

El tamaño de la muestra

Vicente Manzano Arrondo – 2009-2013

Decidir cuál es el mejor tamaño para una muestra es una de las preocupaciones principales relativas al muestreo. Si atiendo a las inquietudes de las muchas personas que han pasado por mi despacho, graduadas en disciplinas diversas y preocupadas por aspectos de la investigación que están llevando, el tamaño de la muestra les interesa más que cualquier otro asunto de la muestra, como su representatividad, por ejemplo.

El primer aviso es que no existe un tamaño bueno para todo. Según el tipo de muestreo que se vaya a realizar, los objetivos que se persigan, las características de la población y las condiciones en las que se van a realizar las estimaciones, serán aconsejables unos tamaños u otros. Podría parecer que una muestra es mejor cuanto más grande. Pues sí, podría parecerlo, pero no tiene por qué ser cierto. Cuanto más grande, las estimaciones serán más precisas y con menos riesgo de error. Pero también saldrán más caras y tal vez se reduzca el control en la recogida de datos, por lo que, repito, no existe un tamaño bueno para todo. Ocurre además que si el

muestreo ha sido malo, la muestra grande será grande pero igualmente mala.

Este documento tampoco es bueno para toda decisión de tamaños de muestra. Se ha pensado para unas situaciones y no para cualesquiera. Entre otras cosas, parte de un único modelo de muestreo: aleatorio simple. Pero no nos preocupemos en exceso. Mi experiencia es que buena parte de las investigaciones científicas publicadas en revistas de prestigio hacen muestreos que dejan bastante que desear. Intentaremos no ser peores sino mejorar un poco el panorama. Por otro lado, aunque no todas las situaciones estén contempladas aquí, sí que lo está la lógica común a todas ellas.

Para entender este documento es necesario que cuentes con conocimientos previos, que se garantizan si has aprovechado antes los monográficos sobre muestreo y estimación estadística.

Cálculo del tamaño de muestra

Algunas consideraciones previas

Ya sabes que hay dos elementos fundamentales de una estimación que guardan una relación inversa entre ellos: la precisión y la seguridad. Son dos objetivos altamente deseables pero que se contraponen: a más precisión menos seguridad, a más

seguridad menos precisión. Ya lo hemos razonado suficientemente en la estimación estadística.

Sin embargo, hay un camino para incrementar la precisión y la seguridad hasta el nivel que queramos: aumentar el tamaño de la muestra. Conforme n se hace mayor, disminuye el error tipo y, por tanto, también el error de precisión, generando un intervalo más estrecho, es decir, más preciso, sin que ello haya requerido tocar el nivel de seguridad. Por si no lo recuerdas, observa la fórmula que permite calcular el error de precisión:

$$e_p = Z_{seg} \frac{\sigma}{\sqrt{n}}$$

La expresión, por un lado, permite observar que al aumentar la distancia estandarizada (Z) que expresa la seguridad, aumenta también el error de precisión, que amplía el intervalo. En otras palabras: a más seguridad, más imprecisión. Pero si no queremos tocar la seguridad (Z permanece constante) al mismo tiempo que deseamos incrementar la precisión (disminuir el error de precisión), entonces tenemos otro recurso: aumentar el tamaño de la muestra, puesto que conforme n se hace más grande, e_p se hará más pequeño.

La cosa no es tan sencilla como aumentar n indefinidamente. Conforme el tamaño de la muestra se hace más grande, también lo hacen el coste

económico, el tiempo necesario y los inconvenientes del trabajo de campo. Así que las soluciones perfectas no existen en esta dimensión y siempre hay que manejar variables que apuntan a direcciones diferentes. El objetivo es conseguir un n razonable en tres sentidos: apunta a un nivel de seguridad razonable, con una precisión razonable y unos recursos razonables.

La lógica, pues, es decidir a qué precisión y a qué seguridad aspiramos, y calcular con ello el tamaño que debería tener la muestra. Si contamos con una estimación aceptable de la desviación tipo poblacional (σ) y hemos decidido qué valores son los razonables para el error de precisión (e_p) y la seguridad (Z_{seg}), bastará con despejar de la fórmula anterior el valor que ha de tener el tamaño de la muestra. Que n dependa de σ nos va a dar algún problema. En el caso del muestreo aleatorio simple y estimación de una media aritmética, esa característica poblacional o parámetro es la desviación tipo. Obviamente no la conocemos. Tampoco podemos acudir a la desviación o cuasidesviación de la muestra, pues no hay muestra todavía. Ya nos enfrentaremos a este problema más adelante. Lo importante aquí es saber que hay soluciones, aunque tampoco perfectas. Del acierto de la solución dependerá que nuestras previsiones sobre la precisión y la seguridad se cumplan cuando estemos realizando los análisis con los datos de la muestra.

Antes de entrar de lleno, unas puntualizaciones concretas:

1. n depende del tipo o modelo de muestro. Cada uno requiere sus propias expresiones de cálculo. Aquí partimos del muestreo aleatorio simple.
2. n depende del objetivo de análisis. Aquí vamos a suponer que estimamos una media aritmética o una proporción.
3. En una investigación se busca responder a varios objetivos. Cada objetivo requiere su propio cálculo de tamaño de muestra. Pero la muestra es única para todos. ¿Qué hacer? Hablaremos de ello más adelante (ver apartado “Tamaño y objetivos”). Pero comenzaremos suponiendo que tenemos un solo objetivo.
4. Aunque cada objetivo y modelo de muestreo cuenta con sus propias expresiones de cálculo, siempre existen cuatro variables básicas: tamaño de la población, variación en la población, precisión de la estimación y confianza o riesgo de error. En cuatro apartados del mismo nombre vamos a tratar estos conceptos.
5. n es una apuesta. Para calcularlo hemos de suponer cosas que no sabemos con exactitud. Podemos haber supuesto bien, menos bien o incluso mal. Por eso solemos ponernos en la peor de las situaciones, como verás en el

ejemplo sobre la influencia de la variación de la población en el tamaño de la muestra.

Expresiones de cálculo

Las expresiones pueden deducirse directamente de las que ya conocemos para el cálculo del error de precisión. En el caso de la estimación de medias, si el tamaño de la población es suficientemente grande como para suponerlo prácticamente infinito, la expresión de cálculo será:

$$e_p = Z_{seg} \frac{\sigma}{\sqrt{n}} \quad \Rightarrow \quad n = \left(\frac{Z_{seg} \sigma}{e_p} \right)^2$$

Si se contempla el tamaño de la población y utilizando la expresión n_{inf} para representar el tamaño de muestra en poblaciones de tamaño prácticamente infinito:

$$e_p = Z_{seg} \sigma \sqrt{\frac{N - n}{n(N - 1)}} \quad \Rightarrow \quad n = \frac{N}{\frac{e_p^2 (N - 1)}{Z_{seg}^2 \sigma^2} + 1} =$$

Desde el monográfico sobre muestreo, estoy haciendo referencias frecuentes a una población “prácticamente infinita”. El infinito no ocurre en

poblaciones reales compuestas por unidades. Ese comportamiento “como si fuera una población de tamaño infinito” se refiere a situaciones en las que da igual utilizar una expresión para poblaciones finitas o para infinitas, puesto que el resultado es el mismo. No hay una cantidad exacta para ello, pues depende de varios aspectos. Así, la característica de población “prácticamente infinita” cuando estamos estimando una media aritmética, depende de los valores del error de precisión y del nivel de seguridad. Cuanto mayor sea el nivel de exigencia (más seguridad y más precisión), más grande tendrá que ser la población para ser considerada prácticamente infinita. En el siguiente apartado, sobre el tamaño de la población, puedes acceder a un ejemplo concreto.

Tamaño de la población

Parece lógico suponer que conforme aumenta el tamaño de la población, también debe hacerlo el de la muestra. En cierta medida es así. Pero la relación entre ambos no es proporcional: la muestra aumenta más despacio hasta que llega un momento en que por muy grande que se haga la población, el tamaño de la muestra no se inmuta.

El objetivo de la muestra es permitirnos realizar estimaciones sobre valores de la población. Es necesario partir de un tamaño mínimo. Estaremos de

acuerdo en que una muestra de $n = 1$ no sirve para generalizar nada cuando medimos variables que tienen alguna variación en la población. Si lo que queremos medir es si la gente se muere al quitarle la cabeza, no hace falta obtener ninguna muestra, pues ya sabemos que la cabeza es fundamental para mantener la vida (la utilicemos mucho o poco). Pero para averiguar qué piensan las personas sobre el precio de la remolacha o en qué medida están de acuerdo con la gestión de un personaje político, ya sabemos que las opiniones difieren y que no hay que jugársela con una muestra de $n = 1$. ¿Y $n = 2$? No hemos mejorado mucho.

Lo siento, no existe un tamaño mínimo. Pero sí un tamaño máximo.

Conforme aumenta el número de elementos de la población (personas, hogares, fábricas o lo que sea que estemos investigando), es lógico que deba ser mayor la muestra, pero cada nuevo elemento en la población añade menos a n . Esto también es lógico. De 1 a 2 hemos duplicado el número, pero de 107 a 108 no hay mucha diferencia, mientras que 9999999995 y 9999999996 son la misma cosa.

Vamos a realizar una estimación por intervalo de una proporción. Suponemos que la varianza poblacional es 0,25. Utilizamos un error de precisión de 0,03. Y el riesgo de error es 0,05. Vamos a ir considerando tamaños de población como potencias de

10. Verás cómo la muestra no se multiplica del mismo modo:

N	10	100	1000	10000	100000	1000000	10000000
n	10	92	517	965	1056	1066	1066
%	100	92	52	10	1	0,1	0,01

He añadido una fila más a la tabla: la relación entre n y N medida en porcentaje. Observa que en el caso de $N=10$, n es el 100% de N , mientras que el % va decreciendo con rapidez conforme aumentamos sensiblemente N al tiempo que n varía cada vez menos.

Esta circunstancia es trascendente. Imagina que queremos realizar estimaciones en una región con diez provincias. Cada provincia cuenta con 10 mil habitantes. Para hacer estimaciones con el error de precisión y el riesgo de error considerados, necesitamos una muestra de 1056 personas para toda la región ($N=$ cien mil habitantes). Pero si queremos hacer estimaciones por provincia, necesitamos 965 para cada una ($N_i=$ diez mil); es decir, 9650 (¡9 veces más!) en total.

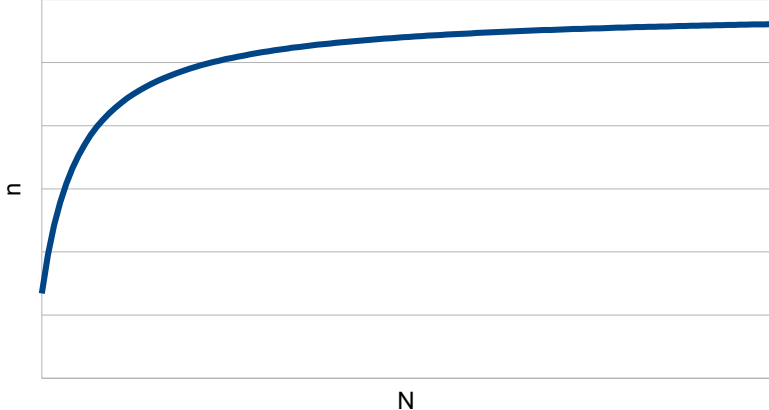


Figura 1. Relación entre n y tamaño de la población.

Variación en la población

También es lógico pensar que conforme varíe más lo que queremos conocer, será necesario indagar en más unidades. Si la gente opina de forma muy similar con respecto a algo, bastará con preguntar a unas pocas. Si queremos hacer un estudio sobre el gasto de los hogares andaluces en alimentación, tal vez necesitemos una muestra grande porque sospechamos que hay muchas formas diferentes de organizar la economía familiar y muchas peculiaridades.

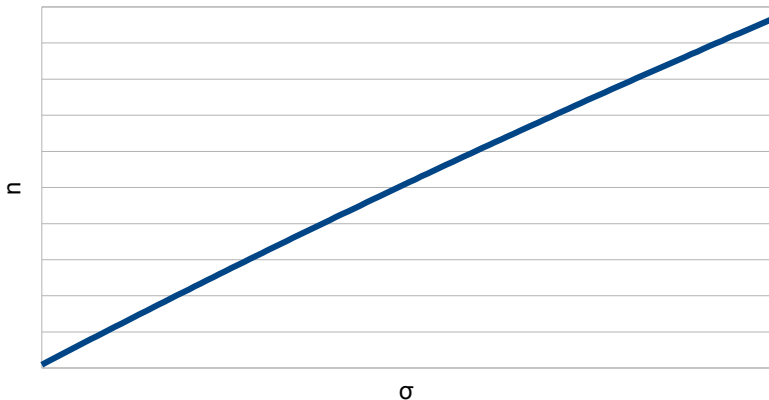


Figura 2. Relación entre n y variación.

Vamos a ver esta influencia con otra tabla, esta vez partimos de un tamaño de población constante (diez mil unidades) para estimar una proporción con un error de precisión de 0,03; riesgo de equivocación de 0,05; y una varianza que vamos a considerar entre 0,10 y 0,25; utilizando intervalos de 0,03.

σ^2	0,10	0,13	0,16	0,19	0,22	0,25
n	410	526	640	751	859	965

Cuando no conocemos la varianza en la población

He dicho que este documento se escribe suponiendo que vamos a realizar estimaciones de medias aritméticas o de proporciones. No he dicho nada de varianzas o desviaciones tipo. Se supone también que si no conocemos la media o la proporción poblacional ¿cómo vamos a conocer la desviación? ¿Qué hacemos entonces? Contamos con tres métodos. Uno es acudir a estudios previos sobre la misma temática que suministren información útil para definir de algún modo la varianza. Otro es llevar a cabo un *estudio piloto*, es decir, un ensayo general donde se recogen datos a partir de una muestra relativamente pequeña. Con esos datos se llevan a cabo unos análisis preliminares que permiten contar con un valor orientativo para la varianza poblacional. El estudio piloto no solo tiene esta función. Es siempre muy recomendable llevarlo a cabo, puesto que permite probar las herramientas de recogida de datos y otros detalles del trabajo de campo.

Cuando ocurre lo habitual, es decir, cuando no sabemos qué valor tiene la variación poblacional, adoptamos lo que se llama una “postura conservadora”. Ya hemos visto que conforme es mayor la variación, también es mayor el tamaño de la muestra. Imagina que infravaloramos la variación, que es 0,22 en lugar

de 0,13 como pensamos. La muestra va a ser más pequeña de lo que debería ser (calcularemos 526 en lugar de 859 que es lo que deberíamos haber hecho). Después, a la hora de hacer las estimaciones a partir de los datos de la muestra que hemos obtenido, nos tropezaremos con las consecuencias del error: por culpa de un tamaño de muestra demasiado pequeño, nuestras estimaciones serán más inseguras o más imprecisas o ambas cosas a un mismo tiempo.

Para evitar ese efecto indeseable, nos ponemos en la peor de las situaciones: la de la máxima varianza que pueda ocurrir teniendo en cuenta la información con que contamos. A esto se le llama “postura conservadora”. Si nos equivocamos será por exceso. Habremos implicado más dinero o más tiempo de lo necesario, pero conseguiremos estimaciones tanto o más precisas como lo deseado y tanto o más seguras de lo diseñado. Mejor eso que lo contrario.

¿Qué valor tiene la varianza entonces, desde una perspectiva conservadora? Hemos de pensar en una situación límite: la máxima varianza que se pueda obtener. En el caso de las proporciones es fácil. La varianza de proporciones es $p(1-p)$. Haz todas las pruebas que quieras con diferentes valores de p y encontrarás que la varianza máxima ocurre cuando $p=0,5$, lo que lleva a $S^2=0,5(1-0,5)=0,25$. Así que si no sabemos qué valor tiene la varianza poblacional

cuando estamos estimando proporciones, vale con la postura conservadora de que $\sigma^2=0,25$.

En las medias aritméticas es más complicado. Es muy difícil que el valor de la desviación tipo llegue a igualar el valor de la media aritmética. Como la varianza es el cuadrado de la desviación tipo, lo que podemos hacer es imaginar un intervalo de valores esperables para la media de la población, tomar el límite superior del intervalo y elevarlo al cuadrado. He ahí una estimación conservadora para la varianza poblacional. Por ejemplo, vamos a estimar el número de peines que compra una persona en una población, cada año. No sé nada sobre ello, pero dudaría que la media sea superior a 3. Entonces, podemos utilizar $3^2=9$ como estimación conservadora de la varianza.

Otro método es considerar la amplitud total de la variable cuya media estamos estimando, es decir, la distancia entre los valores mínimo y máximo. Considerando esa acotación de valores posibles, la máxima cuantía que puede tener la varianza es el cuadrado de semi-amplitud¹.

Error de precisión

1 Tienes una demostración matemática de ello en Manzano, V. (2002). Variación y posición de variables: estrategias para su estudio exploratorio. *Metodología de Encuestas*, 4(1), 4-61.

Como sabes del monográfico sobre la estimación estadística, el error de precisión es, al mismo tiempo:

- La cantidad que se resta y se suma al estadístico o estimador para construir el intervalo de confianza en cuyo interior esperamos que se encuentre el parámetro.
- La máxima diferencia que cabe esperar que exista entre el estadístico y el parámetro, dada una confianza concreta.

Luego, a mayor error de precisión, menos precisión (más imprecisión). Lo ideal es buscar la máxima precisión posible (el mínimo valor para el error de precisión). Pero como ya hemos visto, conforme hacemos el intervalo más pequeño, hay que exigir un tamaño de muestra más grande o correr más riesgo de errar al suponer que el parámetro se encuentra dentro del intervalo. Como el riesgo es algo que preferimos no tocar, está claro que la aspiración de un error de precisión mínimo debe buscarse mediante un tamaño de muestra máximo, en las circunstancias concretas donde se defina tal cosa.

Conforme aumenta el tamaño de la muestra se obtienen valores de los estadísticos que deben parecerse más al valor del parámetro, si todo lo demás sigue igual. La tabla siguiente muestra esta relación, partiendo de una estimación por intervalo de una

proporción, mediante un riesgo de error de 0,05, una población de tamaño 10 mil y una varianza de valor 0,25.

0,01	0,02	0,03	0,04	0,05	0,06	0,07	0,08	0,09
900	1937	965	567	370	260	193	148	118

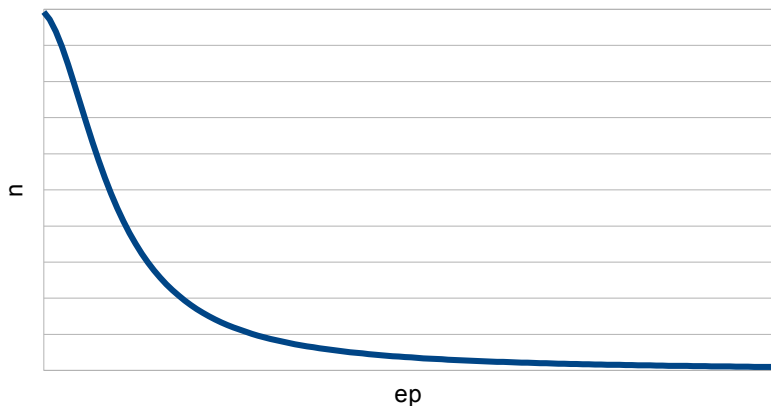


Figura 3. Relación entre n y el error de precisión.

Observa con qué velocidad se reduce el tamaño de la muestra al aumentar el error de precisión. Sin duda es una variable con efecto. Entre otras razones, por esto mismo es la que suele utilizarse para variar las consecuencias de un tamaño de muestra diferente. Imagina que aunque deseamos un error de precisión de valor 0,06, no tenemos dinero para 260 entrevistas.

De hecho, sólo podríamos costear 200. Pues no es muy grave si aceptamos igualmente un error de precisión de valor 0,07. O, al menos, solemos juzgar que sería aun más grave reducir la seguridad en las estimaciones.

¿Cómo decidir un valor concreto para e_p ? En principio hay que considerar la escala en la que se mueve el parámetro. Si vamos a estimar una proporción y creemos que ronda el 50%, podríamos considerar un buen intervalo, por ejemplo, de 45% a 55% (un error de precisión de 5%). Pero si la proporción ronda un 3%, no tiene sentido sumarle y restarle 5%. Algo admisible en este caso podría ser un error de precisión de valor 1%, que construye un intervalo del 2% al 4%, entre otras posibilidades.

Como eso de plantearse “la sensación positiva” de la amplitud de un intervalo es algo que roza lo asombroso para muchas mentes y muchos corazones, podemos pensar en estrategias aparentemente menos sentimentales. Una de ellas es considerar un error de precisión que no vaya más allá de un 10% del valor supuesto para el parámetro. Esto implica aceptar que la amplitud del intervalo de estimación no supere una quinta parte del valor de la media o de la proporción de la muestra a partir de la que estamos realizando la estimación. Por ejemplo, para una proporción del 50% decidimos un error de precisión del 5%; para una proporción del 3%, un error de precisión de 0,3%.

Pongamos que estimamos una media aritmética: el número medio de litros de bebida alcohólica que se consume en una familia española durante las vacaciones de Navidad. Creemos que ese valor debe encontrarse entre los 15 y 25 litros. Utilizamos el centro (20 litros) como referencia y un error de precisión de la décima parte, es decir, de 2 litros.

La sugerencia del párrafo anterior sólo tiene sentido en situaciones de pérdida total de referentes, que es lo que suele ocurrir, por ejemplo, cuando hay que hacer un trabajo para una asignatura. En una situación práctica se utilizan los referentes habituales. Así, por ejemplo, si vamos a estimar el número de escaños que va a ocupar un partido en el Congreso de los Diputados tras unas elecciones generales, lo habitual es procurar un error de precisión no superior a 2 escaños. Muy especialmente cuando la batalla por el poder parlamentario es dura, unos pocos escaños implican un resultado muy distinto, por lo que la amplitud del intervalo es fundamental.

Nivel de confianza o riesgo de error

Como ya sabemos, el nivel de confianza es la probabilidad de acertar al realizar una estimación, mientras que el riesgo de error es su complementario: la probabilidad de errar en la estimación. En las estimaciones suele pensarse más en confianza,

mientras que en una prueba de significación de la hipótesis nula se tiende a pensar en términos de riesgo de error. Es un valor de probabilidad que, como todos los valores de probabilidad, debe encontrarse entre 0 y 1.

Como ocurre con el error de precisión, es el equipo investigador quien decide un valor para el riesgo de error. ¿Cómo escogerlo? Sabemos que a mayor riesgo, menor tamaño de muestra, puesto que somos menos exigentes con la situación. Lo ideal es no equivocarse, lo que es imposible de garantizar. Así que aspiramos al mínimo riesgo en la práctica, lo que aconseja el máximo tamaño de muestra. Observa la relación en la siguiente tabla, para la estimación de una proporción con una población de 10 mil unidades, una varianza poblacional de 0,25 y un error de precisión de 0,03.

α	0,01	0,02	0,03	0,04	0,05	0,06	0,07	0,08	0,09
	1557	1307	1157	1049	965	895	836	785	740

Ya ves que su efecto no parece ser tan drástico como el del error de precisión. Observa no obstante el comportamiento de las figuras 4 y 5. Dado que la distribución muestral muestra una gran acumulación en torno a la media y una dilatada dispersión en los extremos, una misma variación de seguridad o de riesgo cuando nos encontramos cerca de la media

implica una variación muy pequeña en Z , pero esta relación se invierte conforme nos alejamos del centro de la curva. En la fórmula de cálculo para el tamaño de la muestra se contempla el riesgo en términos de distancia estandarizada. La figura 5 expresa cómo las variaciones lineales de Z generan efectos en n . Pero si partimos de valores de probabilidad según una curva normal, pequeñas variaciones del riesgo tendrán un efecto pequeño o pronunciado en Z (y, por tanto, en n) en función de si hablamos respectivamente de valores altos o bajos de riesgo. Ocurre además que la figura 4 es decreciente, mientras que la 5 es creciente. Esto se debe a que la distancia estandarizada es mayor conforme es mayor la confianza, es decir, lo contrario del riesgo.

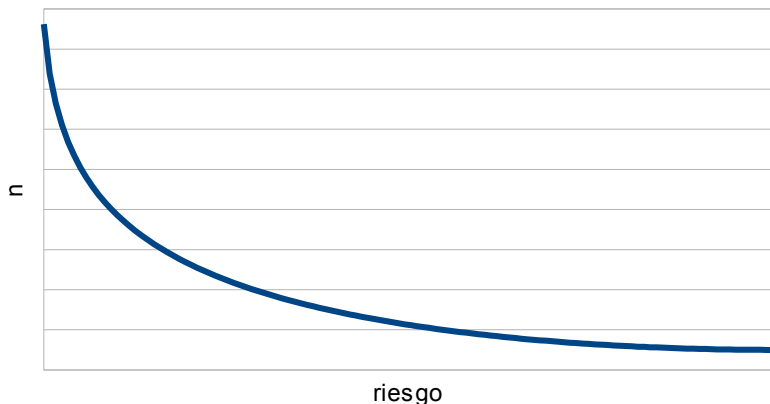


Figura 4. Relación entre n y el riesgo.

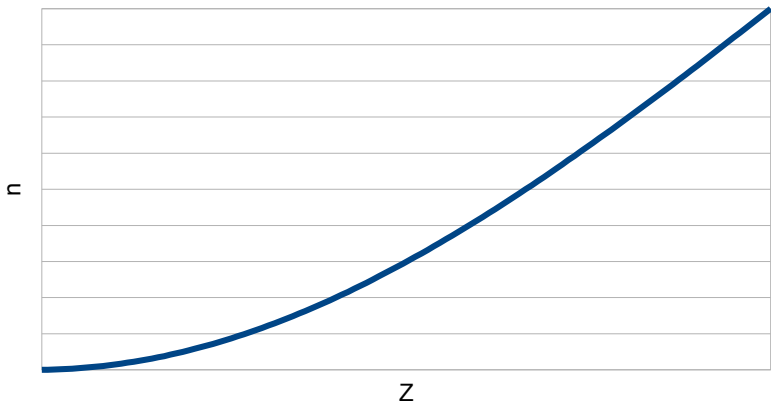


Figura 5. Relación entre n y la seguridad estandarizada.

Pero, volvamos a la pregunta ¿Cómo escoger un valor para el riesgo de error?

Lo suyo sería pensar en cuáles son las consecuencias que se derivan de equivocarse. Cuanto más graves o indeseables sean, menor tendrá que ser el riesgo de llegar a ello. Imagina cuál es el valor de riesgo que considerarías en las siguientes ocasiones:

- 1) si te equivocas mueres pues hablamos del riesgo de tener un accidente de tráfico si aprietas o no ese botón,
- 2) si te equivocas no tendrás el yogurt gratis, pues hablamos de tomar una decisión sobre qué calle tomar

para ir a la universidad sabiendo que en las inmediaciones están repartiendo yogurt.

En muchas ocasiones (por ejemplo, trabajos de asignaturas, como siempre), es difícil ponerse en situación y decidir un valor para el riesgo que sea adecuado a la conciencia de las consecuencias. En estos casos acudimos a Fisher, al que hemos conocido en el monográfico sobre estimación estadística. Desde que sugirió un riesgo del 5%, nadie ha abierto la boca para decir otra cosa. Se considera que un riesgo de 0,05 es válido por defecto. Si alguien osa tomar otra decisión se le preguntará por qué. Como no queremos dar explicaciones, todo el mundo escoge irreflexivamente 0,05. Se admiten dos excepciones: 0,01 para mostrar que uno es muy exigente y ** , es decir, el número de asteriscos que suministra SPSS, un programa de ordenador que se está utilizando como referente desde hace tiempo. Sabemos que muchas personas deciden dejar de pensar cuando se les presenta la oportunidad para ello. SPSS, un programa tonto de ordenador (como todos los programas de ordenador, por definición), suministra esa oportunidad. Si afirmas “me lo dijo SPSS”, suelen dejarte en paz.

Tamaño y objetivos

Los estudios tienen varios objetivos de análisis. En principio, habría que calcular un tamaño de muestra

para cada objetivo. Imagina este caso: tenemos dos proporciones que estimar. Una consideramos que ronda el valor 40% (es decir, una varianza de $0,4 \cdot (1-0,4) = 0,24$; y un error de precisión de $0,4/10 = 0,04$). En una población de tamaño prácticamente infinito y un riesgo de 0,05, el tamaño de muestra es $n = 577$. De la otra proporción suponemos que tiene un valor en torno al 10% (es decir, una varianza de $0,1 \cdot (1-0,1) = 0,09$; y un error de precisión de $0,1/10 = 0,01$). En la misma población de tamaño prácticamente infinito y un riesgo de 0,05, el tamaño de muestra es $n = 3458$. ¿Qué hacemos? Cada objetivo exige su propio tamaño de muestra, mientras que hemos de trabajar con la misma muestra para todos los objetivos.

Una postura conservadora aconseja tomar el tamaño más grande de entre todos los calculados. En este caso es 3458. No obstante, puede ser muy exagerado. En tales situaciones, podemos prescindir de los tamaños excesivos, sabiendo que las estimaciones que realizaremos después para esos objetivos tendrán menos precisión o seguridad. Otra posibilidad es calcular n para cada objetivo y decidir con la media o la mediana (más resistente a las cantidades raras) de todos los tamaños calculados. Puede ser una solución de compromiso. La cuestión es que no hay una norma matemática o estadística para resolver este problema. Lo ideal sería tener delante una tabla con todos los objetivos y los tamaños de muestra

que aconseja cada uno, añadiendo una columna con los errores de precisión que se derivarían en cada uno de los objetivos en función de cada decisión concreta para el tamaño de muestra. Observando esa tabla, tomaremos la decisión que consideremos mejor atendiendo especialmente a los objetivos que se quedan más desprotegidos. No obstante, esta tarea es tediosa y no es habitual llevarla a cabo. En el cálculo de n con RAD tienes un procedimiento para conseguirlo.

En la práctica, lo más recomendable dentro de lo no excesivamente laborioso, es seleccionar un objetivo estrella, el más importante de todos, y calcular el tamaño de muestra según lo necesario para él.

Cálculo de n con MAS

Existen varias utilidades que permiten calcular el tamaño de una muestra. Hay de pago y gratis. Las hay de Internet y de puesto local. Las hay en Windows, en Linux y en otros sistemas. Las hay más claras y más oscuras, con más o menos situaciones de consideración, etc.

Una de las utilidades de libre distribución, pensadas para cubrir todas las situaciones expuestas aquí, es MAS. Se puede bajar de

La siguiente figura muestra una pantalla típica para MAS. Como observarás, se encuentran todas las variables consideradas: varianza de la población, riesgo de error, error de precisión, tamaño de la población y tamaño de la muestra.

Hay dos aspectos interesantes para resaltar. El primero es que MAS sirve también para poner cualquiera de las variables en función de las otras. Observa el recuadro de la derecha de la ventana. Al marcar sobre la variable, esta se transforma en dependiente, su recuadro queda marcado de amarillo y recibirá los resultados de cálculo que se establezcan sobre todas las demás.

02:30:58

Muestreo Aleatorio Simple

Salir

Varianza en la población (V): 0,09

Probabilidad de error (prob): 0,05

Distancia estandarizada (Z): 1,96

Error de precisión (E): 0,01

Tamaño de la muestra (n): 3458

Tamaño de la población (N):

☒ Infinito

☐ Finito

10000

Dependiente

☐ V

☐ prob.

☐ E

☒ n

☐ N

Leer / Escribir

Información

Calcula

Figura 6. Pantalla principal de MAS.

La segunda característica de interés a resaltar es que MAS genera un informe donde consta la situación de partida, los valores de las variables independientes, la fórmula de cálculo y el resultado. Se accede a esta generación pulsando el botón “Leer / Escribir” y escogiendo la opción “Generar”.

En el ejemplo que se muestra en la pantalla, el resultado de la generación del informe es el que sigue:

MAS: Información de contexto y resultados

* Contexto:

Muestreo Aleatorio Simple desde una población de tamaño infinito.
Estimación de una media o proporción.
Supuesta una distribución muestral normal.
Tiempo de interacción: 01:43:14.

* Variables:

Variable dependiente: Tamaño de la muestra
(n)

Variables independientes:

Varianza poblacional (V) :
0,09

Distancia estandarizada (Z) :
1,96

Probabilidad de error (p) :
0,05

Error de precisión (E) :
0,01

* Expresión de cálculo:

$$n = V \cdot Z^2 / E^2$$

* Resultado: 3458

Cálculo de n con RAD

RAD es un programa más ambicioso que MAS. En su versión actual permite resolver varios objetivos relacionados con el muestreo: calcular el número

posible de muestras, decidir un tamaño de muestra para un objetivo y decidirlo también para un conjunto de objetivos.

Las tres opciones las tienes disponibles desde el menú *Muestreo* del principal. Al escoger *Tamaño de muestra* accedes a una ventana similar a la de MAS, pero con una diferencia fundamental: cuenta con botones de ayuda para guiarte en la decisión de valores ante los que dudas.

Muestreo: cálculo del tamaño de la muestra.

Tamaño de la población: ☒ Finito: ☐ Infinito

Varianza de la característica:

Error de precisión:

Riesgo estandarizado:

RAD: cálculo del tamaño de la muestra

Tamaño de la población (N): 12300
Varianza poblacional (V): 100
Error de precisión (E): 2
Riesgo estandarizado (Z): 1,96

Tamaño de la muestra $n = N / E^2 * [N - 1] / Z^2 / V + 1 = 95$

Figura 7. Pantalla de RAD para el cálculo de n .

Si pulsas sobre los botones de ayuda, RAD te orientará siguiendo las mismas explicaciones que tienes en este monográfico. La figura 8 muestra un ejemplo para el caso de la varianza, donde se sugiere $\sigma^2 = 100$. Podrás observar que si acudes a las ayudas, ese proceso queda finalmente registrado en el recuadro de resultados que se genera al pulsar sobre el botón “Calcular”.

RAD también permite resolver problemas de objetivos múltiples. Para probarlo, imaginemos que contamos con tres objetivos: dos estimaciones de proporciones y una estimación de media aritmética. Supongamos que la población es finita, con un tamaño de 13200 unidades y que deseamos hacer las estimaciones con una confianza del 97% (es decir, un riesgo del 3%). Estas dos informaciones (tamaño de la población y riesgo en las estimaciones) van a ser comunes para todos los objetivos. Siguiendo un muestreo aleatorio simple, lo que añade cada objetivo son las dos informaciones que faltan: la varianza poblacional de la característica y el error de precisión para el intervalo. Pues bien, supongamos que estos valores son:

Objetivo	Varianza	Error de precisión
1	0,25	0,03
2	100	2
3	0,12	0,02


Acotación de la varianza


☐ Se va a estimar una proporción

☒ Se va a estimar una media aritmética

☒ La variable se mueve entre los valores:

 mín y máx

☐ Un valor razonable para la media es:

Figura 8. Ayuda para el parámetro varianza.

Al introducir esta información en la opción *Muestreo/Tamaño multiobjetivo* de RAD, la pantalla queda como consta en la figura 9.

Muestreo: múltiples objetivos.

Riesgo: 0,03

Población: ☐ Infinita ☒ Finita: 12300

Obj.	V	Ep	n	Sol. Ep	Sol. n
Obj.1	0,25	0,03			
Obj.2	100	2			
Obj.3	0,12	0,02			

Calcular n's

Nuevo objetivo

Criterios:

☒ n máximo

☐ n mínimo

☐ n medio

☐ n mediano

Aplicar

Volver Limpiar Guardar

Figura 9. Estudio de n común para varios objetivos.

Al pulsar sobre el botón “Calcular”, RAD resuelve los tres tamaños de muestra, que son dispares. Si escogemos la opción “n mediano” y pulsamos “Aplicar”, guardando después el resultado, observa lo que obtenemos:

RAD: registro de objetivos múltiples

- Tamaño de población: 12300
- Riesgo de error: 0,03

Objetivos:

Obj	Var	-----Previo-----		-----Posterior-----	
		Ep	n	Ep	n
1	0,25	0,03	1182	0,030	1182
2	100	2	117	0,600	1182
3	0,12	0,02	1267	0,021	1182

Como puedes ver, la segunda columna Ep contiene los efectos que resultan de forzar el mismo tamaño de muestra para todos los objetivos.

Ensayando diversas soluciones (nuevos errores de precisión, otros criterios...), llegaremos a una buena decisión sobre el tamaño común de muestra para todos los objetivos. En este caso, un tamaño de $n=1182$ para los tres objetivos permite mantener el error de precisión inicial para los objetivos 1 y 3 y mejorarlo sensiblemente para el objetivo 2 (de 2 a 0,6).