P_4(x)$ el polinomio nulo, entonces $P_3(x) = -9.4x + 18.95 = -9.4 \cdot (x - 2)$ nos da el $\text{MCD}(P^*, P^{*'}) = (x - 2)$.

5.3. Calcular x^* una raíz de $M(x)$ (Rutina BIRGE)

El teorema DES-1 nos asegura que cualquier raíz real de $M(x)$ es un óptimo global. Tenemos que resolver pues la ecuación $M(x) = 0$. Lo haremos con el método de Birge-Vieta descrito en el apéndice. En muchos casos $M(x) = 0$ será una ecuación de primer grado y la rutina BIRGE no hará más que un cálculo directo.

EJEMPLO REP-6. Continúa del ejemplo REP-5.

La rutina BIRGE calcula la única raíz de $M(x) = -9.47x + 18.95$ que es $x = 2$.

En consecuencia ya tenemos el óptimo global de $P(x)$ i.e. $x^* = 2$.

```
***** BIRGE  0

OGP/in   mcd      -9.47  18.95

OGP/out  OPTIM                      2.0000

***** BIRGE  1
```

6. EJEMPLOS

6.1. Ejemplo completo

Ahora reunimos el ejemplo que hemos visto de forma fragmentada para dar una visión global del método.

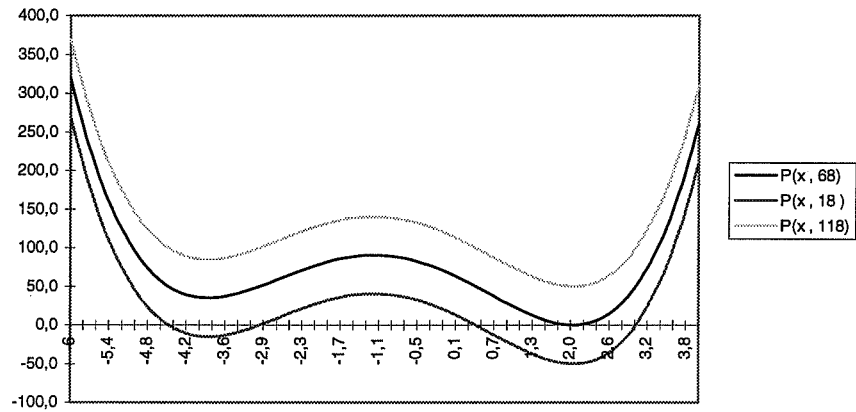
6.1.1. Ideas, gráficos y cálculos

Queremos calcular $\min_x P(x)$ donde $P(x) = x^4 + 4x^3 - 11x^2 - 36x + 18$.

Según el valor del término independiente c , tendremos un polinomio bueno, rico o pobre.

$$P(x, c) := x^4 + 4x^3 - 11x^2 - 36x + c.$$

$P(x, 68)$	polinomio bueno	(1 raíz real).
$P(x, 18)$	polinomio rico	(4 Raíces reales).
$P(x, 118)$	polinomio pobre	(0 Raíces reales).

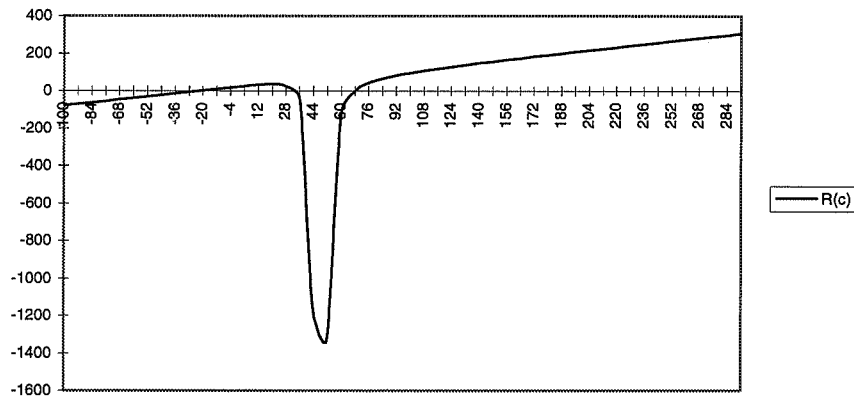


Para obtener $P^*(x)$ aparte de calcular $\#Raíces(P)$ usamos la función residuo $R_P(c)$ que en este caso, dado que $P(x, c) = x^4 + 4x^3 - 11x^2 - 36x + c$ tiene la siguiente expresión:

$$R_P(c) = A(c)/B(c)$$

donde $A(c) = 289 \cdot (16c^3 - 1256c^2 - 483c + 809676)$
 $B(c) = (68c - 3255)^2$

En la siguiente representación de $R_P(c)$ se ve como su mayor raíz real es $c^* = 68$ confirmando el hecho de que $P(x, 68)$ sea un polinomio bueno.



El método del MCD nos conduce a $P^*(x) = x^4 + 4x^3 - 11x^2 - 36x + 68$.

A modo de ejemplo veamos ahora Sturm (P^*) y el cálculo de $\#Raíces(P^*)$:

$$\begin{aligned} P_0(x) &= P^*(x) = x^4 + 4x^3 - 11x^2 - 36x + 68 \\ P_1(x) &= P^{*'}(x) = 4x^3 + 12x^2 - 22x - 36 \\ P_2(x) &= 8.5x^2 - 21.5x - 77 \\ P_3(x) &= -9.4x + 18.8 = -9.4 \cdot (x - 2) \\ P_4(x) &= 0 \end{aligned}$$

$$\begin{aligned} N(-\infty) &:= \text{Cambios de signo en } \{ P_0(-\infty), P_1(-\infty), P_2(-\infty), P_3(-\infty), P_4(-\infty) \} = \\ &= \text{Cambios de signo en } \{ +, -, +, +, 0 \} = \\ &= 2 \end{aligned}$$

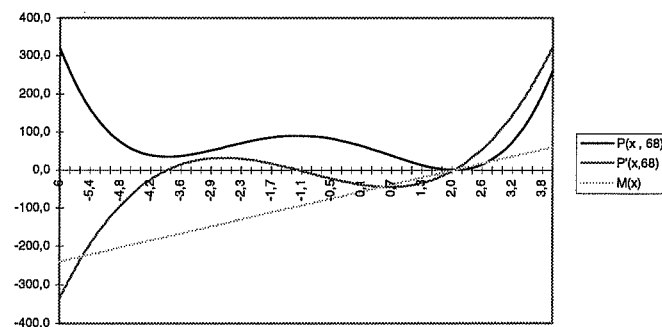
$$\begin{aligned} N(\infty) &:= \text{Cambios de signo en } \{ P_0(\infty), P_1(\infty), P_2(\infty), P_3(\infty), P_4(\infty) \} = \\ &= \text{Cambios de signo en } \{ +, +, +, -, 0 \} = \\ &= 1 \end{aligned}$$

Entonces por el teorema de Sturm tenemos que $\#Raíces(P^*) = N(-\infty) - N(\infty) = 1$ y podemos asegurar que $P^*(x)$ es un polinomio bueno. Al principio de este ejemplo podemos ver la representación de $P^*(x)$ que coincide con $P(x, 68)$ y confirmar que $P^*(x)$ tiene una única raíz real. Nótese que al ser $P_4(x) = 0$, entonces $P_3(x) = -9.4 \cdot (x - 2)$ coincide, excepto por una constante, con $\text{MCD}(P^*, P^{*'}) = (x - 2)$.

Una vez tenemos $P^*(x)$, calculamos los polinomios $P^{*'}(x)$ y $M(x)$ que representamos a continuación:

$$\begin{aligned} P^*(x) &= x^4 + 4x^3 - 11x^2 - 36x + 68 \\ P^{*'}(x) &= 4x^3 + 12x^2 - 22x - 36 \\ M(x) &= \text{MCD}(P^*, P^{*'}) = (x - 2) \end{aligned}$$

En el siguiente gráfico se ve como:



- $x = 2$ es el mínimo global de $P(x, 68)$.
- $x = 2$ es el único punto donde se anulan simultáneamente $P^*(x)$ y $P^{*'}(x)$.
- Por lo tanto $(x - 2)$ es un factor de M (en realidad el único) y así $M(2) = 0$.
- Podemos concluir que $x^* = 2$ es el óptimo global de $P^*(x)$ y por lo tanto de $P(x)$.

NOTA. En lugar de representar $y = M(x)$, hemos representado $y = 30 \cdot M(x)$ pues la representación de $M(x)$ no se distinguiría del eje OX .

6.1.2. Resultados con el programa OGP.EXE

Para correr el programa OCP.EXE necesitamos el fichero de datos OGP.DAT que consta de tres partes:

- Grado del polinomio (entero positivo).
- Precisión exigida (por ejemplo 0.0001).
- El polinomio que queremos minimizar escrito en forma de columna.

EJEMPLO REP-6. Para introducir el polinomio $P(x) = x^4 + 4x^3 - 36x + 18$ crearíamos un fichero OGP.DAT con el siguiente contenido.

4	(Grado del polinomio)
0.0001	(Precisión exigida)
1.1	(a_1)
4.0	(a_2)
0.0	(a_3)
-36.0	(a_4)
18.0	(a_5)

Dado que $P(x, 18)$ es un polinomio rico, después de 2 iteraciones de la rutina «empo-
brint» obtenemos $P(x, 72)$ que es un polinomio pobre. A continuación tenemos Sturm
($P(x, 72)$)

***** EMPOBRINT 0

Polinomio a optimizar $P(x, 18)$

1.00 4.00 -11.00 -36.00 18.00

OGP/out

sturm

1.00	4.00	-11.00	-36.00	<u>72.00</u>	P_0
.00	4.00	12.00	-22.00	-36.00	P_1
.00	.00	8.50	1.50	-81.00	P_2
.00	.00	.00	11.36	18.06	P_3
.00	.00	.00	.00	25.30	P_4

***** EMPOBRINT 2

El intervalo inicial donde se encuentra c^* es (36, 72). La rutina BISSECCIÓ encuentra c^* con 8 iteraciones del método de la secante. Al final tenemos Sturm ($P(x, 68)$):

***** BISSECCIÓ

c_k	$R_P(c_k)$
<u>36.0000</u>	-39.4913
<u>72.0000</u>	25.3022
57.9418	-198.6757
70.4119	16.4875
69.4563	10.4699
67.7938	-1.6308
68.0178	.1392
68.0002	.0017
<u>68.0000</u>	<u>.0000</u>

BIS/SEC	iter	8					
OGP/out	sturm						
	1.00	4.00	-11.00	-36.00	<u>68.00</u>	P_0	
	.00	4.00	12.00	-22.00	-36.00	P_1	
	.00	.00	8.50	21.50	-77.00	P_2	
	.00	.00	.00	-9.47	18.95	P_3	
	.00	.00	.00	.00	<u>.00</u>	P_4	

***** BISSECCIÓ

La rutina BIRGE calcula la única raíz de $M(x) = -9.47x + 18.95$ que es $x = 2$.

En consecuencia ya tenemos el óptimo global de $P(x)$ i.e. $x^* = 2$.

***** BIRGE 0

OGP/in mcd -9.47 18.95

OGP/out OPTIM 2.0000

***** BIRGE 1

6.2. Pruebas con otros polinomios

A continuación definiremos unos polinomios que nos servirán para poner a prueba el método del MCD descrito anteriormente. Dado un polinomio

$$P(x) = a_1 \cdot x^n + a_2 \cdot x^{n-1} + a_3 \cdot x^{n-2} + \cdots + a_n \cdot x + a_{n+1}$$

le haremos corresponder el vector formado por sus coeficientes

$$P(x) = (a_1, a_2, \dots, a_{n+1}).$$

La etiqueta del polinomio nos dirá el grado y el número de identificación del polinomio. Por ejemplo con 4POL1, designamos un polinomio de grado 4 y que identificamos con el número 1.

$$\begin{aligned}
4\text{POL } 1 &= (\begin{array}{cccccc} 1. & 0. & 0. & 0. & -1. & \end{array}) \\
4\text{POL } 2 &= (\begin{array}{cccccc} 1. & 4. & -11. & -36. & 18. & \end{array}) \\
4\text{POL } 3 &= (\begin{array}{cccccc} 1. & -3. & -1.5 & 10. & 0. & \end{array}) \\
4\text{POL } 4 &= (\begin{array}{cccccc} 1. & -4. & 4. & 0. & 0. & \end{array}) \\
6\text{POL } 1 &= (\begin{array}{cccccc} 1. & -6. & 15. & -20. & 15. & -6. & 1. & \end{array}) \\
6\text{POL } 2 &= (\begin{array}{cccccc} 1. & 0. & -14. & 0. & 49. & 0. & -36. & \end{array}) \\
6\text{POL } 3 &= (\begin{array}{cccccc} 0.16666 & -2.08 & 0.4875 & 7.1 & -3.95 & -1. & 0.1 & \end{array}) \\
8\text{POL } 1 &= (\begin{array}{cccccc} 0.001 & 0. & -0.102 & 0.096 & 3.009 & -5.520 & -23.068 & 65.904 \\ & -40.320 & & & & & & \end{array}) \\
8\text{POL } 2 &= (\begin{array}{cccccc} 1. & 0. & 0. & -20.320. & & 12.6 & 1. & 0. \\ & -13. & & & & & & \end{array}) \\
10\text{POL } 2 &= (\begin{array}{cccccc} 0.001 & 0.019 & -0.012 & -1.842 & -4.347 & 60.291 & 142.862 & -869.188 \\ & -864.264 & 5165.28 & -3628.8 & & & & \end{array})
\end{aligned}$$

En las siguientes tablas se muestra la información más relevante obtenida al minimizar los polinomios anteriores con una tolerancia = 0.0001. Hemos de notar que:

- La rutina Birge hace normalmente una sola iteración al tener que resolver una ecuación de primer grado (caso de un único óptimo global).
- La precisión obtenida es buena con unas pocas iteraciones.

	4POL1	1POL2	4POL3	4POL4	6POL1
Óptimos globales exactos	0,0000	2,0000	-1,0000	0,0000	1,0000
			2,0000		
Óptimo calculado	0,0000	2,0000	-0,9999	0,0000	1,0351
Iteraciones «EMPOBRINT»	1	2	1	1	1
Iteraciones «BISSECCIÓ»					
• Primer intento secante	2	8	5	2	2
• Segundo intento secante					
Iteraciones «BIRGE»	1	1	1	1	1290
Comentarios			(1)		(2)
	6POL2	6POL3	8POL1	8POL2	10POL2
Óptimos globales exactos	-2,6457	10,0000	-7,3400	2,2744	6,4347
	0,0000				
	2,6457				
Óptimo calculado	0,0000	10,0004	-7,3416	2,2400	6,4347
Iteraciones «EMPOBRINT»	1	14	4	6	4
Iteraciones «BISSECCIÓ»					
• Primer intento secante	2	7	14	4	27
• Segundo intento secante			11		
Iteraciones «BIRGE»	1	1	1	1	1
Comentarios	(1)		(3)		

Comentarios:

- (1) Hay polinomios con más de un óptimo global.
- (2) Aquí, excepcionalmente, la rutina BIRGE necesita 1290 iteraciones para alcanzar una aproximación al óptimo. Esto es debido a que 6POL1 es $(x-1)^6$ polinomio de curvatura prácticamente nula en un entorno amplio del óptimo $x^* = 1$. El método Birge no es pues adecuado para este tipo de casos.
- (3) Aquí el método del MCD con 14 iteraciones ha calculado una primera raíz de $R_P(c)$ que no es la mayor; seguidamente, con 11 iteraciones más ya ha calculado la raíz máxima de $R_P(c)$.

7. CONCLUSIONES

Hemos presentado un nuevo algoritmo para minimizar globalmente en el conjunto de los reales un polinomio de grado par. La característica más importante del algoritmo es el uso de información global de la función a minimizar (número de raíces reales y función residuo) a diferencia de los métodos clásicos basados en información local (derivada en un punto). El enfoque aquí presentado parece más adecuado que métodos anteriores aunque falta hacer una prueba comparativa. Numéricamente el algoritmo se ha mostrado eficiente y robusto.

8. APÉNDICE

8.1. Algoritmo Euclides para calcular el máximo común divisor de dos polinomios $\text{MCD}(P_0, P_1)$

Dividimos P_0 entre P_1 y denotamos el resto $-P_2$. Seguidamente dividimos P_1 entre P_2 y denotamos el residuo $-P_3$. Continuamos este procedimiento hasta que obtengamos un resto igual a cero. Al final del algoritmo podremos establecer las siguientes relaciones:

$$\begin{aligned}P_0 &= Q_1 \cdot P_1 - P_2 \\P_1 &= Q_2 \cdot P_2 - P_3 \\&\vdots \\P_{m-2} &= Q_{m-1} \cdot P_{m-1} - P_m \\P_{m-1} &= Q_m \cdot P_m\end{aligned}$$

Entonces $P_m = \text{MCD}(P_0, P_1)$

8.2. (División sintética) Generalización de la «Regla de Ruffini» para dividir P_0 entre P_1

Escribimos horizontalmente los coeficientes del dividendo en orden decreciente (como en el método de Ruffini). Seguidamente colocamos verticalmente los coeficientes del divisor cambiados de signo (como en el método de Ruffini) con el coeficiente de máximo grado normalizado (igual a uno). Una referencia más detallada puede encontrarse en Acton (1990) pag. 181.

EJEMPLO APEN-1. Supongamos que queremos dividir P_1 entre P_2 con:

$$\begin{array}{rcl}P_1(x) &= & 2x^4 - 4x^3 + x^2 + 3 \\P_2(x) &= & x^3 - 2x^2 - x + 1\end{array}$$

Escribimos en forma de tabla:

P_2	P_1	2	-4	1	0	3
+2		*	4	0	*	*
+1		*	*	2	0	*
-1		*	*	*	-2	0
Cociente(x)	2	0	3	-2	3	Resto(x)

$$\text{Cociente}(x) = 2x \qquad 3x^2 - 2x + 3 = \text{Resto}(x)$$

La regla de formación es:

1. Sumamos la columna.
2. Multiplicamos el resultado del paso 1 por cada coeficiente de P_2 , empezando por arriba, colocando el producto en sucesivas columnas hacia la derecha y en la misma fila que el coeficiente de P_2 correspondiente.
3. Volvemos al paso 1, pero omitiendo el paso 2 después de haber calculado el cociente.

Al final, tendremos dos triángulos en blanco que en el ejemplo hemos marcado con «*».

8.3. Método de la secante para el cálculo de raíces reales

Queremos resolver $g(x) = 0$.

Este método es una variante del método de Newton, donde en lugar de usar $g'(x_k)$, usamos una aproximación $(g(x_k) - g(x_{k-1})) / (x_k - x_{k-1})$.

0. Partimos del intervalo inicial (a, b) que contiene la raíz buscada.

1. Inicializamos $x_0 = a$ $x_1 = b$.

2. Mientras $(|x_k - x_{k-1}| > \epsilon)$ hacer:

$$\text{paso} = g(x_k) \cdot \{(x_k - x_{k-1}) / (g(x_k) - g(x_{k-1}))\}$$

$$x_{k+1} = x_k - \text{paso}.$$

Final mientras.

8.4. Método de Birge-Vieta para calcular una raíz de un polinomio

Queremos resolver $P(x) = 0$ siendo $P(x)$ un polinomio.

Este método es una especialización del método de Newton, donde calculamos $P(x_k)$ y $P'(x_k)$ de forma muy eficiente, aprovechando las propiedades algebraicas de los polinomios.

Teorema del residuo. Consideramos $P(x)$ y suponemos que expresamos $P(x) = (x - x_0) \cdot Q(x) + R_1$. Entonces:

- a) $P(x_0)$ coincide con el residuo R_1 del cociente $P(x) : (x - x_0)$
- b) $P'(x_0)$ coincide con el residuo R_2 del cociente $Q(x) : (x - x_0)$

ALGORITMO. (Método Birge-Vieta)

- 0. Partimos del intervalo inicial (a, b) .
- 1. Inicializamos $x_0 = a$ $x_1 = b$.
- 2. Mientras $(|x_k - x_{k-1}| / > \epsilon)$ hacer
 - paso = R_1 / R_2
 - $x_{k+1} = x_k - \text{paso}$.
- Final mientras.

BIBLIOGRAFÍA

- [1] **Acton, F.S.** (1990). «Numerical Methods that Work», EEUU, *The Mathematical Association of America*, 2ª Ed., 1990.
- [2] **Beltran, C.** (1997). «Globally Minimizing Even Degree Polynomials on the Real Line». *Research Report 97/01 of the Statistics and Operations Research Department-UPC*. Date 2/97.
- [3] **Bromberg, M. and Chang, T.** (1992). «One Dimensional Global Optimization Using Linear Lower Bonds» from the book *Recent Advances in Global Optimization*, New Jersey, EEUU, Princeton University Press, 1992, pp. 201-220.
- [4] **Floudas, C.A. and Visweswaran, V.** (1992). «Global Optimization of Problems with Polynomial Functions in one Variable» from the book *Recent Advances in Global Optimization*, New Jersey, EEUU, Princeton University Press, 1992, pp. 164-197.
- [5] **Goldstein, A.A. and Price, J.F.** (1971). «On descent from local minima», *Mathematics of Computation*, **25**, 569-574.
- [6] **Henrici, P.** (1974). «Applied and Computational Complex Analysis». Volume I, EEUU, John Wiley & Sons, 1974.

- [7] **Wingo, D.R.** (1985). «Globaly Minimizing Polinomials Without Evaluating Derivatives», *International Journal of Computer Mathematics*, **17**, 287.
- [8] **Wolfram, S.** (1992). «Mathematica: A System for Doing Mathematics by Computer», EEUU, Addison-Wesley Publishing Company, 2^a Ed.

ENGLISH SUMMARY

GLOBALLY MINIMIZING A POLYNOMIAL ON THE REAL LINE

C. BELTRÁN ROYO

Universitat Politècnica de Catalunya*

An algorithm for globally minimizing an even degree polynomial on the real line is proposed and tested. The algorithm is based on the idea of translating the polynomial graph vertically, until the OX axis is tangent to the graph of the translated polynomial. At this point, any root of the translated polynomial is a global minimizer of the original polynomial.

Keywords: Global optimization, greatest common divisor of two polynomials, Euclides algorithm, Sturm sequence, synthetic division of two polynomials.

AMS Classification: 49J05

*Dept. d'Estadística i Investigació Operativa. Secció d'Informàtica. Universitat Politècnica de Catalunya.
Pau Gargallo, 5. 08028 Barcelona (Espanya).

—Received October 1997.

—Accepted July 1998.

The objective of this work is to minimize an even degree polynomial globally i.e. we want to solve the problem:

$$\min\{P(x) : x \in \mathbb{R}\}$$

Traditionally this problem has been solved using local information such as the first derivative. Our approach will use global information (the total number of different real zeros of $P(x)$ in $\mathbb{R}$). This approach avoids an exhaustive search into an interval (a, b) carried out by local information based algorithms.

The algorithm is based on the idea of translating the polynomial graph vertically until the OX axis is tangent to the graph of the translated polynomial. At this point, any root of the translated polynomial is a global minimizer of the original polynomial.

First of all we define a *rich* polynomial as a polynomial that crosses the OX axis, a *good* polynomial as a polynomial that is tangent to the OX axis and a *poor* polynomial as a polynomial that does not intersect the OX axis.

The main results we have proved are:

- a) If $P^*(x)$ is a good polynomial, $P^{*'}(x)$ its derivative and $\text{GCD}(P^*, P^{*'})$ their greatest common divisor, then, the set of global minimizers of $P^*(x)$ is the set of real zeros of the $\text{GCD}(P^*, P^{*'})$.
- b) Given the polynomial $P(x, c) = a_1 \cdot x^n + a_2 \cdot x^{n-1} + \dots + a_n \cdot x + c$ we define an associated función $R_P(c)$ called the residual function. Then $P(x, c^*)$ is a good polynomial if and only if c^* is the greatest real zero of the residual function $R_P(c)$.

The GCD method for finding a global minimizer of $P(x)$ is:

- 1) $P(x) = P(x, a_{n+1})$ is translated vertically until it becomes a good polynomial. We use two sources of information: first, Sturm theorem tells us the total number of different real zeros of $P(x, c)$ in $\mathbb{R}$ and second, solving the equation $R_P(c) = 0$ we find c^* the exact translation that converts $P(x, a_{n+1})$ into a good polynomial $P^*(x) = P(x, c^*)$.
- 2) Solving the equation $M(x) = 0$ where $M(x)$ is the $\text{GCD}(P, P^*)$, we will have a global minimizer of $P(x)$.

Computational experience has shown that the GCD method is an efficient and reliable method for globally minimizing a polynomial.

Estadística Oficial:

Tècniques de projeccions demogràfiques

LE TRAITEMENT DE L'INCERTITUDE DANS LES PERSPECTIVES DÉMOGRAPHIQUES

JOSIANNE DUCHÊNE

Université Catholique de Louvain*

Les tendances prédites sont souvent des extrapolations des tendances passées observées jusqu'à la date d'établissement des projections alors que la réalité peut présenter des retournements de tendances ou des irrégularités inattendues. L'incertitude étant comprise dans le sens d'incapacité à prévoir l'avenir avec exactitude, les erreurs de prévision sont la manifestation de cette incertitude et proviennent de l'imprécision dont l'extrapolation des composantes du changement démographique est affectée et qui s'accroît lorsqu'on s'éloigne de la date de début de la période de projection. Les erreurs ont tendance à croître avec la longueur de l'horizon temporel de la projection. Presque tous les auteurs de prévisions démographiques officielles construisent plusieurs séries basées sur la combinaison de scénarios. Ils ne fournissent toutefois pas la probabilité avec laquelle la réalité se situera dans l'intervalle délimité par les variantes haute et basse. Cet article dresse une liste de quelques mesures utilisées pour évaluer la précision des perspectives et précise leurs limites et avantages respectifs. Il est par ailleurs proposé d'examiner les erreurs ex-post par une méthode qui s'apparente à la décomposition des effets âge-période-cohorte ou d'ajuster un modèle de série chronologique aux erreurs de projection observées sur certaines variables dans les projections antérieures afin d'établir un intervalle de confiance pour ces variables à appliquer à la nouvelle projection. Il est enfin proposé de raisonner sur la base de modèles stochastiques ou d'utiliser une méthode de projections probabilistes de population basée sur des opinions d'experts à propos des tendances futures de la fécondité, de la mortalité et de la mobilité spatiale ainsi qu'à propos du niveau d'incertitude dont ces tendances peuvent être entachées.

Processing uncertainty in demographic projections.

Mots clé: Incertitude, projection démographique.

*J. Duchêne. Institut de Démographie. Université Catholique de Louvain. Place Montesquieu 1 Boîte 17. B-1348 Louvain-La-Neuve (Belgique). E-mail: duchene@demo.ucl.ac.be.

Texte révisé et mis à jour d'une communication préparée pour les Jornades Tècniques sobre Projeccions Demogràfiques de Catalunya, Barcelona, Palau de la Generalitat 26 i 27 de maig del 1997.

— Reçu en septembre de 1998.

— Accepté en janvier de 1999

*«Le futur est-il donné
ou est-il en perpétuelle construction?»*

Ilya Prigogine

1. QUELQUES REFLEXIONS INTRODUCTIVES SUR L'INCERTITUDE

La réponse à la question posée par Prigogine (1996) dans son ouvrage intitulé *«La Fin des Certitudes»*, et reprise en exergue de ce texte, devrait permettre de préciser le type d'incertitude dont il faudrait tenir compte dans les projections démographiques. Si le futur est «donné», il peut en effet être prédit avec un degré relativement élevé d'exactitude et le degré d'incertitude sera plutôt faible. Si, par contre, le futur est «en perpétuelle construction», il ne peut pas être prédit et le degré d'incertitude sera presque maximum. Les tendances prédites sont souvent des extrapolations des tendances passées observées jusqu'à la date d'établissement des projections alors que la réalité peut présenter des retournements de tendances ou des irrégularités inattendues. Les spécialistes des projections voudraient croire que le futur est, dans une certaine mesure, «donné» mais ne serait-il pas plutôt «en perpétuelle construction»? Incertain n'est pas synonyme d'incorrect mais l'incertitude doit être comprise dans le sens d'incapacité à être prédit avec exactitude. Le futur démographique est par nature incertain et plus on s'écarte de la date de début de la période de projection, plus l'incertitude augmente. Par conséquent, on peut poser autant d'hypothèses que l'on veut, et les combiner pour prévoir le devenir des populations. Cependant, il n'y a aucune projection, et donc, entre autres, aucun intervalle de confiance, dont on peut prétendre a priori qu'ils soient «corrects». La réflexion faite par Hajnal en 1955 reste vraie près de 45 ans plus tard: *«Prophecy about the future of human societies is an uncertain business; there is no reason to expect more success in this endeavor than in forecasting other features of historical development»*.

2. SOURCES D'INCERTITUDE DANS LES PROJECTIONS

Même si les probabilités d'occurrence des événements démographiques sont connues avec une relative certitude au niveau de la population sur un horizon temporel assez court, les projections pour de petites populations (sous-populations) peuvent être affectées d'une marge d'incertitude liée au fait que les probabilités d'occurrence des événements sont entachées d'aléa (Sykes, 1969; Pollard, 1973). L'incertitude étant comprise dans le sens d'incapacité à prévoir l'avenir avec exactitude, les erreurs de prévision sont donc la manifestation de cette incertitude et proviennent de l'imprécision, sur l'extrapolation des composantes du changement démographique, qui s'accroît lorsqu'on s'éloigne de la date de début de la période de projection (Long, 1992). Les erreurs ont tendance à croître avec la longueur de l'horizon temporel de la projection (Ascher,

1978; Stoto, 1983; Armstrong, 1985 et Smith, 1987). Keyfitz (1981) écrit qu'au-delà d'un quart de siècle, on est incapable de dire ce que la population sera. En ce qui concerne la structure de la population, les effectifs des cohortes des personnes déjà nées au point de départ de la projection ne seront évidemment affectés que par les éventuels biais d'observation et par la difficulté à prédire la mortalité et la mobilité spatiale de cette population, alors que, la population qui va naître pendant la période de projection sera affectée par, entre autres, les hypothèses émises à propos de la fécondité qui pourraient être les plus influentes sur l'incertitude dans les projections. Alho (1992) écrit en effet que l'incertitude due aux hypothèses relatives à la migration et à la mortalité des populations déjà nées au début de la projection peut être ignorée dans la propagation des erreurs quoiqu'elle ne soit pas négligeable. Contrairement à ce qu'écrit Alho (1992), plusieurs auteurs s'accordent pour écrire que l'un des phénomènes démographiques les plus difficiles à prévoir est probablement la migration, plus sensible que les autres phénomènes aux conditions économiques ou aux mesures politiques prises en matière d'accueil des étrangers.

La qualité des projections démographiques dépend à la fois de la qualité de l'estimation de l'effectif et de la structure de la population au point de départ de ces projections et des hypothèses concernant l'incidence des événements démographiques modifiant cet effectif et cette structure (naissances, décès, migrations). Des erreurs dans les données de base (structure de la population, niveau et/ou tendance des probabilités d'occurrence des événements démographiques) peuvent avoir des répercussions plus ou moins importantes sur les projections démographiques si celles-ci sont calculées par extrapolation des tendances observées. L'incertitude sur le futur est une autre source importante de marge d'erreur dans les projections. Ainsi, Inoue et Yu (1979) ont montré que les projections des Nations Unies établies jusqu'en 1968 étaient principalement influencées par les données initiales sur les niveaux et les taux de croissance des populations. A partir de 1968, la principale source d'erreur était liée aux hypothèses de projection. Keyfitz (1977) distingue cinq types d'incertitudes dans les perspectives: des variations aléatoires autour de chaque paramètre, l'incertitude dans l'estimation des paramètres, la variabilité des paramètres dans le temps, l'utilisation d'un modèle incorrect de projection,¹ des discontinuités entre le passé et le présent.

3. METHODES *EX-POST* POUR QUANTIFIER L'INCERTITUDE

Presque tous les auteurs de prévisions démographiques officielles construisent plusieurs séries. Ces séries sont basées sur la combinaison de scénarios (le plus souvent au

¹ Il nous semble que cette source d'erreur ne doit pas être mise au même niveau que les autres car elle relève du savoir-faire du chercheur et non d'éléments qu'il ne peut pas maîtriser.

nombre de trois: haut, moyen et bas) de fécondité, et parfois de mortalité et de migration qui génèrent des projections sous l'hypothèse que les taux (ou probabilités) suivent exactement les trajectoires hypothétiques.² Certains scénarios, purement hypothétiques, peuvent être proposés à titre de simulation en vue de montrer la sensibilité des projections démographiques à des hypothèses spécifiques ou extrêmes. Si on prend en considération des scénarios qui définissent un couloir de possibles trop large, cette approche apporte une information très floue et peu intéressante. Les utilisateurs de ces projections (et en particulier les planificateurs) sont alors enclins à ne retenir qu'un seul scénario. Ils demandent à ce qu'on leur indique le scénario le plus vraisemblable et peuvent difficilement faire intervenir l'incertitude dans les mesures législatives ou sociales qu'ils prennent. La baisse de la fécondité a été la cause la plus importante de diminution de la proportion de jeunes dans les populations des pays en transition (Coale, 1972), toutefois, étant donné les hauts niveaux d'espérance de vie qui sont actuellement atteints dans un certain nombre de pays industrialisés, la mortalité devient un facteur important de vieillissement de la structure par âge de la population (Duchêne et Wunsch, 1990 et 1991). Les migrations internationales peuvent être fortement affectées par la législation internationale ou, plus encore, par la législation du pays d'accueil.

Les auteurs de prévisions démographiques qui ont recours à des scénarios ne fournissent pas la probabilité avec laquelle la réalité se situera dans l'intervalle délimité par les variantes haute et basse. L'utilisateur de prévisions démographiques n'est pas non plus en mesure de déterminer quelle est la probabilité qu'un scénario se réalise. Cette probabilité dépend à la fois du temps écoulé depuis le point de départ des scénarios et de la robustesse des différentes hypothèses. Or, dans de nombreux domaines de la vie économique, il convient de raisonner en terme de courbe de probabilité ou d'espérance mathématique plutôt qu'à partir d'une seule valeur (un scénario) choisie arbitrairement. Mc Nown, Rogers et Knudsen (1993) proposent de sélectionner des scénarios de projection sur la base de modèles stochastiques mais les degrés d'incertitude associés aux différents éléments de la projection sont les mêmes que pour le modèle des trois scénarios.

Ahlburg (1995) dresse la liste des principales mesures utilisées pour évaluer la précision des perspectives et de leurs limites et avantages respectifs. On peut retenir, entre autres,

- la moyenne des erreurs absolues (MAE en anglais) recommandée si la fonction de coût des erreurs est linéaire en termes d'erreurs absolues;

²Dans les perspectives établies pour la Catalogne sur la période 1995-2030, il est fait référence à trois hypothèses de mortalité et à trois hypothèses de mobilité spatiale ainsi qu'à quatre hypothèses de fécondité (dans ce cas, aux hypothèses basse, moyenne et haute, a été ajoutée une hypothèse de remplacement). La combinaison de toutes ces hypothèses a donné lieu à dix scénarios de projection. (Institut d'Estadística de Catalunya, 1997, p.9-12)

- la moyenne des erreurs pourcentuelles absolues (mean absolute percent error ou MA-PE en anglais) recommandée si la fonction de coût des erreurs est linéaire en termes de pourcentages³;
- la moyenne des carrés des erreurs (MSE en anglais) ou la racine carrée de la moyenne des carrés des erreurs (RMSE en anglais) sont adéquates si les erreurs importantes doivent avoir plus de poids que les erreurs faibles dans la mesure. Quoique ces mesures soient les plus fréquemment utilisées, elles ne permettent pas des comparaisons entre modèles et entre horizons temporels différents;
- la racine carrée de la moyenne des carrés des erreurs pourcentuelles (RMSPE en anglais).

En vue d'introduire des intervalles de confiance dans l'établissement des projections des Nations Unies, Keyfitz (1981) et Stoto (1983) ont procédé pour la première fois à l'analyse *ex-post* des erreurs dans les projections réalisées par les Nations Unies pour les différents pays. Cette approche ne fournit des intervalles que pour l'effectif total de la population: l'écart-type des erreurs sur le taux d'accroissement varie entre 0,3% par an pour les pays développés et 0,5% par an pour les pays en développement. Une autre conclusion des études de Keyfitz (1981) et de Stoto (1983) est que la population croît selon des taux qui sont, dans un cas sur trois, situés à l'extérieur de l'intervalle délimité par les variantes haute et basse. Lee (1990) a complété l'analyse de Keyfitz et de Stoto en étudiant les six projections sur des horizons temporels de cinquante à plusieurs centaines d'années rapportées par Frejka (1981): il obtient une erreur moyenne sur le taux d'accroissement de l'ordre de 1% par an.

Sur la base de projections mises à jour régulièrement en Norvège et aux Pays-Bas, et à partir de nouvelles données, Keilman (1995) examine les erreurs *ex-post* dans les prévisions de la fécondité de la mortalité à l'aide d'un modèle qui distingue les effets de période (année de calendrier), les effets de durée (horizon temporel de la projection) et l'effet (de la date de départ) de la projection. Les détails de la modélisation et de l'estimation ont été présentés dans Keilman (1990). Ce modèle s'apparente techniquement au modèle âge —période— cohorte (modèle APC).

- L'effet de cohorte du modèle APC est interprété par Keilman comme la contribution à l'erreur de projection de la projection particulière dont les données initiales correspondent à une année (ou une période) précise. Cet effet de la projection est considéré comme un indicateur qui exprime, à travers un seul nombre, les circonstances

³Cet indicateur a été proposé à côté de la moyenne des erreurs pourcentuelles algébriques (mean algebraic percent error ou MALPE en Anglais) par Smith (1987), la National Academy of Sciences (1980) et Isserman (1977). Le premier indicateur (MAPE) calcule la moyenne des erreurs sans tenir compte de leur signe (+ ou -) tandis que le deuxième indice (MALPE) prend ce signe en compte.

particulières dans lesquelles la projection a été élaborée (composition de l'équipe, méthodologie, ...).

- L'effet d'âge du modèle APC, que l'on pourrait plutôt appeler un effet de durée, reflète la contribution, à l'erreur, de l'horizon temporel sur lequel la projection a été effectuée: il traduit le fait que l'incertitude augmente lorsque l'horizon temporel de la projection s'accroît.
- L'effet de période de ce modèle recouvre l'effet de période des modèles APC. Il quantifie un ensemble de facteurs qui résument les conditions socio-démographiques observées au cours d'une période de temps, et leur impact sur les erreurs de projection indépendamment des circonstances qui ont entouré la production des projections (résumées dans l'effet de la projection) ou de la longueur de l'horizon temporel (saisie par l'effet de durée).

Pour la projection basée sur les données observées au cours de l'année i , telle qu'observée au cours de l'année (ou de la période) j , avec un horizon de projection de k années, la valeur absolue de l'erreur de prévision d'une variable, définie comme la valeur absolue de la différence entre la valeur prédite et la valeur observée,⁴ s'écrit comme le produit des exponentielles d'un effet de projection $F(i)$, d'un effet de période $P(j)$, d'un effet de durée $D(k)$ et d'un résidu $u(i, j, k)$. En vue de rendre les effets plus faciles à interpréter, Keilman (1995) a donc donné, au modèle de désagrégation des effets, l'expression suivante:

$$|E(i, j, k)| = \exp \{F(i) + P(j) + D(k) + u(i, j, k)\}$$

Si plusieurs scénarios sont proposés à propos d'un phénomène pour une même projection, Keilman suggère de réécrire le modèle de désagrégation des effets, en y introduisant un terme supplémentaire $V(m)$ qui peut s'interpréter comme l'effet de la variante m :

$$|E(i, j, k, m)| = \exp \{F(i) + P(j) + D(k) + V(m) + u(i, j, k, m)\}$$

Les effets de durée sont traduits par Keilman au moyen de l'expression où D est un coefficient estimé à partir des données. Keilman (1995) a appliqué cette méthodologie aux naissances et décès pour la Norvège et les Pays-Bas. Il souligne quelques imperfections de la méthode:

- la fraction de variance exprimée par le modèle est moyenne pour les naissances en Norvège (0,61) et légèrement plus élevée pour les Pays-Bas (0,74); cette fraction s'élève à 0,71 et 0,73 pour les décès respectivement aux Pays-Bas et en Norvège;

⁴Ce que l'on estime ainsi, c'est la difficulté qu'a eu le spécialiste des projections à anticiper d'éventuels retournements ou changements brusques de tendances.

- il y a une forte autocorrélation entre les erreurs qui affectent tant les naissances que les décès;
- pour la Norvège, les erreurs sur les décès sont plus irrégulières que celles sur les naissances.

Malgré ces limites soulignées par l'auteur, cette méthode paraît assez facile à appliquer et, au vu des commentaires de Keilman (1995), elle semble très informative.

Une autre manière de procéder est d'ajuster un modèle de série chronologique aux erreurs de projection observées sur certaines variables dans les projections antérieures et d'établir un intervalle de confiance pour ces variables à appliquer à la nouvelle projection. Par exemple, De Beer (1992) a montré qu'un modèle auto-régressif d'ordre un (AR(1)) conviendrait pour ajuster les erreurs sur le taux de croissance de l'effectif total de la population. Plusieurs auteurs (parmi lesquels Alho (1990), Carter et Lee (1986), Lee et Tuljapurkar (1994), Mc Nown, Rogers et Little (1995) et Pfaumer (1992)) ont développé et appliqué des méthodes, basées sur les séries temporelles, permettant de projeter des distributions de population au niveau national. Les méthodes d'analyse des séries chronologiques produisent, non seulement une projection centrale, mais aussi une distribution de probabilités permettant de localiser le niveau des taux de fécondité, de mortalité, et de mobilité spatiale à tout moment dans le futur sous l'hypothèse d'une distribution normale des erreurs. Quel est l'intérêt d'une projection de la population des Etats-Unis où la fécondité est maintenue, en moyenne, constante au niveau de 2,1 enfants par femme pendant 75 ans (de 1990 à 2065) mais où l'intervalle de confiance à 95% de l'indice synthétique de fécondité contient toutes les valeurs comprises entre 0,7 enfant par femme et plus de 3 enfants par femme (Lee et Tuljapurkar, 1994)? Il est en effet évident qu'un tel intervalle n'a aucune valeur informative car l'indice synthétique de fécondité des Etats-Unis se situera avec certitude entre 0,7 enfant par femme (valeur la plus basse observée dans le monde en Allemagne de l'Est en 1994 selon Conseil de l'Europe (1996)) et 3 enfants par femme (valeur observée en 1990 en Turquie et plus jamais observée depuis dans les pays du conseil de l'Europe selon Conseil de l'Europe (1996)).

Pfaumer (1988) a procédé à une modélisation des taux d'occurrence des événements au moyen de processus stochastiques et a estimé les paramètres du modèle en vue d'établir des simulations stochastiques pour les Etats-Unis. Il a montré qu'au bout d'un siècle, l'erreur type de la projection s'élève à 8% de la valeur attendue. Selon Lee et Tuljapurkar (1994), la distribution utilisée par Pfaumer pour simuler la fécondité est trop optimiste; de plus, et toujours selon Lee et Tuljapurkar (1994), Pfaumer n'introduit pas d'autocorrélation au niveau de la fécondité ce qui conduit à une importante sous-estimation de l'incertitude dans la projection de population.

Lutz, Sanderson et Scherbov (1996) exposent une méthode de projections probabilistes de population basée sur des opinions d'experts à propos des tendances futures

de la fécondité, de la mortalité et de la mobilité spatiale ainsi qu'à propos du niveau d'incertitude dont ces tendances peuvent être entachées. Cette approche rejoint les projections selon la méthode Delphi, appliquées par exemple aux scénarios de la demande d'énergie ou du développement technologique futur (cf. Minrowicz, Chapuy et Louineau, 1990 ainsi que Héraud, Munier et Nanopoulos, 1997). Au vu des résultats fournis par les auteurs, et, entre autres, de l'amplitude de variation des résultats⁵ obtenus à partir de 1000 simulations, on peut se poser quelques questions sur l'intérêt de telles méthodes. Peut-être visent-elles à rendre les spécialistes des projections encore un peu moins certains de ce qu'ils avancent comme résultats mais la subjectivité, dont ils sont bien conscients, nécessaire pour proposer les hypothèses sur lesquelles ils ont établi leurs projections, n'y suffit-elle pas?

4. EN GUISE DE CONCLUSION

En l'absence d'informations sur les tendances à prévoir, les valeurs observées dans le passé sont gardées constantes ou extrapolées. La réalité peut présenter des retournements de tendances ou des irrégularités inattendues. De nouvelles tendances en termes de déclin rapide de la fécondité ou d'augmentation de l'espérance de vie après une période de stagnation n'ont été prises en compte par les auteurs que plusieurs années après que ces tendances se soient manifestées. Des modèles de projection basés sur l'extrapolation des tendances passées ne conviennent que dans le cas où ces tendances sont très régulières. Pour pouvoir prévoir des tendances moins lisses, il faudrait disposer de modèles de comportement d'un niveau explicatif suffisant⁶ ou plus exactement de schémas conceptuels moins entachés de faiblesses. Une autre solution serait de disposer de «*leading indicators*», permettant de prévoir l'évolution de certains comportements à partir d'indicateurs qui anticipent de quelques années cette évolution (Farley, 1988). Par exemple, dans les sociétés où le mariage reste un «pré-requis» à la fécondité, la nuptialité peut être un «*leading indicator*» de la première naissance, une brusque évolution du nombre des mariages s'accompagnant, un ou deux ans plus tard, d'une adaptation du nombre des naissances. On peut aussi se demander si l'évolution du nombre d'enfants désirés par des femmes âgées de 20 ans pourrait expliquer l'évolution de la fécondité des femmes, quelques années plus tard. Nous ne disposons pas de tels schémas qui nous permettraient de mieux comprendre les mécanismes à l'origine de l'évolution de la fécondité qui a une composante comportementale plus importante que la mortalité.

⁵En 2020, l'effectif total de la population varierait dans un intervalle de confiance à 95% qui s'étendrait de 116 à 133 millions en Europe de l'Est, de 209 à 240 millions dans la partie européenne de l'ex Union Soviétique, et de 446 à 512 millions en Europe de l'Ouest (Lutz, Sanderson et Scherbov, 1996).

⁶Et pourtant, selon le titre des entretiens de René Thom (1991) avec Émile Noël, Prédire n'est pas expliquer. Dans cet ouvrage, Thom montre qu'à côté de la science quantitative et prédictive, il existe une approche qualitative dont la valeur explicative est peut-être plus fine et plus décisive pour la connaissance.

Si une erreur importante d'appréciation de l'évolution a été commise dans des perspectives antérieures, la tendance sera de réagir en sens inverse de manière exagérée.

L'application des erreurs *ex-post* aux nouvelles projections repose sur l'hypothèse que les spécialistes des prévisions d'aujourd'hui commettront les mêmes erreurs ou négligeront les mêmes discontinuités que ceux qui ont calculé les projections antérieures. Il serait possible de réduire l'incertitude si on pouvait procéder à un contrôle régulier de l'adéquation entre le futur établi par les projections et la réalité.

5. BIBLIOGRAPHIE

- [1] **Ahlburg, D.A.** (1995). «Simple versus complex models: evaluation, accuracy and combining». *Mathematical Population Studies*, **5**(3), 281-290.
- [2] **Alho, J.M.** (1990). «Stochastic methods in population forecasting». *International Journal of Forecasting*, **6**, 521-530.
- [3] **Alho, J.M.** (1992). «The magnitude of error due to different vital processes in population forecasts». *International Journal of Forecasting*, **8**, 301-314.
- [4] **Armstrong, J.S.** (1985). *Long range forecasting: From crystal ball to computer*, 2nd edition, New York, John Wiley.
- [5] **Ascher, W.** (1978). *Forecasting: An appraisal for policy makers and planners*, Baltimore, The John Hopkins University Press.
- [6] **Carter, L.R.** et **R.D. Lee** (1986). «Joint forecasts of US marital fertility, nuptiality, births and marriages using time series models». *Journal of the American Statistical Association*, **81**(396), 902-911.
- [7] **Coale, A.J.** (1972). *The Growth and Structure of Human Populations. A Mathematical Investigation*, Princeton, New Jersey, Princeton University Press.
- [8] **Conseil de l'Europe** (1996). *Evolution démographique récente en Europe 1996*, Strasbourg, Editions du Conseil de l'Europe.
- [9] **De Beer, J.** (1992). «Uncertainty variants of population forecasts». *Statistical Journal of the United Nations ECE* **9**, 233-253.
- [10] **Duchêne, J.** et **G. Wunsch** (1990). «Les tables de mortalité limite: quand la biologie vient au secours du démographe», in *Population âgée et révolution grise - Chaire Quetelet '86*, M. Loriaux, D. Remy et E. Vilquin (sous la direction de), Bruxelles, Éditions CIACO ARTEL, p.321-332.
- [11] **Duchêne, J.** et **G. Wunsch** (1991). «Population aging and the limits to human life», in *Future Demographic Trends in Europe and North America: What can we assume today?* W. Lutz (ed.), London, Academic Press, p. 27-40.

- [12] **Farley, R.** (1988). «After the starting line: Blacks and women in an uphill race». *Demography*, **25**(4), 477-495.
- [13] **Frejka, T.** (1981). «World population projections: A concise history», in *Proceedings of the IUSSP International Population Conference, Manila*, **3**, 505-528.
- [14] **Hajnal, J.** (1955). «The prospects of population forecasts». *Journal of the American Statistical Association*, **50**, 309-322.
- [15] **Héraud, J.A., F. Munier et K. Nanopoulos** (1997). «Méthode Delphi: une étude de cas sur les technologies du futur». *Futuribles*, **218**, 33-53.
- [16] **Inoue, S. et Y.C. Yu** (1979). «United Nations new population projections analysis of *ex post* factor errors». Communication présentée au *Annual Meeting of the Population Association of America*.
- [17] **Institut d'Estadística de Catalunya** (1997). «Ponència tècnica inicial de l'Institut d'Estadística de Catalunya». *Jornades Tècniques sobre Projeccions Demogràfiques de Catalunya*, Barcelona, Palau de la Generalitat 26 i 27 de maig del 1997.
- [18] **Isserman, A.** (1977). «The accuracy of population projections for subcountry areas». *Journal of the American Institute for Planners*, **43**, 247-259.
- [19] **Keilman, N.** (1990). *Uncertainty in national population forecasting*, Amsterdam, Swets and Zeitlinger.
- [20] **Keilman, N.** (1995). «Accuracy and uncertainty in national population forecasts». Communication présentée à la *Chaire Quetelet 1995 sur «Le défi de l'incertitude. Nouvelles approches en perspective et prospective démographiques»*, Louvain-la-Neuve, Institut de Démographie, 26 p. + tableaux et figures.
- [21] **Keyfitz, N.** (1977). *Applied Mathematical Demography*, New York, Wiley.
- [22] **Keyfitz, N.** (1981). «The limits of population forecasting». *Population and Development Review*, **7**(4), 579-593.
- [23] **Lee, R.D.** (1990). «Long-run global population forecasts: A critical appraisal». *Population and Development Review*, **16**(supplement), 44-71.
- [24] **Lee, R.D., L. Carter et S. Tuljapurkar** (1995). «Disaggregation in population forecasting: Do we need it? and how to do it simply?». *Mathematical Population Studies*, **5**(3), 217-234.
- [25] **Lee, R.D. et S. Tuljapurkar** (1994). «Stochastic population forecasts for the United States: Beyond high, medium and low». *Journal of the American Statistical Association*, **89**(428), 1175-1189.
- [26] **Long, J.F.** (1992). «Accuracy, monitoring, and evaluation of national population projections», in *National Population Forecasting in Industrialized Countries*, Keilman Nico et Cruijsen Harri (eds), Amsterdam/Lisse, Swets & Zeitlinger, 129-145.
- [27] **Long, J.F.** (1995). «Complexity, accuracy, and utility of official population projections». *Mathematical Population Studies*, **5**(3), 203-216.

- [28] **Lutz, W., J. Goldstein et C. Prinz** (1996). «Alternative approaches to population projection», in *The future population of the world. What can we assume today? Revised and updated edition*, Lutz Wolfgang (ed.), Laxenburg (Austria), International Institute for Applied Systems Analysis, 14-44.
- [29] **Lutz, W., W.C. Sanderson et S. Scherbov** (1996). «Probabilistic population projections based on expert opinion», in *The future population of the world. What can we assume today? Revised and updated edition*, Lutz Wolfgang (ed.), Laxenburg (Austria), International Institute for Applied Systems Analysis, 397-428.
- [30] **MC Nown, R., A. Rogers et C. Knudsen** (1993). «Projections of fertility, mortality and the population in the United States; 1990-2050». *Technical Paper WP-93-1, Population Program*, University of Colorado, Boulder.
- [31] **MC Nown, R., A. Rogers et J. Little** (1995). «Simplicity and complexity in extrapolative population forecasting models». *Mathematical Population Studies*, **5(3)**, 235-257.
- [32] **Mirenowicz, P., P. Chapuy et Y. Louineau** (1990). «La méthode Delphi-Abaque. Un exemple d'application: la prospective du bruit». **143**, 49-63.
- [33] **Pfaumer, P.** (1988). «Confidence intervals for population projections based on Monte Carlo methods». *International Journal of Forecasting*, **4**, 135-142.
- [34] **Pfaumer, P.** (1992). «Forecasting US population totals with the Box-Jenkins approach». *International Journal of Forecasting*, **8(3)**, 329-338.
- [35] **Pollard, J.H.** (1973). *Mathematical Models for the Growth of Human Populations*, U.K.: Cambridge University Press.
- [36] **Prigogine, I.** (1996). *La Fin des Certitudes. Temps, Chaos et les Lois de la Nature*, Paris, Éditions Odile Jacobs.
- [37] **Rogers, A.** (1995). «Population forecasting: Do simple models outperform complex models». *Mathematical Population Studies*, **5(3)**, 187-202.
- [38] **Sanderson, W.C.** (1995). «Predictability, complexity and catastrophe in a collapsible model of population, development and environmental interactions». *Mathematical Population Studies*, **5(3)**, 259-279.
- [39] **Sardon, J.P.** (1996). «Les prévisions de fécondité». Communication présentée à la Chaire Quetelet 1995 sur «Le défi de l'incertitude. Nouvelles approches en perspective et prospective démographiques», Louvain-la-Neuve, Institut de Démographie, 16 p. + tableaux et figures.
- [40] **Smith, S.K.** (1987). «Tests of forecast accuracy and bias for country population projections». *Journal of the American Statistical Association*, **82(400)**, 991-1003.
- [41] **Stoto, M.A.** (1983). «The accuracy of population projections». *Journal of the American Statistical Association*, **78**, 13-20.
- [42] **Sykes, Z.M.** (1969). «Some stochastic versions of the matrix model for population dynamics». *Journal of the American Statistical Association*, **44**, 111-130.

- [43] **Tayman, J., E. Schafer et L. Carter** (1998). «The role of population size in the determination and prediction of population forecast errors: An evaluation using confidence intervals for subcounty areas». *Population Research and Policy Review*, **17**(1), 1-20.
- [44] **Thom, R.** (1991). *Prédire n'est pas expliquer*. Entretiens avec Émile Noël. Deuxième édition revue et corrigée, Paris, Flammarion.
- [45] **Wilmoth, J.R.** (1995). «Are mortality projections always more pessimistic when disaggregated by cause of death?», *Mathematical Population Studies*, **5**(4), 293-31.

ENGLISH SUMMARY

PROCESSING UNCERTAINTY IN DEMOGRAPHIC PROJECTIONS

JOSIANNE DUCHÊNE

Université Catholique de Louvain*

Often the predicted trends are extrapolations of trends observed before the projection's initial date, while reality may present an unexpected return of past trends or irregularities. Uncertainty being related to the impossibility to accurately predict the future, errors in the predictions are the manifestation of such uncertainty and come from the lack of accuracy which affects the extrapolation of the components of demographic change. The lack of accuracy increases as the date of the beginning of the projection gets more remote. Errors tend to increase with the length of the projection. Most of the authors of official demographic predictions construct several time series based on a combination of scenarios. However, they do not provide the probability that the real outcome will be inside the interval defined by the high and low scenarios. This paper gives an outline of the measures used to evaluate the accuracy of the perspectives and determine their advantages and limitations. It also proposes to examine the errors exposed by means of a method which breaks down the effects age-period-cohort, or to apply a chronological series model to the projection errors detected in variables from earlier projections so as to establish a confidence interval for these variables and use it in the new projection. Finally, it suggests to work with stochastic models or to use a method of population probabilistic projections based on experts' opinions regarding future trends in fertility, mortality and mobility as well as the level of uncertainty linked to these trends.

Keywords: Uncertainty, demographic projection.

*J. Duchêne. Institut de Démographie. Université Catholique de Louvain. Place Montesquieu 1 Boîte 17. B-1348 Louvain-La-Neuve (Belgique). E-mail: duchene@demo.ucl.ac.be.

Revised and updated from a paper prepared for the Jornades Tècniques sobre Projeccions Demogràfiques de Catalunya, Barcelona, Palau de la Generalitat, 26 i 27 de maig del 1997.

—Received September 1998.

—Accepted January 1999

*«Prophecy about the future of human societies
is an uncertain business;
there is no reason to expect more success
in this endeavor than in forecasting
other features of historical development»*

John Hajnal

Uncertain does not mean the same as incorrect. Moreover, uncertainty should be understood as the impossibility to make an accurate prediction. In demography, the future is essentially uncertain, and the farther from the projection's starting date the more uncertain (Long, 1992). We can, therefore, formulate as many hypotheses as we want and combine them so as to predict the future of populations. However, no projection or confidence interval is ever «correct» beforehand.

Even when we are relatively certain about the probability of certain demographic events regarding population in a short time span, projections of small populations (subpopulations) may be under a margin of uncertainty due to the fact that the probability of each event is subject to random (Sykes, 1969; Pollard, 1973). Regarding the structure of the population, the persons in the cohorts during the observation period previous to the projection's starting date will of course only be affected by an occasional bias in the observation and by the difficulty to predict the mortality and mobility of this population, while those who were born within the span of the projection will be affected by the hypotheses regarding fertility, which might have more influence on the projection's uncertainty than anything else.

Many authors agree that migration is probably one of the most difficult to predict demographic phenomena, since it is more sensitive to economic changes or political decisions regarding immigration than the rest of phenomena.

The quality of demographic projections depends both on the quality of the estimation of the total population and the structure of the population at the start of the projection and on the hypotheses regarding demographic events which affect both the total and the population structure (such as births, deaths and migrations). Errors in initial data (population structure, probability of demographic events) may affect demographic projections if they are calculated by extrapolating the observed trends. Uncertainty about the future is another important source of errors in projections.

Most of the authors of official demographic predictions construct several time series. These series are based in a combination of fertility and, sometimes, mortality and migration scenarios (usually three: high, medium and low), which generate the projections, based on the hypothesis that the rates (or probabilities) will follow exactly the expected trends. Some purely hypothetical scenarios may be proposed as simulations in order to show the sensitivity of demographic projections to very specific or extreme

hypotheses. If a wide range of scenarios is considered, the resulting information is very vague and of little interest. Users of these projections (especially planners) favour the use of a single scenario. They ask to be given the most likely scenario and rarely let uncertainty play a role in their social and legislative decisions.

Authors of demographic predictions who use scenarios do not provide the probability that the real outcome will fall inside the interval defined by the high scenario and the low scenario. A user of demographic predictions can not establish the probability that a given scenario is real either. This probability depends on both the time elapsed from the scenario's starting point and the soundness of the hypotheses. However, in many domains of economics it is advisable to work with probabilities or expected value curves rather than relying on a single randomly chosen value (or scenario). McNown, Rogers and Knudsen (1993) suggest that projection scenarios should be selected on a basis of stochastic models, but the degree of uncertainty for each element of the projection is the same as in the three-scenario model.

Ahlburg (1995) gives a list of the main measures used to evaluate the accuracy of the perspectives and their advantages and limitations. In order to introduce confidence intervals in the United Nations' projections, Keyfitz (1981) and Stoto (1983) have made, for the first time, an *ex post* analysis of the errors in the United Nations' projections for different countries. This approach only provides confidence intervals for the size of the population. Another conclusion of Keyfitz (1981) and Stoto's (1983) studies is that population growth rates are outside the interval defined between the high and low scenarios one out of every three times. Lee (1990) completed Keyfitz and Stoto's analysis by studying projections concerning from fifty to hundreds of years referred to by Frejka (1981), and obtained an average error for the growth rate of 1% each year.

Based on regularly updated projections from Norway and the Netherlands and on new data, Keilman (1995) examines the *ex post* errors in fertility and mortality predictions with the help of a model which differentiates between the period effect (calendar year), the duration effect (time range of the projection) and the effect of the projection's starting date. Modeling and estimating details were presented by Keilman (1990). Technically, this model is related to the age-period-cohort model (APC model).

Another possible procedure is to apply a chronological series model to the projection errors detected in variables from earlier projections and establishing a new confidence interval for these variables, which will be applied to the new projection. For example, De Beer proved that an autoregressive model of order one (AR(1)) would be very helpful in order to adjust errors in the total population growth rate. Several authors (like Alho (1990), Carter and Lee (1986), Lee and Tuljapulkar (1994), McNown, Rogers and Little (1995) and also Pfaumer (1992)) have developed and applied methods based on time series which allow them to project the distribution of the population at a national level. The chronological series analysis methods do not only provide a central projection but they provide also a distribution of probabilities which allows us to locate, at any

time in the future, the level of fertility, mortality and mobility rates, taking into account a normal distribution of errors. Is it interesting to have a projection of population where fertility has stuck to an average of 2.1 children per woman for 75 years but where the 95% confidence interval for fertility rate contains all the values between 0.7 and more than 3 children per woman (Lee and Tuljapurkar, 1994)?!

Pfaumer (1988) came to model the rates by means of stochastic processes and estimated the model's parameters so as to produce stochastic simulations for the United States. He showed that, at the end of the century, the projection's typical error is 8% above the expected value. According to Lee and Tuljapurkar (1994), the distribution that Pfaumer used to simulate the fertility is too optimistic; moreover, and still according to Lee and Tuljapurkar (1994), Pfaumer did not introduce any fertility autocorrelation, which leads to an important underestimation of the uncertainty in the projection of population.

Lutz, Sanderson and Scherbov (1996) put forward a method of population probabilistic projections based on experts' opinions about future trends in fertility, mortality and mobility and their level of uncertainty. This approach puts the projections together under the Delphi method (Minrowicz, Chapuy and Louineau, 1990, and Héraud, Munier and Nanopoulos, 1997). After inspecting the results provided by the authors and the wide range of variation of the results obtained after 1000 simulations, one can question the value of such methods. Perhaps their contribution is to help the experts on projections realize about the uncertainty of the results. However, is not the awareness about the subjectivity inherent to the hypotheses on which the projections are based enough?

MÈTODES I MODELS PER PROJECTAR ELS COMPONENTS DEL CANVI DEMOGRÀFIC EN LES PROJECCIONS DE POBLACIÓ DE L'INSTITUT D'ESTADÍSTICA DE CATALUNYA

M. FARRÉ MALLOFRÉ

J.A. SÁNCHEZ CEPEDA

Institut d'Estadística de Catalunya*

Aquest treball descriu la realització de les projeccions de població 1996-2010-2030 de l'Institut d'Estadística de Catalunya (Idescat) des del punt de vista metodològic. Es descriuen les metodologies particulars desenvolupades per a cada un dels components del creixement demogràfic. Aquestes metodologies no només se centren en els valors dels indicadors demogràfics a projectar sinó que introdueixen comportaments específics i dinàmics per als diferents grups d'edat de la població. En la fecunditat s'ha aplicat un mètode que integra la descendència final de les diferents generacions de dones, l'edat a la maternitat i la descendència segons ordre de naixement. En la mortalitat s'ha aplicat un mètode que permet combinar l'evolució de l'esperança de vida (traç cronològic) amb evolucions no homogènies de les taxes de mortalitat a diferent edats (coeficients de millora). En la migració s'ha treballat amb el model migratori de Rogers, aplicant uns perfils dinàmics en el temps. Finalment, es fa una anàlisi de la incertesa inherent als resultats i la seva quantificació.

Methods and models to project the components of demographic change in the population projections of the Institut d'Estadística de Catalunya.

Paraules clau: Projeccions de població, fecunditat, mortalitat, migració, descendència final, traç cronològic, perfil migratori.

* Institut d'Estadística de Catalunya. Via Laietana, 58. 08003 Barcelona (Espanya).

— Rebut el juliol de 1998.

— Acceptat l'octubre de 1998.

1. INTRODUCCIÓ

Les projeccions de població de Catalunya 1996-2010-2030 portades a terme per l'Institut d'Estadística de Catalunya (Idescat) són el resultat d'un procés de reflexió interdisciplinària que va tenir com a objectiu incloure, en la seva elaboració, les opinions de les persones expertes i tenir en compte les necessitats dels usuaris de les projeccions de població. L'estudi de la població futura catalana abasta un període de 34 anys, del 1996 al 2030, obtenint-se les poblacions per sexe i edat simple any a any. La metodologia de projecció emprada ha estat el mètode dels components, consistent en aplicar a la darre-ra piràmide coneguda la projecció separada de la natalitat, la mortalitat i la migració. S'han elaborat 3 hipòtesis sobre l'evolució futura de les migracions i de la fecunditat i 2 hipòtesis d'evolució de la mortalitat. La combinació selectiva d'aquestes hipòtesis ha donat lloc a 4 escenaris d'evolució de la població fins el 2030, que es diferencien no només per les xifres absolutes de població sinó també per les estructures demogràfiques que en resulten. Els escenaris resultants s'anomenen: baix o vell, central, tendencial i alt o jove.

Aquest article explica la metodologia emprada per realitzar la projecció de la població, per centrar-se posteriorment en cada una de les metodologies específiques desenvolupades per a la projecció dels 3 components demogràfics: fecunditat, mortalitat i migració.

2. EL MÈTODE DELS COMPONENTS

Totes les oficines estadístiques dels països desenvolupats que elaboren projeccions de població per sexe i edat apliquen el mètode dels components per cohorts o generacions. Aquest mètode permet estimar la població futura i la seva distribució per sexe i edat a partir de la piràmide actual de població i de determinats supòsits sobre l'evolució de la mortalitat, la fecunditat i les migracions. Per a la projecció s'ha emprat LIPRO, una aplicació informàtica per a la projecció i anàlisi de models demogràfics multidimensionals.

En els models demogràfics, la població queda classificada en estats o posicions, d'acord amb el seu sexe i edat. La relació entre el nombre d'individus en un estat i el nombre de successos experimentats per ells es pot descriure per l'equació de Markov: la probabilitat que un individu a la posició i salti a la posició j en un interval infinitesimal de temps, és igual a un nombre constant de vegades la longitud de l'interval:

$$\lim_{dt \rightarrow 0} Pr(I(t+dt) = j | I(t) = i) / dt = m_{ij}(t)$$

Aquí, $m_{ij}(t)$ és una constant dependent del temps i de les posicions anterior i posterior al succés (i, j respectivament). Aquesta constant s'anomena intensitat instantània o taxa del salt d' i a j . Quan $m_{ij}(t) = m_{ij}$ al llarg de tot l'interval d'observació, el mo-

del resultant s'anomena model exponencial o model d'intensitats constants. El model exponencial genera expressions complexes que requereixen tècniques d'avaluació iterativa. Per simplicitat computacional, sovint se suposa que els esdeveniments es distribueixen uniformement al llarg del període de projecció. Aquesta suposició s'anomena hipòtesi d'integració lineal i permet fer els càlculs en un sol pas. El model corresponent s'anomena model lineal.

En les projeccions de població de Catalunya 1996-2030 s'ha emprat el model exponencial. Siguin,

I = matriu identitat

i = vector fila unitat

e = edat

s = sexe

t = temps a l'inici de la projecció

h = longitud del període de projecció

$P(s, e, t)$ = població inicial

$M_i(s, e, t)$ = matriu de successos interns. Té 0 a la diagonal.

$M_e(s, e, t)$ = matriu de sortides (emigració i defuncions)

$M_b(s, e, t)$ = matriu de naixements segons la posició de la mare

$M(s, e, t) = M_i(s, e, t) - \text{diag}(M_i(s, e, t) \cdot i^T + M_e(s, e, t) \cdot i^T)$

$O(s, e, t; h)$ = matriu d'entrades exògenes (immigrants). L'edat és en el moment inicial t , no en el moment d'entrada.

La formulació del model exponencial és

$$P(s, e + \tau, t + \tau) = P(s, e, t) \cdot e^{M(s, e, t) \cdot \tau} + (1/h) \cdot i \cdot O(s, e, t; h) \cdot \\ \cdot M^{-1}(s, e, t) \cdot \{e^{M(s, e, t) \cdot \tau} - I\}$$

El vector de persones-any viscudes pels individus és

$$L(s, e, t; h) = \int_0^h P(s, e + \tau, t + \tau) d\tau$$

D'aquí es pot obtenir el nombre d'esdeveniments com

$N_i(s, e, t; h) = \text{Diag}(P(s, e, t; h)) \cdot M_i(s, e, t)$ successos interns (migració interna)

$N_e(s, e, t; h) = \text{Diag}(P(s, e, t; h)) \cdot M_e(s, e, t)$ sortides (emigració i defuncions)

$N_b(s, e, t; h) = \text{Diag}(P(1, e, t; h)) \cdot M_b(s, e, t)$ naixements

Notar que els naixements es calculen a partir de la població de dones només.

Per a l'últim grup d'edat (que és un grup obert), la població al final de període s'obté combinant els supervivents dels 2 últims grups d'edat.

En les projeccions per sexe i edat $M_i(s, e, t) = 0$; en una projecció de la població per comarques, actualment en curs, s'emptra $M_i(s, e, t)$ per a les migracions internes.

3. LA PROJECCIÓ DE LA FECUNDITAT

La fecunditat és, potser, el component que més variabilitat pot introduir en l'evolució de la població. En primer lloc, els seus efectes es concentren en la base de la piràmide, que en el futur constituirà el cor del recanvi generacional; una concentració de naixements (com per exemple el baby-boom dels anys 70) repercuteix en unes generacions plenes que, a mesura que vagin transitant per la piràmide d'edats, anirant generant demandes addicionals de productes i serveis (places escolars, llocs de treball, habitatges, pensions, etc).

En segon lloc, el comportament de la fecunditat no és lineal, sinó més aviat cíclic, amb períodes on hi ha una anticipació de naixements i períodes que coneixen, en canvi, un retard de l'edat de les mares en tenir fills, de manera que després del boom de naixements per anticipació, la baixada subsequent és molt més accentuada si hi ha un retard important del calendari: les dones grans ja n'han tingut i les dones joves s'esperen a ser més grans. D'altra banda, l'edat de constituir la descendència influeix clarament en el ritme de creixement d'una població: si una dona decideix tenir 2 fills, no és el mateix que els tingui als 20 anys que als 40. En el primer cas, pot haver-hi 3 generacions en 40 anys, mentre que en el segon cas només n'hi haurà 2. Els cicles d'augment i disminució dels naixements no només estan determinats pels canvis en el calendari, sinó que també reflecteixen la variació en la descendència que acaben assolint les diferents generacions. Una marcada preferència per famílies de dos fills i una important reducció dels fills d'ordre superior, ha suposat una disminució sensible de la descendència en les generacions recents. L'evolució de la natalitat a Catalunya els darrers vint anys sembla respondre a un model d'anticipació-retard del calendari, acompanyat d'una baixada de la descendència de les dones nascudes a finals dels anys cinquanta i començament dels anys seixanta, respecte les generacions predecessores nascudes els anys quaranta.

El mètode utilitzat projecta la fecunditat segons la generació, l'edat i l'ordre de naixement, i obté les taxes de fecunditat any a any i els indicadors generacionals i transversals associats, com ara l'indicador conjuntural de fecunditat i l'edat mitjana de les mares. Les variables utilitzades en la projecció —generació, edat de la mare i ordre de naixement dels fills— són les rellevants en l'anàlisi de la fecunditat d'un any determinat. En integrar la generació (o cohort) i l'ordre dels naixements, el model està millor especificat, i permet establir supòsits explícits sobre la descendència de les generacions per ordre de naixement.

El mètode FLEM (Fertility Longitudinal Extrapolation Method) correspon a les metodologies edat-cohort-període-ordre i és de caràcter bàsicament extrapolatiu: els desenvolupaments de la fecunditat en cohorts passades s'utilitzen per projectar el comportament de les cohorts recents i futures. Atès que la fecunditat de les generacions és més estable que la fecunditat transversal, és un avantatge incorporar l'enfocament per generacions a l'hora de projectar. En la implementació portada a terme per a les projeccions de Catalunya 1996-2030 s'ha optat per una línia extrapolativa-projectiva: en algunes edats i ordres de naixement s'extrapolen les tendències enregistrades, però en altres es forcen les dades (mitjançant els coeficients regressors) a assolir determinats valors, amb l'objectiu d'ampliar el ventall d'hipòtesis.

La informació s'estructura en les matrius F , de fecunditats, i FA , de fecunditats acumulades, per cohort, edat i ordre de naixement. Siguin

r = ordre de naixement. Valors: 1, 2, 3 i 4 i més.

e = edat. Valors: 15 a 50.

e_0 = última edat amb dada observada en cada cohort.

c = cohort. Valors: 1940 a 1990.

c_0 = última cohort completa (1945)

cf = última cohort a projectar (1980 o 1990 segons la hipòtesi)

cs = cohort estàndard.

y = any del calendari. Valors: 1997 a 2030.

$F(c, e, r)$ = taxa de fecunditat d'ordre « r » a l'edat « e » per la cohort « c »,

$FA(c, e, r)$ = taxa de fecunditat d'ordre « r » acumulada a l'edat « e » per la cohort « c », amb

$$FA(c, e, r) = \sum_{k=15 \dots e} F(c, k, r)$$

$$F(c, e, r) = FA(c, e, r) - FA(c, e - 1, r)$$

D'aquestes 2 matrius tridimensionals, F i FA , es dedueix la fecunditat per a les 4 variables d'estudi: ordre, edat, generació i moment. Efectivament, per a una cohort c i una edat e , el vector $FA(c, e, r)$ indica les fecunditats dels diferents ordres. Per a una cohort c i un ordre r , el vector $FA(c, e, r)$ indica les fecunditats a les diferents edats. Per a una edat e i un ordre r , el vector $FA(c, e, r)$ indica les fecunditats de les diferents generacions. Per últim, la lectura en diagonal de la matriu desagregada, $F(c = y - k, e = k, r)$, $k = 15 \dots 50$ proporciona la informació de la fecunditat d'ordre r per a les diferents generacions en un any donat.

Els indicadors que permeten resumir la informació sobre la fecunditat del moment, la generació i l'edat són l'indicador conjuntural de fecunditat (ICF), la descendència final de les generacions (DF) i l'edat mitjana de la mare al naixement (EM), respecti-

vament. La seva definició a partir de la matriu de fecunditats és

$$ICF(y) = \sum_e \sum_r F(c = y - e, e, r)$$

$$DF(c) = \sum_r FA(c, 50, r)$$

$$EM(y) = \left\{ \sum_e \sum_r e \cdot F(c = y - e, e, r) \right\} / \left\{ \sum_e \sum_r F(c = y - e, e, r) \right\}^1$$

El mètode consisteix a projectar la fecunditat acumulada d'una cohort, emprant els valors de la fecunditat de la mateixa cohort en edats anteriors i els de cohorts anteriors a la mateixa edat. Per exemple, de la cohort de 1965 es coneix la seva fecunditat fins els 30 anys. La fecunditat acumulada per aquesta cohort als 35 anys es calcula a partir de l'acumulada fins els 30 anys i de la fecunditat de les cohorts immediatament anteriors, 1960 a 1964, entre els 30 i 35 anys.

En una primera etapa, per a totes les cohorts a projectar, s'extrapola/projecta la fecunditat acumulada cada 5 anys (als 20, 25, 30, 35, 40, 45 i 50 anys) basant-se en la relació observada entre la fecunditat acumulada a les edats e i $e + 5$. Per fer aquesta extrapolar/projecció FLEM utilitza regressions entre les fecunditats acumulades a l'edat e i $e + 5$ en cohorts anteriors:

$$FA(c, e + 5, r) = \alpha + \beta \cdot FA(c, e, r) + \epsilon$$

Aquestes regressions determinen si una cohort c acumula més, menys o igual fecunditat que les cohorts anteriors en les mateixes edats. D'aquesta manera s'obtenen les fecunditats acumulades cada 5 anys, encara que manquen les intermitges. En aquesta etapa intervenen clarament dos dels indicadors de la fecunditat: la descendència final DF (donat que és la fecunditat acumulada als 50 anys) i l'edat mitjana EM (que ve determinada pel ritme d'acumulació de la fecunditat). En la construcció de les tres hipòtesis de fecunditat un dels punts importants ha estat, dins la situació actual de descens de la fecunditat, determinar la generació que presenta el mínim de fecunditat i quin valor pren aquest mínim, així com fer efectiu un rejuveniment en la corba de fecunditat. En la determinació d'aquests valors ha estat peça important l'estudi per separat dels diferents ordres de fecunditat.

En una segona etapa, es completa la fecunditat de l'última cohort a estudiar, cf , amb un model de Gompertz; aquest mètode transforma la corba de fecunditat en una recta mitjançant la transformació logarítmica $G(p) = -\ln(-\ln(p))$, essent p una proporció.

¹En el càlcul de l'edat mitjana de fecunditat, en la fórmula es pren « e » com l'edat de la generació a final d'any, no en el moment del naixement.

S'aplica la funció G a la cohort incompleta, cf , i a una altra estàndard completa, cs , a les edats de 20, 25, 30, 35 i 40 anys. L'elecció de la cohort estàndard varia segons la hipòtesi, per recollir distribucions més joves o menys joves. La graficació d'una transformació Gompertz contra una altra produeix una línia aproximadament recta, fet que permet aplicar una regressió lineal per passar d'una a l'altra, i completar la cohort incompleta:

$$G\{FA(cf, e, r)/FA(cf, 50, r)\} = \alpha + \beta \cdot G\{FA(cs, e, r)/FA(cs, 50, r)\}$$

Un cop completada cf , es completa la resta de cohorts de la següent manera:

Per a cada generació c es calculen les diferències entre $FA(c, e_0, r)$ i $FA(c, 50, r)$. Es procedeix de la mateixa manera per $c - 1$ i cf . D'aquestes dues últimes cohorts es calcula, per a cada edat $e > e_0$, el seu pes en la diferència,

$$p(c - 1, e, r) = (FA(c - 1, e, r) - FA(c - 1, e_0, r)) / (FA(c - 1, 50, r) - FA(c - 1, e_0, r)).$$

$$p(cf, e, r) = (FA(cf, e, r) - FA(cf, e_0, r)) / (FA(cf, 50, r) - FA(cf, e_0, r)).$$

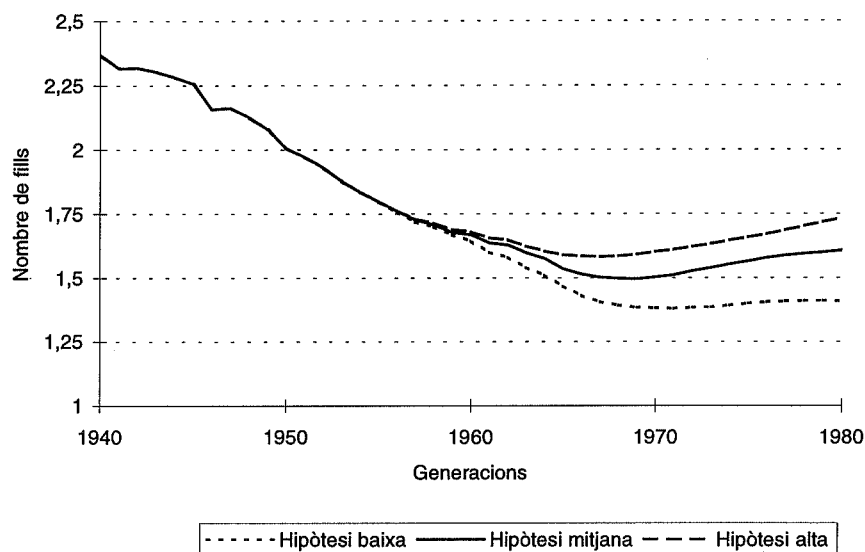
Seguidament es calcula $p(c, e, r)$ com una mitjana ponderada entre $p(c - 1, e, r)$ i $p(cf, e, r)$. Finalment,

$$FA(c, e, r) = FA(c, e_0, r) + p(c, e, r) \cdot (FA(c, 50, r) - FA(c, e_0, r)).$$

D'aquesta manera es fa tendir la fecunditat de les generacions cap a la projectada per a la generació cf . A partir d'aquesta generació se suposa una estabilitat en els nivells projectats.

La projecció de la fecunditat consta de 3 hipòtesis. La hipòtesi alta correspon a una recuperació en la descendència final de les dones, per situar-se en nivells d'1,8 fills per dona, amb una recuperació de la fecunditat en edats joves a mig termini. Predominarien les dones amb 2 fills, produint-se un increment dels fills d'ordre 3 i posterior. En la hipòtesi mitjana la descendència final descendiria fins nivells d'1,5 fills per dona per situar-se finalment en 1,64. El rejuveniment de la fecunditat també es produiria, però a menor nivell; predominarien les dones amb 2 fills. En la hipòtesi baixa, la descendència final cauria fins els 1,4 fills per dona, sense produir-se una recuperació posterior. L'edat mitjana de la fecunditat no descendiria i hi hauria un important augment de la infecunditat i dels fills únics.

És important constatar, però, que en totes 3 hipòtesis, encara que disminueixi la descendència final, es produeix un augment de la fecunditat conjuntural (ICF), que assoleix el valor màxim entre el anys 2015 i 2020, amb valors de 2, 1,77 i 1,45 fills per dona. Aquest augment és efecte de l'actual cicle de retard de la fecunditat. Les tendències de la descendència final i de l' ICF segons les diferents hipòtesis es poden observar en els gràfics 1 i 2 respectivament.

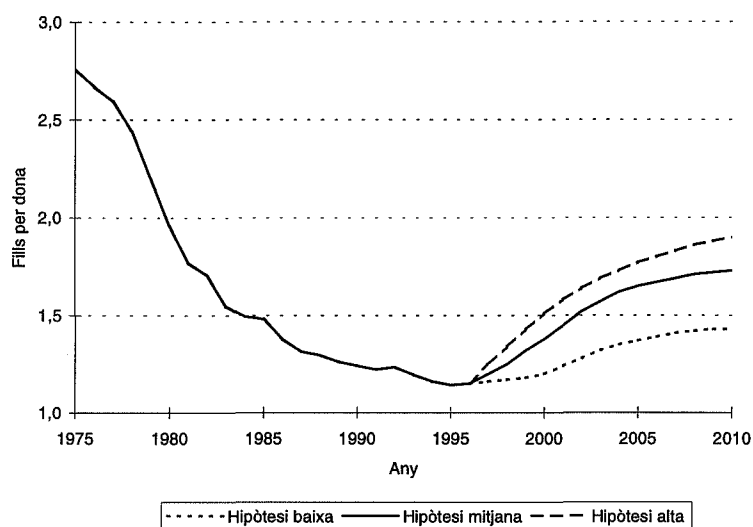


Gràfic 1. Descendència final acumulada als 50 anys per generacions. Catalunya. 1940-1980.

4. LA PROJECCIÓ DE LA MORTALITAT

En els darrers anys la variable mortalitat va prenent més i més importància en les projeccions de població per diferents motius; primerament perquè les previsions d'estabilitat en l'evolució de l'esperança de vida, que era el supòsit més freqüentment contemplat en les projeccions de població, han estat repetidament desmentides per un creixement de la longevitat que ha estat sistemàticament superior al previst.² D'altra banda, l'augment de l'esperança de vida, que històricament s'havia basat en la reducció de les taxes en edats joves i infantils i havia estat un factor de rejuveniment de la piràmide demogràfica, s'ha convertit avui en dia en un factor d'envelliment, perquè són les edats avançades les que registren la disminució més important en la freqüència relativa de defuncions. Com que una de les qüestions que més interessen de la població futura és quines xifres, absolutes i relatives, pot arribar a assolir la població de més edat, les hipòtesis d'evolució de la longevitat han esdevingut crucials.

²Les Nacions Unides han anat revisant a l'alça l'esperança de vida límit en les successives actualitzacions. Les projeccions de l'any 1984 situaven l'edat límit en 75 anys pels homes i en 82,5 anys per les dones, nivells que ja han estat sobrepassats actualment en moltes poblacions. En les projeccions de l'any 1988, l'edat límit era de 82,5 anys per als homes i de 87,5 anys per a les dones.



Gràfic 2. Hipòtesis d'evolució de l'indicador conjuntural de fecunditat. Catalunya. 1975-2010.

En la projecció dels nivells futurs de la mortalitat s'han tractat per separat l'evolució dels nivells de l'esperança de vida i l'evolució de les taxes de mortalitat per edat. Si bé aquests dos factors van íntimament lligats, cal tenir en compte que l'evolució de la mortalitat no és homogènia en totes les edats, de manera que un augment de l'esperança de vida no és necessàriament la traducció d'una millora de les taxes de defunció a totes les edats, podent-se produir patrons oposats d'evolució de la mortalitat per edats, com ho demostra l'evolució recent de l'esperança de vida masculina a Catalunya, que ha registrat una tendència a l'alça compatible amb un augment de la mortalitat en edats joves. Així doncs, en una primera etapa es projecta l'evolució de l'esperança de vida per a cada sexe i cada any, i en una segona etapa es tradueixen les esperances de vida en taxes de mortalitat per edats.

Per projectar els valors futurs de les esperances de vida en néixer, Coale i Guo proposen el mètode del traç cronològic: d'acord amb les seves observacions existeix una relació lineal decreixent entre l'increment mitjà anual de l'esperança de vida en néixer (e_0) i el nivell de la pròpia esperança de vida. Els traços cronològics es deriven de parells de taules de mortalitat femenina i masculina en poblacions amb extenses games de nivells de mortalitat, separades per un temps mínim de 5 anys. Han calculat una regressió lineal entre l'increment anual d' e_0 en relació al nivell d' e_0 . La relació lineal és de:

$$de_0/dt = a - b \cdot e_0$$

i condueix a una equació per a l'esperança de vida en néixer per un moment « t » que segueixi a un moment inicial « t_0 »:

$$e_0(t) = \epsilon_0 - (\epsilon_0 - e_0(t_0)) \cdot \exp(-b \cdot (t - t_0))$$

on

ϵ_0 = esperança de vida màxima,

$e_0(t_0)$ = esperança de vida a l'any inicial t_0 ,

b = descens en el canvi anual d' e_0 (de_0/dt).

El valor d' e_0 s'obté quan e_0 ja no augmenta i per tant és $de_0/dt = 0$, o sigui $\epsilon_0 = a/b$.

D'acord amb les dades estudiades per Coale i Guo, s'obtenen uns valors per les dones d' $\epsilon_0 = 83,25$ anys i $b = 0,03099$. Atès el creixement continuat de l'esperança de vida en els països desenvolupats, els mateixos autors donen un límit alternatiu d'esperança de vida d' $\epsilon_0 = 84,9$ anys. Un pas previ a la determinació dels nivells de mortalitat que s'hauria d'assolir en cada hipòtesi, és l'estudi dels valors enregistrats en els últims anys donat que l'evolució d' e_0 no és del tot contínua, presentant variacions puntuals. De l'estudi de les diverses aproximacions realitzades es conclou que el millor any a prendre com punt inicial del traç cronològic és $t_0 = 1992$.

El càlcul de l'evolució futura de l'esperança de vida pels homes necessita d'un instrument addicional, donada la menor regularitat del seu traç cronològic. Aquest instrument són les taules de vida del model de mortalitat regional de Coale, Demeny i Vaughan. Aquestes taules inclouen quatre models regionals (nord, sud, est i oest) i estan classificades per nivells. Cada nivell té una taula de vida per a cada sexe (per exemple en el nivell 26, e_0 pels homes és de 76,19 anys i per les dones de 82,50 anys). Per calcular el traç cronològic dels homes es procedeix de la següent manera: amb el mateix any de partida, t_0 , es pren $e_0(t_0)$ per a l'home i es calcula el nivell de les taules de vida que li correspon i l'esperança de vida per les dones en aquest nivell, e'_0 . A continuació es calcula el traç cronològic de les dones amb punt de partida e'_0 . Finalment, es tradueixen les esperances de vida del traç cronològic $e'_0(t)$ a termes de nivell de taula de mortalitat i es pren l'esperança de vida dels homes en el mateix nivell.

La hipòtesi mitjana d'esperança de vida s'ha obtingut amb un traç cronològic amb $t_0 = 1992$, $e_0(t_0) = 81,42$ i una esperança de vida límit de 84,9 anys per a les dones. El traç cronològic dels homes es determina com s'ha indicat prèviament.

Alguns experts consideren que l'esperança de vida té un creixement sostingut i que, en qualsevol cas, el límit biològic de 85 anys se situa massa a prop de les dades que ja actualment s'enregistren. És per aquesta raó que la hipòtesi alta d'esperança de vida s'ha obtingut amb un traç cronològic per a les dones amb el mateix punt de partida,

però amb una esperança de vida límit de 91,4 anys que correspon a la taula de vida límit de Duchêne i Wunsch.³ El traç cronològic de l'esperança de vida masculina ha estat obtingut aplicant als homes guanys d'esperança de vida semblants als de les dones, amb la qual cosa s'ha obtingut per aquesta hipòtesi un traç cronològic superior al que s'obtindria segons el mètode de Coale i Guo.

La segona fase de la projecció de la mortalitat correspon al càlcul de les taxes per edat que determinen les esperances de vida del traç cronològic (cal tenir present que una mateixa esperança de vida es pot obtenir amb patrons diferents de mortalitat per edats). El mètode utilitzat és el dels coeficients de millora, metodologia també utilitzada per l'INE (1995), segons el qual la taxa de mortalitat per una edat i un any es referencia a la taxa de la mateixa edat l'any anterior:

$$t_{e,s}(t) = t_{e,s}(t-1) \cdot \gamma_{e,s}(t)$$

on

$t_{e,s}(t)$ és la taxa de mortalitat en un any « t » per una edat « e » i sexe « s »

$\gamma_{e,s}(t)$ és el coeficient de millora anual de la mortalitat a l'edat « e », el sexe « s » i l'any « t »

Plantejant el problema d'aquesta manera, els coeficients de millora $\gamma_{e,s}(t)$ varien cada any i per cada edat, presentant valors inestables.

S'han estudiat els seus valors en el període 1985-1990 i 1990-1995 prenent els coeficients constants al llarg de 5 anys

$$t_{e,s}(t) = t_{e,s}(t-5) \cdot \gamma_{e,s}(t)^5$$

i s'ha observat que la dependència de γ respecte de l'edat « e » es concentra en alguns grups d'edats, de manera que s'ha reescrit el model com

$$t_{e,s}(t) = t_{e,s}(t-1) \cdot \gamma_s(t) \cdot cv_{e,s}(t)$$

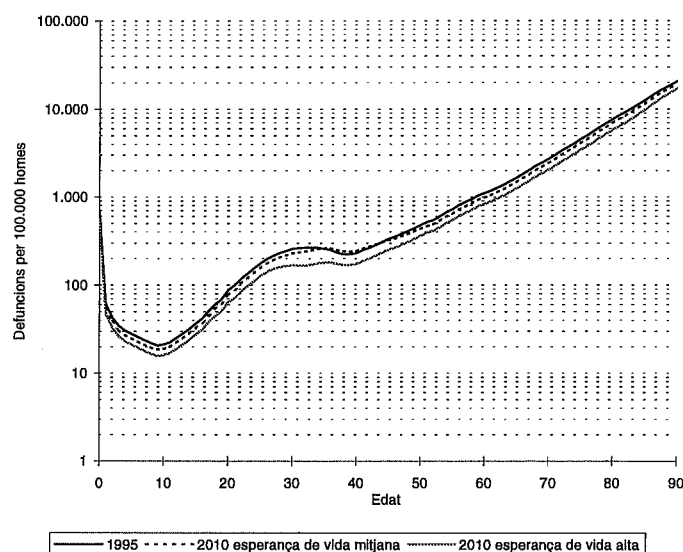
on

$t_{e,s}(t)$ és la taxa de mortalitat en un any « t » per una edat « e » i sexe « s »

$\gamma_s(t)$ és el coeficient de millora anual de la mortalitat del sexe « s » a l'any « t »

$cv_{e,s}(t)$ és el coeficient de variació de l'edat « e » i sexe « s » respecte de la millora anual a l'any « t »

³Aquesta taula de vida límit s'ha calculat suposant l'eliminació de totes les causes de mortalitat i deixant l'envelliment biològic com a única causa d'extinció.



Gràfic 3. Hipòtesis d'evolució de les taxes de mortalitat per edat. Homes. Catalunya. 1995 i 2010.

El model queda reescrit d'aquesta manera per aprofitar el fet que en la majoria de les edats la millora de l'esperança de vida serà homogènia ($cv_{e,s}(t)$ valdrà 1), presentant-se un comportament diferent només en la mortalitat infantil, les edats avançades i els adults joves. L'avantatge d'aquest model és que només es necessita precisar aquelles edats en les quals la millora de la mortalitat estarà per sobre o per sota de la mitjana, indicant la variació sobre la resta d'edats; per exemple, $cv_{e,s} = 1$ equival a no tenir un comportament diferenciat, $cv_{e,s} = 0,9$ equival a una millora del 10% respecte de la resta de les edats o $cv_{e,s} = 1,1$ equival a una taxa de mortalitat un 10% superior a la resta d'edats i equiparable a un empitjorament de la taxa de mortalitat. Per un algorisme iteratiu es determina el valor de $\gamma_s(t)$ que fa correspondre el conjunt de $t_{x,s}(t)$ amb l'esperança de vida marcada pel traç cronològic.

Els valors de $\gamma_s(t)$ obtinguts es troben al voltant de 0,99 per a la hipòtesi mitjana i 0,98 per a la hipòtesi alta. Les hipòtesis sobre els valors de $cv_{e,s}(t)$ s'han realitzat pels anys objectiu, $t=2010$ i $t=2030$, obtenint-se els anys intermitjos per interpolació lineal. Per finalitzar el procés ha estat necessari incloure uns coeficients de correcció amb l'objectiu de modificar les taxes de mortalitat obtingudes en alguna edat. És el cas de la mortalitat infantil, per poder fixar uns valors objectiu, o les edats adolescents, on s'assolien unes taxes de nivells molt baixos. La inclusió d'aquesta correcció introdueix petites variacions en l'esperança de vida del traç cronològic.

S'ha treballat amb 2 hipòtesis de mortalitat: mitjana i alta. En la hipòtesi mitjana les taxes es redueixen en un 9%. Les millores en els adults joves de 30 a 45 són inferiors que a les altres edats, produint-se fins i tot un estancament, amb el resultat d'un cert desplaçament del màxim de sobremortalitat, tal com es pot observar en el gràfic 3. En la hipòtesi alta, en canvi, hi ha una reducció de les taxes de mortalitat del 25% que afecta totes les edats, i especialment les adultes joves, amb el resultat d'una disminució significativa del pic de sobremortalitat dels joves de sexe masculí. En totes dues hipòtesis, la mortalitat en edats avançades coneix reduccions menys importants que en la resta d'edats.

5. LA PROJECCIÓ DE LA MIGRACIÓ

Les dades empíriques mostren que la migració és un fenòmen altament selectiu segons l'edat. Els joves entre els 20 i els 30 anys exhibeixen les taxes de migració més elevades, la qual cosa s'explica perquè la mobilitat residencial d'aquestes edats està relacionada freqüentment amb motius laborals i amb els processos d'emancipació de la llar parental, sovint associats al matrimoni i/o a l'inici de la vida en parella. L'alta mobilitat relativa en edats infantils també reflecteix una mobilitat de tipus familiar, d'adults de més de 30 anys que es desplacen amb els seus fills, potser a residències més àmplies. Les edats adolescents, en canvi, exhibeixen les taxes més baixes, juntament amb els adults madurs de 40 a 55 anys. A les edats entorn la jubilació s'ha detectat també un augment de la mobilitat, que reflecteix canvis de residència lligats a la jubilació. Finalment, a les edats molt avançades, cap als 70 anys o més, també es produeixen canvis de residència, que es poden relacionar amb diferents motius, com ara canvis en l'estat civil, per defunció del cònjuge i/o per canvis en l'estat de salut que representin una pèrdua d'autonomia que fa sovint necessari deixar el domicili propi per ingressar en establiments col·lectius o per anar a viure a casa de parents.

Aquesta propietat de la migració, de ser selectiva segons l'edat, permet matematitzar el fenomen a partir d'un model dependent de l'edat, conegut com model migratori de Rogers. La formulació és la següent:

model general:

$$p(e) = a_1 \exp(-\alpha_1 e) + a_2 \exp\{-\alpha_2(e - \mu_2) - \exp(-\lambda_2(e - \mu_2))\} + c$$

model amb pic de jubilació:

$$p(e) = a_1 \exp(-\alpha_1 e) + a_2 \exp\{-\alpha_2(e - \mu_2) - \exp(-\lambda_2(e - \mu_2))\} + \\ + a_3 \exp\{-\alpha_3(e - \mu_3) - \exp(-\lambda_3(e - \mu_3))\} + c$$

model amb pendent a les edats avançades:

$$p(e) = a_1 \exp(-\alpha_1 e) + a_2 \exp\{-\alpha_2(e - \mu_2) - \exp(-\lambda_2(e - \mu_2))\} + \\ + a_3 \exp(\lambda_3 e) + c$$

L'ajustament d'aquest model a les dades empíriques de la migració a Catalunya és realment bo, per la qual cosa constitueix un instrument de treball especialment adequat a l'hora de projectar la migració. S'han aplicat els 3 models a les dades de la immigració, l'emigració i la migració de l'estranger per a cada sexe i en tots els casos el que millor s'ha aproximat a les dades és el model amb pic de jubilació, si bé en el cas de la immigració de la resta d'Espanya el pic es desplaça cap a edats més avançades.

A diferència de la fecunditat o la mortalitat, on les hipòtesis de treball es fan en termes d'indicadors, com ara l'indicador conjuntural de fecunditat o l'esperança de vida, en la migració les hipòtesis de treball es fan en termes de xifres absolutes, és a dir, de migrants. La distribució del nombre total de migrants per edats es realitza mitjançant el model de Rogers. Efectivament, siguin

y = any

s = sexe

e = edat

$M(y, s, e)$ = migrants

$E(y, s)$ = emigrants cap a la resta d'Espanya

$I(y, s)$ = immigrants de la resta d'Espanya

$O(y, s)$ = migrants de l'estranger

$P(y, s, e)$ = població

$t(y, s, e) = M(y, s, e)/P(y, s, e)$ taxa de migració

$GMR(y, s) = \sum_{e=0,1,\dots,95} t(y, s, e)$ índex sintètic de migració

$p(y, s, e) = t(y, s, e)/GMR(y, s)$ pes d'una edat dins el conjunt de taxes

$(p(y, s))_{e=0,1,\dots,95}$ perfil migratori.

Notar que $\|(p(y, s))_{e=0,1,\dots,95}\|_1 = 1$, fet que permet separar intensitat i perfil. L'estimació del nombre de migrants es realitza per

$$M(y, s, e) = P(y - 1, s, e - 1) \cdot GMR(y, s) \cdot p(y, s, e).$$

Donat que no es fan hipòtesis en termes de GMR , l'indicador sintètic de la migració, sinó de M , el nombre de migrants, cal efectuar una sèrie de càlculs. En el cas de

l'emigració cap a la resta d'Espanya l'objectiu és determinar les taxes de migració per edat $t(y, s, e)$ a partir del perfil d'emigració p i el nombre d'emigrants E . Per fer-ho es calcula

$$GMR(y, s) = E(y, s) / \sum_{e=0,1,\dots,95} P(y-1, s, e-1) \cdot p(y, s, e)$$

$$t(y, s, e) = GMR(y, s) \cdot p(y, s, e)$$

En el cas de la immigració de la resta d'Espanya la situació és lleugerament diferent, donat que es tracta d'una entrada exògena al sistema de projecció i s'ha de proporcionar en xifres absolutes i no en taxes. Aprofitant, però, que l'estructura demogràfica de la resta d'Espanya és molt semblant a la catalana, s'han distribuït els immigrants per edat a partir de la piràmide catalana. Així doncs, a partir del perfil d'immigració p i el nombre total d'immigrants I , es calculen els immigrants per edat com

$$GMR(y, s) = I(y, s) / \left(\sum_{e=0,\dots,95} P(y, s, e) \cdot p(y, s, e) \right)$$

$$I(y, s, e) = P(y-1, s, e-1) \cdot GMR(y, s) \cdot p(y, s, e).$$

El cas del saldo migratori amb l'estranger és semblant al de la immigració de la resta d'Espanya perquè cal donar-lo en xifres absolutes, però aquí no és possible cap analogia de les estructures demogràfiques. En el seu lloc, es treballa amb el perfil dels migrants calculat en dades absolutes, i no en taxes. Les entrades i sortides de l'estranger s'estudien conjuntament i en donar un saldo netament positiu es tracten com entrades; així doncs, els migrants estrangers per edat es calculen a partir del seu total O i el perfil per edats p :

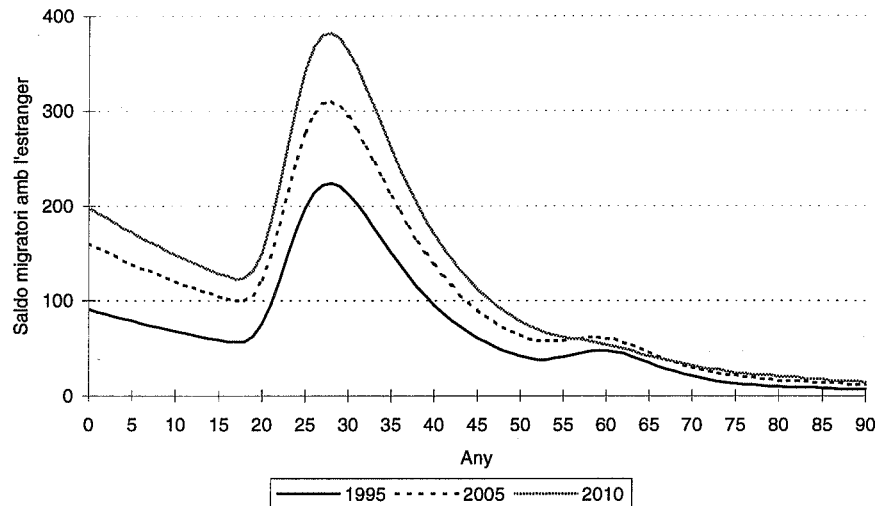
$$O(y, s, e) = O(y, s) \cdot p(y, s, e)$$

El tractament dels perfils migratoris no és estàtic al llarg del temps. S'han introduït els coeficients $k_1(y)$, $k_2(y)$ i $k_3(y)$ per poder modificar la intensitat de la migració infantil, laboral i postlaboral respectivament, de manera que els perfils queden definits com

$$p(y, s, e) = k_1(y) \cdot a_1 \cdot \exp(-\alpha_1 e) + k_2(y) \cdot a_2 \cdot \exp\{-\alpha_2(e - \mu_2) - \exp(-\lambda_2(e - \mu_2))\} + k_3(y) \cdot a_3 \cdot \exp\{-\alpha_3(e - \mu_3) - \exp(-\lambda_3(e - \mu_3))\} + c$$

Els coeficients $k_1(y)$, $k_2(y)$ i $k_3(y)$ valen 1 inicialment i els seus valors oscil·len entre 0,25 i 2. Els coeficients prenen valors segons les tendències i els ritmes marcats a cada hipòtesi per uns valors objectiu. Per exemple, en la hipòtesi mitjana de migració de l'estranger, $k_3(2000) = 1$, $k_3(2005) = 0,75$ i $k_3(2010) = 0,25$ equivalen a una reducció progressiva del pes de la migració de retorn de gent gran entre el 2000 i el 2005

que s'accentua entre el 2005 i el 2010, moment a partir del qual queda constant (gràfic 4). En canvi, a la hipòtesi mitjana d'emigració $k_3(2005) = 1,25$ i $k_3(2015) = 0,5$ determinen un augment a curt termini de la migració de retorn cap a la resta d'Espanya, coincidint amb l'arribada a la jubilació de les generacions que van emigrar cap a Catalunya, per descendir posteriorment a nivells inferiors als actuals.



Gràfic 4. Hipòtesis d'evolució de la migració amb l'estranger. Catalunya. 1995-2010.

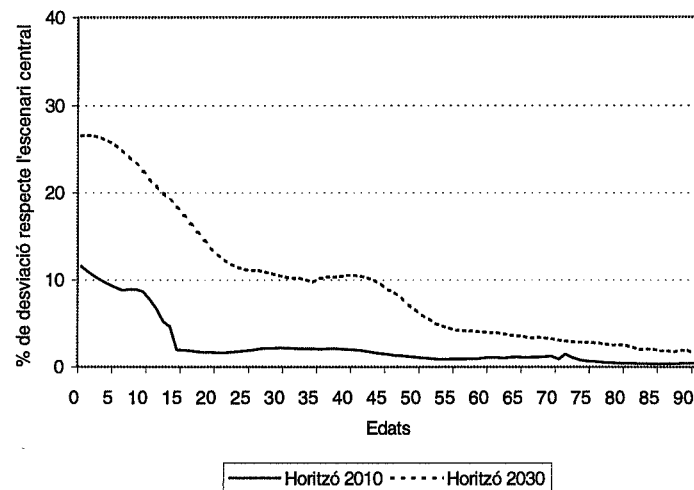
S'ha treballat amb 2 hipòtesis de migració amb la resta d'Espanya (baixa i mitjana) i 2 amb l'estranger (mitjana i alta). En la migració amb la resta d'Espanya, tant la hipòtesi baixa com la mitjana suposen un augment de les sortides, caracteritzat pel major pes de l'emigració de retorn al voltant del 2005 i la seva disminució posterior. En canvi, a les arribades l'augment seria més notable a la hipòtesi mitjana i es caracteritzaria per l'increment del pes de la migració dels joves cap el 2015, coincidint amb la incorporació al mercat laboral de les generacions buides nascudes els anys noranta. En els dos supòsits d'evolució de la migració procedent de l'estranger s'ha considerat una equiparació entre els efectius de cada sexe. Atès que el reagrupament familiar serà previsiblement un dels trets que més caracteritzaran els fluxos externs en el curt termini, s'han intensificat els fluxos migratoris en les edats infantils, tant en la hipòtesi alta com en la mitjana; en aquesta darrera hipòtesi, en canvi, s'ha suposat una disminució de la migració en edats de jubilació en considerar una tendència a la disminució dels fluxos de retorn de migrants espanyols.

6. ELS RESULTATS

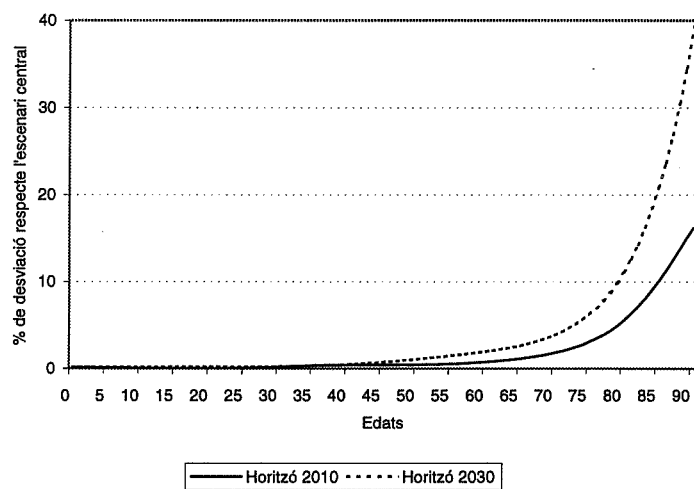
La combinació selectiva de les hipòtesis sobre els components dona lloc a 4 escenaris; cada un d'ells té una caracterització particular i permet la seva utilització per a finalitats específiques. Els dos escenaris intermedis són l'escenari tendencial i l'escenari central, que reflecteixen la visió més àmpliament acceptada sobre el futur de la migració i la fecunditat, amb opcions alternatives pel que fa a l'evolució futura de l'esperança de vida. L'escenari situat en l'extrem superior, l'escenari alt (o jove), serveix per situar tant els màxims efectius de població que es poden atènyer plausiblement com el mínim envelliment demogràfic, ja que combina totes les hipòtesis, que actuen per ralentitzar la tendència a l'envelliment de la població. Simètricament, l'escenari situat en l'extrem inferior o escenari baix (o vell) reflecteix el creixement demogràfic mínim que es pot esperar i representa un extrem en termes de l'envelliment de l'estructura demogràfica, en combinar baixos nivells de fecunditat i migració amb alts nivells d'esperança de vida.

Sovint es posa en dubte la capacitat predictiva de les projeccions. El problema fonamental no rau en les deficiències metodològiques de les tècniques sinó en la dificultat de pronosticar les pautes reproductives futures, l'evolució de l'esperança de vida i, sobretot, el fenomen de la migració, lligat a factors externs a la demografia com són l'economia o les polítiques reguladores. Aquesta incertesa provoca que els usuaris de projeccions acostumin a demanar, especialment quan hi ha diversos escenaris de futur, que s'assigni una probabilitat d'ocurrència a cada projecció. Els científics no han deixat de fer esforços per intentar quantificar la incertesa; els resultats, però, són mínims. Un indicador útil per quantificar la incertesa és l'amplada del ventall de projecció mesurada respecte de la població inicial. En un horitzó de 34 anys de projecció el ventall té una amplada del 17% de la població inicial, és a dir, es produeix un increment mitjà anual del 5 per mil en l'amplada del ventall. De totes maneres, la indeterminació de la població futura no afecta per igual totes les edats: els dos extrems de la piràmide són els que poden patir més variabilitat. En efecte, mentre que la migració es reparteix entre totes les edats, mitigant el seu efecte, no passa el mateix amb els altres 2 components. El grup d'edat de 0 anys està format en la seva pràctica totalitat pels naixements de cada any (l'efecte de la migració i la mortalitat és comparativament molt baix), de manera que el grau d'incertesa és notable. A mida que augmenta l'horitzó de projecció creix la indeterminació del grup de 0 anys, perquè a la incertesa del nivell de fecunditat s'afegeix la dels efectius de dones en edat fèrtil. El gràfic 5 mostra la variació entre la població de 2 escenaris causada per les diferències en la fecunditat i la migració, i com aquesta variació augmenta amb l'horitzó de projecció. La mortalitat, per la seva banda, es concentra en les edats més avançades, amb unes taxes prou importants com per causar variacions significatives en el nombre de supervivents. La incertesa creix amb l'amplada de l'horitzó de projecció, a mida que els efectius de població transiten per les diferents edats. El gràfic 6 mostra la variació entre la població de 2 escenaris

causada per les diferències en la mortalitat, així com el seu augment amb l'horitzó de projecció.



Gràfic 5. Variació de la població de l'escenari alt respecte del central. Horitzons 2010 i 2030.



Gràfic 6. Variació de la població de l'escenari tendencial respecte del central. Horitzons 2010 i 2030.

La millor manera de reduir la incertesa de les projeccions demogràfiques és considerar-les, no com un estudi tancat sinó com un instrument de treball que permeti futures revisions, a mesura que es conegui l'evolució real dels components del creixement demogràfic i la població.

REFERÈNCIES

- [1] **Coale, A., P. Demeny i B. Vaughan** (1983). *Regional Model Life Tables and Stable Populations*. Academic Press.
- [2] **Coale, A. i G. Guo** (1991). «Utilización de nuevas tablas modelo de mortalidad para tasas de mortalidad muy bajas en proyecciones demográficas». *Boletín de Población de las Naciones Unidas*, Nueva York.
- [3] **Crujisen, H. i altres** (1995). *Fertility Longitudinal Extrapolation Method*. Eurostat.
- [4] **Duchêne, J. i G. Wunsch** (1991). *From the demographers Cauldron: Working Paper*, Institut de Démographie, Université Catholique de Louvain.
- [5] **Institut d'Estadística de Catalunya** (1998). *Projeccions de població de Catalunya 2010-2030*. Generalitat de Catalunya.
- [6] **Institut d'Estadística de Catalunya Catalunya i Departament de la Presidència**. (1997). *Jornades tècniques sobre projeccions demogràfiques de Catalunya*. Generalitat de Catalunya.
- [7] **Gill, R.D. and N. Keilman** (1990). «On the estimation of multidimensional demographic models with population registration data». *Mathematical Population Studies*, 2(2), 119-143.
- [8] **Instituto Nacional de Estadística** (1995). *Proyecciones de Población de España calculadas a partir del Censo de Población de 1991*.
- [9] **Lutz, W. (Ed)** (1996). *The future population of the world. What can we assume today?* International institute for applied systems analysis (IIASA). Earthscan Pub.
- [10] **Rogers, A. i F. Willekens (Ed)** (1986). *Migration and settlement, a multiregional comparative study*. Reidel Pu. Co.
- [11] **Van Imhoff, E. i N. Keilman** (1991). *LIPRO 2.0.: An application of dynamic demographic projection model to household structure in the Netherlands*. Swets and Zeitlinger.
- [12] **Van Imhoff, E.** (1994). *LIPRO user's guide. Version 3.0*. Working paper 1994/1A. Netherlands interdisciplinary demographic institute (NIDI).

ENGLISH SUMMARY

METHODS AND MODELS TO PROJECT THE COMPONENTS OF DEMOGRAPHIC CHANGE IN THE POPULATION PROJECTIONS OF THE INSTITUT D'ESTADÍSTICA DE CATALUNYA

M. FARRÉ MALLOFRÉ

J.A. SÁNCHEZ CEPEDA

Institut d'Estadística de Catalunya*

The Catalan statistical office (Idescat) has projected the Catalan population for the period 1996-2030. The methodologies used to project each component of the demographic change (fertility, mortality and migration) are presented in this work. The projection of each component has taken into account not only the synthetic indicators (as life expectancy at birth, total fertility rate or gross migration rate) but also the specific and dynamic evolution for the different ages or cohorts. To project fertility, both cohort and crosssectional perspectives have been included, using a method that integrates mother's age and cohort, and birth order. The chronological trend of life expectancy at birth has been projected and combined with hypotheses about the evolution of mortality rates in specific ages. Rogers' migration model has been used to project a dynamic age migration profile along the projection period. The article concludes with a discussion about the results and attempts to give a measure of the uncertainty associated to different ages.

Keywords: Population projections, fertility, mortality, migration, total fertility rate, chronological trend, migration profile.

* Institut d'Estadística de Catalunya. Via Laietana, 58. 08003 Barcelona (Espanya).

—Received July 1998.

—Accepted October 1998.

The study of future Catalan population covers a period of 34 years, from 1996 to 2030, with the result of yearly populations broken down by sex and age. The projection methodology is the components' method: projected births, deaths and are added to the last population pyramid. LIPRO,¹ a computer program for multidimensional demographic projection, has been used. Three hypotheses about future evolution of migration and fertility, and two for mortality have been defined. The selective combination of these hypotheses has generated 4 scenarios of future population evolution, which are different not only because of the figures of total population but also because of the resulting demographic structures.

Fertility

Fertility has been projected studying not only mother's age but also their year of birth and the birth's order. Explicit hypotheses about number of children per woman have been made. The methodology is called FLEM² (fertility longitudinal extrapolation method) and works mainly with extrapolations: developments in past and present generations fertility at different ages are used to project present and future generations fertility.

Let « c » be a generation, « e » an age and « r » the order of birth. Then $F(c, e, r)$ is the fertility rate of order « r » at age « e » for cohort « c » $FA(c, e, r)$ is the cumulative fertility rate of order « r » at age « e » for cohort « c » where $F(c, e, r) = FA(c, e, r) - FA(c, e - 1, r)$.

The 3 indicators that allow to summarize the information about fertility of a year, a generation and the age of motherhood are

$$ICF(y) = \sum_e \sum_r F(c = y - e, e, r) \text{ total fertility rate}$$

$$DF(c) = \sum_r FA(c, 50, r) \text{ total cumulated fertility of a cohort}$$

$$EM(y) = \left\{ \sum_e \sum_r e \cdot F(c = y - e, e, r) \right\} / \left\{ \sum_e \sum_r F(c = y - e, e, r) \right\}$$

mean age at motherhood

¹Van Imhoff, E. i N. Keilman (1991). *LIPRO 2.0.: An application of dynamic demographic projection model to household structure in the Netherlands*. Swets and Zeitlinger.

²Crujisen, H. et al. (1995). *Fertility Longitudinal Extrapolation Method*. Eurostat.

FLEM calculates regressions between accumulated fertilities at ages 20, 25, 30, 35, 40 and 45 in previous generations

$$FA(c, e + 5, r) = \alpha + \beta \cdot FA(c, e, r) + \epsilon$$

Let « e_0 » be the last age of a generation with observed values and « cf » the target cohort. The way of calculating cumulated fertility at the rest of ages for any cohort « c » is

$$FA(c, e, r) = FA(c, e_0, r) + p(c, e, r) \cdot (FA(c, 50, r) - FA(c, e_0, r))$$

for every $e > e_0$, where

$$p(c - 1, e, r) = (FA(c - 1, e, r) - FA(c - 1, e_0, r)) / (FA(c - 1, 50, r) - FA(c - 1, e_0, r)).$$

$$p(cf, e, r) = (FA(cf, e, r) - FA(cf, e_0, r)) / (FA(cf, 50, r) - FA(cf, e_0, r)).$$

and $p(c, e, r)$ is a weighted average of $p(c - 1, e, r)$ and $p(cf, e, r)$.

Mortality

In the projection of future levels of mortality, the evolutions of life expectancy at birth and mortality rates by age have been analyzed separately. Although these agents are closely related, it must be borne in mind that the evolution of mortality is not homogenous for all ages.

Life expectancy at birth « e_0 » has been projected using the method known as the chronological trend, which provides the equation for past and future values of e_0 as

$$e_0(t) = \epsilon_0 - (\epsilon_0 - e_0(t_0)) \cdot \exp(-b \cdot (t - t_0))$$

where

ϵ_0 = top life expectancy,

$e_0(t_0)$ = life expectancy for a starting year t_0 ,

b = decrease in the annual change of $e_0 (de_0/dt)$.

According to data evaluated by Coale and Guo,³ the parameters' values are $\epsilon_0 = 83,25$ years and $b = 0,03099$, though persisting increases of life expectancy in developed countries made them to give an alternative value of $\epsilon_0 = 84,9$ years. The evolution of mortality rates by age is described by the equation

$$t_{e,s}(t) = t_{e,s}(t - 1) \cdot \gamma_{e,s}(t)$$

³Coale, A. i i G. Guo (1991). *United Nations' population bulletin*. Nr 30. New York.

where

$t_{e,s}(t)$ is the mortality rate depending on year, age and sex

$\gamma_{e,s}(t)$ is the anual improvement coefficient depending on year, age and sex.

As the values of $\gamma_{e,s}(t)$ are homogenous for groups of ages, the model can be written as

$$t_{e,s}(t) = t_{e,s}(t-1) \cdot \gamma_s(t) \cdot cv_{e,s}(t)$$

where $cv_{e,s}(t)$ is the coefficient of variation for an age and sex respect the anual improvement $\gamma_s(t)$.

Hypotheses for values of $cv_{e,s}(t)$ in certain ages are made in target years 2010 and 2030. As the set of $t_{e,s}(t)$ must match the cronological line $e_0(t)$ described previously, the values of $\gamma_s(t)$ are calculated in an iterative way. Values for $\gamma_s(t)$ are between 0.99 for the medium hypotheses and 0.98 for the high one.

Migration

In the study of migration two elements must be analized: the level and the profile.

Let $t(y, s, e)$ be migration rates and $M(y, s)$ the number of migrants. The level and the profile of migration are defined as

$$GMR(y, s) = \sum_{e=0,1,\dots,95} t(y, s, e) \quad \text{gross migration rate}$$

$$p(y, s, e) = t(y, s, e) / GMR(y, s)$$

Usually, hypotheses for migration level are not made in terms of GMR but in number of migrants. Anyway, its calculation is necessary as it is the link between the number of migrants and the profile of migration. Dinamic profiles are defined, adding 3 parameters, $k_1(y)$, $k_2(y)$, $k_3(y)$ to Rogers' migration model. The resulting model is

$$p(y, s, e) = k_1(y) \cdot a_1 \cdot \exp(-\alpha_1 e) + k_2(y) \cdot a_2 \cdot \exp\{-\alpha_2(e - \mu_2) - \exp(-\lambda_2(e - \mu_2))\} + k_3(y) \cdot a_3 \cdot \exp\{-\alpha_3(e - \mu_3) - \exp(-\lambda_3(e - \mu_3))\} + c$$

Rates are then calculated as follows:

$$GMR(y, s) = M(y, s) / \left(\sum_{e=0,1,\dots,95} P(y-1, s, e-1) \cdot p(y, s, e) \right)$$

$$t(y, s, e) = GMR(y, s) \cdot p(y, s, e)$$

where $M(y, s)$ stands both for number of immigrants and for number of outmigrants.

Results

The production of several scenarios and the uncertainty of future population evolution lead people to ask for a probability for each scenario. Scientists have worked on this problem but the results are poor. A useful way for measuring the uncertainty is the range of the interval for future population respect to the initial population. In a 34 years horizon, the interval of future population has a width of 17% the initial population; so, the mean anual increase is 5%. The analysis of the scenarios reveals that uncertainty isn't homogenous for different ages: younger and older ages are the most uncertain.

The best way of reducing uncertainty in demographic projections is not to consider them as something closed but as a tool that will be revised and updated as the evolution of demographic components is known.

Biometria

PAPEL ÉTICO DEL ESTADÍSTICO EN LA EXPERIMENTACIÓN HUMANA

ERIK COBO

Universitat Politècnica de Catalunya*

En diferentes ocasiones, desde la Sociedad de Estadística e Investigación Operativa o desde la Región Española de la International Biometry Society, se nos ha pedido a los estadísticos aplicados que hablemos de aquellas características de nuestro trabajo profesional no contempladas en los planes de estudio. En estas líneas pretendemos recordar el relevante papel que el estadístico desempeña en la investigación y desarrollo de nuevas pautas terapéuticas. Para ello, expondremos los principios éticos básicos que guían la experimentación con seres humanos y sus implicaciones en la tarea profesional del estadístico. También esbozaremos aquellas líneas de investigación metodológica básica que podrían aligerar las tensiones entre principios éticos e investigación con seres humanos. Finalmente, insistimos en la necesidad de la reincorporación inmediata de estadísticos en las comisiones encargadas de valorar la validez científica de las propuestas de investigación clínica.

Statistical ethical issues in human experimentation.

Palabras clave: Ensayos clínicos, ética, código deontológico, estadísticos, causalidad.

Clasificación AMS: 62K99, 92C50

*Erik Cobo. Departament d'Estadística i Investigació Operativa. Secció Informàtica. Universitat Politècnica de Catalunya. Pau Gargallo 5. 08028 Barcelona (Espanya). E-mail: cobo@eio.upc.es

—Recibido en junio de 1998.

—Aceptado en noviembre de 1998.

1. INTRODUCCIÓN

Sir Austin Bradford Hill fue pionero en la aplicación de los principios del diseño experimental de Sir Ronald Arnold Fisher a la investigación de procedimientos terapéuticos. Su participación en el estudio de la eficacia de la Estreptomina en el tratamiento de la tuberculosis (Medical Research Council, 1948) originó el que ha sido reconocido como primer Ensayo Clínico aleatorizado. Mediante asignación al azar, los pacientes fueron incluidos, o bien en el grupo control —que realizaba reposo—, o bien en el grupo tratado —que recibía, además, tratamiento antibiótico. Al cabo de seis meses, dos radiólogos, que desconocían el tratamiento recibido evaluaban la evolución de la enfermedad. Los resultados, estadísticamente significativos, abrieron una nueva época en el tratamiento de esta enfermedad. Por supuesto, los pacientes sometidos al grupo control no pudieron beneficiarse de un efecto que, por entonces, era sólo una hipótesis. Stuart Pocock (1993) ha resumido perfectamente el paradigma que marca la experimentación con seres humanos: las obligaciones de la investigación científica para encontrar los mejores tratamientos para los pacientes futuros se contraponen con los derechos de los pacientes actuales a recibir el mejor tratamiento. Esta contraposición de intereses viene impuesta por un rigor científico que —por lo menos hoy en día— exige grupo control y asignación al azar del tratamiento.

Los primeros ensayos clínicos coincidieron, pues, con el momento en que empezaron a conocerse algunos aspectos de los —mal llamados— experimentos nazis durante la segunda guerra mundial. Todo ello provocó una intensa labor en la búsqueda de referencias de comportamiento ético. A la publicación en 1947 del código de Nuremberg, siguieron en 1958 los principios de conducta estadística profesional de Edward Deming (1965, 1972), donde distingue entre las obligaciones del estadístico y las del investigador «substantivo», cada uno responsable de la correcta aplicación de los métodos científicos de sus áreas de conocimiento. Hacia 1959, Clegg, director del British Medical Journal, empezó el difícil proceso de obtener un consenso en la comunidad médica internacional de lo que era y no era éticamente aceptable en la experimentación humana. El mismo Hill (1963) abordó muy valientemente el tema adoptando una postura que hoy en día sólo sorprende por no ser estricta en la solicitud del consentimiento informado —aspecto incluido en la declaración de ética profesional del International Statistical Institute (1994). Este largo proceso culminó con la aprobación en la 62ª asamblea de la World Medical Association (Gilder, 1964) de la Declaración de Helsinki que, con pequeñas modificaciones, sigue vigente (Real Decreto 561/1993 de 16/4/93). Veamos, a continuación, estos principios éticos.

2. PRINCIPIOS ÉTICOS

Hoy en día (Simón, 1995) se reconocen como principios éticos que afectan toda investigación con seres humanos los siguientes: (1) *El principio de relevancia y solidez*

científica, que establece que toda experimentación humana debe estar dirigida a contestar una pregunta relevante y debe tener un diseño correcto en todos sus apartados metodológicos; (2) *El principio de autonomía*, que establece el derecho individual a decidir y regular el propio destino, requiere el consentimiento informado y libre ante cualquier intervención en nuestro organismo; (3) *Los principios de beneficencia y de no maleficencia*, que imponen la obligación de aplicar únicamente tratamientos que tengan potencial positivo y no dejar de aplicar terapias cuya carencia pueda provocar secuelas irreversibles; (4) *El principio de igualdad* o de justicia distributiva, que requiere que la experimentación se base proporcionalmente en todas las capas sociales; finalmente, también aplica (5) *El principio de confidencialidad*, común a todo trabajo estadístico (García Benavides 1988, Bacaria Martrus 1993), que requiere respetar el anonimato del paciente durante todo el proceso de los datos, pero especialmente en la publicación de los mismos.

Estos principios pueden competir entre sí (García Alonso 1994). Por ejemplo, el principio de autonomía, al imponer la participación voluntaria, desaconseja la experimentación con reclusos, que no pueden decidir libremente, por lo que atenta contra los principios de igualdad —al no soportar esta subpoblación el peso de la investigación— y de beneficencia —al no poder acceder a los hipotéticos efectos de un nuevo tratamiento. Nótese también que el principio de autonomía, para ser ejercido propiamente, presupone mantener la claridad mental en el difícil momento de conocer la enfermedad. Finalmente, hay que decir que el conflicto entre los principios de beneficencia y autonomía es más general y escapa a los límites del Ensayo Clínico. Por ejemplo, cuando la autoridad judicial impone el empleo de transfusiones sanguíneas en comunidades religiosas que las rechazan, está priorizando los principios de beneficencia y de no maleficencia sobre el de autonomía.

Sobre el nivel de seguimiento de estos principios existe disparidad de opiniones. Mientras que De Abajo et al. (1989) aconsejan incluso a los comités de ética inspeccionar si los pacientes reconocen haber firmado el consentimiento informado (Gill y Baber, 1995), en cambio, informan de su cumplimiento en las compañías farmacéuticas con sede en el Reino Unido.

3. PAPEL ÉTICO DEL ESTADÍSTICO

Según la división de responsabilidades de Deming (1965, 1972), es el profesional sanitario o el experto en Ciencias de la Vida quién debe tomar la iniciativa en garantizar el respeto a estos principios, aunque ello no libera al estadístico de su obligación moral de proteger al participante «as far as possible, against potentially harmful effects» (punto 4.4. del código deontológico del ISI).

En nuestra opinión, el estadístico debe liderar la vigilancia de los conflictos éticos en los apartados metodológicos.

En primer lugar, garantizando que el diseño y análisis propuestos permitirán alcanzar el objetivo del estudio.

En segundo lugar, minimizando el «coste» para los pacientes actuales. Así, dentro del marco de la asignación aleatoria, la investigación estadística ha ofrecido aportaciones metodológicas relevantes para aligerar la presión sobre los participantes en el estudio: desde los procedimientos secuenciales —que permiten interrumpir tempranamente el estudio si el grado de evidencia acumulada lo aconseja—, a los diseños de asignación aleatoria dinámica —que aumentan la probabilidad de pertenecer al grupo que obtiene mejores resultados—, pasando por los diseños $n = 1$ con intercambio del tratamiento —que permiten establecer la pauta terapéutica para cada paciente. Como muy bien recuerda Pocock (1993), se debe ser muy prudente en el uso de estas alternativas —a menudo mal recibidas por la comunidad médica— que pueden resultar en una pérdida de la credibilidad del estudio, haciéndolo, a la postre, inútil para su objetivo final: permitir cambiar los hábitos terapéuticos cuando corresponda.

En tercer lugar, y con un papel más próximo al de la actividad notarial, el estadístico debe garantizar la credibilidad en las conclusiones alcanzadas. A un primer nivel ello implica planificar a ciegas y con detalle todo el plan de análisis estadístico, salvando —con la ayuda de toda la información pre-experimental— las dificultades inherentes al cumplimiento de las premisas de los modelos. A un segundo nivel, implica que el estadístico debe colaborar, con el resto del equipo, en la búsqueda de fraude científico. Mike Collins (1993) aconseja levantar la voz de alarma ante las siguientes circunstancias: ausencia casi absoluta de *missing* («datos demasiado buenos»), falta de patrón regular en el ritmo de inclusión de pacientes, menor variabilidad de la esperada, mayor efecto del esperado, sobrepresencia de ciertos números (3, 7,...) y relaciones no previsibles entre variables.

En cuarto lugar, creemos que es papel ético del estadístico potenciar el desarrollo de procedimientos científicos rigurosos que no exijan la asignación aleatoria. En el terreno de la Epidemiología, la asignación aleatoria es muchas veces imposible y no por ello dejan de inferir relaciones de causa-efecto. Tomemos por ejemplo el caso del tabaco, donde es impensable un Ensayo Clínico con asignación aleatoria al grupo fumador y al grupo control. Cuando Fisher (1959), al criticar los primeros estudios observacionales sobre el tabaco y el cáncer de pulmón, apuntaba que un posible condicionante genético podría originar la relación entre ambas variables, Hill —que había sido pionero en la introducción de la asignación aleatoria en el ensayo clínico— acabó proponiendo (1965) una serie de contrastes alternativos que permitían decantar el peso de la evidencia hacia la relación causal.

A otro nivel, Donald Rubin (1974), en el entorno de las ciencias del comportamiento y desde una perspectiva próxima a los planteamientos bayesianos, define la asignación ignorante del tratamiento como aquella situación en la que la respuesta potencial (previa al efecto del tratamiento en estudio) es la misma en los grupos estudiados. Así, caso de observar diferencias en la respuesta final —y de mantenerse el equilibrio entre los grupos—, podrían ser atribuidas al tratamiento en estudio. Es decir, aún existiendo diferencias en la distribución de ciertas variables que influyen en la respuesta, sus efectos podrían compensarse mutuamente, de forma que haya equilibrio en la respuesta potencial. Para poner a prueba esta comparabilidad de los grupos se ha propuesto que un comité de expertos intente «adivinar» el grupo de tratamiento al que pertenece cada caso. Nótese que si la decisión de asignar un tratamiento u otro se ha adoptado mediante algún criterio, este comité podría adivinar el grupo en el que ha sido incluido cada caso. Así, cruzando ambas variables (grupo real y clasificación de los expertos) pondríamos a prueba, de forma indirecta, la asignación ignorante de Rubin. A pesar de sus indudables ventajas éticas y de haber sido descrita en 1974, hasta donde nosotros sabemos, este planteamiento no ha sido ni aplicado ni tan sólo discutido en el terreno de los ensayos clínicos. Remitimos al lector interesado a referencias estadísticas sobre el establecimiento de causalidad (Holland, 1986; Stone, 1993; Senn, 1994) o epidemiológicas sobre comparabilidad de los grupos (Greenland, 1986; Mickey, 1989; Greenland, 1996).

4. RECAPITULACIÓN

Así pues, queremos repetir que, conscientes de las implicaciones éticas que tiene la aplicación del diseño experimental en las Ciencias de la Salud, debemos, siguiendo el ejemplo de Bradford Hill: (1) utilizar el diseño experimental para dilucidar solamente aquellas cuestiones que no pueden ser esclarecidas por otro método, y (2) estimular a nuestros compañeros, los estadísticos teóricos, en el desarrollo de nuevas metodologías que, como la de Rubin, permitan evolucionar hacia procedimientos de similar validez científica pero con menores tensiones éticas.

Finalmente, y parafraseando a Deming, dado que la sola presencia del estadístico en la tabulación de los datos no garantiza la validez científica del experimento, es nuestra obligación incorporarnos de forma temprana en estos estudios. A pesar de lo extendido y aceptado de todos estos principios e incluso de la necesaria presencia del estadístico (Ricoy, 1992), sorprende la ausencia de estadísticos profesionales en muchos organismos con responsabilidad en la autorización de Ensayos Clínicos. A diferencia de sus homólogos americanos (*Food and Drug Administration*) o europeos (*Medicines Community Agency*), que cuentan con conocidos estadísticos como John Lewis o Robert O'Neill, en la Dirección General de Farmacia y Productos Sanitarios del Ministerio de Sanidad y Consumo, no nos consta la existencia de estadísticos oficiales. Asimismo,

no es obligatoria la presencia de un estadístico en los extensos Comités de Ética de Investigación Clínica (CEIC) que deben autorizar previamente cualquier Ensayo Clínico (Dal Ré, 1995; Carné, 1997). Mientras que la reglamentación anterior (Orden 3/8/1982 por la que se desarrollaba el R.D. 944/1978 sobre ensayos clínicos en humanos) especificaba que uno de los miembros permanentes u obligatorios del CEIC debía ser «un bioestadista (sic) o persona que posea amplios conocimientos en Bioestadística», la normativa actual (Ley 25/1990 del Medicamento de 20/12/90, y R.D. 561/1993 de 16/4/93 que desarrolla el título III de Ensayos Clínicos), extensa en cuanto a los colectivos profesionales representados, no menciona al estadístico. A pesar de que ha fijado sus objetivos en, por este orden, «ponderar los aspectos metodológicos, éticos y legales del protocolo propuesto».

Dada nuestra obligación ética de garantizar la calidad de los estudios antes de ser iniciados, debemos utilizar nuestra influencia para acelerar el proceso de incorporación de los estadísticos a todos estos comités y organismos. No debe olvidarse que nunca será ética cualquier investigación en humanos que no optimice sus planteamientos metodológicos, de forma que los pacientes futuros puedan realmente beneficiarse del esfuerzo que se solicita a los pacientes actuales.

5. AGRADECIMIENTO

A un revisor anónimo que, entre otras sugerencias, remarcó la presencia del estadístico en la legislación previa de los comités de ética.

6. REFERENCIAS

- [1] **Bacaria-Martus, J.** (1993). «El secret estadístic: contingut jurídic». *Qüestió*, **17**, 405-412.
- [2] **Carné, X.** (1997). «Dos años de comités éticos de investigación clínica». *Investigación Clínica y Bioética*, **20**, 13-15.
- [3] **Collins, M., Piper, D. and Thomas, P.** (1993). «The detection of fraud in Medical research: the role of statistics». *PSI newsletter*, 18-22.
- [4] **Dal-Ré, R.** (1995). «Comités éticos de investigación clínica: algo más que el cambio de nombre». *Med. Clin. (Barc.)*, **105**, 580-582.
- [5] **De Abajo, F., J, Serrano-Castro, M.A., Galende, I. y Tristán, C.** (1989). «El consentimiento informado y los comités de Ensayos Clínicos». *Med. Clin. (Barc.)*, **93**, 801.

- [6] **Deming, W.E.** (1965). «Principles of professional statistical practice». *Annals of Mathematical Statistics*, **36**, 1883-1990.
- [7] **Deming, W.E.** (1972). «Code of professional conduct». *Int. Statistical Rev.*, **40**, 215-219.
- [8] **Fisher R.A.** (1959). *Smoking and the cancer controversy*. Oliver and Boyd, Edinburgh.
- [9] **García-Alonso, F.** (1994). «Los principios básicos de la bioética». *Investigación Clínica y Bioética*, 5-6.
- [10] **García-Benavides, F.** (1988). «El secreto estadístico y la investigación sanitaria: un delicado, pero necesario, himeneo». *Gaceta Sanitaria*, **8**, 221-223.
- [11] **Gilder, S.** (1964). «World medical association Meets in Helsinki». *B.M.J.*, 299-300.
- [12] **Gill, L.S. and Baber, N.S.** (1995). «Ethical and scientific review procedures in the pharmaceutical industry: a survey of 12 British-based companies». *Pharm. Med.*, **9**, 11-20.
- [13] **Greenland, S. and Robins J.M.** (1986). «Identifiability, exchangeability and epidemiological confounding». *Int. J. Epidemiol.*, **15**, 413-419.
- [14] **Greenlan S.** (1996). «Absence of confounding does not correspond to collapsibility of the rate ratio or rate difference». *Epidemiology*, **7**, 498-501.
- [15] **Hill, A.B.** (1963). «Medical ethics and controlled trials». *B.M.J.*, 1043-1049.
- [16] **Hill, A.B.** (1965). «The environment and disease: association or causation?» *Proc. R. Soc. Med.*, **58**, 295-303.
- [17] **Holland, P.W.** (1986). «Statistics and causal inference (with comments and rejoinder)». *J.A.S.A.*, **81**, 945-970.
- [18] **International Statistical Institute** (1993). «Declaration on Professional Ethics». *Qüestió*, **17**, 3, 373-403.
- [19] **Medical Research Council** (1948). «Streptomycin treatment of pulmonary tuberculosis». *B.M.J.*, **ii**, 769-782.
- [20] **Mickey R.M. and Greenlan S.** (1989). «The impact of confounder selection criteria on effect estimation». *Am. J. Epidemiol.*, **129**, 125-137.
- [21] **Pocock, S.J.** (1993). «Statistical ethical issues in monitoring clinical trials». *Stat. in Med.*, **12**, 1459-1469.
- [22] **Ricoy, J.C.** (1992). «Ética y política científica». *Med. Clin. (Barc.)*, **98**, 423-426.

- [23] **Rubin, D.B.** (1974). «Estimating causal effects of treatment in randomized and non randomized studies». *J. Educ. Psychol.*, **66**, 688-701.
- [24] **Senn, S.** (1994). «Fisher's game with the devil». *Stat. in Med.*, **13**, 217-230.
- [25] **Simón Lorda, P. y Barrio Cantalejo, I.M.** (1995). «Un marco histórico para una nueva disciplina: la bioética». *Med. Clin. (Barc.)*, **105**, 583-597.
- [26] **Stone R.** (1993). «The assumptions on which causal inference rest». *J.R. Statist. Soc. B*, **55**, 455-466.

ENGLISH SUMMARY

STATISTICAL ETHICAL ISSUES IN HUMAN EXPERIMENTATION

ERIK COBO

Universitat Politècnica de Catalunya*

In this paper I seek to highlight, by looking at the ethics behind any human experiment, the significant role that the statistician plays in the research and development of any new therapy. I discuss the conflicts between human rights and the scientific principles that guide clinical trials, and also the desire for progress in methodological developments in order to diminish this conflict. I also point out that in the Spanish regulation of clinical trials, the statistician is not a compulsory member of the scientific and ethical committee who authorise clinical trials.

Keywords: Clinical trials, ethics, code of conduct, statisticians, causality.

AMS Classification: 62K99, 92C50

*Erik Cobo. Departament d'Estadística i Investigació Operativa. Secció Informàtica. Universitat Politècnica de Catalunya. Pau Gargallo 5. 08028 Barcelona (Espanya). E-mail: cobo@eio.upc.es

–Received June 1998.

–Accepted November 1998.

There are at least five ethical principles (Simón 1995; García Benavides 1988, Bacaria Martus 1993) to be considered when planning clinical trials.

- (1) The *scientific relevance and correctness* principle says that any human experimentation should be devoted to answering a relevant scientific question using an appropriate methodology.
- (2) The *autonomy* principle expresses the individual human right to decide his/her own destiny, which involves written and informed consent to any intervention in his/her body.
- (3) The *beneficence and non-maleficence* principle make it compulsory to give only treatments with a positive effect on the body and not interrupt therapies whose absence may have irreversible negative effects.
- (4) The *equality* principle of distributive justice requires experimentation to be done proportionally in all social strata.
- (5) The *confidentiality* principle, common to any gathering of human data, also applies to experimentation.

There are natural clashes between these principles (García Alonso 1994). For example, the autonomy principle precludes experimentation on convicts since they are not free to decide. Therefore, this population cannot receive the benefits of new therapies, as the beneficence principle demands.

Stuart Pocock (1993) defined ethical conflict in clinical trials as «the balance between the patients within the trial —that is, the individual ethics of randomising the next patient— and the longer term interest of obtaining reliable conclusions on sufficient data —that is, the collective ethics of making appropriate treatment policies for future patients». Following Deming's division of responsibilities between the substantive and the statistical aspects of a study (1965, 1972), I believe the statistician should deal with the ethical conflicts in the methodological part of the study.

First, the scientific issue should be addressed by an assessment of the planned design and analysis. Second, the human cost to the patients enrolled in the trial should be minimised, which involves the statistician seeking the «best» design for these patients. Sequential trials, dynamic allocation of patients, n -of- 1 trials, and other methodological alternatives should be studied and applied if possible. Third, the statistician should assess the credibility of the trial results since the only ethical argument for initiating a human experiment is to improve the scientific knowledge to be employed in future patients. Therefore, the statistician should plan quality checks of data and the robustness of the conclusion to any protocol violation, as well as to the planned analysis and its

assumptions. Fourth, I think the statistician should encourage the research and development of new methodologies which would potentially diminish the ethical conflict. Are there alternatives to randomisation? We should always remember that, although there are no randomised clinical trials on the topic, there is no doubt today about the harmful effect of tobacco. By randomising we establish exchangeability (Greenland and Robins 1986) –that is, that the expected value of our estimated causal effect would be the same if the new treatment were given to the other group– but there are alternative assumptions allowing us to make causal inference (Stone 1993). In the field of social science, Rubin proposed the «ignorant treatment assignment» (1974) as the necessary assumption to make causal inferences in the observational context.

Finally, I advocate early incorporation of the statistician into the scientific team, since «the statistician has no magic touch by which he may come in at the stage of tabulation and make something of nothing» (Deming 1965). Whereas the former Spanish medical law (*Order 3/8/1982 developing Royal Decree 944/1978 on clinical trials on humans*) required a member of the ethical committee to be an expert in statistics, the new law (*Law 25/1990 on Medicaments of 20/12/90, and Royal Decree 561/1993 of 16/4/93, which develops the law's Section III on clinical trials*) concerning the ethical committee which «should weigh up all the methodological, ethical and legal views of the proposal» requires the presence of a large number of different professionals, yet say nothing about statisticians. As it is our ethical responsibility to guarantee the quality of a study before it begins, we should ask for the early incorporation of the statistician into any committee or organisation dealing with human experimentation.

Secció Docent i Problemes

SECCIÓ DOCENT I PROBLEMES

La «Secció docent i problemes» té l'objectiu de publicar articles de caire docent, difícilment publicables en revistes de recerca. A cada número de *Qüestió* s'inclouen d'un a tres problemes i les solucions es donen en el número següent.

Els lectors poden 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.

MODELS GRÀFICS D'INDEPENDÈNCIA

J.M. DURAN RÚBIES

Universitat Autònoma de Barcelona*

Els models gràfics d'independència són una eina de l'anàlisi multivariant que utilitza gràfics per representar models. En particular, els grafs d'independència resumeixen i clarifiquen les interaccions entre variables, interaccions no sempre fàcils d'interpretar, especialment quan hi intervenen tres o més variables.

En aquest treball es dona, en clau pedagògica, una introducció a la teoria de grafs d'independència, començant per les nocions d'independència necessàries i arribant a les propietats de Markov, claus per a la interpretació d'aquests grafs.

S'estudia també un exemple d'aplicació amb variables numèriques, que mostra la viabilitat d'aquesta teoria a l'hora de simplificar i aclarir les relacions entre variables.

Independence graphs models.

Paraules clau: Independència, independència condicional, paradoxa de Simpson, vectors aleatoris, teoria de grafs, grafs d'independència condicional, propietats de Markov.

Classificació AMS: 62 H 20

*J.M. Duran Rúbies. Professor associat U.A.B. Departament de Matemàtiques. 08193 Bellaterra (Espanya).

E-mail: duran@mat.uab.es

—Rebut el novembre de 1996.

—Acceptat el desembre de 1998.

1. INTRODUCCIÓ

En anàlisi multivariant, quan hi intervenen tres o més variables, les relacions d'independència condicional no són fàcils de veure ni d'interpretar. Els grafs d'independència condicional resumeixen i clarifiquen aquestes interaccions.

La teoria de grafs d'independència és relativament recent. Fou Speed [8] un dels primers en relacionar els mètodes de la teoria de grafs amb les interaccions entre variables. El treball clau, però, és el de Darroch, Lauritzen i Speed [1], seguit després per altres com Wermuth i Lauritzen [9], Edwards i Kreiner [3] i Whittaker [10], el qual ha estat el motivador principal del treball que segueix, recull de tot aquest material, presentat en clau pedagògica. Més recentment Edwards [2] ha publicat un llibre amb aplicacions dels models gràfics d'independència. Inclou una guia de referència del programa MIM per a PC dissenyat per estudiar aquests models.

L'objectiu d'aquest treball és l'enunciat, demostració i aplicació de la propietat global de Markov, que permet interpretar els grafs d'independència i està dirigit a tots aquells que, sense haver d'anar a llibres com Whittaker [10] i Edwards [2], desitgin tenir una primera idea dels models gràfics.

L'exposició està estructurada de la manera següent: en les seccions 2 a 4 es dona una revisió ràpida de les nocions d'independència i d'independència condicional, fent especial èmfasi en les propietats que s'utilitzaran en la secció 7 per demostrar el teorema de separació. En la secció 5 definim els conceptes de la teoria de grafs necessaris per definir en la secció 6 els grafs d'independència condicional. En la secció 7 es presenta el teorema de separació i en la secció 8 les propietats de Markov. Finalment s'exposa, en la secció 9, un exemple d'aplicació amb variables numèriques.

La majoria de les demostracions estan agrupades fora del text principal, en l'apèndix, per poder seguir així més fàcilment els resultats que són el fil conductor d'aquest treball.

2. CONCEPTES D'INDEPENDÈNCIA

Dos esdeveniments A i B , d'un espai de probabilitat $(\Omega, \mathcal{A}, P)$, són *independents* i ho escriurem $A \perp\!\!\!\perp B$ si:

$$P(A \cap B) = P(A)P(B).$$

Tenint en compte que, suposant $P(B) > 0$, la probabilitat de l'esdeveniment A condicionat per l'esdeveniment B és:

$$P(A/B) = \frac{P(A \cap B)}{P(B)},$$

una altra manera de definir la independència és:

$$A \perp\!\!\!\perp B \iff P(A) = P(A/B).$$

Es pot provar que:

- $A \perp\!\!\!\perp B \iff A \perp\!\!\!\perp B^c$. (vegeu pàg. 186).
- $A \perp\!\!\!\perp B \iff A^c \perp\!\!\!\perp B$, i
- $A \perp\!\!\!\perp B \iff A^c \perp\!\!\!\perp B^c$.

També es pot provar que la relació $\perp\!\!\!\perp$ definida en un mateix espai de probabilitat és simètrica, però no és reflexiva ni transitiva. (vegeu pàg. 186).

En el cas de tenir tres esdeveniments o més, podem definir la ***independència marginal*** com la independència dos a dos. El concepte d'independència però té dues formes de generalització més naturals.

Direm que els esdeveniments A , B i C són ***independents***¹ si i només si compleixen:

- Són marginalment independents
- $P(A \cap B \cap C) = P(A)P(B)P(C)$.

Cal fer notar la necessitat de les dues condicions ja que ni la primera implica la segona, ni viceversa.

Exemple. Per veure que $P(A \cap B \cap C) = P(A)P(B)P(C)$ no implica que A , B i C siguin independents dos a dos, observem el següent exemple:

C	A	A^c	C^c	A	A^c
B	1/16	1/16	B	1/16	5/16
B^c	1/16	1/16	B^c	5/16	1/16

Per una banda tenim que $P(A \cap B \cap C) = 1/16$ i en canvi A i B no són independents, ja que $P(A) = P(B) = 1/2$ i $P(A \cap B) = 1/8$.

¹ Alguns autors prefereixen parlar de esdeveniments ***mútuament independents***.

Per veure que A , B i C independents dos a dos no implica $P(A \cap B \cap C) = P(A)P(B)P(C)$, l'exemple següent és clarificador:

C	A	A^c
B	0	1/4
B^c	1/4	0

C^c	A	A^c
B	1/4	0
B^c	0	1/4

El punt clau d'aquesta última afirmació és de fet que:

$$A \perp\!\!\!\perp B, A \perp\!\!\!\perp C \not\Rightarrow A \perp\!\!\!\perp (B \cap C).$$

En els dos exemples anteriors es pot comprovar. En el primer exemple no es verifica la suficiència i en el segon no es verifica la necessitat.

Tenim, per tant, que la independència implica la independència marginal però no és certa l'afirmació contrària.

La segona possibilitat de generalització és la següent. Un esdeveniment A és **independent respecte l'àlgebra generada pels** esdeveniments B i C , i ho escriurem $A \perp\!\!\!\perp [B, C]$, si i només si:

$$A \perp\!\!\!\perp (B \cap C), A \perp\!\!\!\perp (B \cap C^c), A \perp\!\!\!\perp (B^c \cap C) \text{ i } A \perp\!\!\!\perp (B^c \cap C^c).$$

Amb aquesta definició tenim que:

- $A \perp\!\!\!\perp [B, C] \iff A \perp\!\!\!\perp B, A \perp\!\!\!\perp C \text{ i } A \perp\!\!\!\perp (B \cap C)$. (vegeu pàg. 187).
- $A \perp\!\!\!\perp [B, C] \iff A \perp\!\!\!\perp E$, per tot esdeveniment, E , generat per unions, interseccions i pas al complementari dels esdeveniments B i C . (vegeu pàg. 187).

3. INDEPENDÈNCIA CONDICIONAL

Donarem dues definicions d'esdeveniments condicionalment independents.

Direm que dos esdeveniments A i B són **dèbilment independents condicionats per** C si i només si:

$$P(A \cap B / C) = P(A / C)P(B / C), \text{ on suposem que } P(C) > 0.$$

Escriurem aquest fet de la manera següent:

$$A \perp\!\!\!\perp B/C.$$

Aquesta relació entre esdeveniments és simètrica, $A \perp\!\!\!\perp B/C \iff B \perp\!\!\!\perp A/C$, però no és reflexiva ni transitiva.

Compleix també:

- $A \perp\!\!\!\perp B/C \iff A \perp\!\!\!\perp B^c/C$.
- $A \perp\!\!\!\perp B/C \iff A^c \perp\!\!\!\perp B/C$.
- $A \perp\!\!\!\perp B/C \iff A^c \perp\!\!\!\perp B^c/C$.

Atenció, però, al fet que $A \perp\!\!\!\perp B/C$ no implica $A \perp\!\!\!\perp B/C^c$, com es pot veure en l'exemple de la pàg. 173, ni tampoc, per simetria, $A \perp\!\!\!\perp B/C^c$ implica $A \perp\!\!\!\perp B/C$.

També hem de tenir en compte els següents resultats aparentment paradoxals:

- $A \perp\!\!\!\perp B/C$ i $A \perp\!\!\!\perp B/C^c \nRightarrow A \perp\!\!\!\perp B$.
- $A \perp\!\!\!\perp B \nRightarrow A \perp\!\!\!\perp B/C$ ó $A \perp\!\!\!\perp B/C^c$.

Aquests resultats queden il·lustrats per la paradoxa de Simpson.

Exemple. Dels molts exemples que s'han publicat escollim el del mateix E.H.Simpson [7]. Uns laboratoris volen provar un cert tractament. S'estudien 52 casos donant els resultats de la taula següent:

C	A	A^c
B	4/52	8/52
B^c	3/52	5/52

C^c	A	A^c
B	2/52	12/52
B^c	3/52	15/52

on C i C^c representen home i dona respectivament. A i A^c ens diu si ha rebut tractament o no i B i B^c si s'ha curat o ha mort.

A la vista d'aquestes dades veiem que el nou tractament és beneficiós tant per la població masculina com per la població femenina, mentre que si observéssim els resultats

globals, sense tenir en compte el sexe, veuríem que és independent, pel desenvolupament de la malaltia, el fet d'aplicar o no aquest nou tractament.

La segona definició és una versió forta de l'anterior i està relacionada amb la definició d'independència d' A respecte l'àlgebra generada per B i C .

Dos esdeveniments A i B són **fortament independents condicionats per l'esdeveniment** C , i ho escriurem $A \perp\!\!\!\perp B/[C]$, si i només si:

$$A \perp\!\!\!\perp B/C \text{ i } A \perp\!\!\!\perp B/C^c.$$

La definició $A \perp\!\!\!\perp B/[C]$ és més forta que $A \perp\!\!\!\perp B/C$ en el sentit de que la primera implica la segona però no al contrari.

Més en general, dos esdeveniments A i B són **fortament independents condicionats pels esdeveniments** C i D , i ho escriurem $A \perp\!\!\!\perp B/[C, D]$, si i només si:

$$A \perp\!\!\!\perp B/(C \cap D), A \perp\!\!\!\perp B/(C \cap D^c), A \perp\!\!\!\perp B/(C^c \cap D) \text{ i } A \perp\!\!\!\perp B/(C^c \cap D^c).$$

Per poder arribar a demostrar el teorema de separació, secció 7, necessitem el següent resultat anomenat **independència de blocs**, que ens relaciona la definició forta d'independència condicionada i la definició d'independència entre un esdeveniment i una àlgebra:

Sigui l'espai de probabilitat $(\Omega, \mathcal{A}, P)$. Suposant que P és estrictament positiva en la partició generada pels esdeveniments A , B i C , les propietats següents són equivalents:

1. $A \perp\!\!\!\perp [B, C]$.
2. $A \perp\!\!\!\perp [B]/[C]$ i $A \perp\!\!\!\perp [C]/[B]$. (vegeu pàg. 188).

4. EXTENSIÓ A VECTORS ALEATORIS

Direm que els vectors aleatoris X i Y són **independents**, $X \perp\!\!\!\perp Y$, si i només si la funció de densitat conjunta, f_{XY} , compleix:

$$f_{XY}(x, y) = f_X(x)f_Y(y).$$

O també:

$$X \perp\!\!\!\perp Y \iff f_{X/Y}(x; y) = f_X(x), \forall x.$$

Recordem que $f_{X/Y}(x; y) = \frac{f_{XY}(x, y)}{f_Y(y)}$ i que per tant hem de suposar que $f_Y(y) \neq 0$ quasi-per-tot.

Dos resultats importants per poder demostrar més endavant el teorema de separació són els següents:

1. Criteri de factorització.

Els vectors aleatoris X i Y són independents si i només si existeixen dues funcions g i h tals que:

$$f_{XY}(x, y) = g(x)h(y), \forall x \text{ i } \forall y. \quad (\text{vegeu pàg. 188}).$$

2. Lema de reducció.

Si (X, Y, Z) és un vector aleatori, aleshores:

$$X \perp\!\!\!\perp (Y, Z) \Rightarrow X \perp\!\!\!\perp Y \text{ i } X \perp\!\!\!\perp Z,$$

és a dir, la independència implica la independència marginal. (vegeu pàg. 188).

La implicació contrària no es verifica.

Vegem-ho en el següent exemple ²:

Exemple. Sigui el cub $\{(x, y, z); 0 < x < 1, 0 < y < 1, 0 < z < 1\}$ i la funció de densitat conjunta:

$$f_{XYZ}(x, y, z) = 2(\chi_1(x, y, z) + \chi_2(x, y, z) + \chi_3(x, y, z) + \chi_4(x, y, z))$$

on

$$\chi_1(x, y, z) = \begin{cases} 1 & \text{si } x < 0.5, y < 0.5 \text{ i } z < 0.5 \\ 0 & \text{en cas contrari} \end{cases};$$

²Aquest exemple així com d'altres de semblants tenen el seu origen en l'exemple de S. Bernstein, vegeu per exemple [4] pg. 126. Es pot trobar un exemple en el cas continu quelcom diferent en [5].

$$\chi_2(x, y, z) = \begin{cases} 1 & \text{si } x < 0.5, y > 0.5 \text{ i } z > 0.5 \\ 0 & \text{en cas contrari} \end{cases};$$

$$\chi_3(x, y, z) = \begin{cases} 1 & \text{si } x > 0.5, y < 0.5 \text{ i } z > 0.5 \\ 0 & \text{en cas contrari} \end{cases} \text{ i}$$

$$\chi_4(x, y, z) = \begin{cases} 1 & \text{si } x > 0.5, y > 0.5 \text{ i } z < 0.5 \\ 0 & \text{en cas contrari} \end{cases}$$

són les funcions indicadores de cubs d'aresta 1/2 situats dins del cub gran d'aresta unitat. En aquest cas tenim per $0 < x < 1$ i $0 < y < 1$ que $f_{XY}(x, y) = 1$ que és igual al producte $f_X(x)f_Y(y)$. Per tant $X \perp\!\!\!\perp Y$.

De la mateixa manera $X \perp\!\!\!\perp Z$. Però $f_X(x)f_{YZ}(y, z) = 1$ que no és igual a $f_{XYZ}(x, y, z)$.

El criteri de factorització dóna una caracterització de la independència i el lema de reducció permet escriure-la amb marginals.

Definim ara la independència condicional.

Els vectors aleatoris Y i Z són **independents condicionats pel** vector aleatori X , $Y \perp\!\!\!\perp Z/X$, si i només si la densitat conjunta dels vectors Y i Z condicionada a $X = x$ és igual al producte de les seves marginals:

$$f_{YZ/X}(y, z; x) = f_{Y/X}(y; x)f_{Z/X}(z; x), \forall y, z \text{ i } \forall x \text{ tal que } f_X(x) > 0.$$

O equivalentment:

$$f_{Y/XZ}(y; x, z) = f_{Y/X}(y; x), \text{ o}$$

$$f_{XYZ}(x, y, z) = f_{XY}(x, y) \frac{f_{XZ}(x, z)}{f_X(x)}.$$

De fet és una extensió de la definició forta d'independència condicional d'esdeveniments, $A \perp\!\!\!\perp B/[C]$, ja que és per tot x .

Generalitzem també el criteri de factorització, i els lemes de reducció i independència de blocs.

- Criteri de factorització.

$Y \perp\!\!\!\perp Z/X \iff \exists g, h$ funcions tals que:

$$f_{XYZ}(x, y, z) = g(x, y)h(x, z), \forall y, z \text{ i } \forall x \text{ amb } f_X(x) > 0. \quad (\text{vegeu pàg. ??}).$$

- Lema de reducció.

Siguin X, Y, Z_1, Z_2 vectors aleatoris, aleshores:

$$Y \perp\!\!\!\perp (Z_1, Z_2) / X \Rightarrow Y \perp\!\!\!\perp Z_1 / X \text{ i } Y \perp\!\!\!\perp Z_2 / X. \quad (\text{vegeu pàg. 189}).$$

- Independència de blocs.

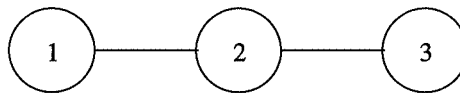
Siguin X, Y, Z_1, Z_2 vectors aleatoris i $f > 0$, aleshores les afirmacions següents són equivalents:

1. $Y \perp\!\!\!\perp (Z_1, Z_2) / X$.
2. $Y \perp\!\!\!\perp Z_1 / (X, Z_2)$, i $Y \perp\!\!\!\perp Z_2 / (X, Z_1)$.
3. $Y \perp\!\!\!\perp Z_2 / (X, Z_1)$, i $Y \perp\!\!\!\perp Z_1 / X$. (vegeu pàg. 189).

5. TEORIA DE GRAFS

Un **graf** és un parell (V, F) on V és un conjunt no buit de **vèrtexs** i F és una col·lecció de parells no ordenats de punts diferents de V , **arestes**.

Per exemple, el graf amb $V = \{1, 2, 3\}$ i $F = \{(1, 2), (2, 3)\}$ es pot representar com:



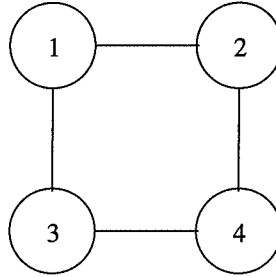
Un parell de vèrtexs, i, j , són **contigus** si $(i, j) \in F$.

Un **camí** és una seqüència de vèrtexs contigus. Un camí és un **cicle** si el primer i l'últim vèrtex coincideixen.

Dos vèrtexs i i j estan **connectats** si existeix un camí amb extrems i i j .

Un conjunt de vèrtexs S **separa** dos vèrtexs i, j si tot camí amb origen en i i final en j conté al menys un vèrtex de S .

Per exemple, en el graf:



el conjunt $S = \{1, 4\}$ separa els vèrtexs 3 i 2.

Un conjunt $S \subset V$ de vèrtexs *separa* dos conjunts V_1 i V_2 de vèrtexs si S separa tot parell de vèrtexs, un de V_1 i l'altre de V_2 .

Frontera d'un conjunt de vèrtexs V és el conjunt de tots els vèrtexs contigus a un vèrtex de V i que no són de V . Per exemple, en el graf anterior el conjunt $\{3, 2\}$ és la frontera del conjunt $V = \{1\}$.

6. GRAFS D'INDEPENDÈNCIA CONDICIONAL

Sigui $X = (X_1, \dots, X_k)$ un vector aleatori i sigui $V = \{1, \dots, k\}$ el corresponent conjunt de vèrtex.

Un graf (V, F) direm que és un *graf d'independència condicional* si

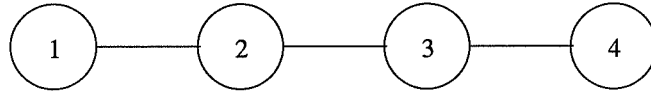
$$X_i \perp\!\!\!\perp X_j / X_{V - \{i, j\}} \iff (i, j) \notin F.$$

Exemple. Sigui $X = (X_1, X_2, X_3, X_4)$ amb les següents independències:

$$X_1 \perp\!\!\!\perp X_3 / \{X_2, X_4\}, X_1 \perp\!\!\!\perp X_4 / \{X_2, X_3\} \text{ i } X_2 \perp\!\!\!\perp X_4 / \{X_1, X_3\}.$$

El graf d'independència condicional vindrà donat per $F = \{(1, 2), (2, 3), (3, 4)\}$.

Gràficament:



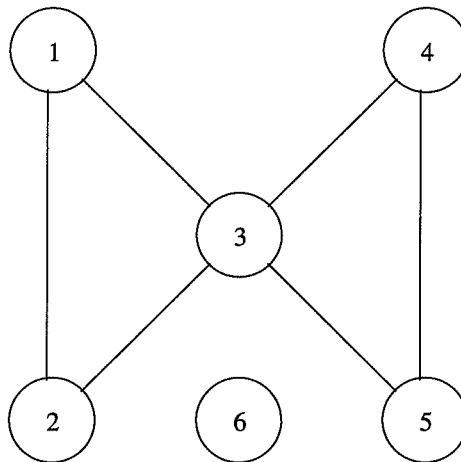
7. TEOREMA DE SEPARACIÓ

Enunciem en primer lloc el teorema.

Sigui $X = (X_1, \dots, X_k)$ un vector aleatori i sigui $V = \{1, \dots, k\}$ el corresponent conjunt de vèrtexs. Siguin a , b i c subconjunts disjunts de V . En aquestes condicions, si a separa b i c en el graf d'independència condicional (V, F) , aleshores $X_b \perp\!\!\!\perp X_c / X_a$.

La demostració d'aquest teorema es basa en l'aplicació dels lemes de reducció i independència de blocs. Per donar una idea del seu mecanisme vegem com aniria en un cas senzill.

Sigui el següent graf d'independència condicional:



on $a = \{3\}$, $b = \{1, 2\}$ i $c = \{4, 5\}$. És clar que a separa b i c .

Per definició de graf d'independència tenim:

$$X_1 \perp\!\!\!\perp X_4 / (X_2 X_3 X_5 X_6) \text{ i } X_1 \perp\!\!\!\perp X_6 / (X_2 X_3 X_5 X_4).$$

Ara pel lema d'independència de blocs amb $Y = X_1$, $Z_1 = X_4$, $Z_2 = X_6$ i $X = (X_2 X_3 X_5)$ tenim

$$X_1 \perp\!\!\!\perp (X_4 X_6) / (X_2 X_3 X_5)$$

i pel teorema de reducció arribem a

$$X_1 \perp\!\!\!\perp X_4 / (X_2 X_3 X_5).$$

De la mateixa manera podríem haver arribat a $X_1 \perp\!\!\!\perp X_6 / (X_2 X_3 X_4)$.

D'aquestes dues últimes afirmacions i aplicant el lema d'independència de blocs trobem que $X_1 \perp\!\!\!\perp (X_4 X_6) / (X_2 X_3)$.

També hauríem pogut trobar de la mateixa manera que $X_2 \perp\!\!\!\perp (X_4 X_6) / (X_1 X_3)$. Tornant aplicar el lema d'independència de blocs tenim que:

$$(X_1 X_2) \perp\!\!\!\perp (X_4 X_6) / X_3$$

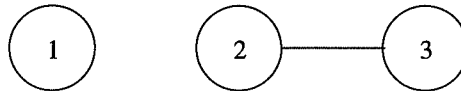
que acabaria la demostració.

Una demostració completa d'aquest teorema es pot trobar a [10].

Una conseqüència del teorema de separació és que alguna de les variables condicionants en una relació d'independència pot ser redundant.

Direm que una independència condicional entre un parell de variables és *minimal* si no es pot aplicar més el teorema de separació per eliminar variables condicionants. En definitiva, quan hem eliminat totes les variables redundants en una relació d'independència.

Exemple. Sigui el graf d'independència condicional:



Totes les afirmacions d'aquest graf són: $X_1 \perp\!\!\!\perp X_2/X_3$ i $X_1 \perp\!\!\!\perp X_3/X_2$. En canvi per aplicació del teorema de separació tenim: $X_1 \perp\!\!\!\perp X_2$ i $X_1 \perp\!\!\!\perp X_3$ que són a la vegada independències condicionals minimal.

8. PROPIETATS DE MARKOV

Presentem a continuació tres propietats que són equivalents, com demostrarem tot seguit, i que, segons la definició donada anteriorment, un graf és d'independència condicional si té la primera d'aquestes propietats. Són les propietats de Markov:

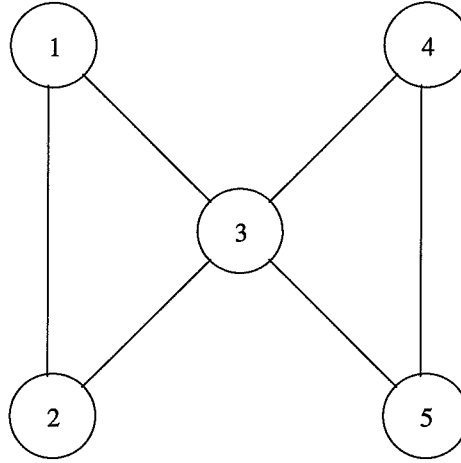
1. Propietat **dos a dos**: $\forall i, j$ no contigus, $X_i \perp\!\!\!\perp X_j/X_a$ on $a = V - \{i, j\}$.
 2. Propietat **global**: $\forall a, b, c$ subconjunts de V disjunts i b i c separats per a implica que $X_b \perp\!\!\!\perp X_c/X_a$.
 3. Propietat **local**: $\forall i$ i $a = \text{frontera}(i)$ i b els restants vèrtexs, tenim que $X_i \perp\!\!\!\perp X_b/X_a$.
- Les tres propietats de Markov són equivalents.

La demostració és com segueix:

- La propietat dos a dos implica la propietat global a causa del teorema de separació.
- La propietat global implica la propietat local pel fet que la frontera sempre és un conjunt separador.
- Falta demostrar que la propietat local implica la propietat dos a dos. Suposem un graf amb vèrtex $V = \{1 \dots k\}$ que satisfà la propietat local. En aquestes condicions, per cada vèrtex i tenim $X_i \perp\!\!\!\perp X_{V - (\{i\} \cup a)} / X_a$ on a és la frontera del vèrtex i .
 Siguin ara un vèrtex $j \in V - (\{i\} \cup a)$ i el conjunt $c = V - (\{i, j\} \cup a)$.
 Podem escriure $X_i \perp\!\!\!\perp (X_j, X_c) / X_a$ i pel lema d'independència de blocs això equival a $X_i \perp\!\!\!\perp X_j / (X_a, X_c)$ i $X_i \perp\!\!\!\perp X_c / (X_a, X_j)$.
 La primera afirmació ja ens dóna la propietat dos a dos ja que $a \cup c = V - \{i, j\}$

Vegem en un exemple aquesta equivalència.

Exemple. Sigui el següent graf d'independència condicional:



• Per la propietat dos a dos tenim:

- $X_1 \perp\!\!\!\perp X_4 / (X_2, X_3, X_5)$
- $X_1 \perp\!\!\!\perp X_5 / (X_2, X_3, X_4)$
- $X_2 \perp\!\!\!\perp X_4 / (X_1, X_3, X_5)$
- $X_2 \perp\!\!\!\perp X_5 / (X_1, X_3, X_5)$.

• Per la propietat global tenim:

- $X_1 \perp\!\!\!\perp X_4 / X_3$
- $X_2 \perp\!\!\!\perp X_4 / X_3$
- $X_1 \perp\!\!\!\perp X_5 / X_3$
- $X_2 \perp\!\!\!\perp X_5 / X_3$ és a dir
- $(X_1, X_2) \perp\!\!\!\perp (X_4, X_5) / X_3$.

• Per la propietat local tenim:

- $X_1 \perp\!\!\!\perp (X_4 X_5) / (X_2 X_3)$
- $X_2 \perp\!\!\!\perp (X_4 X_5) / (X_1 X_3)$
- $X_5 \perp\!\!\!\perp (X_1 X_2) / (X_3 X_4)$
- $X_4 \perp\!\!\!\perp (X_1 X_2) / (X_3 X_5)$.

9. EXEMPLE D'APLICACIÓ

Per acabar vegem un exemple d'aplicació de la teoria de grafs d'independència condicional.

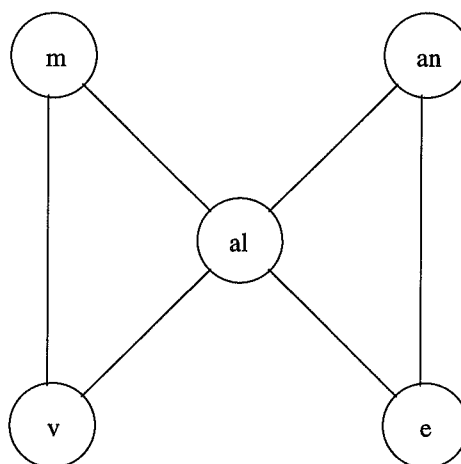
L'exemple està desenvolupat a [10] amb dades extretes de [6]. Les dades corresponen a les notes obtingudes per 88 alumnes en 5 assignatures: mecànica, vectors, àlgebra, anàlisi i estadística, respectivament.

La matriu de correlacions dóna:

$$\begin{pmatrix} 1 & & & & \\ 0.3371 & 1 & & & \\ 0.2294 & 0.2784 & 1 & & \\ \dots & \dots & 0.4336 & 1 & \\ \dots & \dots & 0.3554 & 0.2530 & 1 \end{pmatrix}$$

on els valors marcats amb punts suspensius són les correlacions simples no significatives, és a dir, aproximadament zero.

Una primera interpretació és simplement que cada parell de variables amb correlació parcial zero són independents condicionalment a les altres variables. Això ens porta a construir el següent graf d'independència condicional:



El gràfic és idèntic al de l'exemple de l'apartat 9 on ara, x_1 és mecànica, x_2 és vectors, x_3 és àlgebra, x_4 és anàlisi i x_5 és estadística.

Repasant l'exemple citat trobem que la propietat global de Markov ens permet escriure

$$[\text{mecànica, vectors}] \perp\!\!\!\perp [\text{anàlisi, estadística}]/\text{àlgebra}$$

la qual cosa ens porta a concloure que hi ha dos grups de variables, mecànica, vectors i àlgebra, per una banda, i anàlisi, estadística i àlgebra, per l'altra.

Pel que fa a la predicció de variables, per la propietat local de Markov, tenim:

- àlgebra i vectors són suficients per predir mecànica.
- mecànica i àlgebra són suficients per predir vectors.
- àlgebra i anàlisi són suficients per predir estadística.
- àlgebra i estadística són suficients per predir anàlisi.
- totes les variables són necessàries per predir àlgebra.

10. APÈNDIX

- $A \perp\!\!\!\perp B \iff A \perp\!\!\!\perp B^c$.

secció 2, pàgina 173.

— Cal notar tan sols que:

$$P(A \cap B^c) = P[A - (A \cap B)] = P(A) - P(A)P(B) = P(A)[1 - P(B)] = P(A)P(B^c).$$

- La relació $\perp\!\!\!\perp$ no és transitiva.

secció 2, pàgina 173.

— Tirem un dau. Siguin els esdeveniments següents:

A : «parell», B : «múltiple de tres» i C : «Imparell».

$A \perp\!\!\!\perp B$ i $B \perp\!\!\!\perp C$ però no és cert que $A \perp\!\!\!\perp C$.

- $A \perp\!\!\!\perp [B, C] \iff A \perp\!\!\!\perp B, A \perp\!\!\!\perp C \text{ i } A \perp\!\!\!\perp (B \cap C)$.

secció 2, pàgina 174.

Abans de provar-ho demostrem el següent resultat:

- $A \perp\!\!\!\perp B, A \perp\!\!\!\perp C \text{ i } B \text{ i } C \text{ disjunts} \Rightarrow A \perp\!\!\!\perp (B \cup C)$.

Amb les hipòtesis establertes tenim: $P(A \cap (B \cup C)) = P((A \cap B) \cup (A \cap C)) = P(A \cap B) + P(A \cap C) - P(A \cap B \cap C) = P(A \cap B) + P(A \cap C) = P(A)P(B) + P(A)P(C) = P(A)(P(B) + P(C)) = P(A)P(B \cup C)$.

Ara estem en condicions de demostrar el resultat.

- Per $\Rightarrow$ cal recordar que $B = (B \cap C^c) \cup (B \cap C)$, amb $B \cap C$ i $B \cap C^c$ disjunts. Igual per C.

Per veure $\Leftarrow$ hem de demostrar que

$A \perp\!\!\!\perp (B \cap C), A \perp\!\!\!\perp (B \cap C^c), A \perp\!\!\!\perp (B^c \cap C) \text{ i } A \perp\!\!\!\perp (B^c \cap C^c)$

Per exemple, $P(A)P(B^c \cap C) = P(A)(P(C) - P(B \cap C)) = P(A)P(C) - P(A)P(B \cap C) = P(A \cap C) - P(A \cap B \cap C) = P(A \cap B^c \cap C)$.

O també, $P(A)P(B^c \cap C^c) = P(A)(P(B^c) - P(B^c \cap C)) = P(A)P(B^c) - P(A)P(B^c \cap C) = P(A \cap B^c) - P(A \cap B^c \cap C) = P(A \cap B^c \cap C^c)$.

I així també $A \perp\!\!\!\perp (B \cap C^c)$.

- $A \perp\!\!\!\perp [B, C] \Rightarrow A \perp\!\!\!\perp E$, per tot esdeveniment E generat per unions, interseccions i pas al complementari dels esdeveniments B i C .

secció 2, pàgina 174.

De la definició i de la propietat anterior ens queda tan sols demostrar-ho en el cas $E = B \cup C, E = B \cup C^c, E = B^c \cap C$ i $E = B^c \cap C^c$.

- $A \perp\!\!\!\perp (B \cup C)$.

De $A \perp\!\!\!\perp B, A \perp\!\!\!\perp C$ i $A \perp\!\!\!\perp (B \cap C)$ podem escriure:

$P(A \cap (B \cup C)) = P((A \cap B) \cup (A \cap C)) = P(A \cap B) + P(A \cap C) - P(A \cap B \cap C) = P(A)P(B) + P(A)P(C) - P(A)P(B \cap C) = P(A)P(B \cup C)$.

- $A \perp\!\!\!\perp (B^c \cap C)$.

La mateixa idea però tenint en compte $A \perp\!\!\!\perp B^c, A \perp\!\!\!\perp C$ i $A \perp\!\!\!\perp (B^c \cap C)$.

- Igual per les restants independències.

- *Lema d'independència de blocs.*

secció 3, pàgina 176.

$$- (1) \Rightarrow (2). \text{ Tenim: } P(A \cap B/C) = \frac{P(A \cap B \cap C)}{P(C)} = \frac{P(A)P(B \cap C)}{P(C)} = P(A)P(B/C) = P(A/C)P(B/C), \text{ és a dir } A \perp\!\!\!\perp B/C.$$

De la mateixa manera es demostraria $A \perp\!\!\!\perp B^c/C$, $A \perp\!\!\!\perp B/C^c$ i $A \perp\!\!\!\perp B^c/C^c$, amb el que tindríem $A \perp\!\!\!\perp [B]/[C]$. De manera semblant arribaríem a veure $A \perp\!\!\!\perp [C]/[B]$.

$$- (2) \Rightarrow (1). \text{ De (2) tenim que d'una banda, } P(A \cap B \cap C) = P(A \cap B/C)P(C) = P(A/C)P(B/C)P(C) = P(A/C)P(B \cap C), \text{ i d'altra banda } P(A \cap B \cap C) = P(A \cap C/B)P(B) = P(A/B)P(C/B)P(B) = P(A/B)P(B \cap C).$$

Així doncs $P(A/C) = P(A/B)$.

De la mateixa manera i partint de $P(A \cap B^c \cap C)$ arribem a $P(A/C) = P(A/B^c)$, per tant $P(A/B^c) = P(A/B)$.

Això ens permet veure que $P(A) = P(A/B)P(B) + P(A/B^c)P(B^c) = P(A/B)P(B) + P(A/B)P(B^c) = P(A/B)[P(B) + P(B^c)] = P(A/B)$ per tant $A \perp\!\!\!\perp B$.

De la mateixa manera veuríem $A \perp\!\!\!\perp C$.

Finalment, de $P(A \cap B/C) = P(A/C)P(B/C)$, en ser $A \perp\!\!\!\perp B/C$, tenim $P(A \cap B \cap C) = P(A/C)P(B \cap C) = P(A)P(B \cap C)$, el que demostra que $A \perp\!\!\!\perp (B \cap C)$.

- *Criteri de factorització.*

secció 4, pàgina 177.

- La necessitat passa per fer $f_X = g$ i $f_Y = h$. Provem la suficiència.

De $f_{XY}(x, y) = g(x)h(y)$, integrant els dos costats sobre y , tenim $f_X(x) = g(x)k$. De la mateixa manera, integrant sobre x : $f_Y(y) = h(y)k'$, d'on $f_{XY}(x, y) = k''f_X(x)f_Y(y)$; però $k'' = 1$ en integrar sobre x i y .

- *Lema de reducció.*

secció 4, pàgina 177.

- $X \perp\!\!\!\perp (Y, Z) \Rightarrow f_{XYZ}(x, y, z) = f_X(x)f_{YZ}(y, z)$. Integrant els dos costats sobre z tenim:

$$f_{XY}(x, y) = f_X(x)f_Y(y), \text{ per tant } X \perp\!\!\!\perp Y.$$

De la mateixa manera però integrant sobre y es demostraria $X \perp\!\!\!\perp Z$.

- *Criteri de factorització.*

secció 4, pàgina 178.

- Per veure la necessitat cal fer $g(x, y) = f_{XY}(x, y)$ i $h(x, z) = f_{XZ}(x, z)/f_X(x)$.
- Per veure la suficiència integrem els dos costats de la igualtat $f_{XYZ}(x, y, z) = g(x, y)h(x, z)$ primer sobre y i després sobre z . Així tenim: $f_{XZ}(x, z) = h(x, z)k(x)$ i $f_{XY}(x, y)g(x, y)k'(x)$, per tant:
 $f_{XYZ}(x, y, z) = f_{XZ}(x, z)f_{XY}(x, y)k''(x)$, però integrant aquesta igualtat sobre y i z alhora tenim que $k''(x) = 1/f_X(x)$.

- *Lema de reducció.*

secció 4, pàgina 179.

- De $Y \perp\!\!\!\perp (Z_1, Z_2)/X$ tenim que $f_{XYZ_1Z_2}(x, y, z_1, z_2) = f_{XY}(x, y)f_{XZ_1Z_2}(x, z_1, z_2)/f_X(x)$. Ara integrant sobre z_1 tenim: $f_{XYZ_2}(x, y, z_2) = f_{XY}(x, y)f_{XZ_2}(x, z_2)/f_X(x)$, que implica $Y \perp\!\!\!\perp Z_2/X$. De la mateixa manera sortiria $Y \perp\!\!\!\perp Z_1/X$.

- *Lema d'independència de blocs.*

secció 4, pàgina 179.

- (1) $\Rightarrow$ (2). Pel criteri de factorització sabem que si $Y \perp\!\!\!\perp (Z_1, Z_2)/X$ aleshores tenim $f(x, y, z_1, z_2) = g(x, y)h(x, z_1, z_2)$.
 Cal observar tan sols ara que a la dreta de l'expressió no hi ha cap funció que tingui a la y i a les (z_1, z_2) com arguments. En particular no hi ha cap funció que tingui com arguments y i z_1 a la vegada. Aplicant un altre cop el criteri de factorització deduíem que $Y \perp\!\!\!\perp Z_1/(X, Z_2)$. De la mateixa manera demostrariem que $Y \perp\!\!\!\perp Z_2/(X, Z_1)$.
- (2) $\Rightarrow$ (1). De $Y \perp\!\!\!\perp Z_1/(X, Z_2)$ tenim que
 $f_{XZ_2YZ_1}(x, z_2, y, z_1) = g'(x, z_2, y)h'(x, z_2, z_1)$ però de $Y \perp\!\!\!\perp Z_2/(X, Z_1)$, $Y \perp\!\!\!\perp Z_2/X$ podem escriure $g'(x, z_2, y) = g''(x, y)g'''(z_2, x)$ per tant
 $f_{XZ_2YZ_1}(x, z_2, y, z_1) = g''(x, y)g'''(z_2, x)h'(x, z_2, z_1) = g''(x, y)h''(x, z_1, z_2)$
 i això implica (1).
- (2) $\Rightarrow$ (3). Resulta del lema de reducció.
- (3) $\Rightarrow$ (1). De $Y \perp\!\!\!\perp Z_2/(X, Z_2)$ tenim que
 $f_{XYZ_1Z_2} = f_{XYZ_1} \frac{f_{XZ_1Z_2}}{f_{XZ_1}}$ però com $Y \perp\!\!\!\perp Z_1/X$ tenim $f_{XYZ_1Z_2} = f_{XY}f_{XZ_1}f_{XZ_1Z_2}/f_X$ que vol dir $Y \perp\!\!\!\perp (Z_1, Z_2)/X$.

11. BIBLIOGRAFIA

- [1] **Darroch, J.N., Lauritzen, S.L. i Speed, T.P.** (1980). «Markov fiels and log linear interaction models for contingency tables». *Ann. Stat.*, **8**, 522-539.
- [2] **Edwards, D.** (1995). *Introduction to Graphical Modelling*. Springer-Verlag.
- [3] **Edwards, D.E. i Kreiner, S.** (1983). «The analysis of contingency tables by graphical models». *Biometrika*, **70**, **3**, 553-565.
- [4] **Feller, W.** (1968). *An Introduction to Probability Theory and its Applications*, Volume 1. John Willey & Sons.
- [5] **Geisser, S. i Mantel, N.** (1962). «Pairwise independence of jointly dependent variables». *Ann. Math. Statist.*, **33**, 290.
- [6] **Mardia, K.V., J. Kent and J. Bibby** (1979). *Multivariate Analysis*. Academic Press Limited.
- [7] **Simpson, E.H.** (1951). «The Interpretation of Interaction in Contingency Tables». *Journal of the Royal Statistical Society*, **B-13**, **2**, 238-241.
- [8] **Speed T.P.** (1978). «Graph-teoretic methods in the analysis of interaction». *Lecture Notes*, Institute of Mathematical Statistics, University of Copenhagen.
- [9] **Wermuth, N. i Lauritzen, S.L.** (1983). «Graphical and recursive models for contingency tables». *Biometrika*, **70**, **3**, 537-552.
- [10] **Whittaker, J.** (1990). *Graphical Models in Applied Multivariate Statistics*. John Wiley & Sons.

ENGLISH SUMMARY

INDEPENDENCE GRAPHS MODELS

J.M. DURAN RÚBIES

Universitat Autònoma de Barcelona*

Independence Graphical Models are a tool of the Multivariate Analysis which uses graphs to represent models. Independence graphs specially sum up and clarify the interactions among variables; interactions which are not always easy to interpret, specially when three or more take part.

This work is of pedagogical interest and gives an introduction to the Independence Graphs theory, starting with notions of independence and getting to the Markov's properties, keys to the interpretations of these graphs.

It also studies an application example using numerical variables that shows the feasibility of this theory when it comes to simplifying and clarifying the relations among variables.

Keywords: Independence, conditional independence, Simpson's paradox, random vectors, graph theory, conditional independence graphs, Markov properties.

AMS Classification: 62 H 20

*J.M. Duran Rúbies. Professor associat U.A.B. Departament de Matemàtiques. 08193 Bellaterra (Espanya).
E-mail: duran@mat.uab.es

–Received November 1996.

–Accepted December 1998.

The independence graph theory is quite new. One of the first to relate the graph theory methods with the interactions among variables was Speed [8]. Conditional independence graphs sum up and clarify the relations between three or more variables, relations which are not easy to see nor to interpret.

This work is of pedagogical interest and gives an introduction to the independence graphs theory and it is an approach for all those who, without having to look at books, such as Whittaker [10] and Edwards [2], would like to have a graphical models early idea.

Firstly we study Independence notions (the marginal independence definitions, the mutual independence) and conditional independence (the weak conditional independence, the strong conditional independence, the block independence) which are needed to develop the theory. An example is introduced to show the Simpson's paradox which is useful to clarify some results. To be able to show the separation theorem later, previous concepts are extended to random vectors and important results are given, like factorisation criterion and reduction lemma.

Next we study the basic notions of the graph theory in order to define what it is understood as conditional independence graph. The separation theorem, that allows us to detect and remove redundant variables among conditioning variables in an independent relation, is shown through an example.

Further on, Markov properties are given and their equivalence is demonstrated allowing us to interpret Independence graphs. An application example with numerical variables is shown. In this example and from the correlation matrix a conditional independence graph is built, allowing us, thanks to the global Markov property, to sum up and clarify the conditional independence among variables. In the same way, and thanks to the local Markov property, the prediction of variables also remains clarify.

Demonstrations of some of the results established in the text are assembled in the appendix.

SOLUCIONS ALS PROBLEMES PROPOSATS AL VOLUM 22 N. 3

PROBLEMA N. 72

La función característica de la distribución uniforme en el intervalo $(-1, +1)$ es

$$\int_{-1}^{+1} \frac{1}{2} e^{itx} dx = \frac{1}{2} \left[\frac{e^{itx}}{it} \right]_{-1}^{+1} = \frac{e^{it} - e^{-it}}{2it} = \frac{\text{sen}(t)}{t}$$

Si $\varphi(t)$ es la función característica de X , entonces $\varphi(-t)$ es la función característica de $-Y$ y la función característica de $X - Y$ es

$$\varphi_{X-Y}(t) = \varphi(t) \varphi(-t)$$

por ser X, Y independientes.

Si existe X tal que $X - Y$ es uniforme en $(-1, +1)$, donde Y sigue la misma distribución que X y es independiente de X , entonces necesariamente

$$\alpha < X < \alpha + 1$$

para alguna constante α . Pero entonces $X - \alpha$ tiene la misma propiedad, así que podemos suponer que el soporte de X es

$$0 < X < 1$$

Se verifica pues

$$\varphi(t) \varphi(-t) = \frac{\text{sen}(t)}{t}$$

Pero

$$\begin{aligned} \varphi(t) \varphi(-t) &= (E \cos(Xt) + i \text{sen}(Xt)) (E \cos(Xt) - i \text{sen}(Xt)) \\ &= (E \cos(Xt))^2 + (E \text{sen}(Xt))^2 \\ &= \frac{\text{sen}(t)}{t} \end{aligned}$$

y para $t = \pi$ es $\text{sen}(\pi)/\pi = 0$

$$(E \cos(X\pi))^2 + (E \text{sen}(X\pi))^2 = 0$$

Luego

$$E(\operatorname{sen}(X\pi)) = 0$$

Sin embargo, el soporte de $\operatorname{sen}(X\pi)$ es

$$0 < \operatorname{sen}(X\pi) < 1$$

y su esperanza no puede ser 0. Luego la variable aleatoria X no puede existir.

C.M. Cuadras
Universitat de Barcelona

PROBLEMA N. 73

Si $H(x, y)$ es la función de distribución conjunta de (X, Y) , la distribución de $(X, G^{-1}(1 - G(Y)))$ es

$$\begin{aligned} H^*(x, y) &= P(X \leq x, G^{-1}(1 - G(Y)) \leq y) \\ &= P(X \leq x, Y > G^{-1}(1 - G(y))) \\ &= F(x) - H(x, G^{-1}(1 - G(y))) \end{aligned}$$

pues

$$[X \leq x] = [X \leq x, Y \leq G^{-1}(1 - G(y))] \cup [X \leq x, Y > G^{-1}(1 - G(y))]$$

Si $H = H^- = \max\{F + G - 1\}$ entonces

$$\begin{aligned} H^*(x, y) &= F(x) - H^-(x, G^{-1}(1 - G(y))) \\ &= F(x) - \max\{F(x) + GG^{-1}(1 - G(y)) - 1, 0\} \\ &= F(x) - \max\{F(x) - G(y), 0\} \\ &= \begin{cases} F(x) - (F(x) - G(y)) = G(y) & \text{si } F(x) \geq G(y) \\ F(x) - 0 = F(x) & \text{si } F(x) < G(y) \end{cases} \\ &= \min\{F(x), G(y)\} = H^+(x, y) \end{aligned}$$

La demostración de la segunda relación

$$H^+(x, y) = G(y) - H^-(F^{-1}(1 - F(x)), y)$$

es similar.

C.M. Cuadras
Universitat de Barcelona

PROBLEMES PROPOSATS

PROBLEMA N. 74

Demostrar por inducción en n la siguiente identidad

$$\frac{1}{n} \sum_{i=1}^n (x_i - \bar{x})^2 = \frac{1}{n^2} \sum_{i=1}^{n-1} \sum_{j=i+1}^n (x_i - x_j)^2,$$

donde

$$\bar{x} = \frac{1}{n} \sum_{i=1}^n x_i.$$

M. Ruíz Espejo
U.N.E.D.

PROBLEMA N. 75

Demostrar algebraicamente que el estimador jackknife de la varianza de la media muestral $\bar{x}_n = (1/n) \sum_{i=1}^n x_i$, es:

$$\hat{V}_J(\bar{x}_n) = \frac{1}{n(n-1)} \sum_{i=1}^n (x_i - \bar{x}_n)^2.$$

M. Ruíz Espejo
U.N.E.D.

Novetats de Software

FORCE4/R, A NEW SOFTWARE PRODUCT FOR FORECASTING AND SEASONAL ADJUSTMENT

A. PRAT

I. SOLÉ

J.M. CATOT

Universitat Politècnica de Catalunya*

J. LORÉS

Universitat de Lleida**

In this paper, we briefly describe the functions and the software architecture of FORCE4/R, a time series software that allows easy interfacing with already existing software: TRAMO and SEATS, X12, SCA, and also with statistical software developed at UPC (Polytechnical University of Catalonia) such as IDAUT, IDAFT, CONJUNCTURE and others. In addition, the results of an extensive simulation study to evaluate the quality of FORCE4/R are presented and discussed.

FORCE4/R has also been partially funded by the Spanish government under CICYT Grant TIC96-1310-CE, with the Institut d'Estadística de Catalunya (IDESCAT) acting as «Ente Observador».

* Universitat Politècnica de Catalunya (UPC), Departament d'Estadística i Investigació Operativa, Av. Diagonal 647 6^a planta, 08028 Barcelona.

** Universitat de Lleida (UdL), Departament d'Informàtica i Enginyeria Industrial, Lleida.

1. FORCE4/R AND THE FORECASTING METHODS

FORCE4/R is a forecasting, modelling and signal extraction software which is based mostly on the Box-Jenkins methodology [see Box et al. (1994)]. It is part of FORCE4, a forecasting software system that comprises a wide variety of methods (http://www-tqg.upc.es/seccio_tqg/projectes/force4/), which has been developed over the last two and a half years under ESPRIT IV project number 20.704. FORCE4/R allows for manual or automatic modelling of one or more time series and is capable of dealing with outliers, calendar effects, Easter effects, intervention analysis, model based seasonal adjustment, signal extraction, conjuncture analysis and transfer function models.

2. FUNCTIONS OF THE FORCE4/R SYSTEM

The main goal of the system is to allow the end user to model and forecast time series using the best available methods, and taking advantage of the experience that professional users have incorporated into the system.

In what follows, we describe the main functions of FORCE4/R.

2.1. Data base management

The Data base system groups the data by domains. A domain is composed of a set of time series of the same periodicity. It is possible to update, import, etc., time series, and also to access external databases with ODBC.

2.2. Automatic univariate modelling

The system will automatically find the SARIMA model $(p, d, q) (P_1, D_1, Q_1)_{s_1} \times (p_2, D_2, Q_2)_{s_2} \times (P_3, D_3, Q_3)_{s_3} \times (P_4, D_4, Q_4)_{s_4}$ that best fits a given time series.

This function can be performed with two different methods. One is based on software developed at UPC (IDAUT) which identifies a tentative model within the SARIMA class using a Mahalanobis type of distance between the theoretical ACF and IACF and the sample equivalents together with a search algorithm. The degree of differencing is also identified. The second method uses the automatic modelling capabilities of TRAMO.

Estimation of parameters and diagnostic checking can also be carried out either by using SCA and an expert system with rules implemented by UPC in KAPPA, or using TRAMO. Some of the parameter rules can be modified by the expert user.

2.3. Possibility of the automatic elimination of outliers

Additive Outlier (AO), Temporary Change (TC), and Level Shift (LS) type outliers can be eliminated automatically by SCA or TRAMO.

2.4. Intervention analysis

The impact of external events can be determined. These events may be of different types: labour strike, promotions, price increases, etc. This is done by the introduction of binary variables with various filters that will affect the time series (intervention analysis).

2.5. Transfer function analysis

The system automatically identifies, estimates and verifies the transfer function model between the time series Y_t and X_t .

2.6. Manual modelling

The system allows manual input of any univariate model or intervention model that the user may wish.

2.7. Forecasting

The system performs forecasting using any one of the previously adjusted models, and batch forecasting of a group of series. It also monitors the forecasting errors of the model by means of CUSUMS of the forecast error and EWMA.

2.8. Seasonal adjustment

The system allows the decomposition of a time series into trend, seasonal component and noise using the adjusted model. For this, the user can choose between X12/ARIMA and SEATS.

2.9. Conjuncture analysis

This analysis is performed following the methodology developed by Espasa (1993).

First the system performs seasonal adjustment using X12 or SEATS, and then it calculates relevant statistics for conjuncture analysis, such as underlying trend, conditional and unconditional rate of growth, inertia, etc. and allows the user to visualise the results in different graphic settings.

2.10. Configuration

It offers a wide range of possibilities with regard to definition of the methods and statistical package to be used, the parameters of the rules of the expert system, etc.

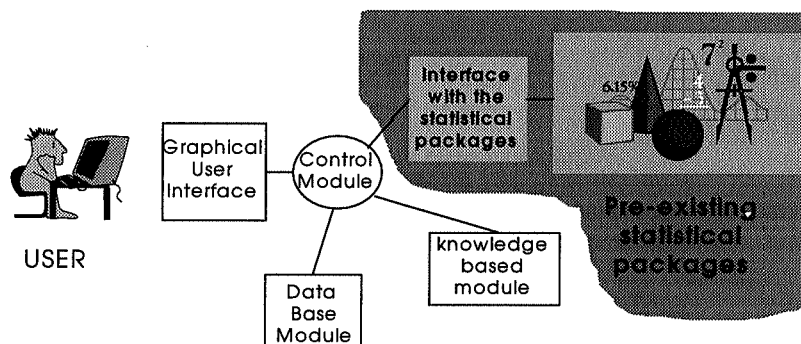
Terminology on time series analysis can be found in Box et al. (1994). Some references on statistical packages are given at the end of this report.

3. ARCHITECTURE OF FORCE4/R

FORCE4/R has been designed with the main objective of separating the user interface from its underlying application subsystem. This provides a flexible framework that allows substitution of some of the statistical packages when desired, together with straightforward, easy introduction of new functions.

With this intention in mind, the ARMAG architecture [see Prat et al. (1994)] was adopted as the basis for the development of the FORCE4/R architecture, and the different components that will compose the complete system were identified and defined.

The following figure shows the FORCE4/R architecture and the components of the system.



In the figure above, we can distinguish:

- The GUI, Graphical User Interface, is the software component responsible for the management of the surface-level human-computer interface, which also includes the strategy of the system.
- The Communications Control and Logging Agent (or Control Module) deals with the management and co-ordination of communications between all the components, which is achieved through a messaging protocol.
- The BEM [Back-End Manager, see Prat et al. (1992)] interfaces the pre-existing statistical packages integrated within the system. The statistical packages used in the FORCE4/R system for performance of the statistical functions are: IDAUT, IDAFT, CONJUNCTURE, SEATS, TRAMO, SCA and X12/ARIMA.
- The knowledge based component contains a set of rules which are used to analyse the results of some of the statistical packages used as back-ends, applying the knowledge of the experts consultants. The KAPPA shell (Intellicorp) is used for this function.
- Data Base Management: developed with ACCESS (Microsoft), although an ODBC connection will be implemented in a second phase.

The link between all the different modules is provided by the nucleus of the ARMAG architecture, which gives the user the impression that he is working with a single system. In other words, FORCE4/R is an example of the creation of a sophisticated system from different and separable elements, using the most appropriate tool available on the market for each component. Developing and/or customising a component is the responsibility of the expert associated with it.

The system is open in the sense that it is possible to specify the desired models instead of using the ones selected by the system, to import/export series, to change certain rule parameters, and in the case of professional users, it is also possible to add or change rules.

Finally, for the «developer»-user, it is possible to add new methods and/or statistical packages, modify interface knowledge with the packages (BEM) and modify the user interface (programmed in Visual Basic).

4. EVALUATION OF THE SYSTEM. SIMULATION STUDY

4.1. Testing and evaluation

Several approaches were followed for the testing and evaluation of the FORCE4/R system:

- In-house testing and evaluation
- Beta-test users
- Presentation to large audiences

The in-house testing and evaluation comprised two different procedures:

- Normal test procedure
- Simulation study

The first procedure involved creation of a test-bed set for the testing of the successive versions of the system. Several time series from the literature and others that were carefully modelled by the experts at UPC were used to check the incremental versions of the system.

The second procedure, described in detail in the following sections, consisted of rigorous work on designing and executing a test involving the simulation of approximately 20,000 series following different ARIMA models, and then checking the results obtained for the different methods included in the system.

The beta-test approach involved search and selection of beta-test users. We tried to find a broad spectrum of different kinds of beta-test users, covering institutions, services and industry organisations. In general, the beta-test users selected analysed a large number of series, and provided us with very useful feedback.

It was necessary to set up a procedure for compiling and recording problems and suggestions, deciding which ones were most important, and re-designing, if necessary, the system. In some cases, re-design was considered appropriate.

4.2. General description of the simulation study

The aim of the study was to compare the two time series modelling methods available in the FORCE4/R software (TRAMO-SEATS and IDAUT-SCA).

The procedure consisted in simulating univariate time series following predetermined ARIMA models, and then assessing the methods' performance by comparing the output model with the original, according to several statistical measures.

The simulation process is carried out by an algorithm which runs the SCA software in order to obtain a vast range of ARIMA models by simulation. The SCA takes as inputs the mean and variance (as a percentage of the mean) of white noise process and the parameter values of a chosen ARIMA model and simulates a time series of a chosen number of time series data.

The time series that are simulated belong to a range of ARIMA models with the following maximum orders: $(p=3, d=2, q=3) \times (P=2, D=1, Q=2)_{12}$. Given a specific ARIMA model, different parameter values are used to simulate several time series: four different values belonging to $|\phi_1| < 1$ for either first-order autoregressive and moving-average processes; for second-order autoregressive and moving-average processes, four pairs of values (ϕ'_1, ϕ'_2) belonging to the admissibility region, two of them corresponding to real roots and the other two to complex ones; and for third-order autoregressive and moving-average processes, the range of values is the result of all the possible combinations of the values above: $(\phi''_1, \phi''_2, \phi''_3) = (\phi_1 \times (\phi'_1, \phi'_2))$, therefore there are sixteen possible values. However, when simulating ARIMA models:

$$\phi_p(B) \Phi_P(B^{12}) \nabla^d \nabla_{12}^P Z_t^\lambda = \theta_0 + \theta_q(B) \Theta_Q(B^{12}) a_t$$

the parameter values of the two terms (autoregressive and moving-average) must not coincide, otherwise the order of the simulated model would become smaller.

On the other hand, the algorithm simulates time series with and without constant term θ_0 and with and without the logarithmic transformation of the original data.

The next step consists of modelling all the simulated time series with the two alternative softwares and obtaining a model for each series. The models are then saved in a data base.

Finally, the accuracy of the different methods is assessed by comparing several measures of forecasting performance, such as the percentage mean absolute error (PMAE), the root mean squared error (RMSE) [García-Ferrer et al. (1997) and Ruiz et al. (1996)] and the distance between the π weights [Piccolo (1990)].

The first results of this study show that:

1. The model obtained by TRAMO and by IDAUT plus SCA is a valid model in approximately 90% of the time series.
2. The PMAE and RMSE obtained by the two methods is at most 5% greater than the corresponding values obtained using the true model in about 65% of the cases (65%-60%).
3. TRAMO execution is five times faster than that of IDAUT + SCA.

5. CONCLUSIONS

If it is to be competitive, European industry needs systems that provide rigorous forecasts of demand. Forecasting researchers and professionals need to optimize their

time in performance of their normal work. The system presented in this paper makes it possible to use sophisticated methods and powerful statistical packages with a high degree of ease, and also to intervene at a low level if desired. Various organisations that acted as beta-testers in a validation study carried out during the last stages of the FORCE4 project (FIAT, HESPERIA, Caixa d'Estalvis de Catalunya, Institut d'Estadística de Catalunya, Instituto Vasco de Estadística and others) agreed with this assertion.

The results of the simulation study comparing the different methods included in the system show it to be reliable.

It is planned that different forecasting systems will evolve from this open expert system. One of these is FORENET [Prat et al. (1997)], a commercial system, having a back-end (an algorithm) based on application of the genetic algorithm and neural network techniques, and distributed in Spain and South America by BBS (Best Business Systems).

Technical details and licences

The system runs on the Windows platform (Windows 3.1, Windows 95 and Windows NT). It requires at least a Pentium 200, 32 Mbytes of RAM memory and 10 Mbytes of hard disk space.

A licence for the use of the system can be obtained from UPC (at the same address as given at the head of this paper). This licence covers the cost of the different packages and authorises use of the system in two different computers.

Statistical package references

IDAUT: Identifies the model using the distance between the ACF and IACF of the theoretical model and those of the given time series [M. Valls i Colom (1983)].

IDAFT: A system that allows identification of the transfer function based on the Corner Table, computation of the standard deviation, and determination of the elements that are statistically significant, by using the recognition algorithms to determine the transfer function. Developed at UPC.

CONJUNCTURE: Developed at UPC following the methodology of A. Espasa (1993).

KAPPA: A shell for the development of a knowledge based system, from Intellicorp.

SCA: Statistical System. Software for Forecasting and Time Series Analysis, from Scientific Computing Associates.

SEATS: Signal Extraction in Arima Time Series, from V. Gómez and A. Maravall (1992).

TRAMO: Time Series Regression with Arima Noise Missing Observations and Outliers, from V. Gómez and A. Maravall (1992).

X12/ARIMA: Decomposition of a time series into trend, seasonal component and noise, from US Bureau of the Census.

REFERENCES

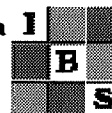
- [1] **Box, G.E., Jenkins, G.M. and Reinsel, G.C.** (1994). *Time series analysis, forecasting and control*. Third Edition. Holden-Day, San Francisco.
- [2] **Espasa, A.** (1993). *Métodos cuantitativos para el análisis de la coyuntura económica*. Alianza Editorial, Jose Ramón Cancelo (Ed.), Madrid.
- [3] **García-Ferrer, A., Del Hoyo, J. and Martín-Arroyo, A.S.** (1997). «Univariate Forecasting Comparisons: The Case of the Spanish Automobile Industry», *Journal of Forecasting*, Jan 1997.
- [4] **Gómez, V. and Maravall, A.** (1992). «Time Series Regression with Arima Noise and Missing Observations - Program TRAMO», *EUI Working Paper ECO No. 92/81*, Dept. of Economics, European University Institute.
- [5] **Piccolo, D.** (1990). «A distance measure for classifying ARIMA models», *Journal of Time Series Analysis Vol. II*, 2, 1990.
- [6] **Prat, A., Catot, J.M., Lorés, J., Galmes, J., Fletcher, P. and Sanjeevan, K.** (1992). «A separable architecture for the construction of knowledge based front ends», *AI Communications*, 5, 4, pp. 184-191.
- [7] **Prat, A., Lorés, J., Sanjeevan, K., Solé, I. and Catot J.M.** (1994). «Application of a multiagent distributed architecture for time series forecasting», Ed. Dutter and Grossmann, Physica-Verlag, Proceedings of Computational Statistics (*COMPSTAT 1994*), Vienna.
- [8] **Prat, A. and Delgado, A.** (1997). «Forenet: a Forecasting System based on Neural Networks», *VIII International Symposium on Applied Stochastic Models and Data Analysis*, Anacapri, June 1997.
- [9] **Ruíz, E. and Lorenzo, F.** (1996), «Which univariate time series model predicts quicker a crisis? The Iberia case». *Working Paper 96-51*, July 1996, Universidad Carlos III de Madrid.
- [10] **Valls i Colom, M.** (1983). *Identificació automàtica de sèries temporals*, PhD thesis summary, UPC.

Ressenyes d'activitats institucionals

Sociedad Española de Biometría



Región Española
de la Sociedad Internacional de Biometría
Spanish Region
of the International Biometric Society



<http://www.iata.csic.es/ibsresp>

La Sociedad Española de Biometría/Región Española de la Sociedad Internacional de Biometría (abreviadamente **SEB** o **REsp**) tiene como objetivos promover, impulsar y difundir el desarrollo y la aplicación de los métodos matemáticos y estadísticos a la biología, medicina, psicología, farmacología, agricultura y otras ciencias afines (ciencias relacionadas con los seres vivos). Cualquier profesional o alumno de estas disciplinas puede ser miembro de la SEB.

Consejo Directivo

<i>Presidente:</i>	Guadalupe Gómez i Melis (Biología)
<i>Vicepresidente:</i>	Emilio A. Carbonell Guevara (Agronomía)
<i>Secretario y Tesorero:</i>	Fernando López Santoveña (Agronomía)
<i>Vocal en calidad de</i>	
<i>Miembro del Consejo de la IBS:</i>	Carles M. Cuadras Avellana (Biología)
<i>Vocales:</i>	María Jesús Bayarri García (Medicina)
	Juan Luis Chorro Gascó (Psicología)
	Rosa Estarrelles Rodríguez (Psicología)
	José Luis González Andújar (Agronomía)
	Martín Ríos Alcolea (Medicina)
	Alex Sánchez Pla (Biología)
<i>Corresponsal de la REsp en el</i>	
<i>«Biometric Bulletin» de la IBS:</i>	María Luz Calle Rosingana

La SEB promueve directamente o participa en la promoción de cursos monográficos sobre distintas técnicas de Análisis Estadístico. Durante 1998 se celebró un curso sobre «Regresión Logística» y otro sobre «Modelos Mixtos con S-Plus» (en colaboración con la Universidad Politécnica de Valencia).

Próximamente, del 13 al 17 de Mayo de 1999, se celebrará un curso sobre Análisis de Supervivencia con S-Plus en Barcelona, en colaboración con la Fundació Politècnica de Catalunya. Más información en <http://www.iata.csic.es/ibsresp> o en <http://www.fpc.upc.es/CURSOS/598400.html> o llamando al tfno. 93 401 57 31.



The TES Institute Training of European Statisticians

TES OR THE ENHANCEMENT OF STATISTICS

The TES Institute is a non-profit association created in November 1996 by the Directors General of the National Statistical Institutes of ten Member States of the European Union and of the four Member States of the European Free Trade Association.

As an international post-graduate vocational training institute for statisticians, the mission of the TES Institute is to create truly European vocational training and staff development opportunities at post-graduate level. The annual training programmes are designed for target groups ranging from young statisticians to executives of National Statistical Institutes.

Over the last five years, the objectives of the TES Institute were (a) to provide training courses and seminars of short duration at post-graduate level, thus offering statisticians opportunities to perfect their skills, (b) to provide a forum for mutual consultation on vocational training for statisticians, (c) to give assistance in the training for new European statistical projects, thus disseminating standards, European methods and classifications, and to help future members to prepare for their integration into the European Statistical System, and (d) to promote exchanges of skills and experiences.

TYPES OF TRAINING

The TES Institute offers two types of training: the *Core Programme* and the *Special Courses Programme*. Both of them are covering the following specialisation areas:

(1) *Data Collection and Survey Methodology*, (2) *Economic Statistics*, (3) *Social Statistics* and (4) *Publication, Dissemination and Use of Statistics*, complemented by the *Common Courses* (general «statistical» culture) and the *Support Courses* (statistical analysis, computing techniques and statistical management).

The Core Programme

The annual training programmes are designed for public sector statisticians in the Member States of the EU and EFTA. However, the programmes are also open to the National Statistical Institutes of countries in Central and Eastern Europe, the Mediterranean Basin, some other selected countries and private sector statisticians.

The leading principle for the definition of the annual programmes and the detailed content of each of the courses is that choices are made on the basis of the training needs in the Member States of the EU and the EFTA. The *Core Programme* is subsidised by the Statistical Office of European Commission (Eurostat).

The Special Courses Programme

The *Special Courses Programme* is designed for public sector statisticians from countries outside the EU and the EFTA. At this moment its main clients are statisticians from the Central European countries and the countries of the Mediterranean Basin. The *Special Courses Programme* is not subsidised.

The *Special Courses Programme* consists of courses which are either repeats of existing courses in the *Core Programme* or tailor-made courses at the request of a country or a group of countries. Whenever a course from the *Core Programme* is repeated, the content is reviewed and adapted to the specific needs of the participants.

TRAINING ORGANISATION

Obviously, vocational training programmes of this kind imply the involvement of different groups of people.

- The training staff of the TES courses consists of *an international pool of trainers* (highly qualified practitioners and university professors). This pool is responsible for the design and the execution of the courses and the development of the course material, which are annually reassessed and updated where necessary. Each year the organisation of the programme requires the intervention of about hundred trainers.
- The National Statistical Institutes of the European Union (EU), the European Free Trade Association (EFTA), the Central and Eastern European Countries (ECO¹) and the Mediterranean Basin countries² have nominated *TES Correspondents* for the TES programme. This network of TES Correspondents is responsible for the communication between their Institute and the TES Institute, for the co-ordination of the registration of candidates and for the dissemination of information about TES courses.
- The main task of the staff of the TES Institute is to assist in the design, to organise and to give logistic support to the international vocational training programme and to provide information about existing training possibilities in the various countries in Europe. To carry out these activities, the TES Institute comprises a permanent team of eleven people.

CURRENT ACTIVITIES

Training Core and Special Courses Programmes

Over the last five year, the TES Institute organised some 110 courses in the framework of the Core Programme open to EU, EFTA and ECO statisticians. Besides, Special Courses have also been organised especially for statisticians from ECO countries. On the whole, about 600 participants have been trained every year through these two kind of training programme.

Moreover, since 21st May 1997, when the first Eurostat Task Force on Training for MEDSTAT countries was organised, the TES Institute officially extended its training activities to the Mediterranean Basin countries.

¹ Albania, Bulgaria, Czech Republic, Estonia, Hungary, FYROM (Former Yugoslav Republic of Macedonia), Latvia, Lithuania, Poland, Romania, Slovakia and Slovenia.

² Algeria, Cyprus, Egypt, Israel, Jordan, Lebanon, Malta, Morocco, Palestine, Syria, Tunisia, Turkey.

The TES Institute will be very active between March and July 1999. Indeed besides the courses planned in the framework of the *Core Programme*, a number of additional activities will be organised, such as:

- A *Summer School* on Social Statistics planned from 5 to 10 July 1999 in Siena
- Several *Special Courses* for the Central European countries
- Consulting activities for the Russian Federation

As far as the *Core Programme* is concerned, the following courses will be held:

Course Title	Course leader	Location	Dates
Seasonal Adjustment Methods	Maravall	Luxembourg	22-26.3.99
Estimation for Small Areas	Smith & Sörndal	Vienna	12-14.4.99
Systems of Health Accounts	Bonte	Luxembourg	12-15.4.99
Concepts and Measurement of Inequality and Poverty	Martín & Guzmán	Madrid	19-23.4.99
The Revised European Systems of Accounts (ESA 95) - Financial Accounts	Coin	Luxembourg	26-28.4.99
Geographical information systems	Painho	Lisbon	26-29.4.99
Techniques of Electronic Data Dissemination	Argüeso	Sevilla	3-7.5.99
Variance Estimation and Discrete Data Analysis for Complex Surveys	Lehtonen	Helsinki	17-19.5.99
Measurement in Surveys	Munck	Stockholm	17-21.5.99
Statistical Presentation with Graphic Software	Wailgren	Stockholm	31.5-2.6.99
Analysis of European Labour Market Information	Zighera	Libourne	7-11.6.99
Effective Presentation of Official Statistics to News Media	Wright	London	7-11.6.99
Theory and Application of Household Panel Surveys	Duncan	Porto	7-17.6.99
Enterprise Statistics	Lock	Libourne	14-18.6.99
Sampling Techniques and Practice (E)	Smith	Southampton	14-25.6.99

Please note that all the courses offered by the TES Institute are divided between both theoretical and practical sessions. The lectures will provide the theoretical background necessary for further developments and the practical sessions will offer you the opportunity to discuss and exchange experience with colleagues from other countries and other working areas.

Ad hoc projects

In a near future, the TES Institute will carry out:

Two seminars on 1/2 Statistics and Policy Making+ at the request of the Palestinian Statistical Training Centre (Gaza, Westbank). These seminars will be financed through UN funds and are carried out in co-operation with UNITAR, the UN Institute for Training and Research, Geneva. Again this seminar may be introduced, in a modified form, in the future regular programme.

The first part of a «Summer School on Social Statistics» in July 1998 (5 1/2 days) in Siena. The other two parts are planned for summer 1999 and 2000.

Specific training will be provided to statisticians from the Ukrainian Statistical Office. The training will cover *Business Registers, Enterprise Statistics and National Accounts*.

Consulting

Apart from the well-known training activities, the TES Institute has recently set up a procedure for consulting activities.

In this context, the TES Institute is requested to provide support to the set-up of a training centre for statisticians in a specific country in Ukraine. This support will consist of (a) training of future trainers, as above mentioned and (b) consulting for curriculum development for their future training programmes.

PUBLICATIONS

Articles

The TES Institute is regularly publishing articles on its current activities in periodic Newsletter of several Statistical Institutes and some statistical journals as *Qüestió* (Quaderns d'Estadística i Investigació Operativa), edited by the Institut d'Estadística de Catalunya.

Manuals

The TES Institute decided to set up a TES Manual series for the manuals to be developed in the specialisation areas covered by the vocational training programme (i.e. *Data Collection and Survey Methodology, Economic Statistics, Social Statistics and Publication, Dissemination and Use of Statistics*).

The following publications are available at the TES Institute:

- «*Introduction to Demographic Data Analysis*» by Mr Otto Andersen from Statistics Denmark. This manual will cover subjects like: data in demography, lexis diagram, fertility, mortality, etc. As the manual will be available at the TES Institute.
- «*The Role of Statistics in a Democracy*» (it can be ordered at the TES Institute for free).
- «*Index Number Theory and Practice*» by Prof. Marco Martini from the University of Milano. (to appear)

Please inform the TES Institute if you are interested in the above mentioned publications (articles and/or manuals) and detailed information will be forwarded to you.

Other activities

The TES Institute has been associated with the Maastrich School of Management as training and advice provider in the framework of their MBA.

In order to receive an informative leaflet on the MBA, please address your request to the TES Institute (see below for contact person). Further requests will be treated by MSM directly.

GENERAL INFORMATION

For further information about courses, new brochure, publications, specific projects and so on, please contact Ms Valerie Vandewalle co-ordinator for *Evaluation and Information matters*:

phone: +352/29 85 85 34
fax: +352/29 85 29/30
e-mail: vvandewalle@tes-institute.lu

In order to get any information on courses, latest activities, publications, etc., you are always welcome to visit our webside at **www.tes-institute.lu**.



COURSE «SYSTEMS OF SOCIAL STATISTICS»

COURSE LEADER

Pieter EVERAERS

OBJECTIVE

To get an insight knowledge in the systems and models on social statistics in use in Eurostat and in the Member States of the EU, in the possible data sources, based on the interdependency of the themes in social statistics. Special attention is given on the methodology of harmonization and integration of data sources and results of this processes for the relevant social themes in the national and international context. New developments in harmonisation (e.g. accounting, integration of surveys) and integration techniques and the outcomes are illustrated.

TRAINING METHODS

Lectures, discussion and working groups, practical examples, case studies, exercises, role play.

TARGET GROUP

Statisticians working (or starting to work) in one of the domains in social statistics or in middle management with a task in preparing or implementing harmonisation and integration, working on the crossroads of statistical domains or preparing output/publications covering several data sources and/or statistical domains.

ENTRY QUALIFICATIONS

- Academic or comparable level.
- At least three years of experience in statistical work.
- Sound knowledge of the teaching language (see «Practical Information»).

EXPECTED OUTPUT

Direct usable knowledge of the interdependency between themes in social statistics, methods and levels of harmonisation and integration of data from different sources, insight and experience in working with systems/models of social statistics and the availability and characteristics of European 1/2comparable+ data sources for social statistics.

CONTENTS

Overview of several systems in social statistics, in the process of development or in use, their contents and characteristics, methodological and technical aspects, possible use, available data sources. Relations between the different social themes in statistics, European data sources, policy use, and experiences in NSI's with the systems/models of social statistics.

The course is structured along three main lines:

SYSTEMS OF SOCIAL STATISTICS

- Social accounting and social demographic accounting.
- Supply and demand of services (health, education).
- Situations relating to systems of living/life cycle approaches.
- Time accounting/time use systems.
- The indication approach.

LEVELS AND METHODS FOR HARMONIZATION AND INTEGRATION

- Accounting, macro linkage.
- Imputation, weighting correcting on the table level.
- Combination of information from tables, reconciliation of data.
- Integration on the micro level, micro linkage.

GENERAL BACKGROUNDS/RELEVANCE

- Principles of harmonisation, co-ordination and integration.
- Theoretical bases, methods and techniques, data sources for social statistics.
- The European system of social statistics.
- National experiences.

REQUIRED READING

C.A. van Bochove and P.C.J. Everaers. Micro-macro and micro-micro linkage in social statistics (in Netherlands Official Statistics, 1996).

SUGGESTED READING

Reading of about 50 pages of separate articles required as well as preparing presentation on one or two specific articles. Extra reading of about 100 pages of separate articles suggested.

Description of main characteristics of own national statistical systems. To be brought to the course.

PRACTICAL INFORMATION

Code: SOC-001-E

Duration: 5 days

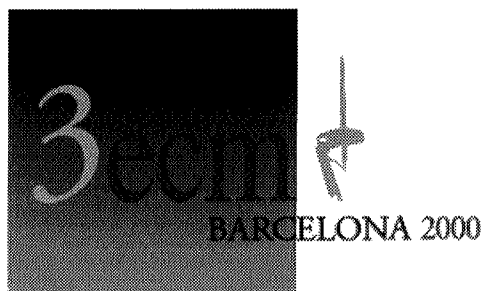
Timing: 4-8 October 1999

Training site: Institut d'Estadística de Catalunya (Idescat)
Via Laietana, 58. 08003 Barcelona. SPAIN

Language: English

Training Fee: 1.240 Euro

Deadline for registration: 31 August 1999



Tercer Congrés Europeu de Matemàtiques

Barcelona, del 10 al 14 de juliol del 2000

Primer Anunci

El Comitè Organitzador es complau a anunciar que el **Tercer Congrés Europeu de Matemàtiques (3ecm)** tindrà lloc a Barcelona del 10 al 14 de juliol de l'any 2000. L'organitza la Societat Catalana de Matemàtiques (SCM), sota els auspicis de la Societat Matemàtica Europea (EMS).

Conferenciants plenaris

- **Robbert Dijkgraaf** (Universitat d'Amsterdam, Holanda)
- **Hans Föllmer** (Universitat Humboldt de Berlín, Alemanya)
- **Hendrik W. Lenstra, Jr.** (Universitat de Califòrnia a Berkeley, Estats Units, i Universitat de Leiden, Holanda)
- **Yuri I. Manin** (Institut Max Planck de Matemàtiques, Bonn, Alemanya)
- **Yves Meyer** (Escola Normal Superior de Cachan, França)
- **Carles Simó** (Universitat de Barcelona)
- **Marie-France Vignéras** (Universitat de París 7, França)
- **Oleg Viro** (Universitat d'Uppsala, Suècia, i POMI de Sant Petersburg, Rússia)
- **Andrew J. Wiles** (Universitat de Princeton, Estats Units)

Programa científic

El programa del congrés inclourà nou conferències plenàries, trenta conferències invitades en sessions paral·leles, conferències impartides pels guardonats amb els premis de l'EMS, minisimposis, taules rodones i sessions de pòsters. Igual com es va fer en els congressos europeus anteriors, s'atorgarà un cert nombre de premis a investigadors/res joves en matemàtiques, de menys de trenta-dos anys d'edat. Els minisimposis són una de les novetats del **3ecm**; el Comitè Científic escollirà una llista de temes actuals i interdisciplinaris. El programa complet de conferències, minisimposis i taules rodones s'especificarà en el segon anunci. Tots els participants podran presentar comunicacions en forma de pòsters. També està previst que s'organitzin demostracions de programari matemàtic, vídeo i material multimèdia.

Amb finalitats organitzatives, s'ha establert la següent llista numerada de temes científics:

1. Lògica i fonaments
2. Àlgebra. Teoria de nombres
3. Geometria algebraica i analítica
4. Geometria diferencial
5. Topologia
6. Matemàtica discreta i informàtica
7. Modelització i simulació
8. Equacions diferencials ordinàries i sistemes dinàmics
9. Equacions en derivades parcials
10. Anàlisi funcional
11. Anàlisi complexa
12. Probabilitat i estadística
13. Anàlisi real
14. Física matemàtica

Comitès

- El Comitè Científic és presidit per **Sir Michael Atiyah** (Universitat d'Edimburg).
- El Comitè de Premis és presidit per **Jacques-Louis Lions** (Col·legi de França).
- El Comitè de Taules Rodones és presidit per **Miguel de Guzmán** (Universitat Complutense de Madrid).
- El Comitè Organitzador és presidit per **Sebastià Xambó Descamps** (Universitat Politècnica de Catalunya).

Presentació de pòsters

Tots els participants inscrits al **3ecm** podran presentar treballs en forma de pòsters. El Comitè Organitzador decidirà quins pòsters s'accepten a partir de resums que haurà d'haver rebut abans de l'1 d'abril del 2000. Us demanem que envieu el vostre resum preferiblement fent servir el programa que hi haurà a la pàgina web <http://www.iec.es/3ecm/posters.htm>. També es pot enviar per correu electrònic a posters.3ecm@upc.es, posant com a *subject* només el número de la secció escaient (vegeu la llista de temes indicada més amunt). Si no el podeu enviar electrònicament, utilitzeu l'adreça següent: *Pòsters 3ecm (Prof. Josep M. Font), Facultat de Matemàtiques, Universitat de Barcelona, Gran Via 585, 08007 Barcelona.*

Presentacions de programari matemàtic

Durant el congrés tindrà lloc una sessió de programari matemàtic, en la qual es podran presentar programes relacionats amb tots els camps de les matemàtiques i aplicables a objectius diversos. El Comitè Organitzador avaluarà les propostes i en seleccionarà un cert nombre, aplicant criteris d'originalitat matemàtica, novetat i possibilitats d'aplicació, i tenint en compte l'equilibri temàtic de la sessió.

Les propostes han d'arribar als organitzadors abans de l'1 de febrer del 2000. Es poden enviar electrònicament, fent servir el full que hi ha a la web <http://www.iec.es/3ecm/mathsoft.htm>, o bé a l'adreça mathsoft.3ecm@upc.es, posant com a *subject* només la paraula *mathsoft*. L'adreça següent també es pot utilitzar per enviar material complementari: *Mathsoft 3ecm (Prof. Santiago Zarzuela), Facultat de Matemàtiques, Universitat de Barcelona, Gran Via 585, 08007 Barcelona.*

Presentacions de vídeo i multimèdia

Durant el **3ecm** hi haurà un seguit d'activitats complementàries i actes culturals. Una d'aquestes activitats serà la producció d'un DVD amb vídeos i material multimèdia amb contingut matemàtic. Aquest DVD s'exhibirà en sessions públiques i també serà accessible en diversos llocs de la seu del **3ecm**. Es podran enviar contribucions per a aquest DVD des de totes les àrees de les matemàtiques.

Els treballs hauran d'arribar als organitzadors abans de l'1 de febrer del 2000. S'han d'enviar per correu ordinari a: *Video 3ecm (Prof. Santiago Zarzuela), Facultat de Matemàtiques, Universitat de Barcelona, Gran Via 585, 08007 Barcelona*. Trobareu més informació a la web <http://www.iec.es/3ecm/video.htm>. Podeu utilitzar l'adreça video.3ecm@upc.es per contactar amb els organitzadors d'aquesta activitat.

Activitats satèl·lit

Els congressos i les altres activitats de la llista següent han estat acceptats com a satèl·lits del **3ecm** pel Comitè Executiu abans del mes de febrer de 1999. Us encoratgem que feu la llista més llarga. Les propostes s'han de fer arribar al president del Comitè Organitzador abans de l'1 de febrer del 2000, per correu electrònic a 3ecm@iec.es o bé per carta a la SCM.

- **Summer School on Interactions between Algebraic Topology and Invariant Theory.** Ioannina, Grècia, del 26 de juny a l'1 de juliol del 2000. *Contacteu amb:* Nondas Kechagias (Universitat de Ioannina), nkechag@cc.uoi.gr.
- **Functional Analysis Valencia 2000, an International Functional Analysis Meeting on the Occasion of the 70th Birthday of Professor Manuel Valdivia.** València, del 3 al 7 de juliol del 2000. *Contacteu amb:* José Bonet (Universitat de València), vlc2000@mat.upv.es, o bé Klaus D. Bierstedt (Universitat de Paderborn), vlc2000@uni-paderborn.de.
- **6th International Conference on Harmonic Analysis and Partial Differential Equations.** El Escorial, Madrid, del 3 al 7 de juliol del 2000. *Contacteu amb:* Eugenio Hernández (Universitat Autònoma de Madrid), eugenio.hernandez@uam.es.
- **Alhambra 2000, a Joint Mathematical European-Arabic Conference.** Granada, del 3 al 7 de juliol del 2000. *Contacteu amb:* Ceferino Ruiz (Universitat de Granada), alhambra2000@ugr.es.
- **First Euro-Mediterranean Topology Meeting.** Bellaterra, del 4 al 7 de juliol del 2000. *Contacteu amb:* Carlos Broto (Universitat Autònoma de Barcelona), broto@mat.uab.es.
- **cem 2000, Congrés d'Educació Matemàtica, I Jornades d'Educació Matemàtica a Catalunya.** Mataró, del 3 al 5 o del 5 al 7 de juliol del 2000. *Contacteu amb:* Xavier Vilella (FEEMCAT), xvilella@pie.xtec.es.
- **Distributions with Given Marginals and Statistical Modelling.** Barcelona, del 17 al 19 de juliol del 2000. *Contacteu amb:* Carles M. Cuadras (Universitat de Barcelona), carlesm@porthos.bio.ub.es.

Preinscripcions

Si desitgeu rebre el segon anunci i més informació per correu electrònic sobre el **3ecm**, us podeu preinscriure a través de la web <http://www.iec.es/3ecm> (si encara no ho heu fet). La preinscripció no costa diners ni us obliga a res. Per tal d'esdevenir participants del **3ecm**, caldrà que formalitzeu la inscripció quan s'obri el termini per fer-ho i pagueu la quota corresponent. També us podeu preinscriure per correu electrònic a l'adreça **3ecm@iec.es**, o bé enviant una carta a la SCM. Cal que indiqueu el vostre nom, la vostra institució, l'adreça postal completa, l'adreça de correu electrònic (si en teniu) i els camps científics que us interessin (en podeu escollir un o més d'un de la llista indicada més amunt).

Patrocinadors (relació actualitzada a febrer de 1999)

Generalitat de Catalunya, Comissionat per a Universitats i Recerca
Generalitat de Catalunya, Departament d'Ensenyament
Ministerio de Educación y Cultura, S.E.U.I.D.
Fundació Catalana per a la Recerca
Ajuntament de Barcelona
Institut d'Estudis Catalans
Universitat de Barcelona
Universitat Autònoma de Barcelona
Universitat Politècnica de Catalunya
Institut d'Estadística de Catalunya
International Mathematical Union
Real Sociedad Matemática Española
Sociedad Española de Matemática Aplicada
Fundación Retevisión
Fundació «la Caixa», Museu de la Ciència
Borsa de Barcelona
Port de Barcelona
Fundació Caixa Catalunya
Fundació Banc Sabadell
Fundació Caixa de Sabadell
Logic Control
Springer-Verlag

Adreces de contacte

Correu electrònic: **3ecm@iec.es**

Web: **<http://www.iec.es/3ecm/>** o també **<http://www.si.upc.es/3ecm/>**

Correu ordinari: Societat Catalana de Matemàtiques
Institut d'Estudis Catalans
Carrer del Carme, 47
08001 Barcelona

Telèfon: +34 93 270 16 20

Fax: +34 93 270 11 80

Informació per als autors i lectors

NORMES PER A LA PRESENTACIÓ D'ARTICLES A QÜESTIÓ

La revista accepta, per a la seva publicació, articles originals no sotmesos a consideració en cap altra revista dins els àmbits de l'Estadística, la Investigació Operativa, l'Estadística Oficial i la Biometria. Els articles poden ser teòrics o aplicats, incloent aspectes computacionals i/o de caire docent, i poden presentar-se en anglès, francès, català o qualsevol altra llengua oficial a l'Estat espanyol.

Tots els originals destinats a les esmentades seccions temàtiques de *Qüestió*, incloent-hi els articles per a la «Secció docent i problemes», seran sotmesos sistemàticament a una avaluació prèvia a càrrec d'especialistes independents i/o membres del Consell Editorial, llevat dels articles convidats per la revista i les reimpressions d'articles. El resultat de l'avaluació serà comunicat a l'autor principal als efectes d'eventuals correccions formals o dels seus continguts.

Per a totes les trameses d'originals, la revista emetrà un acusament de recepció la data del qual figurarà com a «data de rebuda» en la publicació de l'article. Per la seva banda, la «data d'acceptació» de l'article serà la data de recepció de la versió definitiva.

Per a la presentació d'articles, l'autor trametrà a la Secretaria de *Qüestió* (Institut d'Estadística de Catalunya) dues còpies del treball mecanografiat en DIN A4, a una sola cara, a doble espai i amb marges amplis. Cada article ha d'incloure el títol, el nom de l'autor o autors, la seva afiliació i l'adreça completa, així com un resum de 75-100 paraules al principi de l'article, seguit de les principals paraules clau (en l'idioma original) i la seva adscripció a la classificació AMS. Abans de sotmetre els articles a la revista, s'aconsella els autors que revisin la correcció lingüística de textos d'acord amb l'idioma original i les eventuals traduccions a l'anglès.

Les referències bibliogràfiques es faran indicant el cognom de l'autor seguit de l'any de la publicació entre parèntesi [i.e.: Mahalanobis (1936), Rao (1982b)] i seran llistades alfabèticament al final de l'article; les referències múltiples d'un mateix autor s'ordenaran cronològicament. Les notes explicatives es numeraran correlativament i han d'aparèixer al peu de la pàgina corresponent. Les taules i figures també es numeraran correlativament en el text i seran reproduïdes directament dels originals tramesos en cas que no sigui possible la seva autoedició.

Una vegada avaluat satisfactòriament l'article cal que, a més de la versió impresa, l'autor el trameti en disquet de 3.5 polsades i en format MS-DOS, on han de constar de forma clara els noms dels autors i el títol de l'article. Aquesta versió final s'ha de trametre preferiblement en el processador de textos $\text{\LaTeX} 2_{\epsilon}$ [subsidiàriament, es poden trametre els textos i les taules en Word Perfect —versió 6.0A o anterior— o ASCII]; en el cas de figures, diagrames o gràfics es recomanen els formats adients per als programes editors PS, EPS o PCX. Els autors han de garantir la correspondència exacta entre la versió impresa i la còpia electrònica. D'altra banda, si l'article no està escrit en llengua anglesa s'haurà d'adjuntar la traducció del títol original, de l'abstract i de les paraules clau, així com un ampli resum en anglès (amb una extensió d'entre 2 i 5 pàgines i amb la mateixa estructura de l'article original).

La Secretaria de *Qüestió* posa a disposició dels autors que ho sol·licitin plantilles en format $\text{\LaTeX} 2_{\epsilon}$ de *Qüestió* per a la seva edició i les referències adients de la classificació AMS.

QÜESTIÓ
Institut d'Estadística de Catalunya
Via Laietana, 58
08003 Barcelona
Tel: +34-93 412 15 36
Fax: +34-93 412 31 45
E-mail: questio@idescat.es

GUIDELINES FOR THE SUBMISSION OF ARTICLES FOR *QÜESTIÓ*

The journal well comes submission of articles and contributions 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 other official languages in Spain.

All originals assigned to the thematic sections of *Qüestió*, including articles for the «Teaching section and problems» will be systematically reviewed by independent referees and/or members of the Editorial Board, who will send a report to the main author of the article in order to correct, if necessary, any formal or content aspects. The articles invited by the journal and articles reprinted will be excluded from this evaluation process.

For all submissions, the journal will issue a receipt corresponding to the submission date, which will appear as «date received» in the final publication of the article. 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 copies of the paper typed on A4 sheets, one side of the paper only, double spaced and with wide margins. Each article should include the title, the name of the author or authors, their affiliation, full address and also an abstract of the paper (75-100 words) at the beginning of the article, followed by the main keywords (in the original language) and its assignation in the AMS classification. Before submitting their papers, authors are advised to seek assistance in the writing of their articles for the correct use of English and/or of original language.

Bibliographical references should state the author's name followed by the year of publication in brackets [e.g.: Mahalanobis (1936), Rao (1982b)] and they should be listed at the end of the article in alphabetical order; multiple references to the same author should be given in chronological order. Footnotes should be numbered in the article and appear at the foot of the corresponding page. Figures and tables are to be numbered in consecutive order in the text using Arabic numerals and will be directly reproduced from the originals submitted if it is not impossible to print them electronically.

Once the evaluation has been passed, the author is required to provide the article on a diskette (a 3.5-inch disk in MS-DOS format) together with its paper copy; it must be a new diskette and must bear very clearly the names of the authors and the title of the article. This final version should be processed by $\text{\LaTeX} 2_{\epsilon}$, preferably, or, failing that, by Word Perfect (6.0A or earlier) or ASCII for text and tables; for figures, diagrams or graphs, the appropriate formats of PS, EPS or PCX software tools are strongly recommended. Authors must ensure that the version of the electronic copy is exactly the same as the paper copy which accompanies it. Furthermore, if the article is not written in English, the translation of its original title, short abstract and keywords should be enclosed, as well as a full summary of the article in English (that is, 2-5 pages with the same structure as the original).

The Secretary of *Qüestió* can send, by request of the authors, the $\text{\LaTeX} 2_{\epsilon}$ style of *Qüestió* for manuscript preparation and the appropriate AMS classification references.

QÜESTIÓ
Institut d'Estadística de Catalunya
Via Laietana, 58
08003 Barcelona
Tel: +34-93 412 15 36
Fax: +34-93 412 31 45
E-mail: questio@idescat.es

NORMES PER A LA PUBLICACIÓ D'ANUNCIS INSTITUCIONALS A *QÜESTIÓ*

Qüestió convida les entitats patrocinadores, les institucions col·laboradores, els organismes públics i privats, i tota la comunitat científica vinculada a l'estadística o la investigació operativa, a la publicació d'anuncis institucionals sobre cursos, seminaris, congressos i activitats similars que, preferentment, tinguin lloc en el nostre país. Els textos poden presentar-se en anglès, francès, català o en qualsevol altra llengua oficial a l'Estat espanyol. Les iniciatives per a una possible publicació sempre són a instància de les entitats interessades, de manera que *Qüestió* no fa una cerca sistemàtica d'esdeveniments d'aquesta naturalesa, ni té cap ànim d'exhaustivitat en les ressenyes d'activitats finalment publicades.

Una vegada aprovada la inclusió dels anuncis sol·licitats es procedirà a la seva publicació, i es reproduirà directament dels originals tramesos amb les mides adequades i la màxima qualitat tipogràfica possible; en aquest cas, *Qüestió* no procedeix a cap mena de procés d'autoedició de la versió impresa que l'anunciant hagi tramès. Si els originals es trameten en els mateixos termes electrònics exigits per als articles (vegeu «Normes per a la presentació d'articles a *Qüestió*»), la revista procedirà a la seva autoedició. Si es desitja una qualitat superior a la reproducció simple o l'autoedició, o bé la seva publicació en color, els sol·licitants hauran de posar-se en contacte amb la Secretaria de *Qüestió* per tal de trametre els fotolits dels textos originals corresponents.

La disposició dels textos i les figures adjuntes dels anuncis han de procurar la màxima intel·ligibilitat i claredat expositiva, sense atapeir la informació ni utilitzar formats o fonts de lletres excessivament petites. D'altra banda, la publicitat ha de ser fidedigna, exempta d'engany i respectuosa amb les persones i institucions. En qualsevol cas, la direcció de *Qüestió* es reserva la decisió final pel que fa a la seva publicació.

L'anunciant es compromet a lliurar els textos/materials amb l'antelació que se li indiqui per a la inserció en els números/volums de *Qüestió* que prèviament s'hagi establert. La revista no es fa responsable dels retards, per part de l'anunciant, que impedeixin la publicació de l'anunci en els termes previstos.

Mònica M. Jaime
Secretaria de *Qüestió*
Institut d'Estadística de Catalunya
Via Laietana, 58
08003 Barcelona
Tel: +34-93 412 15 36
Fax: +34-93 412 31 45
E-mail: questio@idescat.es

GUIDELINES FOR INSTITUTIONAL ADVERTISEMENTS IN *QÜESTIÓ*

Qüestió invites all sponsor entities, collaborating institutions, other public and private bodies and the entire scientific community related to Statistics or Operations Research to submit institutional advertisements on courses, seminars, congress and similar activities that will be held, preferably in our country. These will be accepted in English, French, Catalan or any of the other official languages in Spain. The initiative should always come from the entities interested in advertising them so that *Qüestió*'s aim is not to do a systematic search of these events and therefore does not publish a comprehensive list of such activities.

Once their insertion is approved the advertisements will be reproduced from the most accurate photocopy of the originals sent by the advertiser to *Qüestió* in paper copy, with the appropriate size and at the best possible typographic quality. Therefore, in this case the journal does not elaborate any further editing process to the printed version that the advertiser has sent. If the original advertisements are sent in the same electronic format requested by the articles (please see «Guidelines for the submission of articles for *Qüestió*») the journal will print it directly from the file. If a better quality than the simple reproduction or automatic printing or a colour version of the adverts is desired, the authors should contact the Secretary of *Qüestió* in order to negotiate this.

The typesetting of texts and figures in the advertisement should have maximum intelligibility and clearness, neither compressing the information too much nor using formats or letter fonts that are too small. Furthermore, the information has to be reliable, without errors and respectful of the people and institutions. The management of *Qüestió* has the right to a final decision concerning the insertion of the advertisement.

Advertisers commit themselves to give the text/materials on request in order to insert them in the issues of *Qüestió* that have been previously agreed. The journal is not responsible for any delay from the announcer that could prevent the advertisement from been published on the agreed terms.

Mònica M. Jaime
Secretaria de *Qüestió*
Institut d'Estadística de Catalunya
Via Laietana, 58
08003 Barcelona
Tel: +34-93 412 15 36
Fax: +34-93 412 31 45
E-mail: questio@idescat.es

NORMES PER A LA PUBLICACIÓ D'ANUNCIS PRIVATS O AMB FINALITAT COMERCIAL A *QÜESTIÓ*

Qüestió accepta la publicació d'anuncis privats o amb finalitat comercial sobre productes, serveis o altres eines promocionals a l'entorn de l'estadística o la investigació operativa. Els textos poden presentar-se en anglès, francès, català o en qualsevol altra llengua oficial a l'Estat espanyol. Les iniciatives per a una possible publicació sempre són a instància de les organitzacions que hi estiguin interessades, de manera que *Qüestió* no fa una cerca sistemàtica de novetats o productes d'aquesta naturalesa ni té cap ànim d'exhaustivitat en els anuncis finalment publicats.

Els anuncis en **blanc i negre** s'elaboren a partir de la fotocòpia més acurada possible dels originals que trameti l'anunciant en versió impresa, amb les mides adequades i la màxima qualitat tipogràfica. Per tant, en aquest cas la revista no efectua cap procés d'edició ulterior respecte de la versió impresa que l'anunciant hagi tramès. Alternativament, si els anuncis originals es trameten en els mateixos termes formals exigits per als articles (vegeu «Normes per a la presentació d'articles a *Qüestió*»), la revista procedirà a la seva autoedició. Igualment, si es desitja una qualitat superior a la reproducció simple, els sol·licitants hauran de trametre els fotolits dels originals corresponents o encarregar-los a *Qüestió*, que els facturarà separatament.

Els anuncis en **color** requereixen els fotolits dels textos originals, que poden ser subministrats directament per l'anunciant o bé encarregats per la revista a compte de l'anunciant; en el segon cas, l'anunciant ha de trametre a la revista els originals impresos en color amb la màxima qualitat, per tal de filmar-los amb les millors garanties i condicions. El cost dels fotolits realitzats per *Qüestió* serà sempre a càrrec de l'anunciant, a qui se li repercutirà l'import i l'IVA d'aquests, juntament amb les tarifes que corresponen a la modalitat d'anunci per la qual hagi optat.

La disposició dels textos i figures adjuntes dels anuncis ha de procurar la màxima intel·ligibilitat i claredat expositiva, sense atapeir la informació ni utilitzar formats o fonts de lletres excessivament petites. D'altra banda, la publicitat ha de ser fidedigna, exempta d'enganys i respectuosa amb les persones i institucions. En qualsevol cas, la direcció de *Qüestió* es reserva la decisió final de la seva inclusió.

L'anunciant es compromet a lliurar els textos/materials amb l'antelació que se li indiqui per a la seva inserció en el(s) número(s)/volum(s) de *Qüestió* que prèviament s'hagi establert. La revista no es fa responsable dels retards per part de l'anunciant que impedeixin la publicació de l'anunci en els termes previstos.

Imports:

1 pàgina en color (un número aïllat):	125.000 PTA + IVA
1 pàgina en color (tres números consecutius):	200.000 PTA + IVA
1 pàgina en blanc i negre (un número aïllat):	30.000 PTA + IVA
1 pàgina en blanc i negre (tres números consecutius):	50.000 PTA + IVA
1/2 pàgina en blanc i negre (un número aïllat):	20.000 PTA + IVA
1/2 pàgina en blanc i negre (tres números consecutius):	35.000 PTA + IVA

Mides opcionals dels anuncis:

1 pàgina sencera (espai intern):	19.0 cm. x 12.3 cm.
1 pàgina sencera (espai extern):	23.8 cm. x 17.0 cm.
1/2 pàgina (espai intern):	9.5 cm. x 12.3 cm.
1/2 pàgina (espai extern):	11.9 cm. x 17.0 cm.

Forma de Pagament:

- Transferència bancària al compte: 2013-0100-53-0200698577
- Xec bancari nominatiu a l'Institut d'Estadística de Catalunya
- Pagament amb targeta de crèdit

El pagament serà per l'import total de la factura corresponent, on hi figurarà el cost dels fotolits en el cas que l'edició de l'anunci hagi estat a càrrec de l'Institut. En el cas que l'anunciant necessiti una factura proforma, només cal que ho faci saber amb l'antelació suficient.

Correspondència:

Mònica M. Jaime
Secretaria de *Qüestió*
Institut d'Estadística de Catalunya
Via Laietana, 58
08003 Barcelona
Tel: +34-93 412 15 36
Fax: +34-93 412 31 45
E-mail: questio@idescat.es

GUIDELINES FOR THE PRIVATE OR COMMERCIAL ADVERTISEMENTS IN *QÜESTIÓ*

Qüestió accepts for their publication both private and commercial advertisements on products, services or other promotional tools related to statistics or operational research and will be accepted in English, French, Catalan or any of the official languages in Spain. The initiatives should always come from entities interested in advertising them so that *Qüestió*'s aim is not to do a systematic search of news and therefore does not publish a comprehensive list of such private or profit activities.

The **black and white** advertisements are made out from the most accurate photocopy of the originals sent by the advertiser to *Qüestió* in paper copy with the appropriate size and at the best possible typographic quality. Therefore, in this case the journal does not elaborate any further editorial process to the printed version that the advertiser has sent. Alternatively, if the original advertisements are sent in the same formal terms required by the articles (please see «Guidelines for the submission of articles for *Qüestió*»), the journal will proceed to its autoedition. In the same way, if a better quality than the simple reproduction is wanted, the authors should send the photolits of the corresponding original texts or, on the other hand, order to *Qüestió* their fulfilment, which will be invoiced separately from the rates charged as advertisements.

The advertisements **in colour** need the photolits of the original texts, which can be provided directly by the advertiser or requested by *Qüestió* to the advertiser charge; in the second case, the advertiser must sent to the journal the originals printed in colour with the best possible quality, so that they can be filmed at the best conditions and guarantees. The cost of the photolits made by *Qüestió* will always be charged to the advertiser together with the VAT derived from it, plus the prices corresponding to the type of the advertisement that has been chosen.

The set up of texts and figures of the advertisement should provide the maximum intelligibility and clearness, neither squeezing together the information nor using set ups or letter types that are too small. On the other hand the publicity has to be reliable, without fraud and respectful to the persons and institutions. The direction of *Qüestió* has the right of the last decision concerning the insertion of the advertisement.

The advertiser commits himself to give the texts/materials on request, in order to insert them in the issue(s) of *Qüestió* that had been previously agreed. The journal is not responsible for any delay from the announcer that could prevent the advertisement from been published in the agreed terms.

Rates:

1 colour page (only one issue):	125.000 PTA + VAT
1 colour page (three consecutive issues):	200.000 PTA + VAT
1 black and white page (only one issue):	30.000 PTA + VAT
1 black and white page (three consecutive issues):	50.000 PTA + VAT
1/2 black and white page (only one issue):	20.000 PTA + VAT
1/2 black and white page (three consecutive issues):	35.000 PTA + VAT

Advertisement sizes (optional):

1 full page (internal space):	19.0 cm. x 12.3 cm.
1 full page (external space):	23.8 cm. x 17.0 cm.
1/2 page (internal space):	9.5 cm. x 12.3 cm.
1/2 page (external space):	11.9 cm. x 17.0 cm.

Payment:

- A bank transfer to account number: 2013-0100-53-0200698577
- A bank cheque to Institut d'Estadística de Catalunya
- Charge on a credit card

The payment should be for the amount shown at the invoice, where it will be shown the total cost of the photolits, in case that *Qüestió* would be in charge of the filmation of the advertisement. If advertiser need a pro-forma invoice, he should let us know some time in advance so that *Qüestió* could send it to the proper address.

Mail address:

Mònica M. Jaime
Secretaria de *Qüestió*
Institut d'Estadística de Catalunya
Via Laietana, 58
08003 Barcelona
Tel: +34-93 412 15 36
Fax: +34-93 412 31 45
E-mail: questio@idescat.es

Butlleta de subscripció a la revista **Qüestió**

Nom i cognoms _____	

Empresa/Institució _____	

Adreça _____	

Codi postal _____	Ciutat _____
Tel. _____	Fax _____ NIF _____
Data _____	
Signatura	

Desitjo subscriure'm a **Qüestió** per a l'any 1999.

El preu de la subscripció és de 3.000 PTA (18,03 euros) (IVA inclòs).

Forma de pagament

☐ Transferència al compte 2013-0100-53-0200698577

☐ Domiciliació bancària al compte número

--	--	--	--	--	--	--	--	--	--	--	--	--	--	--	--	--	--	--	--

☐ Xec nominatiu a l'Institut d'Estadística de Catalunya

☐ Gir postal

☐ En efectiu

Retorneu aquesta butlleta (o una fotocòpia) a:

Qüestió:

Institut d'Estadística de Catalunya

Via Laietana, 58

08003 Barcelona

Exemplar per a l'entitat bancària

Autorització de domiciliació bancària per al pagament de les subscripcions anuals de la revista **Qüestió**

El sotassinat _____

autoritza el Banc/Caixa _____

Adreça _____

Codi postal _____ Ciutat _____

a abonar les subscripcions a la revista **Qüestió** amb càrrec al seu compte

número

--	--	--	--

--	--	--	--	--

--	--

--	--	--	--	--	--	--	--	--	--	--	--

Data _____

Signatura

Qüestió:
Institut d'Estadística de Catalunya
 Via Laietana, 58
 08003 Barcelona

SOLAS l'anàlisi de les dades no disponibles

COM TRACTAR LES DADES NO DISPONIBLES?

Són molts els estudis avançats que han estat perjudicats per l'ús de mètodes puntuals en el tractament de dades no disponibles, com ara per exemple les "anàlisis de cas global" o les "anàlisis de cas disponible". Els organismes de control consideren aquests mètodes cada cop més esbiaixats. D'acord amb l'índole de dades que no es troben disponibles (ja siguin acotades, univariants o multivariants), dins d'aquest conjunt de dades, Solas permet escollir entre les tècniques més adients per aplicar a un subconjunt de les dades o a totes en conjunt.

SOLAS, l'anàlisi de les dades no disponibles, és un paquet de programari estadístic completament integrat dins l'entorn Windows i inclou les últimes tècniques reconegudes per la indústria en els tractaments dels valors no disponibles. Les tècniques d'imputació inclouen:

Imputació múltiple

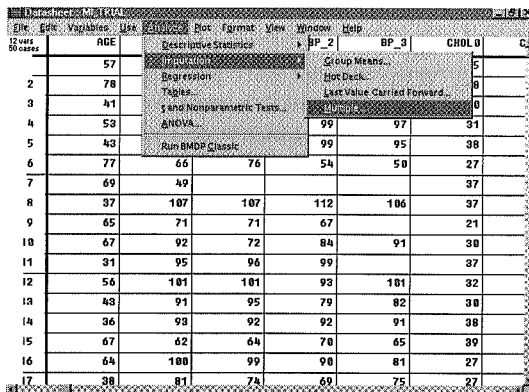
- Basada en tècniques desenvolupades per Rubin et al.
- Genera imputacions múltiples en utilitzar unes puntuacions de propensió, així com la rutina bayesiana d'inicialització més adient.
- Permet l'anàlisi de les dades longitudinals.

Imputació senzilla

- Últim valor transferit: permet a l'usuari imputar el valor que falta en funció de l'últim valor observat.
- Mitjana de grup: permet la imputació d'una mesura de grup (o d'una modalitat en cas que les dades siguin categòriques).
- Hot decking: permet a l'usuari especificar i prioritzar totes les combinacions de variable per a una imputació del tipus "correspondència més propera".

Patró SOLAS de dades no disponibles

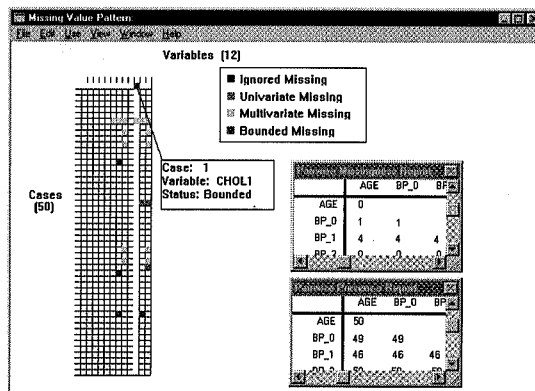
El patró de dades no disponibles és una exclusiva de Solas i ofereix a l'usuari una visualització clara i global de la quantitat, la localització i els tipus de dades no disponibles a cadascuna de les cel·les. Aquesta prestació ofereix la possibilitat d'aïllar, encara més, les cel·les individuals d'un conjunt de dades i identificar-ne les configuracions per defecte d'observació i d'estat. Un cop acabada la imputació aquest dispositiu gràfic, codificat amb colors, identifica amb facilitat el mètode de tractament seleccionat.



	AGE	BP_2	BP_3	CHOL_0
1	57			
2	78			
3	41			
4	53			
5	43			
6	77	66	76	
7	69	49		
8	37	107	107	
9	65	71	71	
10	67	92	72	
11	31	95	96	
12	56	101	101	
13	43	91	95	
14	36	93	92	
15	67	62	64	
16	64	100	99	
17	38	81	74	



Seleccioneu un dels quatre mètodes d'imputació disponibles amb SOLAS per a l'anàlisi de les dades no disponibles.



Variables (12)

- Ignored Missing
- Univariate Missing
- Multivariate Missing
- Bounded Missing

Case: 1
Variable: CHOL_0
Status: Bounded

AGE	BP_0	BP_1
0	1	1
4	4	4
0	0	0

AGE	BP_0	BP_1
50	49	49
46	46	46
50	50	50



Determineu a quin tipus pertanyen els valors no disponibles utilitzant l'Informe de Patró de Dades No Disponibles.

Statistical Solutions Ltd.
8 South Bank
Crosse's Green, Cork, Irland

Tel. + 353 21 319629 • Fax + 353 21 319630
Email: sales@statsol.ie
Website: <http://www.statsol.ie>*

Statistical Solutions
Stonehill Corporate Center Suite 104
999 Broadway, Saugus, MA 01906

Tel. (781) 2317680 • Fax (781) 2317684
Email: info@statsolusa.com
Website: <http://www.statsolusa.com>



solas For Missing Data Analysis 1.0

HOW DO YOU HANDLE MISSING DATA?

Many advanced studies have suffered from the use of ad-hoc methods of handling missing data such as "complete-case-analysis" or "available-case analysis". These methods are becoming increasingly viewed as biased by regulatory agencies. According to the types of missing data (bounded, univariate or multivariate) in your dataset SOLAS allows you to choose the most appropriate techniques to apply to a subset or all of your data.

solas For Missing Data Analysis, a fully Windows based statistical software package, incorporates the latest industry accepted techniques for treating missing values. Imputation techniques include:

Multiple Imputation

- Based on techniques developed by Rubin et Al.
- Generates multiple imputations using propensity scores and the approximate Bayesian bootstrap.
- Allows for analysis of longitudinal data.

Single Imputation

- Last Value Carried Forward - allows the user to impute Missing Value based on last value observed
- Group Means - Imputation of group mean (or mode in the case of categorical data)
- Hot Decking - Allows the user to specify and prioritize all combinations of variables for 'closest match' imputation

Solas Missing Data Pattern

The unique missing data pattern in Solas provides the user with a clear over-view of the quantity, positioning and types of missing values in each cell. This feature allows you to further isolate particular cells in a data set and identify observation and status defaults. Following imputation this colour coded graphics feature easily identifies the treatment method selected.

Case	Var	TreatGrp	Imputation	height	SYMPOUR	AGE	CVC_0	CVC_1	CVC_2
1	1.0000000								
2	1								
3	1								
4	1								
5	1								
6	1								
7	1								
8	1								
9	1								
10	1								
11	1								
12	1								
13	1								
14	1								
15	1								
16	1								
17	1								
18	1								
19	1								
20	1								
21	1								
22	1								
23	1								
24	1								
25	1								
26	1								
27	1								
28	1								
29	1								
30	1								
31	1								
32	1								
33	1								
34	1								
35	1								
36	1								
37	1								
38	1								
39	1								
40	1								
41	1								
42	1								
43	1								
44	1								
45	1								
46	1								
47	1								
48	1								
49	1								
50	1								

Select one of the four imputation methods available in SOLAS for Missing Data Analysis

Variables (13)

Case 1
Var FRS1
Status: Bounded

	SYMPOUR	AGE	CVC_0	CVC_1	CVC_2
CVC_0	0	0	0		
CVC_1	1	1	1	1	1
CVC_2	4	4	4	4	4

Frequency Percent Report

	OBS	Pattern	SYMPOUR	AGE
CVC_0	50	50	50	50
CVC_1	49	49	49	49
CVC_2	46	46	46	46

Identify the type of missing values using the Missing Pattern Report

Statistical Solutions Ltd.
8 South Bank
Crosse's Green, Cork, Ireland

Tel. + 353 21 319629 • Fax + 353 21 319630
Email: sales@statsol.ie
Website: <http://www.statsol.ie>*

Statistical Solutions
Stonehill Corporate Center Suite 104
999 Broadway, Saugus, MA 01906

Tel. (781) 2317680 • Fax (781) 2317684
Email: info@statsolusa.com
Website: <http://www.statsolusa.com>





nQuery Advisor® Versió 2.0

Programari de Planificació d'Estudis

Determinació de la dimensió i la potència de la mostra

nQuery Advisor és una eina molt útil per a investigadors i estadístics a l'hora de planificar els estudis d'investigació i de garantir un ús eficaç dels recursos.

La correcta determinació de la dimensió i de la potència de la mostra representa un aspecte important dins de la planificació dels estudis.

nQuery Advisor proporciona uns mètodes simples, fiables i eficaços per determinar aquests valors, a més d'una experta assistència en computació i una àmplia cobertura quant als problemes relacionats amb la dimensió i la potència de les mostres, així com les respostes sobre la dimensió de les mostres per a les anàlisis de l'interval de confiança i d'equivalència. En la major part dels problemes que es presenten permet que els grups posseïxin una dimensió igual o desigual.

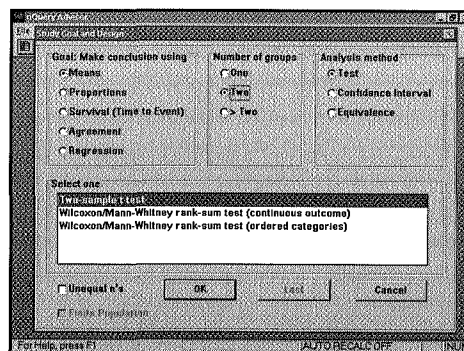
La versió 2.0 de nQuery Advisor compta amb notables millores dissenyades per tal de facilitar, encara més, la planificació dels seus estudis d'investigació. S'han inclòs noves taules amb opcions expandides per a la planificació de les dimensions de les mostres en els estudis de supervivència, el mostreig de poblacions finites, el mostreig per grups i, també, les anàlisis que utilitzen proves no paramètriques. Tant el manual com l'ajuda en pantalla s'han ampliat per proporcionar més detalls sobre l'ús de nQuery Advisor en mesures repetides i dissenys encreuats, demostracions d'equivalència i anàlisis de supervivència.

Per aquells investigadors que estan planificant assajos clínics o altres estudis amb supervivència del temps a un esdeveniment com a punt final, nQuery Advisor versió 2.0 introdueix una nova taula de dimensions de mostres per a les anàlisis de supervivència. Aquesta taula permet a l'usuari especificar unes corbes de supervivència que posseïxin una forma no exponencial o bé ràtios de risc que no es mantenen sempre constants, a més d'especificar uns patrons complexos d'acumulació o d'abandonament.

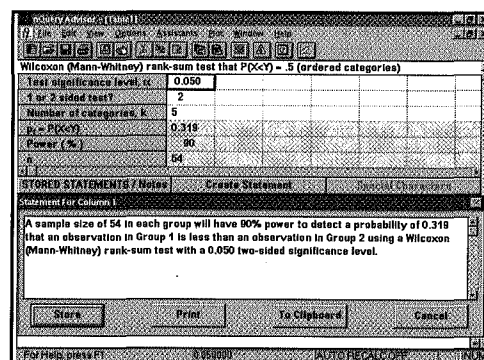
Els investigadors de ciències socials sabran, sens dubte, apreciar les noves opcions per a la planificació dels estudis basats en una població finita. La versió 2.0 de nQuery Advisor ajuda a estimar la desviació estàndard de grups per poder determinar el nombre de grups sobre els quals cal fer el mostreig.

La versió 2.0 de nQuery Advisor incorpora a més computacions de dimensions de mostra en pretendre utilitzar la prova de rang-suma de Wilcoxon/Mann-Whitney per analitzar els resultats d'un estudi de dos grups.

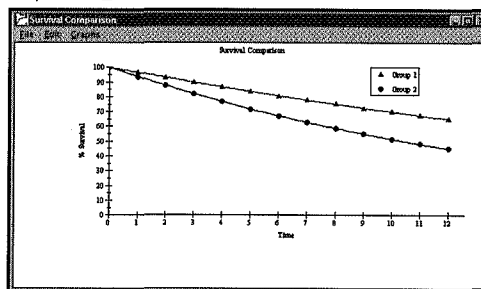
Quan s'utilitza en un entorn Windows nQuery Advisor, proporciona: Un índex estructurat per tal de facilitar l'especificació de l'objectiu de la investigació; Entrades i visualització de la dimensió i potència de mostres realitzades similars a les dels fulls de càlcul; Fitxes guia per ajudar els usuaris amb les seves preguntes estadístiques i pràctiques sobre les entrades; Eines de computació per especificar les dimensions d'efecte i les estimacions de variabilitat; Traces que mostren la relació entre la potència i la dimensió de les mostres. Quant als resultats computats proporciona informes de justificació de les dimensions de la mostra preparats per a la seva protocolització; facilitat en comparar els resultats sota diferents condicions o anàlisis; fórmules i algorismes seleccionats d'acord amb les avaluacions comparatives i accés a les funcions de distribució per als usuaris experts.



Les noves prestacions introduïdes en la versió 2.0 inclouen proves no paramètriques i mostreigs en funció de les poblacions finites.



Per als resultats computats, s'han preparat informes de justificació de la dimensió de la mostra per a la seva protocolització.



nQuery Advisor versió 2.0 permet a l'usuari especificar ja les corbes de supervivència.

nQuery Advisor és un producte desenvolupat per la Dra. Janet Elashoff de Dixon Statistical Associates, Los Angeles, EE.UU.



Statistical Solutions Ltd.
8 South Bank
Crosse's Green, Cork, Ireland

Tel. + 353 21 319629 • Fax + 353 21 319630
Email: sales@statsol.ie
Website: <http://www.statsol.ie>*

Statistical Solutions
Stonehill Corporate Center Suite 104
999 Broadway, Saugus, MA 01906

Tel. (781) 2317680 • Fax (781) 2317684
Email: info@statsolusa.com
Website: <http://www.statsolusa.com>



nQuery Advisor® Release 2.0

STUDY PLANNING SOFTWARE

Sample size and Power determination

nQuery Advisor is a very useful aid to investigators and statisticians in planning research studies to ensure efficient use of resources. Determination of the correct sample size and power are important aspects of study planning.

nQuery Advisor provides simple, reliable, and efficient methods for determining these values. It provides expert computational assistance and extensive coverage of sample size and power determination problems as well as sample size answers for confidence interval and equivalence analyses. It allows for either equal or unequal group sample sizes for most problems.

nQuery Advisor Release 2.0 features major new enhancements designed to make planning your research studies even easier. New tables have been included with expanded options for planning sample sizes for survival studies, sampling from finite populations, cluster sampling, and for analyses using nonparametric tests. The manual and on-line help have been expanded providing more details on the use of nQuery Advisor for repeated measures and crossover designs, equivalence demonstrations, and survival analysis.

For researchers planning clinical trials or other studies with survival or time to an event as an endpoint, nQuery Advisor Release 2.0 introduces a new sample size table for Survival Analysis. This table allows user specification of survival curves which are not exponential in shape, or hazard ratios which are not constant throughout, and user specification of complex patterns of accrual or dropout.

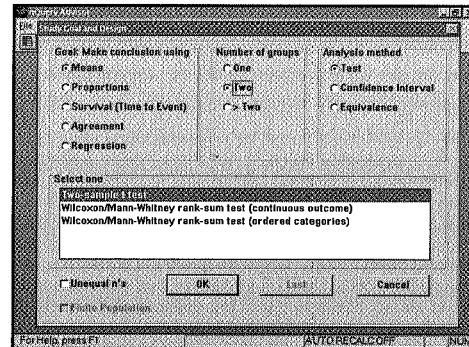
Social science researchers will appreciate the new options for planning studies for sampling from a finite population. nQuery Advisor Release 2.0 now provides sample size tables with a finite population correction for a range of study designs with testing and confidence interval goals.

The design of community intervention studies often requires randomization of clusters of subjects. nQuery Advisor Release 2.0 assists in estimation of the cluster standard deviation to use in determining the number of clusters which need to be sampled.

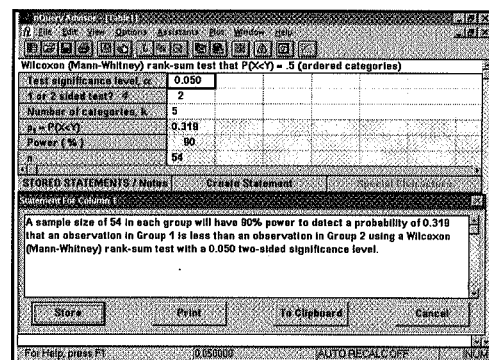
nQuery Advisor Release 2.0 also includes sample size computations when the Wilcoxon/Mann-Whitney Rank-Sum test will be used to analyze results of a two-group study.

Operating in a Windows environment nQuery Advisor® provides:

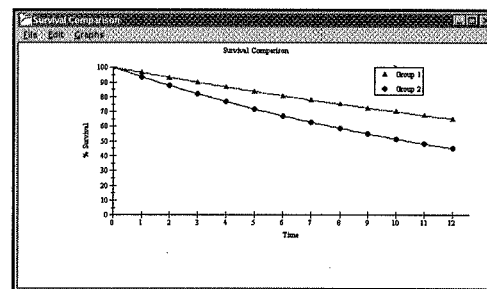
- A structured index to assist in specifying the research goal
- Spreadsheet-style tabular entry and display of sample size and power
- Guide cards to assist users with statistical and practical questions on input
- Computational aids for specifying effect sizes and estimates of variability
- Plots showing the relationship between power and sample size
- Protocol-ready sample size justification statements for the computed results
- Ease of comparison of results under different conditions or analyses
- Formulas and algorithms chosen on the basis of comparative evaluations
- Access to distribution functions for expert users



Among the new features introduced in version 2.0 are Non-Parametric tests, and sampling from a finite population.



Protocol-ready sample size justification statements for the computed results.



nQuery Advisor® Release 2.0 now allows user specification of survival curves.

nQuery Advisor® is developed by Dr. Janet Elashoff Ph.D. of Dixon Statistical Associates, Los Angeles, U.S.A.



Statistical Solutions Ltd.
8 South Bank
Crosse's Green, Cork, Ireland

Tel. + 353 21 319629 • Fax + 353 21 319630
Email: sales@statsol.ie
Website: <http://www.statsol.ie>

Statistical Solutions
Stonehill Corporate Center Suite 104
999 Broadway, Saugus, MA 01906

Tel. (781) 2317680 • Fax (781) 2317684
Email: info@statsolusa.com
Website: <http://www.statsolusa.com>