

ANÁLISIS BAYESIANO

ÍNDICE

ANÁLISIS BAYESIANO.....	3
CONCEPTOS GENERALES.....	3
UN EJEMPLO CLÁSICO.....	4
PROBABILIDAD SUBJETIVA.....	6
INTERVALO DE PROBABILIDAD.....	10
EXPLICACIÓN DE LOS SUBMÓDULOS.....	11
1. UNA POBLACIÓN. ESTIMACIÓN DE UNA PROPORCIÓN.....	12
1.1. Solución convencional.....	12
1.2. Solución bayesiana.....	13
1.3. El enfoque bayesiano bajo supuestos artificialmente simplificados.....	13
1.4. El paso a la continuidad.....	17
1.5. Intervalo de probabilidad.....	20
2. UNA POBLACIÓN. VALORACIÓN DE HIPÓTESIS SOBRE UNA PROPORCIÓN.....	23
2.1. Proporción igual a una constante.....	23
2.2. Proporción dentro de un intervalo.....	25
3. DOS POBLACIONES. ESTIMACIÓN DE UNA DIFERENCIA DE PROPORCIONES.....	26
4. DOS POBLACIONES. VALORACIÓN DE HIPÓTESIS SOBRE UNA DIFERENCIA DE PROPORCIONES.....	32
4.1. Igualdad de proporciones.....	32
4.2. Diferencia de proporciones dentro de un intervalo.....	33
5. UNA POBLACIÓN. ESTIMACIÓN DE UNA MEDIA.....	35
6. UNA POBLACIÓN. VALORACIÓN DE HIPÓTESIS SOBRE UNA MEDIA.....	40
6.1. Media igual a una constante.....	40
6.2. Media dentro de un intervalo.....	42
7. DOS POBLACIONES. ESTIMACIÓN DE UNA DIFERENCIA DE MEDIAS. MÉTODO EXACTO.....	43
8. DOS POBLACIONES. ESTIMACIÓN DE UNA DIFERENCIA DE MEDIAS. MÉTODO APROXIMADO.....	45
9. DOS POBLACIONES. VALORACIÓN DE HIPÓTESIS SOBRE UNA DIFERENCIA DE MEDIAS.....	51
10. TABLAS DE CONTINGENCIA. VALORACIÓN DE HIPÓTESIS DE INDEPENDENCIA.....	52
11. VALORACIÓN BAYESIANA DE PRUEBAS CONVENCIONALES.....	54
BIBLIOGRAFÍA.....	58

ANÁLISIS BAYESIANO

CONCEPTOS GENERALES

La inferencia estadística devino en un recurso extremadamente útil para los editores de revistas y responsables administrativos, cuando a principios del siglo XX la ausencia de una herramienta que aquilatara cuantitativamente el significado de los hallazgos produjo que las anécdotas clínicas poblaran las revistas médicas. Se hacía necesario, por tanto, usar procedimientos que cuantificaran la evidencia y complementaran los razonamientos verbales, de modo que los protegiera de la subjetividad.

Sin embargo, lo cierto es que, aunque la objetividad es un deseo natural y legítimo, lamentablemente resulta inalcanzable en estado puro. La estadística no puede resolver este conflicto, pues todo proceso inferencial, incluso cuando se lleva adelante con su concurso, tendrá siempre un componente subjetivo. Si bien las técnicas estadísticas suelen ser muy útiles, ocasionalmente pueden defraudar al usuario. En efecto, pueden despertar expectativas que a la postre no se cumplan, especialmente cuando el investigador renuncia a examinar la realidad a través de un pensamiento integral y deja todo en manos del veredicto formal de los procedimientos estadísticos. Los argumentos aportados en este sentido por Berger y Berry¹ son sumamente persuasivos.

Los trabajos desarrollados por Fisher en los años 20 y por los matemáticos Neyman y Pearson en la década del 30, dieron lugar al método actualmente conocido como *prueba de significación*, que responde al llamado “paradigma frecuentista”. Este método engloba un índice para medir la fuerza de la evidencia, el llamado *valor p* , y un procedimiento de elección entre hipótesis, llamado *prueba de hipótesis* (PH).

Sin embargo, la metodología de las pruebas de hipótesis vive en la actualidad, según algunos especialistas, cierta crisis; desde esa perspectiva, por tanto, resulta atractiva la idea de manejar un nuevo paradigma o de corregir el actual. En este sentido el enfoque bayesiano se perfila como una alternativa altamente promisorio.

Las críticas más reiteradas al paradigma clásico han sido que obliga a una decisión dicotómica y que constituye una medida que depende vitalmente de un elemento exógeno a los datos: el tamaño de muestra. A un efecto pequeño observado en un estudio con un tamaño de muestra grande puede corresponder el mismo valor p que a un gran efecto registrado a través de una muestra pequeña. Por otra parte, basta que haya una diferencia mínima, intrascendente (o el más mínimo sesgo), para que tal diferencia pueda ser declarada “significativa”, siempre que haya recursos como para tomar una muestra suficientemente grande. Esta problemática se discute e ilustra más adelante.

Se supone que, al aplicar cierto procedimiento estadístico a un conjunto de datos, lo que se procura es que el análisis gane en objetividad; es decir, que los puntos de vista del investigador no puedan modificar sustancialmente las conclusiones. Pero la verdad es que los métodos estadísticos convencionales están lejos de garantizar automáticamente tal *desideratum*. Es bien sabido que la forma de operacionalizar las variables, los puntos de corte que se eligen, el nivel de significación empleado, las escalas de medición adoptadas, las pruebas de significación seleccionadas, son solo algunos ejemplos de la larga lista de instrumentos estadísticos que irremediablemente han de elegirse según un punto de vista que varía entre investigadores. Sin

embargo, donde tal carencia de normas uniformes es más acusada es en el punto culminante del proceso: a la hora de realizar inferencias una vez examinados los resultados².

UN EJEMPLO CLÁSICO

En Benavides y Silva³ se desarrolla un ejemplo que contribuye a exponer más claramente el modus operandi de la técnica clásica y que se retoma más adelante para ilustrar el enfoque bayesiano.

Supóngase que hay motivos teóricos e indicios empíricos nacidos del trabajo de enfermería que hacen pensar que los pacientes afectados por quemaduras se recuperan más rápidamente cuando el tratamiento combina cierta crema antiséptica con un apósito hidrocoloide que cuando solo se utiliza la crema antiséptica. Se diseña, entonces, un experimento con la esperanza de rechazar la hipótesis nula que afirma que el tratamiento simple es tan efectivo como el combinado.

Imagínese que se tienen $n=80$ pacientes; aleatoriamente se eligen 40 de ellos para ser atendidos con el tratamiento experimental (combinación de crema antiséptica y apósito hidrocoloide) en tanto que a los 40 restantes se les aplica el tratamiento convencional (crema únicamente).

Una vez obtenidos los datos (porcentajes de recuperación en el grupo experimental y en el control, p_e y p_c , y su diferencia, $d_0=p_e-p_c$), se calcula la probabilidad asociada a ese resultado bajo el supuesto de que se cumple H_0 . Supóngase que el 75% ($p_e=0,75$) de los pacientes a los cuales se les aplicó el tratamiento experimental mejora apreciablemente a los 5 días, mientras que para los pacientes tratados de manera convencional, esta tasa de recuperación fue del 60% ($p_c=0,60$). La Tabla 1 recoge la información relevante de este ejemplo.

Tabla 1. Distribución de una muestra de 80 pacientes según tratamiento asignado y según se recuperaran o no.

Tratamiento	Se recuperan		Total
	Sí	No	
Experimental	30	10	40
Convencional	24	16	40
Total	54	26	80

Según la práctica regular, ahora procede aplicar una prueba estadística; la más usada para valorar la diferencia entre porcentajes es la prueba Ji-cuadrado (enteramente equivalente a la prueba basada en el estadístico Z que opera con la diferencia entre los porcentajes y se distribuye aproximadamente según una normal estándar). Es fácil constatar que en este ejemplo $\chi^2(\text{obs})=2,05$ y que el correspondiente valor p es igual a 0,15. Puesto que dicho valor no es suficientemente pequeño como para considerar que “hay significación” a ninguno de los niveles habituales (0,10; 0,05 y 0,01) y a pesar de que esta diferencia objetivamente observada es notable, según la práctica al uso, el investigador tiene que concluir (aunque casi con seguridad, y con razón, a regañadientes) que no tiene suficiente evidencia muestral como para afirmar que el tratamiento con crema y apósito sea más efectivo que el tratamiento con crema solamente. Es decir, debe actuar como si el experimento no arrojará información adicional alguna para pronunciarse, lo cual es obviamente inexacto.

Por otra parte, este método, las decisiones se adoptan sin considerar la información externa a las observaciones o al experimento; de ahí que una de las objeciones más connotadas que se hace al

método es que no toma en cuenta de manera formal en el modelo de análisis la información anterior a los datos, proveniente de estudios previos o de la experiencia empírica informalmente acumulada, que **siempre** se tiene sobre el problema que se examina.

Para ilustrar la dependencia que tiene el análisis del tamaño muestral, bastará reproducir la Tabla 1 pero usando mayores valores de muestra total y conservando la misma distribución en su interior. El valor de p se puede reducir tanto como se desee. Por otra parte, supóngase ahora que el estudio se realizó con $n=400$ pacientes y que arrojó los resultados de la Tabla 2.

Tabla 2. Distribución de una muestra de 400 pacientes según tratamiento aplicado y según se recuperaran o no.

Tratamiento	Se recuperan		Total
	Sí	No	
Experimental	103	97	200
Convencional	120	80	200
Total	223	177	400

Las estimaciones son $p_e=0,52$ y $p_c=0,60$. Como se ve, los resultados están en clara colisión con las expectativas del investigador. A juzgar por estas tasas, habría que pensar en principio que el apósito podría ser dañino. Al realizar la prueba de hipótesis formal se obtienen $\chi^2(\text{obs})=2,93$ y $p=0,09$. En este caso, si el investigador usara umbrales fijos, no podría rechazar la hipótesis H_0 al nivel más socorrido ($\alpha=0,05$) pero sí al nivel $\alpha=0,1$. Si acude al recurso más usual (consignar el valor exacto de p), podría declarar algo como lo siguiente: “las tasas observadas difieren significativamente ($p=0,09$)”. En cualquier variante, el rechazo de la hipótesis de igualdad entre tratamientos debería ser a favor de que el que utiliza solamente crema es más efectivo que el que incluye el apósito.

Esta conclusión, sin embargo, **contradice los conocimientos previos, la expectativa racional y la experiencia del investigador**, todo lo cual lo colocaría probablemente en un conflicto: ¿matiza verbalmente el resultado hasta hacerle perder de hecho todo valor? ¿se escuda en que no se llegó al umbral “mágico” de 0,05 y dice simplemente que la diferencia no es significativa? ¿se abstiene de comunicar los resultados y actúa como si no se hubieran producido? A juicio de los críticos del frecuentismo, cualquiera de estas tres variantes sería inconsistente con la “obligación metodológica” que contrajo el investigador al elegir las PH como medio valorativo del procedimiento terapéutico en estudio.

El atractivo de contar con un enfoque alternativo puede ser fácilmente comprendido: resulta natural que se aspire a contar con un procedimiento inferencial libre de las serias impugnaciones que se hacen a las pruebas de significación.

Para conjurar la orfandad en que quedaría el investigador tras el abandono de las PH, se manejan varias alternativas. Una de ellas, sin duda la más sencilla de todas, es simplemente no usar la prueba de hipótesis y circunscribirse a la construcción de intervalos de confianza (IC). De manera informal, un intervalo de confianza para un parámetro P se define como una pareja de números $\hat{P}_1$ y $\hat{P}_2$ entre los cuales se puede “estar confiado” que se halla el parámetro en cuestión. Si bien los IC se inscriben en la órbita de la misma vertiente frecuentista que las pruebas de hipótesis, se apartan de la interpretación automática de los valores p y constituyen un recurso para aquilatar, justamente, el grado en que el conocimiento de la verdadera

diferencia es adecuado. Una ventaja obvia de los intervalos de confianza, sin embargo, es que los resultados se expresan en las mismas unidades en las cuales se hizo la medición y, por tanto, permiten al lector considerar críticamente la relevancia clínica de los resultados.

La otra alternativa es la inferencia bayesiana, objeto del presente módulo. Según comentan detalladamente Benavides y Silva³, se trata de una aproximación metodológica exenta de casi todas las críticas que se le hacen a las pruebas de significación y que goza del atractivo de incorporar las evidencias aportadas por experiencias previas dentro del proceso analítico y las contempla, por ende, en las conclusiones.

Adviértase que, típicamente, el clínico admite con naturalidad que tiene un criterio a priori sobre un paciente; realiza exámenes complementarios y actualiza su visión inicial sobre el diagnóstico que corresponde a ese paciente al conjugar las dos cosas (visión inicial e información complementaria) en lo que constituye un proceso de inducción integral. Parece bastante natural la aspiración de que un investigador se conduzca de manera similar; esa es exactamente la forma en que opera el pensamiento bayesiano.

Los recursos para hacer inferencia estadística bayesiana se conocen desde hace más de 200 años. El reverendo Thomas Bayes resolvió cuantitativamente por entonces el problema de determinar cuál de varias hipótesis es más probable sobre la base de los datos. Su descubrimiento básico se conoce como el Teorema de Bayes.

Bayes nació en Londres en 1702 y falleció el 17 de abril de 1761 en Tunbridge Wells, Kent. Fue distinguido como *Fellow* de la Royal Society en 1742, aunque hasta ese momento no había dado publicidad a trabajo alguno bajo su nombre. Su artículo más emblemático⁴ se titulaba *Ensayo hacia la solución de un problema en la doctrina del azar* (*Essay towards solving a problem in the doctrine of chances*) y fue publicado póstumamente.

El pensamiento bayesiano tiene más similitud que el frecuentista con el tipo de situaciones en que se ve el científico habitualmente: lo que tiene son datos (pacientes con ciertos rasgos, medias muestrales, series de datos) y lo que quiere es descubrir qué circunstancias determinaron que los datos fueran esos y no otros (es decir, quiere hacer juicios acerca de las leyes que gobiernan el proceso que produjo los datos que observa). La diferencia esencial entre el pensamiento clásico y el bayesiano radica en que aquél se pronuncia probabilísticamente sobre los datos a partir de supuestos (la “p” no es otra cosa que eso); en tanto que éste se pronuncia (también probabilísticamente) sobre los supuestos partiendo de los datos.

Aunque las bases de este enfoque datan de hace más de dos siglos, es ahora que empieza a asistirse a un uso apreciable del mismo en la investigación biomédica. Una de las razones que explican tal realidad y que a la vez augura un prominente futuro, es que algunos de los problemas de cierta complejidad que posee este método exigen el uso de recursos computacionales accesibles sólo muy recientemente para el común de los investigadores.

PROBABILIDAD SUBJETIVA

La inferencia bayesiana constituye un enfoque alternativo para el análisis estadístico de datos que contrasta con los métodos convencionales de inferencia, entre otras cosas, por la forma en que asume y maneja la probabilidad⁵.

Existen dos definiciones de la noción de probabilidad: objetiva y subjetiva. La definición frecuentista considera que la probabilidad es el límite de frecuencias relativas (o proporciones) de

eventos observables. Pero la probabilidad puede también ser el resultado de una construcción mental del observador, que corresponde al grado de “certeza racional” que tenga acerca de una afirmación, donde el término “probabilidad” significa únicamente que dicha certidumbre está obligada a seguir los axiomas a partir de los que se erige la teoría de probabilidades; como en este marco las probabilidades pueden variar de una persona a otra, se le llaman *certidumbres personales, credibilidad, probabilidades personales o grados de creencia*^{6,7}. La alternativa bayesiana está abierta a manejarse con esta última variante, aunque no excluye la primera.

Greenland⁶ alerta acerca de que los términos “objetivo” y “subjetivo” tienen una connotación que favorece el prejuicio y tiende a alejar a lectores ingenuos de la perspectiva subjetiva. Pero detrás de esto se esconde cierta perversión semántica. La palabra “subjetivo” sugiere una impronta de arbitrariedad o irracionalidad; sin embargo, esta es una idea simplemente errónea. El hecho de que diferentes sujetos pudieran atribuir probabilidades diferentes al mismo evento no significa que se conduzcan arbitrariamente. La palabra “objetivo” sugiere observabilidad directa; sin embargo, los límites de las frecuencias relativas se definen en términos de secuencias infinitas, las cuales no son observables directamente.

Con frecuencia se admite que se hable de cuán probable es que, por ejemplo, un condiscípulo apruebe un examen; esto se considera una pregunta razonable y para emitir una respuesta nadie exigiría someterlo 1.000 veces a ese examen para poder contar cuántas veces aprobó y poder computar el porcentaje de éxitos. Igualmente, nadie considera insensato afirmar “es muy probable que Juan gane las próximas elecciones”; ni se tilda de irracional a un médico que considere como “improbable” que un paciente en particular sobreviva (de hecho, los médicos actúan cotidianamente de esta manera). Dado que resulta imposible su cuantificación formal, las afirmaciones de este tipo no tienen sentido en el marco frecuentista; sin embargo, son muy usadas en el lenguaje común y, de hecho, desempeñan un papel real (aunque informal) en la toma de decisiones.

La inducción bayesiana consiste en usar recursos probabilísticos para actualizar (cambiar) nuestra asignación probabilística inicial o previa (haya sido ésta “objetiva” o “subjetivamente establecida”) a la luz de nuevas observaciones; es decir, computar nuevas asignaciones condicionadas por nuevas observaciones. El teorema de Bayes es el puente para pasar de una probabilidad a priori o inicial, $P(H)$, de una hipótesis H a una probabilidad a posteriori o actualizada, $P(H|D)$, basado en una nueva observación D . Produce una probabilidad conformada a partir de dos componentes: una que con frecuencia se delimita subjetivamente, conocida como “probabilidad a priori”, y otra objetiva, la llamada verosimilitud, basada exclusivamente en los datos. A través de la combinación de ambas, el analista conforma entonces un juicio de probabilidad que sintetiza su nuevo grado de convicción al respecto. Esta probabilidad a priori, una vez incorporada la evidencia que aportan los datos, se transforma así en una probabilidad a posteriori.

La regla, axioma o teorema de Bayes (se le ha denominado de todas esas formas) es en extremo simple, y se deriva de manera inmediata a partir de la definición de *probabilidad condicional*. Esta sería:

Si se tienen dos sucesos A y B (donde A y B son ambos sucesos posibles, es decir, con probabilidad no nula), entonces la *probabilidad condicional* de A dado B , como es bien conocido, se define del modo siguiente:

$$P(A | B) = \frac{P(A \cap B)}{P(B)} \quad [1]$$

Análogamente,

$$P(B | A) = \frac{P(A \cap B)}{P(A)}$$

de modo que, sustituyendo en [1] la expresión $P(A \cap B) = P(B | A)P(A)$, se llega a la forma más simple de expresar la regla de Bayes:

$$P(A | B) = \frac{P(B | A)P(A)}{P(B)} \quad [2]$$

Ahora bien, supóngase que $A_1, A_2, \dots, A_k$ son k sucesos mutuamente excluyentes, uno de los cuales ha de ocurrir necesariamente; entonces la conocida (e intuitiva) *ley de la probabilidad total* establece que:

$$P(B) = \sum_{i=1}^k P(B | A_i) P(A_i)$$

De modo que, tomando el suceso A_j en lugar de A en la fórmula [2] y aplicando al denominador la mencionada ley, se tiene:

$$P(A_j | B) = \frac{P(B | A_j)P(A_j)}{\sum_{i=1}^k P(B | A_i) P(A_i)} \quad [3]$$

que es otra forma en que suele expresarse la regla de Bayes.

Para el caso particular en que $k=2$, la expresión [3] se reduce a:

$$P(A|B) = \frac{P(B|A)P(A)}{P(B|A)P(A) + P(B|\bar{A})P(\bar{A})}$$

donde $\bar{A}$ representa el complemento de A , para dar lugar a la forma más conocida del teorema de Bayes, ya mencionada antes.

La regla de Bayes produce *probabilidades inversas*, en el sentido de que expresa $P(A | B)$ en términos de $P(B | A)$.

La terminología convencional para $P(A | B)$ es la probabilidad a posteriori de A dado B y para $P(A)$ es la probabilidad a priori de A , (debe su nombre a que se aplica antes, sin estar condicionada por la información de que B ocurrió).

Una forma alternativa de expresar el teorema en el caso en que se trabaja con un evento y su complemento es especialmente útil cuando se trabaja con *odds*⁸; en este caso la expresión sería:

$$\frac{P(H | \text{datos})}{P(K | \text{datos})} = \frac{P(\text{datos} | H)}{P(\text{datos} | K)} \frac{P(H)}{P(K)}$$

donde H representa que una hipótesis dada es válida y K que lo es la hipótesis complementaria. Bajo el enfoque bayesiano, se suele valorar el apoyo relativo que dan los datos a cada una de estas dos hipótesis en términos del llamado *Factor de Bayes* (BF), el cual compara las probabilidades de los datos observados bajo las hipótesis:

$$BF = \frac{P(\text{datos} | H)}{P(\text{datos} | K)}$$

Sintetizando, lo que proclama el enfoque bayesiano es que resulta útil, además de ser consistente con las demandas de la intuición, poder contar con un método que combine las evidencias subjetivamente acumuladas con la información objetiva obtenida de un experimento en particular.

Por ejemplo, al estimar un parámetro, su valor, naturalmente, se desconoce; pero es absurdo suponer un vacío total de ideas o de presunciones en la cabeza de quienes tienen esa ignorancia, y una persona puede expresar su opinión acerca de su conocimiento e incertidumbre por medio de una distribución probabilística de un conjunto de valores posibles del parámetro. La distribución *a priori* es la distribución probabilística que refleja y sintetiza la idea acerca de lo que dicha persona razonablemente piensa antes de observar los datos. Está claro que la realidad objetiva (muchas veces representada a través de lo que los estadísticos llaman “parámetros”) es de determinada manera (ajena a las impresiones que uno tenga de ella), y también es obvio que no se comporta aleatoriamente; pero nuestra visión sobre ella sí, y en gran medida depende de nuestra comprensión del problema y sobre todo, del conocimiento general prevaleciente en el momento del estudio.

Procede, sin embargo, aclarar que es posible (y ocasionalmente recomendable) emplear información empírica para conformar esta opinión “subjetiva”. En esto se basa el **análisis bayesiano empírico**, consistente en la determinación de la distribución *a priori* a partir de la evidencia empírica y el uso de esta misma evidencia para los cálculos bayesianos formales⁹.

En este punto, y a modo de síntesis, cabe preguntar ¿qué se supone que representa la distribución *a priori*?. Una posible respuesta es: la distribución probabilística *a priori* representa el estado de incertidumbre de un individuo particular acerca de ciertos parámetros. Por tanto, si dos individuos tienen diferentes creencias acerca de un parámetro, y estas creencias son representadas honestamente por medio de diferentes distribuciones probabilísticas *a priori*, entonces ambas distribuciones son, en cierto sentido y de momento, igualmente válidas.

⁸ Recordemos que los “odds” asociados a cierto suceso se definen como la razón entre la probabilidad de que dicho suceso ocurra y la probabilidad de que no ocurra⁸.

INTERVALO DE PROBABILIDAD

Recuérdese que en el marco frecuentista un intervalo de confianza se interpreta en términos de repeticiones hipotéticas del estudio. Según este enfoque clásico: “si se repitiera muchas veces el muestreo, se aplicara el mismo procedimiento y se calcularan respectivos intervalos de confianza al 95% según las fórmulas conocidas, 95 de cada 100 intervalos incluirían al verdadero parámetro que está siendo estimado”. Esta afirmación es coherente con la definición de probabilidad como frecuencia “a la larga” que usa esta escuela.

Procede señalar en este contexto que la interpretación de los intervalos de confianza también se presta a confusión. Cuando un investigador calcula un intervalo de confianza al 95%, suele pensar que este intervalo obtenido a partir de una muestra particular contiene al parámetro de interés con una alta probabilidad. Esta interpretación, aunque intuitiva y muy extendida, es errónea. Una vez seleccionada la muestra, es el intervalo de confianza el resultado de la experiencia aleatoria; lo que puede afirmarse es que, aproximadamente para el 95% de las muestras posibles, el intervalo resultante contendría al parámetro desconocido. Con este enfoque, solo se puede *confiar* en el procedimiento seguido: si este intervalo se calculara repetidamente, entonces se sabe que aproximadamente en el 95% de los casos se obtendrán intervalos que contendrán el valor desconocido; si se hace una sola vez (que es lo que se produce siempre), entonces solo puede “confiarse” en que **este** intervalo haya sido uno de los “exitosos”.

Para los bayesianos, que admiten que se considere a la probabilidad como un grado de creencia, el enfoque correcto es diferente: se considera la curva que representa la función de densidad que se obtiene a posteriori, y si el área bajo dicha curva entre los valores X e Y es igual a 95%, entonces se puede hablar de que el verdadero valor esté entre X e Y con cierta probabilidad (por ejemplo, del 95%). Se dice entonces que (X, Y) constituye un *intervalo de credibilidad* al 95% o un *intervalo de confianza bayesiano*¹⁰. Al intervalo construido bajo estos supuestos también se le llama *intervalo probabilístico*.

Nota: En el desarrollo de la explicación sobre la estimación de una proporción se da una definición formal de este intervalo.

La literatura referente a los métodos bayesianos crece a diario. Un solo ejemplo claramente elocuente del auge que han tomado dichos métodos lo proporciona la Revista *Annals of Internal Medicine*, una de las de mayor prestigio internacional. En un editorial firmado por Davidoff¹¹ se plantea: “Convencidos de que la inferencia inductiva es útil y factible para la interpretación de estudios científicos, en 1997 comenzamos a alentar a los autores que envían sus manuscritos a *Annals* a que incluyan la interpretación bayesiana en sus resultados”. Un hecho notable es que en febrero del 2005, una búsqueda en *Google*, arrojó que hay 31.500 sitios en cuyo nombre está la palabra *Bayes*, y 173.000 en que aparece *bayesian*; además hay 2.450.000 sitios WEB donde se menciona en algún lugar dicha palabra *bayesian*. Una medida del crecimiento de la presencia de los métodos bayesianos en los artículos médicos puede apreciarse consultando la base *Pubmed*; operando con la clave “*bayes OR bayesian*” se hallaron los resultados que recoge la Tabla 3.

Tabla 3. Número de artículos registrados en PUBMED en cuyos títulos y resúmenes aparece bayes o bayesian según tres quinquenios consecutivos.

Artículos en PUBMED	Período		
	1990-94	1995-99	2000-04
Títulos	232	397	975
Resúmenes	617	1.114	2.787

En lo que sigue dentro de este texto de ayuda, han de tenerse en cuenta algunas advertencias:

- En aquellos ejemplos en que el algoritmo transita por técnicas de simulación, los resultados, naturalmente, difieren en las aplicaciones sucesivas debido, precisamente, a que el ordenador simula cada vez datos diferentes; ello explica asimismo que si el usuario prueba el programa usando los ejemplos aquí incluidos, no va a obtener exactamente los mismos resultados que dichos ejemplos ofrecen.
- El número de simulaciones que han de hacerse debe ser decidido por el usuario; la única restricción que pone Epidat 3.1 es que no sea inferior a 500; en muchos casos no tiene sentido poner más de 5.000 o 10.000 porque a partir de esos valores los resultados prácticamente no difieren. En cualquier caso, no existe un número “óptimo”, pues nunca va a afectar negativamente que se decida un número mayor.
- Ciertas zonas del texto contienen formulaciones complejas, algunas matemáticamente avanzadas; el usuario que no tenga suficiente versación matemática, debe pasar por alto tales segmentos, que se han incluido en beneficio de quienes deseen acceder a una fundamentación formal de naturaleza más técnica, ya que no se trata de desarrollos fáciles de hallar en la literatura convencional.
- El “tono” con que en algunos puntos se exaltan las virtudes del enfoque bayesiano no equivale a que los autores del programa suscriban enteramente dicho enfoque, el cual también es objeto de algunas valoraciones críticas. Es reflejo de los puntos de vista que típicamente pueden hallarse en la literatura correspondiente y se emplea para que el usuario conozca cabalmente dichas opiniones.
- En algunos submódulos, los datos pueden entrarse según dos modalidades. Típicamente, la forma escogida es irrelevante (el resultado será el mismo), pero a veces pueden producirse pequeñas diferencias en dependencia de cuál se escoja. Dichas diferencias son solo atribuibles a cuestiones de redondeo y no afectan lo esencial de los resultados.

EXPLICACIÓN DE LOS SUBMÓDULOS

A continuación se da una explicación de las bases teóricas de cada uno de los submódulos incluidos en Epidat 3.1 bajo el rubro de *Análisis bayesiano*. En cada caso se dará una explicación teórica y luego se pone un ejemplo. Tal explicación es más extensa en algunos casos que en otros (por ejemplo el de la estimación de una proporción o la diferencia de proporciones) ya que se aprovechan para dar explicaciones de valor general. No obstante, las bases teóricas de los métodos hay que buscarlos en las referencias que se dan al final de esta ayuda.

1. UNA POBLACIÓN. ESTIMACIÓN DE UNA PROPORCIÓN

El problema posiblemente más simple de la inferencia (aunque a la vez uno de los que se presenta con mayor frecuencia) es estimar el valor desconocido de una proporción poblacional θ . La explicación correspondiente a este sencillo problema, permitirá captar algunas ideas generales del enfoque bayesiano (y por ende, válidas para los demás submódulos). Tal circunstancia determina que este punto será bastante más extenso que el resto.

El proceso para resolver la tarea mencionada exige en cualquier variante obtener una muestra de n sujetos de la población, en la que se evalúa una variable dicotómica. Los resultados se dividen en e “éxitos” y $f=n-e$ “fracasos”.

Supóngase entonces que θ es un porcentaje poblacional desconocido (por ejemplo, la tasa de prevalencia de asma en una comunidad) y que se ha concluido una indagación o encuesta para recoger información que permita estimarlo.

1.1. Solución convencional

El procedimiento regular que suele aplicarse es bien conocido: consiste en realizar una estimación puntual y luego estimar el grado de error atribuible a dicha estimación y, con esos datos, construir un intervalo de confianza.

Supóngase que se ha hecho un estudio con una muestra de tamaño $n=200$, entre cuyos miembros se detectaron $e=11$ casos positivos. El procedimiento convencional conduce a construir un intervalo de confianza para la tasa de asmáticos (θ).

Las fórmulas bien conocidas para los límites inferior y superior (I y S , respectivamente) serían:

$$I = p - Z_{1-\frac{\alpha}{2}} \sqrt{\frac{p(1-p)}{n}} \quad \text{y} \quad S = p + Z_{1-\frac{\alpha}{2}} \sqrt{\frac{p(1-p)}{n}}$$

En este caso,

$$n=200, \quad p = \frac{11}{200} = 0,055 \quad \text{y} \quad Z_{1-\frac{\alpha}{2}} = 1,96 \quad (\text{suponiendo que se trabaja con el consabido } \alpha=0,05),$$

de modo que un intervalo de confianza al 95% en este marco convencional vendría delimitado por los valores: $I=0,023$ y $S=0,087$.

Este procedimiento permite “confiar” en que el verdadero porcentaje se halla entre 2,3% y 8,7%. Esto quiere decir que el intervalo en cuestión ha sido obtenido por un método que el 95% de las veces da lugar a intervalos que contienen al verdadero parámetro; y de ahí dimana la confianza que se tiene en que *esta vez* haya ocurrido así. No se tiene derecho, sin embargo, a hablar de la *probabilidad* de que el intervalo contenga o no a la tasa desconocida.

El enfoque convencional parte de un supuesto implícito: no se tiene información alguna sobre la tasa que se quiere estimar. Desde luego, dicha tasa no se conoce exactamente. El objetivo es, justamente, incrementar el conocimiento que se tiene sobre el tema con anticipación; pero eso no equivale a que no se posea idea alguna. Esa es la distinción crucial entre esta solución y la bayesiana.

1.2. Solución bayesiana

La inferencia bayesiana exige tener *a priori* una opinión acerca de la proporción θ (reflejo de nuestro conocimiento antes de realizar observación alguna); una vez que ésta se ha formulado, se modifica nuestro conocimiento acerca de θ , y el mecanismo utilizado para la “actualización” es el teorema de Bayes. De modo que, una vez aplicado el teorema, toda la información acerca de la proporción desconocida está contenida en la distribución probabilística que se obtiene para θ , la llamada *distribución a posteriori*. A partir de ella se pueden realizar diferentes tipos de afirmaciones, como se ilustra de inmediato.

Como se ha dicho, la solución clásica parte de un desconocimiento total sobre el porcentaje de interés. Es razonable suponer, sin embargo, que de antemano se tengan ciertas ideas acerca de cuál puede ser ese valor desconocido; por ejemplo, podría considerarse altamente verosímil que la tasa de prevalencia de asma en la comunidad (θ), no esté muy distante de 9%. Desde luego, tal criterio inicial, nacido quizás de la experiencia propia y ajena, así como del encuadre teórico del problema, se produce en un marco de incertidumbre. Así, tal vez se piensa que 7% ó 13% son también valores posibles, aunque acaso menos verosímiles que 9%. Y también se pudiera estar convencidos de que 55% sería sumamente improbable, a la vez que se pudiera estar persuadidos de que cifras como 1% u 80% serían virtualmente imposibles. Tener de antemano puntos de vista como estos no solo es posible sino casi inevitable para alguien versado en la materia.

En la valoración realizada en el párrafo precedente se ha empleado una terminología probabilística informal, típica en las expresiones cotidianas; pero bien podría llevarse a un plano formal y el lenguaje evocativo de las probabilidades (“sumamente improbable”, “altamente verosímil”, “virtualmente imposible”, etc.) podría ser convertido en números concretos.

De modo que el problema que se aborda es el de ganar conocimiento sobre *una tasa de prevalencia*; se sabe que se necesita información procedente de la realidad como parte del proceso inferencial, pero se está a la vez consciente de que se poseen criterios racionales potencialmente valiosos a los efectos de resolver eficientemente dicho problema. El empleo de estos criterios en el proceso formal de estimación es posible pero, como se verá, exige transitar por el ejercicio de traducir esos criterios, en una u otra medida subjetivos, a un marco cuantitativo.

1.3. El enfoque bayesiano bajo supuestos artificialmente simplificados

Para utilizar el método bayesiano, por tanto, el primer paso consiste en fijar las probabilidades *a priori* $P(\theta)$, por cuyo conducto se resume la información disponible antes de realizar el experimento o la observación.

Supóngase que se conforma una distribución *a priori* a partir de una selección de un conjunto finito de posibles valores de θ , a cada uno de los cuales se le asigna una probabilidad (una distribución *a priori* bajo este supuesto se conoce como “discreta”). En algunos casos se sabe, o se sospecha fuertemente, que el valor de θ debe caer en un intervalo particular; siendo así, se puede elegir un conjunto de valores para θ dentro de este intervalo, o, si se tiene poco conocimiento acerca de la localización de la proporción, entonces se puede elegir una red de valores igualmente espaciados entre 0 y 1.

Acéptese de momento que hay solo k posibles valores para θ : $\theta_1, \theta_2, \dots, \theta_k$ a los que se han atribuido respectivas probabilidades que reflejan formal y cuantitativamente nuestros puntos de

vista de partida (probabilidades *a priori*). Llámense $P(\theta_1), P(\theta_2), \dots, P(\theta_k)$ a dichas probabilidades, donde:

$$\sum_{i=1}^k P(\theta_i) = 1$$

Se debe consignar en este punto que Epidat 3.1 no aborda la solución del problema bajo este supuesto de que la distribución de θ es discreta. Aquí se introduce para allanar el proceso de comprensión del enfoque bayesiano, pero el presente software encara el problema directamente en una situación realista, como se explica e ilustra más adelante.

Nótese que, en la medida que se atribuyen probabilidades a cada uno de los valores posibles, el parámetro *se convierte* en una variable aleatoria con cierta distribución explícita. El hecho de que se hable de que “el parámetro sea una variable aleatoria” constituye una especie de contrasentido para el estadístico clásico. Solo la comprensión de que se trata de un recurso artificial por cuyo conducto se expresan nuestras certidumbres e incertidumbres previas puede conferir sentido a la idea. Superado ese escollo, ya se puede hablar, por ejemplo, del valor esperado de θ y computarlo con acuerdo a la fórmula conocida:

$$E(\theta) = \sum_{i=1}^k \theta_i P(\theta_i) \quad [4]$$

Supóngase también que se ha hecho una experiencia empírica consistente en observar n individuos, y que e de ellos resultan tener la condición de interés.

Lo que permite el teorema de Bayes es reconstruir nuestra visión inicial con los datos que arroja dicha información empírica. Por su conducto se establece la probabilidad de que el verdadero valor de θ sea igual a θ_j . Se trata de la llamada *probabilidad a posteriori*, ya que es la que se obtiene *luego* de haber observado los datos.

Anteriormente se expresó el teorema de Bayes a través de la fórmula [3]:

$$P(A_j | B) = \frac{P(B | A_j) P(A_j)}{\sum_{i=1}^k P(B | A_i) P(A_i)}$$

Si en dicha fórmula, los θ_j ocupan el lugar de A_j y ciertos *datos* (en este caso que e de los n individuos encuestados son asmáticos) ocupan el de B , entonces ésta adopta la forma siguiente:

$$P(\theta_j | \text{datos}) = \frac{P(\text{datos} | \theta_j) P(\theta_j)}{\sum_{i=1}^k P(\text{datos} | \theta_i) P(\theta_i)} \quad [5]$$

La probabilidad $P(\text{datos} | \theta_j)$ es la llamada *verosimilitud* correspondiente a θ_j ; se trata de un número que cuantifica cuán verosímil es que se hayan obtenido esos datos si la verdadera tasa fuera θ_j .

Se aplicará en principio la expresión [5] para resolver el problema que se intenta resolver. Por ejemplo, podría postularse que la tasa de asmáticos pudiera corresponderse con solo uno de los

siguientes $k=5$ porcentajes (modelos posibles): 5%, 7%, 9%, 11% y 13%. Un investigador podría considerar que las probabilidades que se deben atribuir respectivamente a las distintas opciones son las que se recogen en la Tabla 4.

Tabla 4. Distribución de probabilidad *a priori* para los 5 modelos posibles.

	Modelos posibles				
	5%	7%	9%	11%	13%
Probabilidades <i>a priori</i> atribuidas	0,10	0,20	0,40	0,20	0,10

El valor esperado de θ es, de acuerdo con [4]:

$$E(\theta)=(0,05)(0,10)+(0,07)(0,20)+(0,09)(0,40)+(0,11)(0,20)+(0,13)(0,10)=0,09$$

Nótese que este juego de probabilidades *a priori* no solo refleja el punto de vista que se tiene sobre el valor (9%) en torno al cual podría hallarse la tasa desconocida, sino también nuestro grado de certidumbre (e incertidumbre) al respecto. Cuanto mayor sea la dispersión de las probabilidades *a priori* atribuidas a los valores posibles para θ , mayores dudas se estarían reflejando, y viceversa. Por ejemplo, si en lugar de las probabilidades consignadas en la Tabla 4, un segundo investigador confiriese una probabilidad más baja al valor de 9% (dígase, una probabilidad de 0,30) y redistribuyese la probabilidad restante, asignando 0,20 para los valores de 7% y 11%, y 0,15 a los valores de 5% y 13%, este segundo investigador estaría reflejando el mismo punto de vista en torno al posible valor de θ pero también que el grado de convicción sobre esta posibilidad es mucho menor.

En efecto, en este caso, se seguiría teniendo $E(\theta)=0,09$, pero se estaría comunicando una convicción mucho más débil en el sentido de que θ sea igual a 9%. El segundo investigador estaría expresando un grado de incertidumbre mayor; (ya que estaría otorgando más credibilidad en principio a la posibilidad de que la tasa se aleje de tal valor en una u otra dirección).

Sin embargo, los resultados finales del análisis pudieran no ser demasiado diferentes. Es lo que suele ocurrir salvo que se trate de visiones drásticamente distintas. El examen de las posibles modificaciones debidas a respectivas visiones *a priori* es lo que se conoce como “análisis de sensibilidad”, y se trata de un paso metodológicamente recomendado.

Ocasionalmente, por otra parte, el investigador pudiera preferir no anticipar criterio alguno; en tal caso, optará por los llamados “*priors* no informativos” (es decir, que elegirá la distribución uniforme como información *a priori*, lo cual equivale en esencia a dejar “que los datos hablen por sí mismos”; en este ejemplo, equivaldría a atribuir probabilidad 0,2 a cada una de las 5 posibilidades).

El siguiente paso, por tanto, consiste en calcular las verosimilitudes (*likelihoods*). En general, la verosimilitud de un modelo poblacional es la probabilidad de que se hayan producido los datos observados d (o una función de ellos, por ejemplo, una tasa calculada usando dichos datos) supuesto que ese modelo es válido.

Ahora bien, la probabilidad de tener estos datos (en este caso, de que se tengan e éxitos en n intentos independientes), si la probabilidad de cada intento es igual a θ , no es otra cosa que el valor de una densidad binomial con parámetros n y θ evaluada en e (e : 0, 1, ..., n):

$$P(e|n, \theta) = \binom{n}{e} \theta^e (1 - \theta)^{n-e} \quad [6]$$

donde $\binom{n}{e} = \frac{n!}{e!(n-e)!}$

Supóngase ahora que se hace una encuesta sobre 200 sujetos elegidos al azar y que se observa que 11 de ellos son asmáticos. Es decir, en ese caso $n=200$ y $e=11$, lo cual arroja una estimación puntual de 5,5%.

Consecuentemente, la verosimilitud de cada uno de los modelos posibles depende del valor de θ que lo define; por ejemplo, para el modelo que afirma que la tasa de asmáticos es 0,05 ($\theta=0,05$), ésta es:

$$\binom{200}{11} 0,05^{11} 0,95^{189} = 0,1167$$

En la Tabla 5 se muestran todos los cálculos necesarios para este ejemplo. La columna 3 resulta de aplicar la expresión anterior a cada uno de los modelos. A continuación se multiplican las probabilidades *a priori* por sus correspondientes verosimilitudes, se suman todos estos productos (0,0381), y se divide cada producto por dicha suma para obtener las probabilidades posteriores (quinta columna). Por ejemplo, para el modelo que establece que la probabilidad es igual a 0,05, al cual se atribuye una probabilidad *a priori* de 0,1 y al que corresponde una verosimilitud de 0,1167, se multiplica la probabilidad *a priori* por la verosimilitud y se divide el resultado 0,0117 por la suma total de los productos (0,0381); el valor que se obtiene, 0,306, es la probabilidad posterior de ese modelo.

Tabla 5. Resultados de los cálculos del ejemplo.

Modelo	Prior	Verosimilitud	Prior x Verosimilitud	Posterior
0,05	0,10	0,1167	0,0117	0,306
0,07	0,20	0,0847	0,0169	0,445
0,09	0,40	0,0221	0,0088	0,232
0,11	0,20	0,0030	0,0006	0,016
0,13	0,10	0,0003	0,00003	0,001
Suma	1		0,0381	1

Denótese por $P(5\% | d)$ a la probabilidad de que sea correcto el modelo que afirma que $\theta = 0,05$ supuestos los datos de que una muestra de tamaño 200 arrojó una tasa de 5,5% (que se denota por d). Según [5], la probabilidad *a posteriori* $P(5\% | d)$ es igual a:

$$P(0,05 | d) = \frac{(0,1167)(0,1)}{(0,1167)(0,1) + (0,0847)(0,2) + (0,0221)(0,4) + (0,0030)(0,2) + (0,0003)(0,1)} = 0,306$$

Como se puede apreciar, los datos han dado lugar a una redistribución de las probabilidades iniciales. El nuevo valor esperado es:

$$E(\theta|d)=(0,05)(0,306)+(0,07)(0,445)+(0,09)(0,232)+(0,11)(0,016)+(0,13)(0,001)=0,069$$

que resulta ser ahora muy próximo a 7%.

Por otra parte, el grado de incertidumbre es considerablemente menor. De hecho, los modelos $\theta=0,11$ y $\theta=0,13$ podrían virtualmente ser descartados a la luz de los datos empíricos.

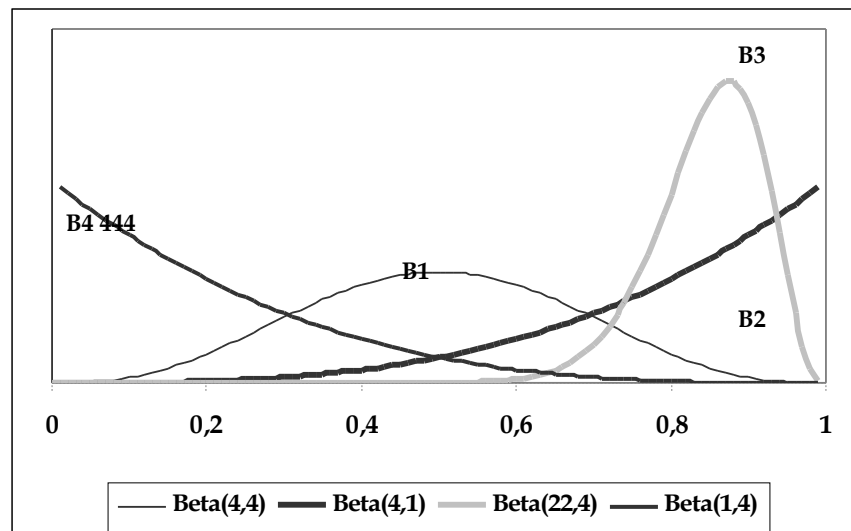
1.4. El paso a la continuidad

Colocando ahora las cosas en un marco realista, considérese el caso en que θ se puede ubicar en cualquier punto del intervalo $(0,1)$. Ahora las probabilidades *a priori* no se pueden enumerar, pues son infinitas. El modo de representarlas es a través de una función de densidad correspondiente a una distribución continua. Cualquier función de densidad continua definida en el intervalo $(0,1)$ puede ser útil (siempre que refleje el punto de vista del investigador), pero la más usual en el caso en que se estima una proporción es la *distribución beta*, que depende de dos parámetros ($a>0$ y $b>0$). Dicha distribución se extiende solo sobre el intervalo $(0,1)$, de modo que cumple con la condición óptima para representar nuestra percepción acerca de cuán probables son unos u otros valores de un parámetro que también está constreñido a tales límites. La función de densidad beta y sus propiedades pueden verse en cualquier libro de probabilidades y, aunque más abajo se exploran sus rasgos con cierta extensión, procede adelantar que, en particular, puede corroborarse que su media y su desviación estándar son, respectivamente, las siguientes:

$$\mu = \frac{a}{a+b} \quad \text{y} \quad \sigma = \sqrt{\frac{ab}{(a+b)^2(a+b+1)}}$$

La Figura 1 permite apreciar parcialmente la diversidad de formas que puede asumir la distribución beta en dependencia de los valores que adopten a y b . Así, se ve que la curva asume muy diferentes formas para cuatro juegos de parámetros: la curva **B1** corresponde a la pareja $a=4, b=4$, para $a=4, b=1$ se obtiene la curva **B2**, $a=22, b=4$ origina la curva **B3** y la pareja $a=1, b=4$, la curva **B4**.

Figura 1. Funciones de densidad beta para cuatro pares de parámetros.



Un caso de particular interés es aquel en que $a=1$ y $b=1$ (no representada en la Figura 1); con estos parámetros, la distribución beta coincide con la uniforme en $(0,1)$. Si se eligiera ésta como distribución *a priori*, ello significa que el investigador está adoptando una posición totalmente “neutra” (sin escorarse en sentido alguno, pues estaría otorgando la misma probabilidad a todos los valores posibles). El análisis bayesiano es en tal caso esencialmente coincidente con el que se obtiene mediante el enfoque clásico.

Nuevamente, supóngase que en una muestra de tamaño n se registran e asmáticos. Al número de no asmáticos se le llamará $f=n-e$. En este caso, la fórmula [3] involucra a una integral en lugar de a una sumatoria (de hecho –expresado de un modo informal- la integral en un intervalo no es otra cosa que la suma de los infinitos valores de la función integrada en dicho intervalo). Nótese que la ley de probabilidad total, que expresaba que:

$$P(B) = \sum_{i=1}^k P(B | A_i) P(A_i)$$

se traduce en el caso continuo a:

$$P(B) = \int_{\Omega} P(B | x) f(x) dx$$

donde $f(x)$ es la función de densidad de una variable aleatoria X evaluada en x , $P(B | x)$ es la probabilidad de B supuesto que $X=x$, y Ω es el recorrido continuo de valores posibles para X .

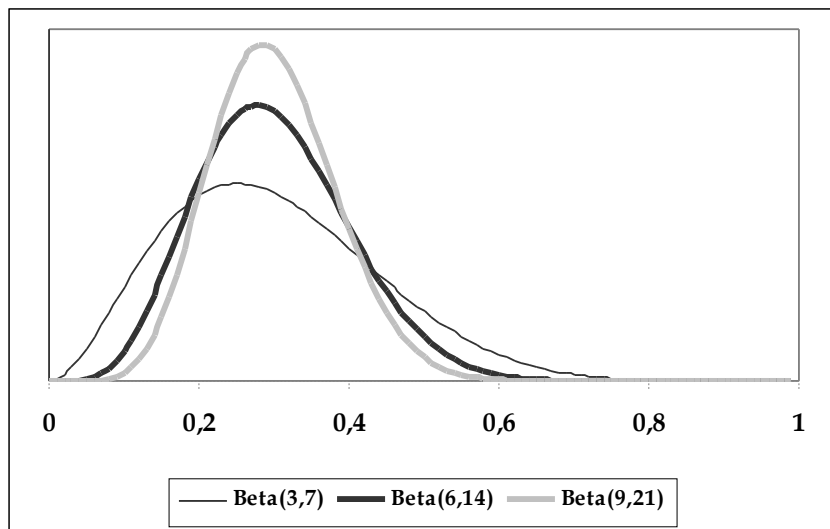
No es difícil probar que si la distribución *a priori* es $\text{Beta}(a,b)$ y los datos siguen una distribución binomial, para la que se han producido e “éxitos” y f “fracasos” entonces la distribución *a posteriori* es $\text{Beta}(a+e,b+f)$. Se dice que la distribución beta es conjugada para la binomial.

Si, por ejemplo, se tuviera cierta convicción *a priori* de que $\theta \approx 0,30$, teniendo en cuenta las fórmulas de la media y la varianza de la distribución Beta, se podría fijar a y b de manera que:

$$\frac{a}{a+b} = 0,30$$

Naturalmente, infinitas parejas cumplen tal condición. Si se estuviera poco convencido de la visión *a priori* sobre el valor en torno al cual pudiera hallarse θ , se podría fijar, por ejemplo, $a=3$ y $b=7$; si se concediera más credibilidad al valor medio postulado, una expresión paramétrica acorde con ese menor grado de incertidumbre podría conducir a fijar $a=6$ y $b=14$ ó, por ejemplo, $a=9$ y $b=21$, o $a=22,5$ y $b=52,5$.

Figura 2. Funciones de densidad Beta para 3 variantes en que la media es 0,30.



La Figura 2 refleja las tres primeras posibilidades. Epidat 3.1 cuenta con una interfase gráfica: cuando el usuario comunica valores para a y b , se produce un gráfico que le permite valorar si la forma de la función de densidad es acorde con su visión. Si quiere enmendar dicha visión y cambiarla por otra que considere más adecuada, bastará cambiar los parámetros y podrá examinar una nueva curva antes de pasar a dar los datos que habrán de emplearse para la “actualización”.

Como puede apreciarse, la forma de la función de densidad es muy parecida para las tres curvas (la media es la misma), pero las curvas varían en el grado de concentración en torno a esa media). Las varianzas de las tres distribuciones son respectivamente: 0,019, 0,010 y 0,007. Supóngase para lo que sigue que se acepta la primera de las tres distribuciones consideradas ($a=3$ y $b=7$).

Supóngase nuevamente que la muestra de tamaño $n=200$ arroja 11 asmáticos ($e=11$ y $f=189$); siendo así, la densidad *a posteriori* sería la correspondiente a una distribución beta con parámetros $a+e=14$ y $b+f=196$, que tiene una media de 0,067 y una desviación estándar de 0,017.

Como se observa, la distribución *a posteriori* se escora notablemente a la izquierda en relación a la distribución inicial, que estaba centrada en 30% y ahora lo está en 6,7%. Este resultado es coherente con el hecho de que la tasa observada, 5,5%, es marcadamente inferior a tal previsión. Por otra parte, la nueva distribución, *actualizada* con la información empírica, es apreciablemente más estrecha, lo que resulta de que el tamaño muestral es bastante grande. La nueva distribución pone de manifiesto que los datos corrigen al supuesto inicial y que, si bien éste tiene cierto peso, la verosimilitud domina. Este es un rasgo típico del enfoque bayesiano: cuando el tamaño muestral es muy grande, la distribución *a priori* que se elija tiende a ser irrelevante. Esa misma circunstancia, dicho sea de paso, subraya que la máxima utilidad del bayesianismo se produce cuando los tamaños muestrales no son muy grandes.

Ya en el marco bayesiano, una vez construida la distribución *a posteriori*, procede sacar conclusiones por su conducto. Lo singular es que, con esa herramienta, tales conclusiones pueden ser directamente expresadas en términos de probabilidad, algo imposible, como se ha dicho, en el marco de la estadística frecuentista clásica. Por ejemplo, se puede decir que la probabilidad de que θ supere a 7% es igual a 0,39; o que la de que esté entre 4% y 8% es igual a 0,75. En efecto, al contar con la distribución del parámetro, se trata en este caso de computar

áreas bajo la función de densidad correspondiente (resolver ciertas integrales); es decir, hallar las áreas bajo la curva de densidad en los intervalos que sean de interés, algo que resulta sencillo con las posibilidades computacionales actuales y que Epidat 3.1, naturalmente, incorpora entre sus posibilidades.

1.5. Intervalo de probabilidad

Un procedimiento de análisis de gran interés se relaciona con los llamados *intervalos de probabilidad*. Ellos evocan claramente a los intervalos de confianza clásicos, pero tienen un rasgo distintivo fundamental: a diferencia de estos últimos, el intervalo de probabilidad (también llamado *de credibilidad*) es un intervalo dentro del cual se hallaría el parámetro con cierta probabilidad especificada. Para fijar las ideas, imagínese que se quiere identificar un intervalo que contenga a θ con probabilidad $(1-\alpha)100\%$ donde $0 < \alpha < 1$. Supóngase que $\alpha=0,05$ (como suele elegirse en la inferencia clásica), de modo que en ese caso la probabilidad con la cual estaría el parámetro dentro del intervalo sería 95%.

Ahora bien, típicamente existirán infinitos intervalos que cumplan esa condición. ¿Con cuál de ellos operar? Hay varias maneras de elegirlo, pero el interés se centra básicamente en el llamado *intervalo de máxima densidad*.

El *intervalo de máxima densidad* se define como aquel intervalo (I,S) para el cual la función de densidad f cumple la siguiente condición: $f(x) \geq f(y)$ cualquiera sea $x \in (I,S)$ y cualquiera sea $y \notin (I,S)$. Este intervalo coincide, salvo situaciones excepcionales, con el intervalo más corto de entre los que cumplen la condición de tener bajo la curva un área de magnitud $(1-\alpha)100\%$; es decir, el intervalo tal que el valor $S-I$ sea mínimo entre aquellos que cumplen $P(I \leq \theta \leq S) = 1-\alpha$.

En el ejemplo, los percentiles más relevantes de la distribución a posteriori Beta(14, 196) son:

Percentil	Valor
2,5%	0,037
5%	0,040
10%	0,045
25%	0,054
50%	0,065
75%	0,078
90%	0,090
95%	0,098
97,5%	0,105

El intervalo de máxima densidad hay que hallarlo a través de programas especializados tales como Epidat 3.1. Como puede comprobarse, dicho intervalo es $[0,035 ; 0,101]$.

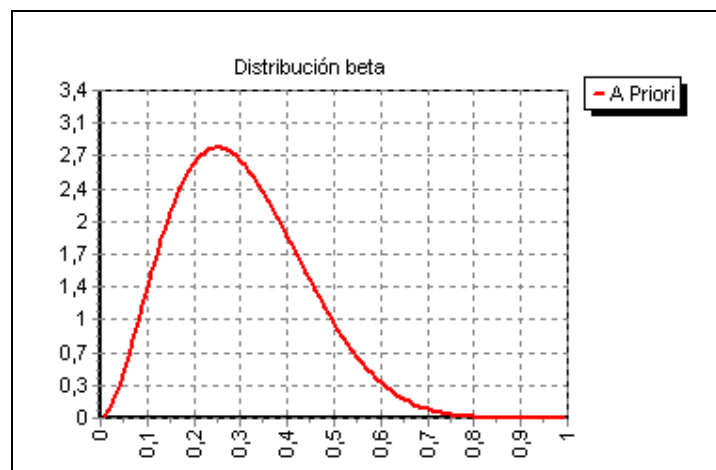
Por último, repárese en la diferencia que existe entre el intervalo de confianza calculado según los métodos tradicionales y el intervalo probabilístico que se obtiene a partir de los métodos bayesianos. Para los datos del ejemplo, se tiene un 95% de confianza en que la proporción de asmáticos de la comunidad se encuentra en el intervalo que tiene como límite inferior a 0,023 y superior a 0,087 (intervalo de confianza tradicional); mientras que, según los métodos bayesianos, se tiene un 95% de probabilidad de que el intervalo $[0,037 ; 0,105]$ contenga el valor del parámetro (proporción de asmáticos). Como se ve, los intervalos son bastante similares

(desde el punto de vista matemático, porque la interpretación es diferente como ya se sabe), lo cual se debe al tamaño de muestra tan grande que se usó; en efecto, si el tamaño de muestra es grande el enfoque bayesiano puede no diferir demasiado del frecuentista, porque la verosimilitud tiene mayor peso que la información *a priori*. Lo cierto es que un tamaño de muestra grande nunca es “malo”, si bien podría perderse parte del atractivo del enfoque bayesiano; pero cuando no es el caso, que es lo que suele ocurrir, entonces el enfoque bayesiano sí juega un papel cardinal, y ahí radica una de sus ventajas con respecto al método frecuentista.

Ejemplo

Considérese el ejemplo de estimación arriba planteado usando Epidat 3.1.

Para definir una distribución Beta(a, b) a priori con $a=3$ y $b=7$, se dan estos datos al programa y se marca PRESENTAR. Se tendría la curva siguiente:



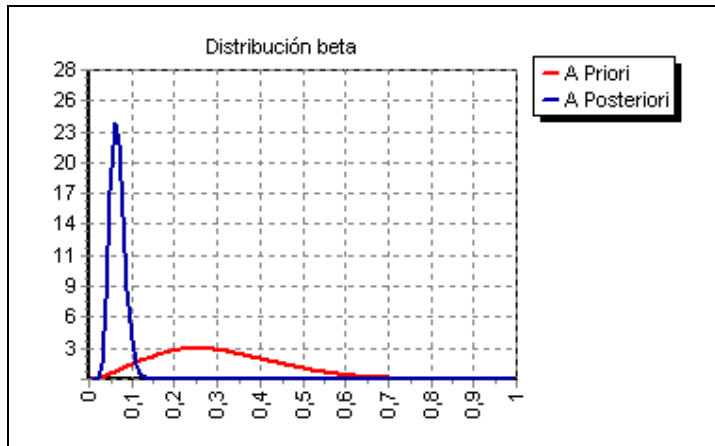
Si no se tiene ninguna idea acerca del valor que deben tomar los parámetros a y b para que la distribución beta correspondiente adopte la forma que represente nuestra convicción sobre la proporción P , Epidat permite, alternativamente, especificar un valor para la media (entre 0 y 1) y para la desviación estándar (mayor que 0) de dicha distribución. A partir de estos datos, el programa calcula los valores de a y b que definen una distribución beta con la media y desviación establecidas. Como estos parámetros tienen que ser mayores que 0, la desviación estándar no puede superar un determinado valor, que depende de la media, y que Epidat indica a través de un mensaje durante la entrada de datos.

Ahora se comunican los éxitos y fracasos de una experiencia concreta; éxitos=11 y fracasos=189, así como que se quiere un intervalo de probabilidad al 95% y conocer el área a la izquierda de 0,12. Los resultados son los siguientes:

Resultados con Epidat 3.1:

Análisis bayesiano. Estimación de una proporción	
Datos muestrales	Número
-----	-----
Éxitos	11
Fracasos	189

-----		-----	
Total		200	
-----		-----	-----
Distribución beta		A priori	A posteriori
-----		-----	-----
Parámetro a		3,0	14,0
Parámetro b		7,0	196,0
-----		-----	-----
Media		0,30	0,07
Desviación estándar		0,14	0,02



Área a la izquierda de los puntos elegidos

Punto	Área
-----	-----
0,120	0,996

Percentiles relevantes

Área	Percentil
-----	-----
0,025	0,037
0,050	0,041
0,100	0,046
0,250	0,054
0,500	0,065
0,750	0,077
0,900	0,089
0,950	0,097
0,975	0,104

Intervalo de máxima densidad que acumula el 95,0%
L.Inferior L.Superior

-----	-----
0,035	0,101

Obsérvese que en este ejemplo se ha pedido el área a la izquierda de 0,12. El resultado fue 0,996. Esto quiere decir que, a la luz de los datos, y hecha con ellos la actualización respecto de nuestra idea original, se puede asegurar que la probabilidad de que la tasa de interés sea superior a 12% es virtualmente nula (0,4%), ya que $1-0,996=0,004$, el área a la derecha, que es el complemento del que queda a la izquierda. El área entre dos puntos, naturalmente, se calcula restándole al área a la izquierda del mayor el área a la izquierda del menor.

2. UNA POBLACIÓN. VALORACIÓN DE HIPÓTESIS SOBRE UNA PROPORCIÓN

2.1. Proporción igual a una constante

Se valora la hipótesis $H: P=P_0$ frente a la hipótesis $K: P \neq P_0$

El primer dato que ha de entrarse es el valor de P_0 que se quiere valorar. Luego hay que dar la distribución a priori que se supone para el parámetro P . Como se explicó en el submódulo de estimación, se operará con la distribución $Beta(a, b)$, y los parámetros a y b pueden introducirse directamente, o bien estimarse a partir de la media y la desviación estándar de la distribución beta. En cualquiera de los dos casos, la media debe ser aproximadamente igual a P_0 , de modo que debe cumplirse la condición:

$$\frac{a}{a+b} \approx P_0 \text{ ó } media \approx P_0$$

lo cual es enteramente razonable, ya que si el usuario pensara a priori que la media del parámetro fuera diferente de P_0 , no tendría sentido para él hacer la prueba de hipótesis. Si no se verifica esta condición, o más específicamente no se cumple que:

$$|P_0 - \frac{a}{a+b}| < 0,015$$

Epidat 3.1 muestra un mensaje de error para informar de que los parámetros a y b no son compatibles con la proporción p_0 y se le da opción al usuario de volver a introducirlos.

Debe advertirse que los parámetros a y b solo pueden tomar valores mayores que 0; asimismo, la media debe ser un valor entre 0 y 1 y la desviación estándar será positiva y acotada superiormente por un valor del que Epidat 3.1 informa, para garantizar que los parámetros a y b calculados no sean negativos.

Definida la distribución a priori se podrá visualizar su representación gráfica, por si el usuario desea cambiar su posicionamiento a priori luego de ver el gráfico.

El procedimiento exige que se establezca también una probabilidad a priori (q) de la validez de H . En ese punto hay que comunicar e y f , donde e es el número de éxitos y f el de fracasos en $n=e+f$ experiencias.

Epidat 3.1 computa (véase Albert, 1996)¹² entonces el Factor de Bayes a favor de H (BF):

$$BF = \frac{P_0^e (1 - P_0)^f B(a, b)}{B(a + e, b + f)},$$

donde $B(a, b)$ es la función Beta evaluada en (a, b) , la cual se define del modo siguiente:

$$B(a, b) = \int_0^1 u^{a-1} (1 - u)^{b-1} du = \frac{\Gamma(a) \Gamma(b)}{\Gamma(a + b)}, \text{ y donde } \Gamma(x) = \int_0^\infty u^{x-1} e^{-u} du$$

Computa también el factor de Bayes en contra de H (BC):

$$BC = \frac{1}{BF},$$

y la probabilidad a posteriori de la veracidad de H:

$$P(H | \text{datos}) = \frac{q BF}{q BF + (1 - q)}$$

Debe tenerse en cuenta que si el usuario no quiere “tomar partido” a priori sobre la validez de H, puede optar por informar que $q=0,5$, lo cual hace que la probabilidad a posteriori solo quede en función de BF:

$$P(H | \text{datos}) = \frac{BF}{BF + 1}$$

Debe advertirse que en el marco bayesiano, no se trata de “elegir” entre una hipótesis y otra, de rechazar o no la hipótesis nula, sino de identificar en qué medida una es más razonable que la otra; tanto el BF como la probabilidad a posteriori sirven a ese fin, aunque esta última sea más expresiva para la mayor parte de los usuarios. Sobre esas bases se pueden tomar decisiones prácticas o incluso hacer otros desarrollos o análisis (tales como la optimización de decisiones con la incorporación de funciones de utilidad, no incluidos en el presente programa).

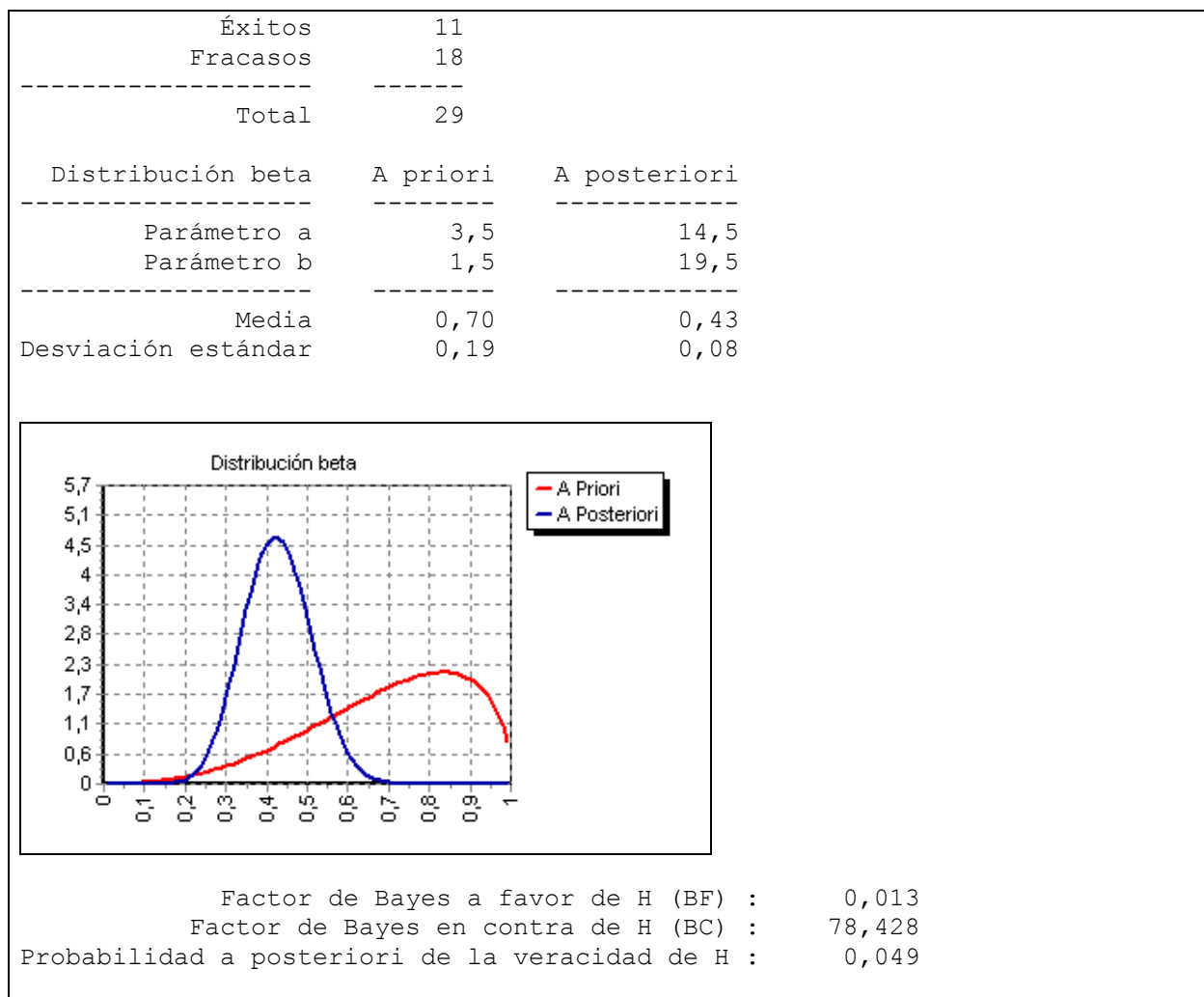
Ejemplo

Supóngase que se discute la validez de una hipótesis que afirma que la tasa de prevalencia de cierta enfermedad es igual a 0,7. Se fija entonces $P_0=0,7$.

Para representar la curva a priori se eligen $a=3,5$ y $b=1,5$. Se otorga una probabilidad bastante alta a priori de que tal hipótesis es cierta: $q=0,8$, y se informa que el resultado de un estudio por muestreo sobre 29 sujetos dio lugar a $e=11$ y $f=18$. Entonces, los resultados que se obtienen del programa son:

Resultados con Epidat 3.1:

Análisis bayesiano. Valoración de hipótesis sobre una proporción	
Hipótesis: Proporción = 0,70	
Probabilidad a priori de la validez de H: 0,80	
Datos muestrales	Número
-----	-----



Esto significa que es 78 veces más probable que sea cierta la hipótesis complementaria a que lo sea H. Si se incorpora la probabilidad a priori que se había dado para este hecho, se obtiene la probabilidad de que H sea cierta; un número sumamente bajo: =0,049.

2.2. Proporción dentro de un intervalo

En este caso se valora la hipótesis H: $P \in [P_1, P_2]$ frente a la hipótesis complementaria K: $P < P_1$ ó $P > P_2$.

Además de los valores de P_1 y P_2 , deben proveerse al programa los parámetros de la distribución beta (a y b), que caracterizan la apreciación a priori del investigador sobre la distribución del parámetro de interés. Los valores de a y b pueden introducirse directamente, o bien estimarse a partir de la media y la desviación estándar de la distribución beta. Definida la distribución a priori, se podrá visualizar su representación gráfica, por si el usuario desea cambiar su posicionamiento a priori luego de ver el gráfico.

El procedimiento exige que se establezca también una probabilidad a priori (q) de la validez de H. En ese punto hay que comunicar e y f , donde e es el número de éxitos y f el de fracasos en $n=e+f$ experiencias.

Epidat 3.1 computa, entonces, la probabilidad de la hipótesis:

$$P(H) = F(P_2) - F(P_1)$$

donde $F(P)$ es la función de distribución a posteriori evaluada en P (área bajo la curva de densidad a posteriori que queda a la izquierda de P),

el Factor de Bayes a favor de H (BF):

$$BF = \frac{P(H)}{P(K)}$$

y la probabilidad a posteriori de la veracidad de H :

$$PP = \frac{qBF}{qBF + 1 - q}$$

Quiere esto decir que la idea inicial que se tenía acerca de cuán probablemente cierta era la hipótesis se transforma, y la probabilidad, luego de haber observado los datos, pasa a ser PP .

Debe tenerse en cuenta que si el usuario no quiere “tomar partido” a priori sobre la validez de H , puede optar por informar que $q=0,5$, lo cual hace que la probabilidad a posteriori solo quede en función de BF :

$$PP = \frac{BF}{BF + 1}$$

lo cual, a su vez, equivale a que $PP=P(H)$.

Ejemplo

Supóngase que se valora la hipótesis correspondiente al caso $P_1=0,5$ y $P_2=0,8$. Para representar la curva a priori se eligen $a=3,5$ y $b=1,5$. Se otorga una probabilidad bastante alta a priori de que tal hipótesis es cierta ($q=0,8$) y se informa que el resultado de un estudio por muestreo sobre 29 sujetos dio lugar a $e=11$ y $f=18$.

Entonces, los resultados que se obtienen del programa son:

Probabilidad de la hipótesis H :	0,192
Factor de Bayes a favor de H (BF) :	0,238
Probabilidad a posteriori de la veracidad de H :	0,488

Nótese que aunque la fracción de “éxitos” ($11/29=0,38$) es considerablemente inferior al extremo inferior del intervalo $[0,50-0,80]$, la probabilidad de que el porcentaje poblacional esté dentro de dicho intervalo dista de ser pequeña. Ello se debe al alto valor que a priori se atribuye a la veracidad de H_0 .

3. DOS POBLACIONES. ESTIMACIÓN DE UNA DIFERENCIA DE PROPORCIONES

En la parte inicial de esta ayuda se resolvió un problema típico de la investigación médica a través de una prueba de hipótesis. A continuación, se expone la solución de este problema bajo

el enfoque bayesiano. Dicha solución, como se verá, transita precisamente por la estimación de una diferencia de proporciones.

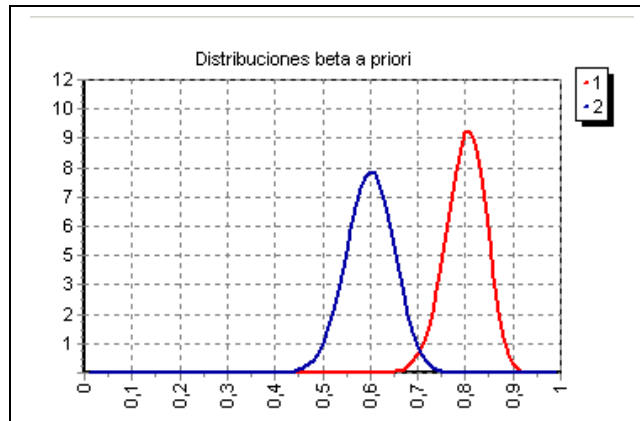
Recuérdese que se trataba de un ensayo clínico para valorar si los pacientes afectados por quemaduras hipodérmicas se recuperan más rápidamente cuando el tratamiento combina cierta crema antiséptica con un apósito hidrocoloide que cuando se utiliza solamente la crema antiséptica. Se conoce con bastante certeza que solo aproximadamente el 60% de los pacientes se recupera con este último recurso. El equipo investigador tiene, por otra parte, motivos teóricos e indicios empíricos surgidos de la literatura y del trabajo cotidiano de enfermería que hacen pensar con bastante optimismo que el tratamiento combinado es más efectivo que el tratamiento simple³.

Se había organizado un experimento con n pacientes, $n/2$ de los cuales se eligen aleatoriamente para ser atendidos con el tratamiento novedoso (combinación de crema antiséptica y apósito hidrocoloide), en tanto que a los $n/2$ restantes se les aplicará el tratamiento convencional (crema solamente).

En el ejemplo interesa ganar conocimiento acerca de dos proporciones: las tasas P_e y P_c de recuperación en cierto lapso con los tratamientos experimental y convencional respectivamente. Cada una de ellas es un valor que puede coincidir con cualquiera de los infinitos números reales que se hallan entre 0 y 1. Desde luego, se ignora cuáles son esos dos valores; pero ello no quiere decir que no se sepa absolutamente nada sobre ellos. Como se ha dicho, el primer paso de este enfoque exige expresar nuestras ideas iniciales (a priori) mediante una distribución de probabilidades para cada proporción. Recuérdese que los investigadores consideran de antemano, sobre bases racionales, teóricas e incluso empíricas, que es muy verosímil que el tratamiento que combina la crema antiséptica con el apósito hidrocoloide sea más efectivo que el tratamiento con crema solamente (de lo contrario no planificarían el ensayo). En este caso, el enfoque bayesiano exige que el investigador consiga expresar en términos probabilísticos su conocimiento (y su ignorancia) sobre esta proporción; es decir, que declare las zonas del intervalo (0,1) en las que resulta virtualmente imposible que se halle el valor, así como el grado en que considera probable que éste se ubique en otros intervalos de las zonas complementarias. La asignación de las probabilidades a priori, por ejemplo, puede dar cuenta de la convicción anticipada de que P_c se encuentra casi con seguridad entre 0,4 y 0,8, con alta probabilidad en una vecindad de 0,6, y que se reduce rápidamente cuando se aleja de ese punto; y en cuanto a P_e , que se halla en un entorno de 0,8, con muy escasa probabilidad de estar fuera del intervalo $[0,7 ; 0,9]$ ².

Como se había visto en el caso de la estimación de proporciones, la forma funcional más empleada para una densidad a priori es la que se asocia a la distribución beta (véanse la definición y diversas propiedades de esta distribución en la ayuda destinada a la estimación de una proporción). En este caso, en que se tienen dos proporciones, procede que los investigadores especifiquen valores de a y b para cada una de ellas con acuerdo a la visión arriba descrita. El investigador puede “jugar” con los valores de a y b , e ir observando los gráficos hasta arribar al que resulte más apropiado de acuerdo con sus creencias a priori sobre cada una de las proporciones de interés. Epidat 3.1 ofrece exactamente esta interfase gráfica con el usuario. En el ejemplo, la distribución beta a priori para P_e que se ajusta a la visión descrita podría ser la que tiene parámetros $a_e=72$ y $b_e=18$, y en el caso de P_c , la distribución beta con parámetros $a_c=57$ y

$b_c=38$. La figura siguiente muestra estas funciones de densidad seleccionadas a priori para las proporciones P_e y P_c tal y como se obtienen a través de Epidat 3.1 (1-Tratamiento experimental, 2- Tratamiento convencional):



Epidat 3.1 también ofrece la posibilidad de especificar la media y la desviación estándar para cada una de las distribuciones beta a priori sobre las dos proporciones. A partir de estos datos, el programa calcula, en cada caso, los valores de a y b que definen una distribución beta con la media y desviación establecidas. Como estos parámetros tienen que ser mayores que 0, la desviación estándar no puede superar un determinado valor, que depende de la media, y que Epidat indica a través de un mensaje durante la entrada de datos.

En el ejemplo, para los datos de la Tabla 1 se tenía: $e_e=30$ y $f_e=10$ por una parte, y $e_c=24$ y $f_c=16$ por otra; como $a_e=72$, $b_e=18$, $a_c=57$ y $b_c=38$, ahora las distribuciones a posteriori para P_e y P_c serían, respectivamente, $Beta_e(102, 28)$ y $Beta_c(81, 54)$.

Análogamente, si los resultados experimentales hubieran sido los de la Tabla 2 para la cual $e_e=103$, $f_e=97$, $e_c=120$ y $f_c=80$, las distribuciones a posteriori para P_e y P_c serían, respectivamente, $Beta_e(175, 115)$ y $Beta_c(177, 118)$.

Epidat 3.1 pide los resultados de la experiencia, genera las distribuciones actualizadas y luego, para cada una de tales distribuciones a posteriori, se genera cierto número s de datos simulados (el mismo número para ambas) y se calculan las s diferencias. Con ellas, finalmente, el programa construye una distribución empírica que será empleada posteriormente para hacer juicios probabilísticos concernientes a la diferencia entre los porcentajes.

En nuestro ejemplo, además de los parámetros que definen las distribuciones a priori y de los resultados del ensayo (primero de los dos escenarios considerados), el usuario debe comunicar el número de simulaciones que ha de realizarse (póngase por caso que fueran 10.000) y también (si lo desea) los puntos para los cuales se quiere el área a la izquierda bajo la distribución de las diferencias entre las proporciones.

Por ejemplo, puede pedirse la probabilidad de que la diferencia esté ubicada a la izquierda de 0,10. El resultado del análisis bayesiano sería como sigue:

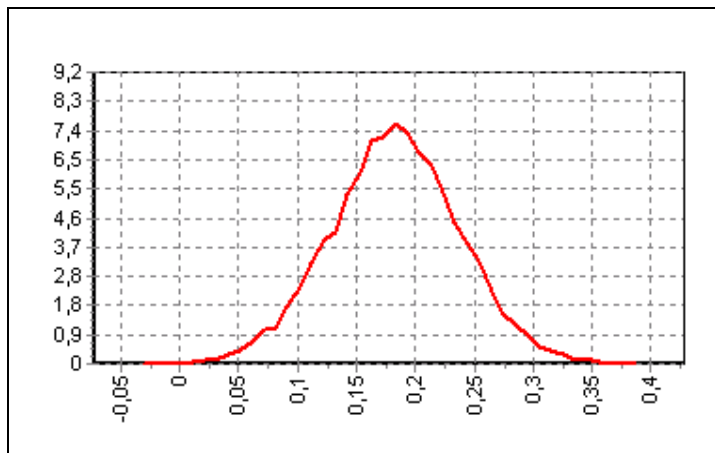
Resultados con Epidat 3.1

Análisis bayesiano. Estimación de una diferencia de proporciones
--

Datos muestrales	Población 1	Población 2
Éxitos	30	24
Fracasos	10	16
Total	40	40

Distribuciones beta a priori	Población 1	Población 2
Parámetro a	72,0	57,0
Parámetro b	18,0	38,0
Media	0,80	0,60
Desviación estándar	0,04	0,05

Distribución empírica a posteriori de las diferencias
Número de simulaciones: 10000



Media	
0,184	
Área a la izquierda de los puntos elegidos	
Punto	Área
0,100	0,063
Percentiles relevantes	
Percentil	Punto
0,025	0,073
0,050	0,094
0,100	0,113
0,250	0,148
0,500	0,185
0,750	0,221
0,900	0,255
0,950	0,275

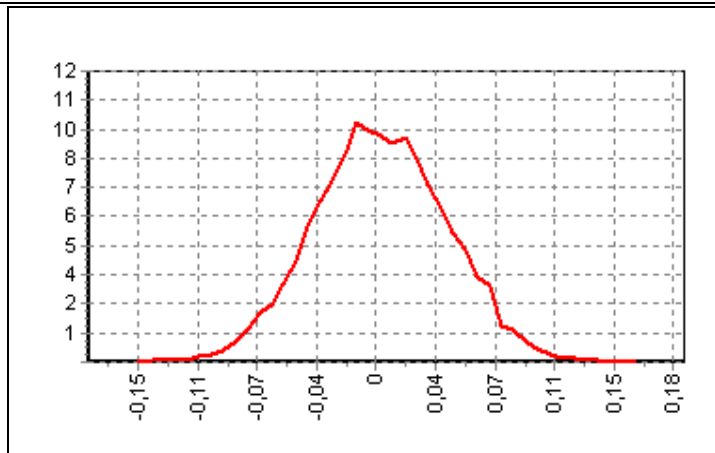
0,975	0,293
-------	-------

La figura representa la distribución empírica de la diferencia entre las proporciones de curación usando los hipotéticos datos experimentales de la Tabla 1. Debe tenerse en cuenta que si se repitiera esta solicitud, Epidat 3.1 produciría percentiles diferentes a estos, ya que se construyen con el resultado de las 10.000 simulaciones, las cuales son diferentes cada vez.

Como se ve, en el caso en que se utiliza la información de la Tabla 1, la zona más probable de la diferencia se ubica en un entorno de 0,18. Se recordará que, bajo el enfoque frecuentista, el investigador no podía arribar a conclusión alguna una vez realizada la prueba χ^2 , ya que la falta de significación estadística solo le permitía afirmar que no existía evidencia muestral suficiente como para declarar la superioridad del tratamiento novedoso. Sin embargo, el enfoque bayesiano permite computar probabilidades con las que se pueden hacer afirmaciones razonables y precisas: por ejemplo, que la probabilidad de que la diferencia entre las proporciones de curación entre estos tratamientos sea como mínimo de un 10%, asciende a $1-0,063=0,937$ (un 93,7%, lo que vale decir que es virtualmente seguro que el tratamiento nuevo hace un aporte sustancial).

Un proceso similar para los datos de la Tabla 2, produce los siguientes resultados:

Análisis bayesiano. Estimación de una diferencia de proporciones			
Datos muestrales		Población 1	Población 2
-----		-----	-----
Éxitos		103	120
Fracasos		97	80
-----		-----	-----
Total		200	200
Distribuciones beta a priori		Población 1	Población 2
-----		-----	-----
Parámetro a		72,0	57,0
Parámetro b		18,0	38,0
-----		-----	-----
Media		0,80	0,60
Desviación estándar		0,04	0,05
Distribución empírica a posteriori de las diferencias			
Número de simulaciones: 10000			



Media

0,003

Área a la izquierda de los puntos elegidos

Punto	Área
0,100	0,993

Percentiles relevantes

Percentil	Punto
0,025	-0,075
0,050	-0,063
0,100	-0,049
0,250	-0,024
0,500	0,002
0,750	0,030
0,900	0,055
0,950	0,068
0,975	0,081

0,025 -0,075

0,050 -0,063

0,100 -0,049

0,250 -0,024

0,500 0,002

0,750 0,030

0,900 0,055

0,950 0,068

0,975 0,081

Al considerar bayesianamente el caso en que los resultados del experimento contradicen las convicciones previas del investigador, con los mismos parámetros de la distribución beta a priori y los resultados de la Tabla 2, la diferencia entre los tratamientos se “mueve” en torno al cero, como se observa en la figura. Si bien, como se recordará, bajo el marco frecuentista se rechazaría (al nivel $\alpha=0,1$) la hipótesis de que el tratamiento que utiliza crema y apósito es equivalente al que prescinde del apósito (y que tal rechazo se haría en favor del segundo), con este enfoque se obtiene que la probabilidad de que la diferencia entre los porcentajes de recuperación sea mayor que un 10% es prácticamente despreciable (apenas un 0,7%), en tanto que la probabilidad de que el porcentaje de recuperación del tratamiento novedoso supere al del convencional en menos de un 5% es enorme (85,7%). En fin, puesto que la diferencia se halla en un entorno de 0 con una masa grande en vecindades estrechas de ese número, habría que concluir que los tratamientos son en esencia equivalentes. En este caso, la asignación de las probabilidades a priori que hizo el investigador desempeña un papel singularmente atractivo: nótese que, aunque su teoría a favor del tratamiento experimental está enfrentando un notable embate empírico, no es suplantada por una afirmación radicalmente opuesta a sus

conocimientos y experiencias anteriores, como habría que hacer bajo el enfoque frecuentista, a la vez que el resultado de la experiencia tiene peso suficiente como para cuestionar aquella convicción.

Lo cierto es que los métodos bayesianos siempre permiten arribar a alguna conclusión, que por otra parte resulta menos rígida, y en ese sentido más cercana al sentido común, que la que dictan los métodos frecuentistas.

4. DOS POBLACIONES. VALORACIÓN DE HIPÓTESIS SOBRE UNA DIFERENCIA DE PROPORCIONES

4.1. Igualdad de proporciones

Se valora la hipótesis $H: P_1 = P_2$ contra la hipótesis $K: P_1 \neq P_2$.

Debe comenzarse dando la probabilidad a priori (q) de la validez de H . Luego hay que definir la distribución $Beta(a_1, b_1)$ a priori en el supuesto de que valga la hipótesis H . Esto, como en los casos anteriormente explicados, se puede hacer dando directamente los valores de los parámetros a_1 y b_1 , o especificando la media y la desviación estándar de la distribución beta. Asimismo hay que definir, del mismo modo, la distribución $Beta(a_2, b_2)$ a priori de la segunda proporción en el supuesto de que NO vale H . Es decir, expresar nuestra convicción acerca del posible valor de P_2 , supuesto que no coincide con P_1 . Finalmente, se informarán los valores e_1 y f_1 , donde e_1 es el número de éxitos y f_1 el de fracasos en $n_1 = e_1 + f_1$ experiencias llevadas adelante en el primer caso, así como e_2 y f_2 , donde e_2 es el número de éxitos y f_2 el de fracasos en $n_2 = e_2 + f_2$ experiencias llevadas adelante en el segundo caso.

La salida es el factor de Bayes a favor de H (véanse las definiciones de esta noción y de los conceptos afines en la ayuda correspondiente a pruebas de hipótesis con una sola proporción), que Epidat 3.1 calcula usando la siguiente expresión:

$$BF = \frac{B(a_1 + e_1 + e_2, b_1 + f_1 + f_2)B(a_2, b_2)}{B(a_1 + e_1, b_1 + f_1)B(a_2 + e_2, b_2 + f_2)}$$

donde $B(x, y)$ denota la función beta evaluada en (x, y) .

Se obtienen, asimismo, el factor de Bayes en contra de H :

$$BC = \frac{1}{BF},$$

y la probabilidad a posteriori de la veracidad de H :

$$P(H | \text{datos}) = \frac{q BF}{q BF + (1 - q)}.$$

Ejemplo

Para el ejemplo del ensayo clínico sobre tratamientos a quemados expuesto en secciones anteriores de esta ayuda, y admitiendo $q=0,1$, se tendría:

$$\begin{array}{cccc} a_1=72 & b_1=18 & a_2=57 & b_2=38 \\ e_1=30 & f_1=10 & e_2=24 & f_2=16 \end{array}$$

Los resultados con Epidat 3.1 son:

Análisis bayesiano. Valoración de hipótesis sobre una diferencia de proporciones			
Hipótesis: Igualdad de proporciones			
Probabilidad a priori de la validez de H: 0,1			
Datos muestrales	Población 1	Población 2	
-----	-----	-----	
Éxitos	30	24	
Fracasos	10	16	
-----	-----	-----	
Total	40	40	
Distribuciones beta a priori	Población 1	Población 2	
-----	-----	-----	
Parámetro a	72,0	57,0	
Parámetro b	18,0	38,0	
-----	-----	-----	
Media	0,80	0,60	
Desviación estándar	0,04	0,05	
Factor de Bayes a favor de H (BF) :	0,076		
Factor de Bayes en contra de H (BC) :	13,210		
Probabilidad a posteriori de la veracidad de H :	0,008		

Los resultados son claros: la hipótesis de nulidad tiene un respaldo muy bajo si se compara con la que afirma que los tratamientos difieren. Luego de haber observado los datos y teniendo en cuenta el conocimiento previo, la probabilidad de que los procedimientos terapéuticos puedan considerarse equivalentes es de 8 entre 1.000.

4.2. Diferencia de proporciones dentro de un intervalo

Se valora la hipótesis $H: P_1-P_2 \in [P_3, P_4]$ contra la hipótesis complementaria:

$K: P_1-P_2 < P_3$ ó $P_1-P_2 > P_4$.

Debe comenzarse especificando los valores de P_3 y P_4 , y definiendo las dos distribuciones a priori, $Beta(a_1, b_1)$ y $Beta(a_2, b_2)$. Esto, como en casos previos, se puede hacer dando directamente los valores de los parámetros a y b , o especificando la media y la desviación estándar de la distribución beta. Finalmente, se informarán los valores e_1 y f_1 , donde e_1 es el número de éxitos y f_1 el de fracasos en $n_1=e_1+f_1$ experiencias llevadas adelante en el primer caso, así como e_2 y f_2 , donde e_2 es el número de éxitos y f_2 el de fracasos en $n_2=e_2+f_2$ experiencias llevadas adelante en el

segundo caso. También hay que indicar la probabilidad a priori (q) de la validez de la hipótesis H.

El método opera por simulación, de modo que hay que decidir el número n de simulaciones a realizar (no menor de 500).

Con esos datos, Epidat 3.1 procede a generar n valores de la distribución $\text{Beta}(a_1+e_1, b_1+f_1)$: $y_{11}, y_{12}, \dots, y_{1n}$ y n valores de la distribución $\text{Beta}(a_2+e_2, b_2+f_2)$: $y_{21}, y_{22}, \dots, y_{2n}$ para luego formar las n diferencias: $d_i = y_{1i} - y_{2i}$ con las que opera para producir las salidas. Estas son:

La probabilidad de la hipótesis:

$$P(H) = F(P_4) - F(P_3),$$

el factor de Bayes a favor de H:

$$BF = \frac{P(H)}{P(K)},$$

y la probabilidad a posteriori de la veracidad de H:

$$PP = \frac{qBF}{qBF + 1 - q}.$$

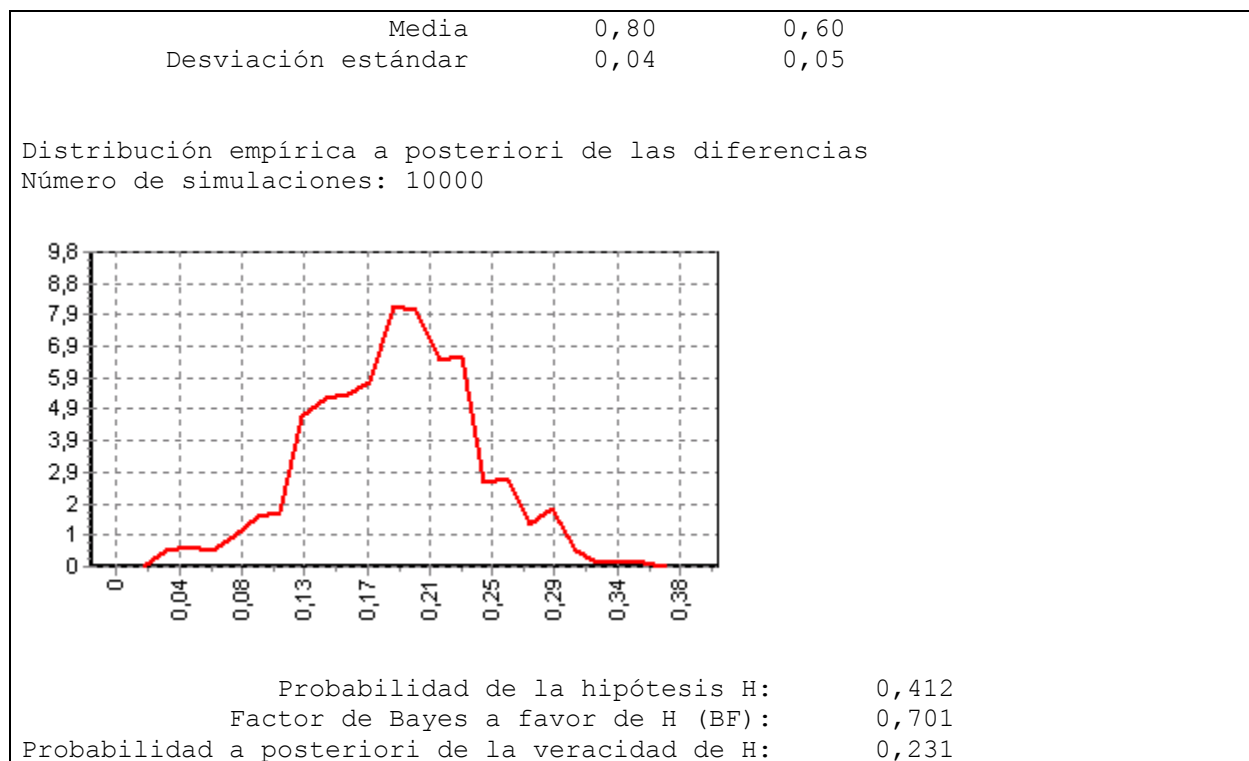
Ejemplo

Para el ejemplo del ensayo clínico sobre tratamientos a quemados expuesto en secciones anteriores de esta ayuda, y admitiendo $q=0,3$; $n=10.000$; $P_3=0,20$ y $P_4=0,35$ con los datos ya conocidos:

$$\begin{array}{cccc} a_1=72 & b_1=18 & a_2=57 & b_2=38 \\ e_1=30 & f_1=10 & e_2=24 & f_2=16 \end{array}$$

los resultados con Epidat 3.1 son:

Análisis bayesiano. Valoración de hipótesis sobre una diferencia de proporciones		
Hipótesis: Diferencia de proporciones [0,20 , 0,35]		
Probabilidad a priori de la validez de H: 0,30		
Datos muestrales	Población 1	Población 2
-----	-----	-----
Éxitos	30	24
Fracasos	10	16
-----	-----	-----
Total	40	40
Distribuciones beta a priori	Población 1	Población 2
-----	-----	-----
Parámetro a	72,0	57,0
Parámetro b	18,0	38,0
-----	-----	-----



5. UNA POBLACIÓN. ESTIMACIÓN DE UNA MEDIA

Tareas muy frecuentes en el mundo investigativo biomédico están relacionadas con datos correspondientes a variables continuas, para los cuales no es válida la división en “éxitos” y “fracasos”. Como las distribuciones poblacionales generalmente tienen la forma de la función de densidad normal, para estimar una media suele suponerse que los datos poblacionales siguen tal distribución con una media M y una desviación estándar S .

Para ilustrar el procedimiento bayesiano orientado a conocer acerca de la media poblacional, considérese el siguiente ejemplo. Supóngase que se realiza un estudio en una escuela primaria enclavada en un barrio marginal con el fin de examinar la situación en materia estomatológica; básicamente, se quiere conocer el promedio de caries por niño. Puesto que el estomatólogo ha estudiado profundamente la enfermedad, y tiene conocimientos y experiencias anteriores obtenidas de su trabajo en esa escuela con los niños, está persuadido, por poner un ejemplo burdo, de que el promedio de caries es, casi con seguridad, mayor que 0,5.

Supóngase que se le realizó el examen bucal a 30 niños, y que los resultados arrojan que el promedio de caries fue de 3,5 y la desviación estándar de 1,5.

La solución convencional acude a la obtención de un intervalo de confianza. La técnica estadística para el caso de una media utiliza un estadístico que sigue una distribución t de Student. Al aplicarla se obtiene que, para un nivel de confianza del 95%, el intervalo es $[2,94 ; 4,06]$.

Veamos la solución bayesiana. En esta situación no es razonable asumir priors uniformes, pues, como se había dicho, la enfermedad se conoce suficientemente como para establecer valores a priori de su frecuencia. De modo que el investigador puede asumir, por ejemplo, que el promedio de caries se encuentra alrededor de 3, pero no resultaría muy sorprendente si fuera 1

ó 5. Supóngase que se decide a priori que el promedio de caries es 3 ($m_0=3$) y la desviación estándar 1,56 ($s_0=1,56$). Berry¹³ proporciona un método para calcular la desviación estándar. Supóngase que se tiene un 10% de probabilidad a priori de que la media es mayor que 5. De acuerdo con la tabla de la distribución normal estándar, el valor de z que tiene un 10% de probabilidad a su derecha es 1,28. De modo que, suponiendo que nuestra percepción a priori es que la media muestral sigue una distribución normal, se tiene:

$$z = \frac{5 - 3}{s} = 1,28 \text{ y, despejando, se obtiene: } s = \frac{2}{1,28} = 1,56.$$

A continuación, solo resta conjugar este supuesto con los resultados que arrojó el examen bucal en una muestra (información empírica) y de ese modo obtener una distribución normal actualizada. Recuerdese que se examinaron 30 niños ($n=30$), y que el promedio de caries fue de 3,5 ($\bar{x} = 3,5$) y la desviación estándar de 1,5 ($s=1,5$).

Cuando la distribución a priori es normal, como ocurre en este caso, se calculan:

$$c_0 = \frac{1}{s_0^2}; \quad c = \frac{n}{s^2 \left(1 + \frac{20}{n^2}\right)^2}; \quad c_1 = c_0 + c.$$

Con los datos muestrales y la información a priori se calculan la media y la desviación de la distribución a posteriori, que también será normal:

$$m_p = \frac{c_0 m_0 + c \bar{x}}{c_1}; \quad s_p = \frac{1}{\sqrt{c_1}}.$$

La media de la distribución a posteriori es 3,48 y la desviación estándar 0,28. La curva posterior, resultante de la unión de ambas informaciones a través de métodos bayesianos, es más estrecha que la densidad a priori.

Se puede conocer ahora la probabilidad de que sean válidos ciertos valores; en este caso, la probabilidad de que el promedio de caries por niño sea mayor o menor que un valor dado. Por ejemplo, la probabilidad de que sea menor que 3 (valor fijado a priori por el investigador) es igual a 0,04.

Por último, también pueden calcularse intervalos de probabilidad para el promedio de caries por niño. Un intervalo de probabilidad al 95%, que es a su vez el de máxima densidad por ser simétrica la curva de densidad normal, tiene como límite inferior a 2,94 y como límite superior a 4,02. El resultado es, en este caso, virtualmente igual al que arroja el enfoque frecuentista, aunque la interpretación, desde luego, no.

Se ha desarrollado un ejemplo en el cual el investigador anticipa una distribución normal para la variable de trabajo, pero Epidat 3.1 contempla la posibilidad de que se trabaje con priors no informativos (con la distribución uniforme en lugar de la normal). De modo que lo primero que se pide al usuario es que responda si la distribución a priori que usará es la uniforme o la normal. En cualquiera de los dos casos, la distribución a posteriori será normal.

Si responde que es la normal, se le hará otra pregunta: ¿va a dar directamente los parámetros (media y una desviación estándar a priori: m_0 y s_0) o desea calcularlos a partir de dos valores posibles de la media (dígase X_1 y X_2) de X , y respectivas probabilidades (áreas a la izquierda de

esos puntos)? En el primer caso, han de teclearse m_0 y s_0 . En el segundo deben pedírsele los dos valores X_1 y X_2 , y las probabilidades asociadas P_1 y P_2 . Con esos cuatro datos, Epidat 3.1 calcula:

$$s_0 = \frac{X_1 - X_2}{\psi_1 - \psi_2} \text{ y } m_0 = X_1 - s_0 \psi_2$$

donde $\psi_i = \phi^{-1}(P_i)$, es decir, la inversa de la distribución normal estándar aplicada a la probabilidad P_i ($i=1, 2$).

En cualquiera de los dos casos, se mostrará el gráfico de la densidad resultante, $\text{Normal}(m_0, s_0)$ y se pedirá al usuario que confirme los datos que ha dado. Si no los confirma, se regresa al punto inicial.

Luego, en cualquier variante, hay que comunicar el tamaño muestral n y los resultados de la experiencia empírica: la media muestral y la desviación estándar muestral: $\bar{x}$ y s .

En ese punto el programa calcula:

$$c = \frac{n}{s^2(1 + \frac{20}{n^2})^2} \text{ y } c_1 = c_0 + c,$$

donde:

$c_0 = 0$, si la distribución a priori es la uniforme,

$c_0 = \frac{1}{s_0^2}$ si la distribución a priori es normal.

Las salidas que ofrece Epidat 3.1 son:

- Media de la distribución a posteriori: $m_p = \frac{c_0 m_0 + c \bar{x}}{c_1}$.
- Desviación de la distribución a posteriori: $s_p = \frac{1}{\sqrt{c_1}}$.
- Estimación paramétrica de percentiles de la distribución normal (a posteriori) más relevantes: 2,5; 5; 10; 25; 50; 75; 90; 95; 97,5.
- Estimación paramétrica de probabilidades (áreas a la izquierda) correspondientes a valores comunicados por el usuario.
- Gráfico de la curva normal a posteriori (en el caso de la normal, se incluye también la distribución a priori junto con esta curva; si se optó por la uniforme, entonces solo aparecerá la densidad a posteriori).
- Intervalo de probabilidad al $(1-\alpha)100\%$: $(X_{\frac{\alpha}{2}}; X_{1-\frac{\alpha}{2}})$.

Ejemplo

Si se usa la distribución uniforme y se informan los datos siguientes: $\bar{x}=176$; $s=3,16$ y $n=10$, así como que se quiere conocer la probabilidad de que el valor de la variable sea menor que 177, entonces el resultado es el siguiente:

Análisis bayesiano. Estimación de una media			
Datos muestrales		Valor	
-----		-----	
Media		176,00	
Desviación estándar		3,16	
Tamaño de muestra		10	
Distribución a priori: Uniforme			
Distribución normal a posteriori			Valor
-----			-----
Media			176,00
Desviación estándar			1,20
Área a la izquierda de los puntos elegidos			
Punto		Área	
-----		-----	
177,000		0,798	
Percentiles relevantes			
Área		Percentil	
-----		-----	
0,025		173,650	
0,050		174,028	
0,100		174,463	
0,250		175,191	
0,500		176,000	
0,750		176,809	
0,900		177,537	
0,950		177,972	
0,975		178,350	
Intervalo de probabilidad (95%)			
L.Inferior		L.Superior	
-----		-----	
173,650		178,350	

Si se elige la distribución a priori normal con m_0 y s_0 conocidas y se informa que $m_0=174$; $s_0=4,69$; $\bar{x}=176$; $s=3,16$ y $n=10$, así como que se quiere conocer la probabilidad de que el valor de la variable sea menor que 177, se obtiene:

Análisis bayesiano. Estimación de una media	
Datos muestrales	Valor

-----	-----	
Media	176,00	
Desviación estándar	3,16	
Tamaño de muestra	10	
Distribución a priori:	Normal	
Distribución normal	A priori	A posteriori
-----	-----	-----
Media	174,00	175,88
Desviación estándar	4,69	1,16
Área a la izquierda de los puntos elegidos		
Punto	Área	
-----	-----	
177,000	0,833	
Percentiles relevantes		
Área	Percentil	
-----	-----	
0,025	173,600	
0,050	173,966	
0,100	174,388	
0,250	175,094	
0,500	175,877	
0,750	176,661	
0,900	177,366	
0,950	177,788	
0,975	178,154	
Intervalo de probabilidad (95%)		
L.Inferior	L.Superior	
-----	-----	
173,600	178,154	

Obsérvese la notable consistencia entre los resultados, independientemente de que se fije una distribución a priori normal o uniforme. Esto no tiene porqué ocurrir siempre, pero cuando sí pasa, como ahora, ello ofrece especial respaldo a las conclusiones.

Finalmente, si se elige una distribución a priori normal pero a partir de los puntos X_1 y X_2 , y las probabilidades asociadas P_1 y P_2 , y si se comunicaran los valores para $X_1=174$; $P_1=0,5$; $X_2=180$; $P_2=0,9$; $\bar{x}=176$; $s=3,16$ y $n=10$, así como que se quiere conocer la probabilidad de que el valor de la variable sea menor que 177 se obtiene:

Análisis bayesiano. Estimación de una media		
Datos muestrales	Valor	
-----	-----	
Media	176,00	
Desviación estándar	3,16	
Tamaño de muestra	10	
Distribución a priori:	Normal	
Parámetros de entrada	Valor 1	Valor 2
-----	-----	-----

Punto	174,00	180,00
Probabilidad	0,50	0,90
Distribución normal	A priori	A posteriori
-----	-----	-----
Media	174,00	175,88
Desviación estándar	4,68	1,16
Área a la izquierda de los puntos elegidos		
Punto	Área	
-----	-----	
177,000	0,833	
Percentiles relevantes		
Área	Percentil	
-----	-----	
0,025	173,600	
0,050	173,966	
0,100	174,388	
0,250	175,093	
0,500	175,877	
0,750	176,660	
0,900	177,366	
0,950	177,788	
0,975	178,154	
Intervalo de probabilidad (95%)		
L.Inferior	L.Superior	
-----	-----	
173,600	178,154	

Los resultados son los mismos que en el caso anterior porque la información que se había dado antes para m_0 y s_0 es equivalente a la que se ha dado ahora para X_1, X_2, P_1 y P_2 (los dos primeros se deducen de estos 4 y viceversa).

6. UNA POBLACIÓN. VALORACIÓN DE HIPÓTESIS SOBRE UNA MEDIA

6.1. Media igual a una constante

Se trata de valorar la hipótesis H: $m = m_0$ frente a la hipótesis K: $m \neq m_0$.

Los datos de entrada son el valor de m_0 y la desviación estándar poblacional σ que se presuponen, la probabilidad a priori de la validez de H (q) y la desviación estándar de la distribución a priori bajo la hipótesis K (t). Finalmente, han de comunicarse los datos empíricos: media muestral y tamaño muestral: $\bar{x}$ y n .

Epidat 3.1 procede a computar el factor de Bayes (véase Albert, 1996)¹² a favor de la hipótesis:

$$BF = \frac{\frac{\sqrt{n}}{\sigma} \exp \left[-\frac{n}{2\sigma^2} (\bar{x} - m_0)^2 \right]}{\left(\frac{\sigma^2}{n} + t^2 \right)^{-\frac{1}{2}} \exp \left[-\frac{(\bar{x} - m_0)^2}{2 \left(\frac{\sigma^2}{n} + t^2 \right)} \right]},$$

el factor de Bayes en contra de H:

$$BC = \frac{1}{BF},$$

y la probabilidad a posteriori de la veracidad de H:

$$P(H | \text{datos}) = \frac{q BF}{q BF + (1 - q)}.$$

Ejemplo

Supóngase que las entradas son: $m_0=170$; $\sigma=3$; $q=0,5$; $t=0,5$; $\bar{x}=176$ y $n=10$.

Entonces, los resultados que se obtienen son:

Análisis bayesiano. Valoración de hipótesis sobre una media	
Hipótesis : Media = 170,00	
Probabilidad a priori de la validez de H: 0,50	
Desviación estándar	Valor
-----	-----
Poblacional	3,00
A priori	0,50
Datos muestrales	Valor
-----	-----
Media	176,00
Tamaño de muestra	10
Factor de Bayes a favor de H (BF):	0,015
Factor de Bayes en contra de H (BC):	68,393
Probabilidad a posteriori de la veracidad de H:	0,014

6.2. Media dentro de un intervalo

Se valora la hipótesis H: $m \in [m_1, m_2]$ frente a la hipótesis K: $m < m_1$ ó $m > m_2$.

Este procedimiento asume distribución a priori uniforme (es decir, priors no informativos). Las entradas son los valores de m_1 y m_2 , la probabilidad a priori (q) de la validez de H y, finalmente, los resultados empíricos: media, desviación estándar y tamaño muestrales: $\bar{x}$, s y n .

Epidat 3.1 procede como cuando se trataba de estimar una media para el caso uniforme; en ese caso, la distribución a posteriori será normal con media m_p y desviación estándar s_p .

Las salidas que se producen son:

- Probabilidad de la hipótesis: $P(H) = F(m_2) - F(m_1)$, donde $F(m)$ es la función de distribución evaluada en m , es decir, el área bajo la curva de densidad a posteriori (normal con media y desviación estándar m_p y s_p) que queda a la izquierda de m .
- Factor de Bayes a favor de H: $BF = \frac{P(H)}{P(K)}$.
- Probabilidad a posteriori de la veracidad de H: $PP = \frac{qBF}{qBF + 1 - q}$.

Ejemplo

Supóngase que las entradas son: $m_1=170$; $m_2=178$; $q=0,8$; $\bar{x}=176$; $s=3$ y $n=10$.

Entonces se obtiene:

Análisis bayesiano. Valoración de hipótesis sobre una media	
Hipótesis : Media dentro de un intervalo [170,00 , 178,00]	
Probabilidad a priori de la validez de H: 0,80	
Datos muestrales	Valor
-----	-----
Media	176,00
Desviación estándar	3,00
Tamaño de muestra	10
Probabilidad de la hipótesis H:	0,961
Factor de Bayes a favor de H (BF):	24,333
Probabilidad a posteriori de la veracidad de H:	0,990

7. DOS POBLACIONES. ESTIMACIÓN DE UNA DIFERENCIA DE MEDIAS. MÉTODO EXACTO

El enfoque convencional para el análisis de la diferencia de dos medias, como se sabe, transita por la formulación de una hipótesis nula, elegir el nivel de significación α , utilizar los datos para calcular el valor del estadístico, para el cual se utiliza la diferencia de las medias de los dos grupos ($d=m_1-m_2$) y el hecho de que ella sigue una distribución t de Student, determinar el valor p asociado a ese estadístico y, por último, rechazar o no la hipótesis nula en dependencia del valor de p . Se ha sugerido fuertemente, sin embargo, el empleo de intervalos de confianza para estimar la diferencia.

Supóngase que se realiza un estudio para evaluar los cambios en la glicemia en pacientes diabéticos controlados y en pacientes sanos. Se estudian 20 pacientes diabéticos controlados y 8 libres de esa enfermedad. La información recogida arroja los siguientes datos:

$$\overline{x}_1 = 6,78; s_1^2 = 1,11; n_1 = 20 \quad \overline{x}_2 = 6,80; s_2^2 = 0,59; n_2 = 8$$

El intervalo de confianza al $(1-\alpha)100\%$ se obtiene del modo siguiente:

$$LimInf = d - t_{1-\frac{\alpha}{2}}(n_1 + n_2 - 2)se(d)$$

$$LimSup = d + t_{1-\frac{\alpha}{2}}(n_1 + n_2 - 2)se(d)$$

donde,

$$d = \overline{x}_1 - \overline{x}_2,$$

$$se(d) = s \sqrt{\frac{1}{n_1} + \frac{1}{n_2}},$$

$$s = \sqrt{\frac{(n_1 - 1)s_1^2 + (n_2 - 1)s_2^2}{n_1 + n_2 - 2}},$$

$t_{1-\frac{\alpha}{2}}(n_1 + n_2 - 2)$ es el percentil $(1 - \frac{\alpha}{2})100\%$ de la distribución t de Student con n_1+n_2-2 grados de libertad.

Este es el procedimiento “usual” que se emplea en el marco frecuentista estándar, tal y como se hace en el módulo de inferencia de Epidat 3.1.

Es fácil corroborar que, con los datos del ejemplo, se tiene: $d=-0,02$; $s=0,985$; $se(d)=0,412$ y $t_{0,975}(26)=2,06$, de modo que el intervalo de confianza al 95% es $[-0,867; 0,827]$.

Bajo el enfoque bayesiano, cuando interesa conocer acerca de la diferencia de dos medias, en Epidat 3.1 se parte del supuesto de que las medias son independientes y que los datos son recogidos a partir de dos muestras independientes.

En este caso, hay cuatro parámetros desconocidos: las medias M_1 y M_2 y las desviaciones estándares S_1 y S_2 de las dos poblaciones. Dado que, en general, es difícil establecer una distribución conjunta para cuatro parámetros, se asume que los parámetros tienen una distribución a priori no informativa estándar proporcional a $1/(S_1 S_2)$ (**Nota:** este planteamiento

exhibe dificultades mayores que lo típico en el presente texto; por su complejidad, para conocer más detalles, véase Albert¹²). Entonces, las medias M_1 y M_2 tienen distribuciones t de Student independientes. La distribución a posteriori de la diferencia entre las medias $M_2 - M_1$ tiene una forma funcional no estándar, por lo que se acude a la simulación. Concretamente: se simulan muestras independientes de un tamaño fijo acordes con las distribuciones posteriores marginales de M_1 y M_2 , después se emparejan las observaciones de las dos muestras simuladas y se toman las diferencias; entonces se obtienen los valores simulados de la distribución posterior de la diferencia de medias.

El proceso detallado es como sigue:

- Se computan: $S_1^2 = n_1 s_1^2$ y $S_2^2 = n_2 s_2^2$.
- Se generan:
 - n valores $y_{11}, y_{12}, \dots, y_{1n}$ con distribución χ^2 con $n_1 - 1$ grados de libertad,
 - n valores $y_{21}, y_{22}, \dots, y_{2n}$ con distribución χ^2 con $n_2 - 1$ grados de libertad.
- Se obtienen dos juegos de n realizaciones $z_{11}, z_{12}, \dots, z_{1n}$ y $z_{21}, z_{22}, \dots, z_{2n}$ con distribución normal estándar.
- Para cada i entre 1 y n se calculan:

$$m_{1i} = \sqrt{\frac{S_1^2 z_{1i}^2}{y_{1i} n_2}} + \bar{x}_1 \quad \text{y} \quad m_{2i} = \sqrt{\frac{S_2^2 z_{2i}^2}{y_{2i} n_2}} + \bar{x}_2$$

- Finalmente, se obtienen las n diferencias: $d_i = m_{1i} - m_{2i}$.

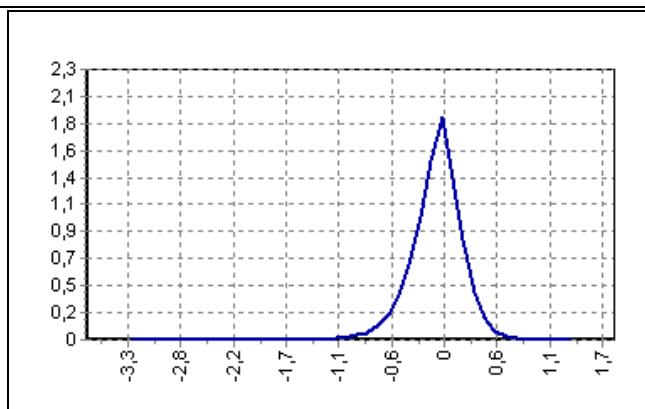
Ejemplo

Aplicado este procedimiento con $n=10.000$ simulaciones, y dando los datos arriba expuestos:

$$\bar{x}_1 = 6,78; s_1^2 = 1,11 \text{ } (s_1=1,05) \text{ y } n_1 = 20 \quad \bar{x}_2 = 6,80; s_2^2 = 0,59 \text{ } (s_2=0,77) \text{ y } n_2 = 8$$

se obtuvo:

Análisis bayesiano. Estimación de una diferencia de medias por método exacto		
Datos muestrales	Población 1	Población 2
-----	-----	-----
Media	6,78	6,80
Desviación estándar	1,05	0,77
Tamaño de muestra	20	8
Distribución empírica a posteriori de las diferencias		
Número de simulaciones:	10000	



Media

-0,082

Percentiles relevantes

Área	Percentil
0,025	-0,694
0,050	-0,550
0,100	-0,414
0,250	-0,223
0,500	-0,056
0,750	0,084
0,900	0,228
0,950	0,325
0,975	0,406

Vale destacar que este valor tan pequeño (-0,082) de la diferencia media entre pacientes diabéticos controlados y libres de la enfermedad no es sorprendente; obsérvese que el valor medio de la glicemia es muy parecido en ambos grupos.

A partir de esta distribución empírica pueden hacerse estimaciones (no paramétricas) de interés, tales como la mediana (-0,056).

8. DOS POBLACIONES. ESTIMACIÓN DE UNA DIFERENCIA DE MEDIAS. MÉTODO APROXIMADO

En este caso Epidat 3.1 procede exactamente como en el caso de la estimación de una media, solo que dos veces en lugar de una: una para cada población. Si se responde que la distribución a priori es la normal, se le hará otra pregunta: ¿va a dar directamente los parámetros (media y desviación estándar a priori para cada población: m_{01} , m_{02} , s_{01} y s_{02}) o desea calcularlos a partir de dos valores de X , y respectivas probabilidades (áreas a la izquierda de esos puntos)?

En el primer caso, han de teclearse m_{01} , m_{02} , s_{01} y s_{02} . En el segundo, deben pedírsele los dos valores X_1 y X_2 para cada población, y las probabilidades asociadas P_1 y P_2 . Con esos datos, se calculan, tal y como se explicó en el caso de una sola media, m_{p1} , s_{p1} primero y m_{p2} , s_{p2} después. Como la distribución a posteriori de cada media es normal, y éstas se suponen independientes, la distribución a posteriori de la diferencia también será normal.

Finalmente, se pide media, desviación estándar y tamaño muestrales para cada población:

$$\bar{x}_1, s_1, n_1 \text{ y } \bar{x}_2, s_2, n_2.$$

Las salidas en este submódulo son:

- Media de la distribución a posteriori de la diferencia: $m = m_{p1} - m_{p2}$.
- Desviación de la distribución a posteriori de la diferencia $s = \sqrt{s_{p1}^2 + s_{p2}^2}$.
- Estimación paramétrica de percentiles de la distribución a posteriori más relevantes: 2,5; 5; 10; 25; 50; 75; 90; 95; 97,5.
- Estimación paramétrica de probabilidades (áreas a la izquierda) correspondientes a valores comunicados por el usuario.
- Gráfico de la densidad normal a posteriori de la diferencia con media m y varianza s^2 .
- Intervalo de probabilidad al $(1-\alpha)100\%$ de la diferencia: $\left(X_{\alpha/2}; X_{1-\alpha/2} \right)$.

Ejemplo

El siguiente ejemplo comienza con el caso en que se comunica que ha de emplearse la distribución uniforme. En tal caso, no hay que dar otra información adicional fuera de la que procede de la muestra. Supóngase que las entradas son:

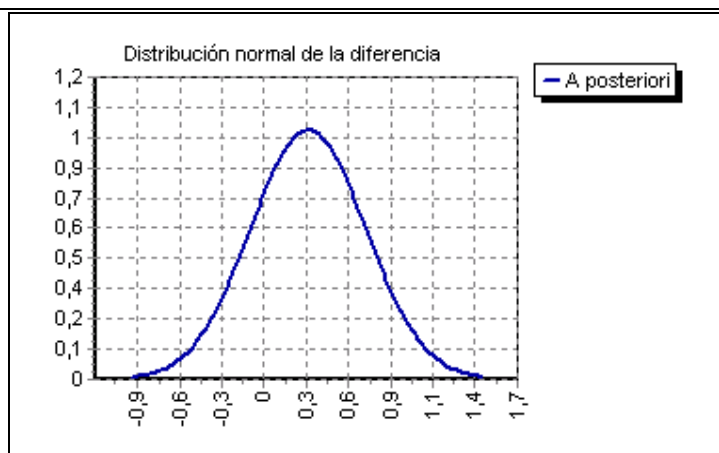
$$\bar{x}_1 = 7,1 \quad s_1=1,12 \quad n_1=21$$

$$\bar{x}_2 = 6,80 \quad s_2=0,63 \quad n_2=8$$

así como que se quiere el área a la izquierda de 0 y de 1,2.

Las salidas que se obtienen son:

Análisis bayesiano. Estimación de una diferencia de medias por método aproximado		
Datos muestrales	Población 1	Población 2
-----	-----	-----
Media	7,10	6,80
Desviación estándar	1,12	0,63
Tamaño de muestra	21	8
Distribución a priori: Uniforme		



Distribución a posteriori	Valor
Media	0,300
Desviación estándar	0,388

Área a la izquierda de los puntos elegidos

Punto	Área
0,000	0,220
1,200	0,990

Percentiles relevantes

Área	Percentil
0,025	-0,461
0,050	-0,339
0,100	-0,198
0,250	0,038
0,500	0,300
0,750	0,562
0,900	0,798
0,950	0,939
0,975	1,061

Intervalo de probabilidad del 95%

L.Inferior	L.Superior
-0,461	1,061

En el caso en que se comunica que ha de emplearse la distribución normal con medias y varianzas conocidas, hay que especificar estos datos. Supóngase que las entradas son:

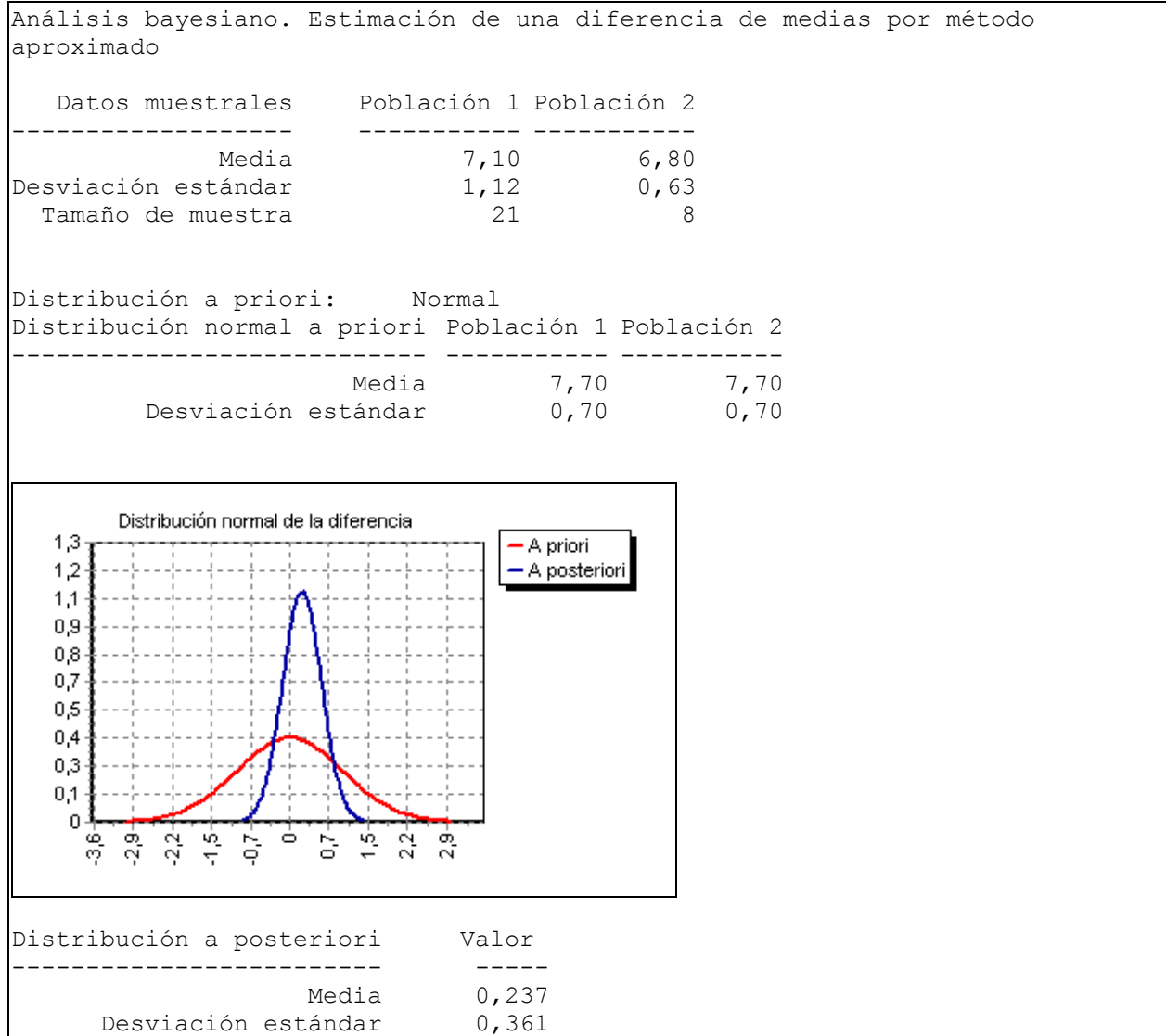
$$m_{01}=7,7 \quad s_{01}=0,7 \quad m_{02}=7,7 \quad s_{02}=0,7$$

$$\bar{x}_1 = 7,1 \quad s_1=1,12 \quad n_1=21$$

$$\bar{x}_2 = 6,80 \quad s_2=0,63 \quad n_2=8$$

y que se quiere el área a la izquierda de 0 y de 1,2.

Los resultados son:



Área a la izquierda de los puntos elegidos	
Punto	Área
0,000	0,256
1,200	0,996
Área a la izquierda de los puntos elegidos	
Punto	Área
0,000	0,256

1,200	0,996
Percentiles relevantes	
Área	Percentil
-----	-----
0,025	-0,471
0,050	-0,357
0,100	-0,226
0,250	-0,007
0,500	0,237
0,750	0,480
0,900	0,700
0,950	0,831
0,975	0,945
Intervalo de probabilidad del 95%	
L.Inferior	L.Superior
-----	-----
-0,471	0,945

Finalmente, en el caso en que se comunica que ha de emplearse la distribución normal pero a partir de los puntos X_{11} , X_{12} , X_{21} y X_{22} y las probabilidades asociadas P_{11} , P_{12} , P_{21} y P_{22} , supóngase que las entradas son:

$$X_{11}=0 \quad X_{12}=1,42 \quad X_{21}=0 \quad X_{22}=1,52$$

$$P_{11}=0,5 \quad P_{12}=0,8 \quad P_{21}=0,5 \quad P_{22}=0,8$$

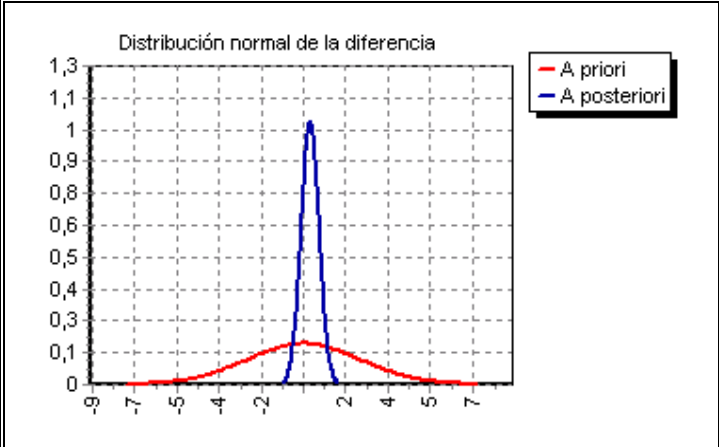
$$\bar{x}_1 = 7,1 \quad s_1=1,12 \quad n_1=21$$

$$\bar{x}_2 = 6,80 \quad s_2=0,63 \quad n_2=8$$

y que se quiere el área a la izquierda de 0 y de 1,2.

Se obtiene lo siguiente:

Análisis bayesiano. Estimación de una diferencia de medias por método aproximado				
Datos muestrales	Población 1		Población 2	
-----	-----		-----	
Media	7,10		6,80	
Desviación estándar	1,12		0,63	
Tamaño de muestra	21		8	
Distribución a priori: Normal				
Parámetros de entrada	Población 1		Población 2	
	Valor 1	Valor 2	Valor 1	Valor 2
-----	-----		-----	
Punto	0,00	1,42	0,00	1,52

Probabilidad	0,50	0,80	0,50	0,80
Distribución normal a priori Población 1 Población 2				
Media	0,00	0,00		
Desviación estándar	1,69	1,81		
Distribución normal de la diferencia 				
Distribución a posteriori		Valor		
Media	0,314			
Desviación estándar	0,384			
Área a la izquierda de los puntos elegidos				
Punto	Área			
0,000	0,206			
1,200	0,990			
Percentiles relevantes				
Área	Percentil			
0,025	-0,437			
0,050	-0,316			
0,100	-0,177			
0,250	0,056			
0,500	0,314			
0,750	0,573			
0,900	0,806			
0,950	0,945			
0,975	1,066			
Intervalo de probabilidad del 95%				
L.Inferior	L.Superior			
-0,437	1,066			

9. DOS POBLACIONES. VALORACIÓN DE HIPÓTESIS SOBRE UNA DIFERENCIA DE MEDIAS

Se considera la diferencia D entre dos medias y se valora la hipótesis $H: D \in [d_1, d_2]$ frente a la hipótesis $K: D < d_1$ ó $D > d_2$. Epidat 3.1 asume distribución a priori uniforme.

Los datos de entrada de este submódulo son:

- Media, desviación estándar y tamaño muestrales en el grupo 1: $\bar{x}_1, s_1$ y n_1
- Media, desviación estándar y tamaño muestrales en el grupo 2: $\bar{x}_2, s_2$ y n_2
- Número de simulaciones: n

(Con estos datos se procede como se hizo en el caso de la estimación de la diferencia de dos medias, método exacto).

- Valores de los extremos involucrados en la hipótesis: d_1 y d_2
- Probabilidad a priori de la validez de H : q

Con esta información, Epidat 3.1 produce las salidas siguientes:

- Probabilidad de la hipótesis, $P(H)$: se computa la estimación no paramétrica de esta probabilidad a partir de la distribución empírica de las diferencias entre las dos medias.
- Factor de Bayes a favor de H : $BF = \frac{P(H)}{P(K)}$.
- Probabilidad a posteriori de la veracidad de H : $PP = \frac{qBF}{qBF + 1 - q}$.

Ejemplo

Supóngase que se dan los siguientes datos de entrada:

$$\bar{x}_1 = 68 \quad s_1 = 8 \quad n_1 = 50$$

$$\bar{x}_2 = 65 \quad s_2 = 7 \quad n_2 = 80$$

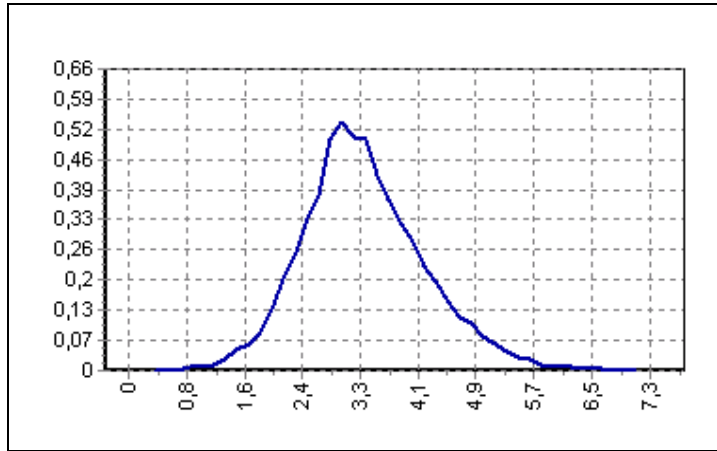
$$d_1 = 2 \quad d_2 = 4 \quad q = 0,6$$

Número de simulaciones: $n = 10.000$

Los resultados que arroja Epidat 3.1 son:

Análisis bayesiano. Valoración de hipótesis sobre una diferencia de medias		
Hipótesis : Diferencia de medias en un intervalo [2,00 , 4,00]		
Probabilidad a priori de la validez de H: 0,600		
Datos muestrales	Población 1	Población 2
-----	-----	-----
Media	68,00	65,00
Desviación estándar	8,00	7,00
Tamaño de muestra	50	80

Distribución empírica a posteriori de las diferencias
Número de simulaciones: 10000



Probabilidad de la hipótesis H: 0,757
Factor de Bayes a favor de H (BF): 3,114
Probabilidad a posteriori de la veracidad de H: 0,824

10. TABLAS DE CONTINGENCIA. VALORACIÓN DE HIPÓTESIS DE INDEPENDENCIA

La exploración de la relación entre mediciones categóricas es un problema básico de la inferencia. Para ilustrar esta situación considérese el siguiente ejemplo. Supóngase que se realiza un estudio para evaluar la relación entre el hábito de fumar de la madre y el peso al nacer de sus hijos. Los resultados se recogen en una tabla de contingencia de 2x3 (Tabla 6).

Tabla 6. Distribución de los niños según peso y edad de la madre.

Peso del niño	Hábito de fumar de la madre			Total
	Nunca	Antes del embarazo	Continuó fumando	
Normopeso	12	69	4	85
Bajo peso	10	24	6	40
Total	22	93	10	125

Según los métodos frecuentistas, para evaluar la posible relación entre el hábito de fumar de la madre y el peso del niño al nacer (valorado en este caso dicotómicamente), se valora la hipótesis de independencia, la cual afirma que los niveles de peso de los niños son los mismos para los tres grupos considerados en materia de hábito de fumar de la madre. El enfoque comúnmente usado comienza por ajustar el modelo de independencia a la tabla y calcula los valores “esperados” bajo este modelo para todas las celdas de la tabla. Entonces, se aplica una prueba estadística, la cual mide cuán alejados están los valores esperados de los observados.

La prueba estadística convencionalmente usada es la Ji-cuadrado de Pearson, y se rechaza la hipótesis de independencia si el valor del estadístico es suficientemente grande. Típicamente, se juzga el tamaño del valor del estadístico a través del valor p , probabilidad (bajo la hipótesis de independencia) de observar un valor del Ji-cuadrado con $(R-1) \times (C-1)$ grados de libertad al menos tan grande como el que se obtuvo, donde R es el número de filas de la tabla y C el de columnas. Si el valor de p es suficientemente pequeño, se rechaza la hipótesis de independencia.

Para estos datos, el valor del estadístico, que sigue una distribución Ji-cuadrado con 2 grados de libertad, es 7,07, al cual se le asocia un valor de p igual a 0,03. Dado que este valor es menor que 0,05, desde el punto de vista frecuentista se considera como una evidencia significativa de que el hábito de fumar de la madre y el peso al nacer no son independientes.

La prueba bayesiana de independencia maneja dos hipótesis, la hipótesis H_0 que afirma que las dos variables son independientes, y la hipótesis complementaria, que plantea que las dos variables son dependientes o están de alguna manera relacionadas. Bajo el enfoque bayesiano, se valora el apoyo relativo que dan los datos a cada una de estas dos hipótesis en términos del factor de Bayes, el cual compara las probabilidades de los datos observados bajo las hipótesis:

$$FB = \frac{P(D = d_0 | H_0)}{P(D = d_0 | \bar{H}_0)}$$

Naturalmente, mientras H_0 es una afirmación concreta, la falsedad de H_0 puede asumir infinitas expresiones puntuales. Sin embargo, en condiciones bastante generales, se pueden hallar cotas inferiores para BF . Así, para el caso en que se valora una asociación a través de la prueba Ji-cuadrado, se tiene una cota inferior para el factor de Bayes¹⁹ en función del valor observado χ^2 :

$$\sqrt{\chi^2} \exp\left(\frac{1 - \chi^2}{2}\right) \leq FB$$

Esto da lugar a la siguiente desigualdad:

$$\left[1 + \frac{1 - p(H_0)}{p(H_0) \sqrt{\chi^2} \exp\left(\frac{1 - \chi^2}{2}\right)} \right]^{-1} \leq P(H_0 | D = d_0)$$

Ejemplo

Si la tabla de contingencia fuera la que se expuso al comienzo de esta sección, los datos iniciales de entrada serían:

12	69	4
10	24	6

Supóngase que se anticipa el valor $q=0,9$. Entonces los resultados son:

aprobación de conductas terapéuticas que poco tiempo después han sido fatalmente desacreditadas; la razón de estos escándalos, simplemente es, a su juicio, estadística. Según este autor, la razón más persuasiva para usar la inferencia bayesiana es la capacidad que tiene de proveer un nivel de protección mucho mayor que el enfoque frecuentista contra la posibilidad de “ver significación en hallazgos procedentes de la investigación científica que son enteramente espurios”.

Muy por el contrario, por lo general p puede llegar a ser sustantivamente menor que $P(H_0 | D=d_0)$. Eso es en esencia lo que se advierte mediante la famosa paradoja de Lindley¹⁵.

La disparidad entre p y $P(H_0 | D=d_0)$ se incrementa en la medida que $P(H_0)$ crece, supuesto d_0 fijo; de modo que p es un indicador más engañoso cuanto menos plausible resulte ser la teoría bajo investigación. Esta diferencia entre la probabilidad que usualmente se computa (p) y la que verdaderamente interesa, también se hace más acusada, por otra parte, en la medida que se incrementa el tamaño muestral¹⁶: quiere esto decir que para tamaños de muestra grandes, el riesgo de atribuir incorrectamente una baja probabilidad a la hipótesis nula debido a que es menor el “ p value”, es mayor.

El enfoque bayesiano es matemáticamente coherente con el hecho de que, si bien el apoyo que puede otorgarse a cierta teoría puede inicialmente ser vago y subjetivo, el cúmulo de datos conduce progresivamente hacia una única y objetiva realidad acerca de la cual todos coincidirán¹⁷.

Como se ha dicho, lamentablemente se puede tener un valor pequeño de p sin que ello signifique que el de $P(H | D)$ lo sea. ¿En qué medida se produce tal discrepancia? Eso depende de la prueba empleada y de las circunstancias propias de cada caso; además, es imposible calcular el valor $P(H_0 | D=d_0)$ exactamente. Pero sí es posible, sin embargo, obtener cotas inferiores para dicha magnitud.

Por conducto del Teorema de Bayes es fácil corroborar que:

$$P(H_0 | D = d_0) = \frac{1}{1 + \frac{1 - P(H_0)}{P(H_0) BF}}$$

donde $P(H_0)$ denota la probabilidad a priori que cabe otorgar a la hipótesis nula de ser cierta y BF es el llamado “factor de Bayes”, definido como:

$$BF = \frac{P(D = d_0 | H_0)}{P(D = d_0 | \bar{H}_0)}$$

Naturalmente, mientras H_0 es una afirmación concreta (típicamente, que dos tratamientos son idénticos), la falsedad de H_0 puede asumir infinitas expresiones puntuales. Sin embargo, en condiciones bastante generales, se pueden hallar cotas inferiores para BF .

Comparación de medias

Cuando se comparan dos medias, si se llama Z al estadístico por conducto del cual se obtiene el valor p , se puede demostrar¹⁸ que se cumple:

$$\exp\left(-\frac{Z^2}{2}\right) \leq FB$$

Consecuentemente, se tiene:

$$\left[1 + \frac{1 - P(H_0)}{P(H_0)} \exp\left(\frac{Z^2}{2}\right)\right]^{-1} \leq P(H_0 | D = d_0).$$

Prueba Ji-cuadrado

Análogamente, para el caso en que se valora una asociación a través de la prueba Ji-cuadrado, se tiene una cota inferior para el factor de Bayes¹⁹ en función del valor observado χ^2 :

$$\sqrt{\chi^2} \exp\left(\frac{1 - \chi^2}{2}\right) \leq FB$$

Esto da lugar a la siguiente desigualdad:

$$\left[1 + \frac{1 - P(H_0)}{P(H_0) \sqrt{\chi^2} \exp\left(\frac{1 - \chi^2}{2}\right)}\right]^{-1} \leq P(H_0 | D = d_0)$$

Ejemplo

Considérese un estudio publicado en una prestigiosa revista debido a Harris y col²⁰. Allí se da cuenta de una valoración de los efectos a distancia de la intercesión mediante plegarias religiosas como recurso terapéutico a través de un ensayo clínico controlado en pacientes ingresados en una unidad de cuidados coronarios.

Se seleccionaron todos los pacientes que ingresaron en la Unidad de Cuidados Coronarios (UCC) del Instituto Americano de Cardiología, Ciudad de Kansas, a lo largo de un año, salvo los que habrían de recibir un transplante de corazón, para quienes se anticipaba una estadía prolongada en el servicio.

Los pacientes, registrados diariamente, fueron asignados a los grupos de tratamiento usando el último dígito del número con que fue recepcionado (número del registro médico): los pares eran asignados al grupo experimental (466 pacientes) y los impares al grupo control (524 pacientes). Se procuraba valorar así el efecto de las oraciones como recurso para disminuir las complicaciones que sobrevienen a los pacientes de enfermedades coronarias.

Una vez que el paciente era asignado al grupo experimental, se informaba su nombre (sin apellidos) por vía telefónica al responsable de uno de los grupos de intercesores para que se comenzara a orar por espacio de 28 días exactos. Para cuantificar los resultados, se creó al efecto un índice de complicaciones (IDC). A cada una de las posibles complicaciones (o acciones

derivadas de ellas) se le atribuyó un número (ponderación) según la gravedad que ellas entrañaban y de ese modo se asignaba un valor a cada paciente. Para valorar estadísticamente la hipótesis de nulidad se compararon las medias de IDC mediante la prueba t de Student de dos colas. Los resultados finales obtenidos por Harris y colaboradores se sintetizan en la Tabla 7.

Tabla 7. Efectos de la intercesión, mediante la oración, sobre el IDC de pacientes en la Unidad de Cuidados Coronarios.

Datos relevantes	Grupo control	Grupo experimental
Índice de Complicaciones (IDC) promedio	7,13	6,35
Tamaño muestral	524	466
Desviación estándar	5,95	5,82

El valor de z en este caso viene dado por la expresión:

$$z = \frac{\bar{x}_c - \bar{x}_e}{\sqrt{\frac{s_c^2}{n_c} + \frac{s_e^2}{n_e}}} = \frac{7,13 - 6,35}{\sqrt{\frac{5,95^2}{524} + \frac{5,82^2}{466}}} = 2,08$$

Debe advertirse que, en rigor, este estadístico sigue una distribución t de Student; sin embargo, la condición se cumple en esencia aunque no se trate exactamente de una normal. Por otra parte, para los tamaños muestrales de este ejemplo, la similitud entre la t y la normal es virtualmente total.

Es fácil ver que a este valor de z corresponde uno de p igual a 0,04. El valor de la cota inferior para $P(H_0 | D=d_0)$ depende de la probabilidad a priori que se atribuya a la hipótesis de que es lo mismo ser motivo de los rezos que no serlo. Si se atribuyera, por ejemplo, el valor 0,5 a dicha probabilidad, se tendría:

$$\left[1 + \frac{1 - P(H_0)}{P(H_0)} \exp\left(\frac{Z^2}{2}\right) \right]^{-1} = 0,10$$

El escollo más importante que se presenta al aplicar los procedimientos bayesianos es, precisamente, el de fijar distribuciones de probabilidad a priori, lo que en este caso se reduce a dar un valor anticipado para $P(H_0)$.

Las opiniones sobre el poder real de los rezos en la salud pueden ser diversas, pero es, como mínimo, muy razonable partir de un posicionamiento neutro (agnóstico) que se traduciría en no escorarse ni a favor ni en contra de la medida y, por ende, en admitir que $P(H_0)=0,5$. Pero en algunos casos (y el que se analiza es un buen ejemplo), podría ser más razonable aun adoptar un punto de vista escéptico (por ejemplo $P(H_0)=0,8$ ó $P(H_0)=0,9$).

La tabla siguiente recoge los valores de probabilidad a posteriori que, como mínimo, se puede atribuir a $P(H_0 | D=d_0)$ para diferentes valoraciones a priori de los datos obtenidos en el estudio.

Niveles	Probabilidad a priori	Valor mínimo de $P(H_0 D=d_0)$
Agnóstico	0,5	0,10

Escéptico	0,8	0,31
Muy escéptico	0,9	0,51

Como se ve, la probabilidad de que las plegarias sean totalmente estériles (H_0) es mucho mayor que p incluso en caso de que no se adopte una posición escéptica sino neutra.

Supóngase ahora en el mismo contexto del estudio anterior que se hubiera obtenido que en el grupo control murieron 124 pacientes y en el experimental 138. Con esos datos se podría configurar la siguiente tabla de 2x2:

Muerte	Grupos	
	Control	Experimental
Sí	a=124	b=138
No	c=400	d=328
	$n_c=524$	$n_e=466$

El valor observado de Ji-cuadrado es igual a:

$$\chi^2 = \frac{n(ad - bc)^2}{n_c n_e (a + b)(c + d)} = 4,49$$

y a este resultado se asocia un valor de $p=0,034$.

Los valores de probabilidad a posteriori que, como mínimo, se puede atribuir a $P(H_0 | D=d_0)$ para diferentes valoraciones a priori de los datos obtenidos en el estudio se muestra a continuación.

Niveles	Probabilidad a priori	Valor mínimo de $P(H_0 D=d_0)$
Agnóstico	0,5	0,27
Escéptico	0,8	0,60
Muy escéptico	0,9	0,77

Esto quiere decir que, incluso en el caso en que se adopte un punto de vista inicial que no se pronuncie a priori ni a favor ni en contra de la validez de la hipótesis de nulidad (que establece que la probabilidad de muerte no varía si se adicionan rezos a los cuidados normales) la probabilidad de que tal hipótesis sea correcta es sumamente alta (como mínimo, asciende a 0,27) a diferencia de lo que podría hacer pensar el valor de p (0,034), si se interpretara incorrectamente, como suele ocurrir.

BIBLIOGRAFÍA

- Berger JO, Berry DA. Statistical analysis and the illusion of objectivity. *American Scientist* 1988; 76: 159-65.
- Silva LC, Benavides A. El enfoque bayesiano: otra manera de inferir. *Gac Sanit* 2001; 15: 341-6.

3. Benavides A, Silva LC. Contra la sumisión estadística: un apunte sobre las pruebas de significación. *Metas de Enfermería* 2000; 3: 35-40.
4. Bayes T. Essay towards solving a problem in the doctrine of chances. [Reproduced from *Phil Trans Roy Soc* 1763; 53:370-418]. Studies in the history of probability and statistics. IX. With a bibliographical note by G.A. Barnard. *Biometrika* 1958; 45: 299.
5. Silva LC, Muñoz A. Debate sobre métodos frecuentistas vs bayesianos. *Gac Sanit* 2000; 14: 482-94.
6. Greenland S. Probability logic and probabilistic induction. *Epidemiology* 1998; 9: 322-32.
7. Lilford RJ, Braunholtz D. The statistical basis of public policy: a paradigm shift is overdue. *Br Med J* 1996; 313: 603-7.
8. Silva LC. *Excursión a la regresión logística en ciencias de la salud*. Madrid: Díaz de Santos; 1995.
9. Carlin BP, Louis TA. *Bayes and empirical Bayes methods for data analysis*. 2nd ed. New York: Chapman & Hall/CRC; 2000.
10. Burton PR. Helping doctors to draw appropriate inferences from the analysis of medical studies. *Stat Med* 1994; 13: 1699-713.
11. Davidoff F. Standing statistics right side up. *Ann Intern Med* 1999; 130: 1019-21.
12. Albert J. *Bayesian computation using Minitab*. Belmont, California: Duxbury Press, Wadsworth Publishing Company; 1996.
13. Berry DA. *Statistics: A Bayesian perspective*. Belmont, California: Duxbury Press; 1996.
14. Matthews RA. Significance levels for the assessment of anomalous phenomena. *Journal of Scientific Exploration* 1999; 13(1): 1-7.
15. Lindley DV. The analysis of experimental data: The appreciation of tea and wine. *Teaching Statistics* 1993; 15: 22-5.
16. Johnson DH. The Insignificance of statistical significance testing. *Journal of Wildlife Management* 1999; 63(3): 1629-32.
17. Matthews RA. Facts versus Factions: the use and abuse of subjectivity in scientific research. European Science and Environment Forum Working Paper; reprinted in *Rethinking Risk and the Precautionary Principle*. Ed: Morris, J. Oxford: Butterworth; 2000.
18. Lee PM. *Bayesian Statistics: an introduction*. 2nd ed. London: Arnold; 1997.
19. Berger J, Sellke T. Testing a point null hypothesis: the irreconcilability of P-values and evidence. *Journal of the American Statistical Association* 1987; 82: 112.

20. Harris WS, Gowda M, Kolb JW, Strychacz CP, Vacek JL, Jones PG et al. A randomized, controlled trial of effects of remote, intercessory prayer on outcomes in patients admitted to the coronary care unit. *Arch Intern Med* 1999; 159: 2273-8.