

ESTIMACION DEL ERROR TOTAL DE UNA CONTABILIDAD INCORPORANDO LA OPINION DEL EXPERTO.*

F.J. Vázquez Polo[†]

A. Hernández Bastida[‡]

RESUMEN

Este trabajo presenta algunas posibilidades de aprovechamiento de la opinión cualitativa de un auditor. El trabajo se desarrolla en torno a un caso concreto que contiene las ideas básicas sobre la metodología expuesta. La utilización práctica está ligada al empleo de un ordenador personal.

Palabras Clave: Densidad a priori conjugada, Cuantiles a priori, Clases de densidades.

ABSTRACT

This paper presents how an auditor may draw his qualitative opinions. This work is performed around a specific auditing situation which contains the basic ideas of the model. The objectives of this report are to present the model BU and two ways for studying the sensitivity of the model.

Keywords: Conjugate prior, Prior Quantiles, Classes of Priors.

1. INTRODUCCION

Una de las cuestiones a las que más frecuentemente tiene que dar respuesta un auditor en ejercicio cuando se dispone a inspeccionar o validar el establecimiento financiero de una empresa, corporación, etc... es la de establecer su opinión sobre la *exactitud* de la cantidad (monetaria) de error que posee el valor registrado (o valor de libro), que obviamente es un agregado de valores registrados y que deben de ser auditados.

La efectividad del sistema de control para la prevención y detección de errores puede hacerse mediante la solicitud de informes sobre el sistema de control, la observación del sistema de trabajo..., es lo que se conoce como **evaluación del sistema**; el muestreo de ítems para comprobar si se les ha aplicado apropiadamente o no un determinado procedimiento (**tests de cumplimiento**); la comparación de varias fuentes de información (**revisión analítica**) y el muestreo de ítems para obtener la cantidad de error monetario de la población (**test sustantivos**), permitirán al auditor (o auditores) pronunciarse sobre el estado financiero del sistema. Estas diferentes fuentes de evidencia suelen conformar el *proceso de auditoría*. La utilización del análisis bayesiano permite que estas fuentes de evidencia puedan ser modeladas probabilísticamente. El modelo que

presentamos, denominado BU (Godfrey y Neter (1984)), puede ser simplificado desde el punto de vista analítico para hacerse más tratable (Cox y Snell (1979), modelo GU, Godfrey y Neter (1984)) se ha optado por uno más sofisticado (poco desarrollado) que incluso puede hacerse aún más complejo. Aún así, el objetivo es mostrar cómo incorpora flexibilidad al proceso de auditoría un enfoque bayesiano; aunque cierto es que exige un *entrenamiento* estadístico del auditor. Steele (1992) recoge varios métodos interactivos de especificación subjetiva a través de preguntas que debe contestar el auditor.

Bajo la hipótesis, usual en auditoría, de que los errores aparecidos son de sobrevaloración (es decir, los valores registrados son mayores que los auditados) se expone un modelo mediante el que obtener, con una probabilidad específica, una cota superior para la proporción del valor de libro que tiene error o equivalentemente la cantidad de error total de la población.

En la Sección 2 se describen los elementos del problema que planteamos y en la Sección 3 la solución propuesta por el modelo desarrollado llevando simultáneamente una aplicación práctica. Por último en la Sección 4 se plantea el problema de la sensibilidad a la hora de reflejar probabilísticamente las creencias del auditor.

2. ELEMENTOS DEL PROBLEMA.

Siguiendo un esquema usual (Cox y Snell (1979)), consideraremos una población de N items (usualmente, N es grande) cada uno de los cuales tiene un valor registrado o valor de libro: $X_i > 0$. Al valor total registrado lo notaremos por

$$T_x = \sum_{i=1}^N X_i \text{ Si inspeccionamos el } i\text{-ésimo item su valor será: } X_i - Y_i \text{ (} Y_i \text{ es la cantidad}$$

de error. En el supuesto anterior de errores de sobrevaloración: $0 \leq Y_i \leq X_i$). El error

$$\text{total de la población es } T_y = \sum_{i=1}^N Y_i. \text{ La metodología de muestreo utilizada}$$

usualmente en las técnicas bayesianas para auditoría es la conocida como D.U.S. (Dollar Unit Sampling) o M.U.S. (Monetary Unit Sampling), en la terminología británica, en la que la unidad muestral es la unidad monetaria y en la que la probabilidad de seleccionar cualquier item es proporcional a su valor. Así, una vez que se ha seleccionado una unidad monetaria en la muestra, se identifica la unidad de auditoría a la que pertenece y se audita. Entonces el error que se encuentra en esa unidad de auditoría se promedia entre todas las unidades monetarias que

están en el apunte y la unidad monetaria muestreada recibe esta cantidad de error promedio o prorrateado.

Utilizando la variable binaria d_i para detectar si un ítem tiene o no error, el valor total registrado de aquellos ítems que contienen error será:

$$T_{xd} = \sum_{i=1}^N d_i \cdot X_i$$

Llamaremos Z_i a la fracción de error de cada ítem:

$$Z_i = \frac{Y_i}{X_i} \quad (i = 1, \dots, N)$$

de la hipótesis de errores de sobrevaloración, se tiene que: $0 \leq Z \leq 1$.

Sea M el número total de errores $M = \sum_{i=1}^N d_i$. Observados $M = m$ ítems con

error, supondremos que $Z_1, Z_2, \dots, Z_M$ son i.i.d. con densidad $g(z \mid d=1)$ y con media $\mu = E\{Z \mid d=1\}$.

Suponiendo que la tasa de error es constante para cada ítem e independiente de su valor, la notaremos por:

$$\varnothing = \text{Prob}\{Y_i > 0\} = \text{Prob}\{Y_i > 0 \mid X_i\}$$

Obviamente, en un ambiente de auditoría, es esperable que $\varnothing$ sea pequeña.

Observemos que hemos pasado de magnitudes discretas a dos cantidades continuas, μ y $\varnothing$. La relación entre ellas puede obtenerse fácilmente (ya que frecuentemente N es grande) de:

$$T_y = T_x \cdot \frac{T_{xd}}{T_x} \cdot \frac{T_y}{T_{xd}}$$

donde, T_y es el error total de la población, T_x es el valor total registrado, $\frac{T_{xd}}{T_x}$ es la proporción de error o tasa de error de la población y $\frac{T_y}{T_{xd}}$ es la fracción error

en items con error. Así pues (y cuanto $N \rightarrow +\infty$):

$$\frac{T_{xd}}{T_x} \rightarrow \varnothing$$

$$\frac{T_y}{T_{xd}} \rightarrow \mu$$

de donde:

$$T_y = T_x \cdot \varnothing \cdot \mu$$

como T_x es conocido, el parámetro de interés será:

$$\psi = \varnothing \cdot \mu$$

que es la proporción del valor de libro o registrado que tiene error.

Será pues sobre el parámetro ψ sobre el que el auditor deba concentrar toda la evidencia anteriormente expuesta, vía $\varnothing$ y μ , para determinar en cada etapa del proceso la situación de la empresa auditada.

Para ilustrar cada uno de los pasos a seguir, vamos a suponer el caso ficticio de un cliente fabricante de componentes electrónicos que al final de un año tiene acumulada un total de 700.000 u.m. a cobrar, distribuidas en 2.000 clientes activos. Suponemos que el rango de las cantidades registradas varía de 10 a 3.000 u.m. El propósito será comprobar la exactitud de esa cantidad. Se trata pues de establecer una cota superior, que con una probabilidad fija, nos informe sobre que las facturas a cobrar no están sobrevaloradas más que una cantidad determinada.

3. EL MODELO BU.

El problema y sus elementos están planteados. Se trata pues de modelizar mediante distribuciones de probabilidad el comportamiento de la tasa de error de la población y del valor medio de la cantidad de error de acuerdo con las creencias que el auditor va obteniendo en cada uno de los pasos del proceso de auditoría.

La hipótesis de que $\varnothing$ y μ son independientes es esencial en todo lo que sigue.

Kaplan (1973) postuló que el proceso generador de que una unidad de auditoría contenga o no error es independiente del proceso que genera la cantidad de error. Este hecho es a veces discutible en el contexto que nos ocupa, sin embargo ha venido aceptándose en los trabajos relacionados con el tema. Analíticamente el problema es muy complejo. Nos encontramos con un espacio paramétrico bidimensional (θ, μ) donde nos interesa $\psi = \theta \cdot \mu$. Especificar la densidad a priori para ψ es extremadamente difícil aunque sí que es posible especificar densidades a priori para θ y μ , en definitiva, marginales. En Rachev (1985) puede verse que sólo obtener conclusiones sobre la media a priori cuando se dispone únicamente de marginales es una cuestión bastante complicada (de hecho, nos encontramos con un caso particular del conocido como problema de Monge-Kantorovich).

3.1. La evaluación del sistema y el análisis a priori.

Será mediante la solicitud de informes técnicos, la observación de sistema de trabajo, etc... como el auditor pueda tener una idea más o menos fiel del sistema que va a auditar. Esta información la debe recoger en densidades que reflejen dichas creencias. Así el esquema que tenemos será:

$$f(\theta, \mu) = f_1(\theta) \cdot f_2(\mu)$$

De entre las distribuciones utilizadas para la modelización anterior, una aproximación *realista* (realismo que paga a cambio de ser menos tratable analíticamente) del problema se encuentra en el modelo BU en el que:

1. θ es beta: $B(\theta_0, c)$:

$$f_1(\theta) = K \cdot \theta^{\theta_0-1} (1-\theta)^{(1-\theta_0)c-1}, 0 \leq \theta \leq 1$$

$$\text{siendo } K = \frac{\Gamma(c)}{\Gamma(\theta_0 c) \cdot \Gamma((1-\theta_0)c)} \text{ y de donde:}$$

$$E\{\theta\} = \theta_0 \text{ y } \text{Var}\{\theta\} = \frac{\theta_0(1-\theta_0)}{c+1}$$

2. μ es uniforme: $\mu(0,1)$:

$$f_2(\mu) = 1, 0 \leq \mu \leq 1$$

Blocher y Robertson (1976) han descrito un método para elegir un miembro de la familia beta que permite cuantificar los juicios a priori del auditor que está interesado en la tasa de error a priori, utilizando la relación ente los parámetros y los momentos de orden uno y dos. *Steele (1992)* también describe métodos alternativos, en los que es posible incorporar, con el mismo fin, otras características usuales θ como la unimodalidad.

Así pues mediante la evaluación del sistema, el auditor ha podido identificar aproximadamente el comportamiento de los parámetros θ y μ . La distribución de ψ es:

$$f(\psi) = \int_0^1 f(\theta, \frac{\psi}{\theta}) \cdot \frac{1}{\theta} \cdot d\theta = \int_0^1 f_1(\theta) \cdot f_2(\frac{\psi}{\theta}) \cdot \frac{1}{\theta} \cdot d\theta$$

luego:

$$f(\psi) = K \cdot \int_0^1 \theta^{c_0-2} \cdot (1-\theta)^{(1-\theta_0)^{c-1}} \cdot d\theta, 0 \leq \psi \leq 1 \quad (1)$$

Volviendo a nuestro ejemplo, supongamos que al auditor no le ha causado buena impresión el sistema de trabajo y asigna un valor alto a la media de la tasa de fallo:

$$\theta_0 = E\{\theta\} = 0.375, \text{Var}\{\theta\} = 0.06 \text{ (c} \approx 3\text{)}$$

Resulta que:

$$f(\psi) = 2.22742 \cdot \int_0^1 \left(\frac{1-\theta}{\theta} \right)^{0.875} \cdot d\theta, 0 \leq \psi \leq 1$$

Esta densidad es evaluada para un número grande de valores de ψ y entonces se aproxima su distribución acumulada mediante algún algoritmo de integración numérico de tal manera que podemos obtener los cuantiles de ψ .

El cuantil 95 es: $Prob\{\psi \leq 0.543635\} = 0.95$ valor bastante grande y obviamente inaceptable. El auditor necesitará de la evidencia que aporten nuevas fases del proceso para poder emitir un juicio válido. Obviamente, la distribución de ψ también podría obtenerse de manera aproximada mediante simulación.

3.2. Incorporación de la evidencia muestral

Visto lo anterior, el auditor desea profundizar algo más y decide muestrear algunas cuentas individuales y algunos procedimientos de control (por ejemplo, podría comprobar si existía autorización apropiada para crédito y cotejar ciertas cantidades, anotando un error cuando alguno de estos controles fallase). Es la fase que comprenden los **tests de cumplimiento**. Obsérvese que un test de cumplimiento repercutirá únicamente sobre θ , ya que no se comprueban o verifican cantidades.

El número total de errores M en una muestra de tamaño n puede representarse mediante una distribución binomial: $\text{Bin}(n, \theta)$:

$$h(m | \theta) = \text{Prob}\{M = m\} = \binom{n}{m} \cdot \theta^m \cdot (1 - \theta)^{n-m}, \quad m = 0, 1, \dots, n$$

(la utilización de una verosimilitud binomial permite un cálculo fácil de la distribución a posteriori de θ , ya que son familias conjugadas).

Una vez pasados los tests de cumplimiento resulta:

$$f(\theta, \mu | m) \propto f_1(\theta) \cdot f_2(\mu) \cdot h(m | \theta) \propto f_1(\theta | m) \cdot f_2(\mu)$$

siendo $f_1(\theta | m)$ una densidad $\beta(\theta_0 c + m, (1 - \theta_0)c + n - m)$.

En nuestro ejemplo, si suponemos que en una muestra de tamaño $n = 100$ se detectaron $m = 30$, tenemos una densidad a posteriori para $\theta \sim \beta(31.125, 71.875)$. En consecuencia, una vez incorporada la evaluación del sistema y los tests de cumplimiento, resulta que ψ se distribuye según:

$$f(\psi | m) = K \cdot \int_{\psi}^1 \theta^{29.125} \cdot (1 - \theta)^{70.875} \cdot d\theta, \quad 0 \leq \psi \leq 1$$

De nuevo mediante integración numérica podemos calcular los cuantiles. El nuevo cuantil 95 es: $\text{Prob}\{\psi \leq 0.3024 | m\} = 0.95$. Cota superior sensiblemente más pequeña que la anterior pero que sigue siendo demasiado grande.

El auditor puede ahora además de medir el número de errores en la muestra, medir el valor de la fracción de error del ítem erróneo. Esta nueva etapa es conocido como fase de **tests sustantivos**.

Si un ítem auditado i es erróneo, su fracción de error Z_i se anota. Es decir, en los

tests sustantivos se actualizan la tasa de error $\varnothing$ y la media de la fracción de error, μ .

Supondremos que $Z \sim \exp \left(\frac{1}{\mu} \right)$:

$$g(z | \mu, d = 1) = \frac{1}{\mu} \cdot e^{-z/\mu}, z > 0$$

conocido el número de errores m , el estadístico $Z = \frac{1}{m} \cdot \sum_{i=1}^m Z_i \sim \text{Gamma} \left(m, \frac{m}{\mu} \right)$.

El modelo puede realizarse considerando la distribución truncada en (0,1). Para evitar unos cálculos que ya se han complicado suficientemente, se ha optado por no truncar. La verosimilitud asociada es:

$$l_z(\bar{z} | \mu) = \frac{(m/\mu)^m}{\Gamma(m)} \cdot \bar{z}^{m-1} \cdot e^{-m\bar{z}/\mu}$$

la distribución a posteriori de μ puede calcularse fácilmente el teorema de Bayes:

$$f_2(\mu | \bar{z}) \propto f_2(\mu) \cdot l_z(\mu | \bar{z}).$$

La actualización final de la proporción del valor registrado que tiene error, después de los tests de cumplimiento y sustantivos es:

$$f(\psi | m, z) = \frac{A}{B} \cdot \int_{\psi}^1 \varnothing^{\varnothing_0 c + 2(m-1)} \cdot (1 - \varnothing)^{(1 - \varnothing_0)c + n - m - 1} \cdot e^{-mz/\varnothing} \cdot d\varnothing \quad (2)$$

donde A y B valen:

$$A = \frac{\Gamma(c + n)}{\Gamma(\varnothing_0 c + m) \cdot \Gamma((1 - \varnothing_0)c + n - m)}$$

$$B = \int_0^1 e^{-mz/\mu} \cdot d\mu$$

Suponiendo en nuestro ejemplo que se ha observado $z = 21$, se obtiene:

$$\text{Prob} \{ \psi \leq 0.1782 \mid m, \bar{z} \} = 0.95$$

Debido a los resultados del test sustantivo la cota superior ha sufrido una reducción sustancial. Aunque posiblemente insuficiente, ya que el auditor puede concluir que con una probabilidad de 0.95 el establecimiento financiero no está sobrevalorado más que: $0.1782 \cdot T_x = 124740$ u.m., obviamente una cantidad muy alta.

4. CONSIDERACIONES SOBRE LAS ENTRADAS EN UN ANALISIS BAYESIANO

Ya hemos comentado que el auditor puede encontrar dificultades para poder expresar una determinada densidad a priori (continua o discreta) sobre los parámetros de interés en auditoría.

Mucho más concreto, a ningún auditor que utilice una metodología bayesiana escapa que su asignación de densidad a priori, por ejemplo para la tasa de error, no puede especificarse de manera exacta y que a lo sumo (y después de un proceso bastante laborioso) podrá obtener una buena aproximación de sus creencias a priori.

Inmediatamente, el auditor puede (y debe) plantearse las siguientes cuestiones: ¿qué ocurrirá con las cantidades calculadas para la (verdadera) densidad a priori (densidad que debe de ser muy parecida a la asignada. La idea de proximidad entre densidades no será tratada aquí)?, ¿tomaré una decisión correcta si utiliza mis creencias a priori?, ¿habrá una diferencia significativa si se toma otra densidad *próxima*?

Una vía natural para este estudio es la siguiente: Si las creencias a priori del auditor tienen una forma menos elaborada que una distribución de probabilidad a priori, ¿por qué no plantear metodologías que permitan la utilización de esas informaciones menos elaboradas? Si la especificación completa de una distribución de probabilidad a priori es, con frecuencia, una petición excesiva en la práctica, ¿podríamos ampliar las "*entradas*" del análisis bayesiano, permitiendo que la especificación a priori fuese una clase (o familia) de distribuciones en vez de una sola? Esa clase o familia estaría definida por aquellos aspectos que son claros para el auditor. Por ejemplo:

- Conocimiento de algunos cuantiles.
- Unimodalidad.

- Porcentaje de seguridad en una distribución concreta y porcentaje de inseguridad, etc.

Evidentemente, sobre esa clase el auditor debe calcular las cantidades de interés para cada uno de los elementos y luego comprobar si existe mucha diferencia entre esas cantidades (dicha diferencia puede ser medida por el extremo inferior y el superior de las cantidades sobre la clase: medida que suele denominarse *rango de variación*). Cuando la diferencia es grande el auditor deberá tomar muchas precauciones en sus decisiones finales puesto que densidades muy *parecidas* no producen cantidades próximas. Suele decirse que en el modelo se ha detectado una falta de *robustez*. Cuando la diferencia sea pequeña el auditor tiene cierta seguridad de que sus decisiones no serán sustancialmente peores con un elemento u otro de la clase. Se dice entonces, que el modelo es *robusto*.

Concretando, para varias de las clases de distribuciones de probabilidad antes mencionadas se dispone de algunas conclusiones:

4.1. Conocimiento de algunos cuantiles

Consideremos, como antes, que la verosimilitud asociada al número de errores m en una muestra de tamaño n es una binomial, $\text{Bin}(n, \theta)$, que aproximamos por una distribución de Poisson, $P(n \cdot \theta)$. Además suponemos que μ toma un valor fijo μ_0 , esta hipótesis no es excesivamente restrictiva (ver Vázquez Polo (1992), para una discusión detallada de esta hipótesis). En este caso la información a priori sobre el parámetro desconocido se refiere a la información θ .

Suponemos que la información a priori del auditor consiste en algunos cuantiles de la distribución de θ , es decir, una información del tipo:

$$p_i = \text{Prob} \{ \theta \notin \theta_i \}, i = 1, 2, \dots, k \quad (3)$$

donde,

$$(0, 1) = \bigcup_{i=1}^k \theta_i; \theta_i \cap \theta_j = \emptyset; i \neq j \quad (4)$$

$$\theta_i = (a_{i-1}, a_i], i = 1, 2, \dots, k-1, a_0 = 0$$

$$\theta_k = (a_{k-1}, a_k), a_k = 1$$

En la práctica los θ_i serán intervalos definidos por cuantiles, deciles, etc. Obsérvese que la expresión anterior induce la siguiente clase de densidades a priori que modeliza las creencias del auditor:

$$\Gamma = \left\{ \xi(\theta) = \sum_{i=1}^k p_i \cdot q_i(\theta) / q_i \text{ dens. a priori con soporte en } \theta_i \right\} \quad (5)$$

Dependiendo del número de cuantiles especificados, la información a priori será más o menos rica, en definitiva la clase Γ será más o menos amplia. Un número muy pequeño de cuantiles induce una clase Γ muy amplia, tal vez excesivamente amplia.

Por lo tanto, dependiendo del número de cuantiles, el rango de variación de las cantidades de interés será más o menos amplio, en definitiva se concluirá con un grado u otro de robustez.

Podemos plantearnos la adición, a la información a priori especificada, de la condición de unimodalidad. Esta condición es habitual en el contexto que nos ocupa y por tanto es verosímil presuponer su existencia. La incorporación de unimodalidad constituye un enriquecimiento de la información a priori disponible y por tanto conduce a rangos de variación de las cantidades de interés más pequeños. Para una discusión más detallada de los aspectos descritos en este apartado ver Vázquez-Polo y Hernández (1993).

4.2. Porcentaje de seguridad en una distribución concreta

Este apartado, conocido como *clases de contaminación*, consiste en suponer que el experto especifica una distribución a priori pero desea protegerse frente a una especificación errónea. Esta protección se concreta determinando un grado de seguridad entre 0 y 1 en la distribución especificada y permitiendo, en el grado complementario, que la distribución fuese otra cualquiera.

Más preciso, la idea es suponer que la distribución a priori de θ está dentro de una clase de distribuciones a priori de la forma:

$$\Gamma = \{ \xi(\theta) = (1 - \varepsilon) \cdot g_0(\theta) + \varepsilon \cdot q(\theta) / q \in D \} \quad (6)$$

donde $\varepsilon \in [0,1]$ expresa el grado de incertidumbre (contaminación) y D es una clase de distribuciones a priori cuya especificación dependerá del tipo de problema que se aborde.

Habitualmente D es una clase muy amplia de distribuciones pudiendo ser la clase de todas las distribuciones de probabilidad. Es evidente que un aumento de la cantidad ε , es decir una mayor contaminación conducirá a una menor robustez, o lo que es igual un mayor rango de variación de las cantidades de interés. El análisis del rango de variación en función de ε nos permite concluir sobre el nivel de contaminación admisible a partir del cual el modelo deja de ser robusto.

Como antes se ha indicado la clase D es habitualmente muy amplia y sólo debe restringirse con aquellos aspectos extraordinariamente claros de nuestras creencias a priori y por tanto irrenunciables. Por ejemplo, la condición de unimodalidad ya tantas veces comentada.

Un estudio más detallado de este apartado puede verse en *Hernández y Vázquez-Polo (1992)* para el objeto que nos ha ocupado en este trabajo y en *Hernández y Vázquez-Polo (1991)* para el problema del tamaño muestral.

BIBLIOGRAFIA

- (1) Blocher, E. y Robertson, J.C. (1976). *Bayesian Sampling Procedures for Auditors: Computer-Assisted Instruction*. The Accounting Review. April. 359-363.
- (2) Cox, D.R. y Snell, E.J. (1979). *On Sampling and the Estimation of Rare Errors*. Biometrika. 66-1. 125-132.
- (3) Godfrey, J.T. y Neter, J. (1984). *Bayesian Bounds for Monetary-Unit-Sampling in Accounting and Auditing*. Journal of Accounting Reserach. Vol. 22, # 2, 497-525.
- (4) Hernández Bastida, A. y Vázquez Polo, F.J. (1991). *Tamaño muestral óptimo en auditoría contable. Contaminaciones en la opinión del auditor*. Actas V Jornadas ASEPELT-España. Las Palmas de G.C.
- (5) Hernández Bastida, A. y Vázquez Polo, F.J. (1992). *Sobre el Modelo Gamma-Gamma inversa de Cox y Snell para la Determinación de una Cota Superior para el Error Total en Auditoría Contable*. Estudios de Economía Aplicada. VI-Jornadas ASEPELT-España, 97-106.

(6) Kaplan, R.S. (1973). *A. Stochastic Model for Auditing*. Journal of Accounting Research. Spring, 38-46.

(7) Rachev, S.T. (1985). *The Monge-Kantorovich Mass Transference Problem and Its Stochastic Applications*. Theory of Probability and Its Applications. 29, 647-671.

(8) Steele, A. (1992). *Audit Risk and Audit Evidence. The Bayesian Approach to Statistical Auditing*. Academic Press.

(9) Vázquez Polo, F.J. (1992). *Técnicas Estadísticas Bayesianas en Auditoría. Un Análisis de Robustez*. Tesis Doctoral. Universidad de Las Palmas de G.C.

(10) Vázquez Polo, F.J. y Hernández Bastida, A. (1993). *Sensibilidad en un Análisis Estadístico Bayesiano con Cuantiles Conocidos para Auditoría*. Estudios de Economía Aplicada. VII Jornadas ASEPELT-España, 311 - 319.

NOTAS DE LOS AUTORES Y FINANCIACION

(*) Este trabajo ha sido realizado con el apoyo financiero del proyecto PB91-0952 de la DGICYT. El orden de los autores es arbitrario. Los autores quieren agradecer los valiosos comentarios de un evaluador anónimo que, sin duda, han contribuido a clarificar este trabajo.

(†) Dpto. de Economía Aplicada. Secc. Matemáticas. Univ. de Las Palmas de G.C. Facultad de CC. Económicas y Empresariales. 35017 Las Palmas de G.C. Tfno.: (34) 28 451 800 Fax.: (34) 28 451 829 e-mail: vazquez@empres.ulpgc.es

(‡) Dpto. de Economía Aplicada. Univ. de Granada. Facultad de CC. Económicas y Empresariales. 18011 Granada. Tfno.: (34) 58 243 711 Fax: (34) 58 243 728.