

mientras que otros se oponen. Para dar un paso adelante, digamos que usted tiene la impresión de que los estudiantes de humanidades y ciencias sociales parecen en general más inclinados a sostener la idea del derecho a abortar que los alumnos de ciencias naturales (estos razonamientos suelen llevar a la gente a diseñar y realizar investigaciones sociales).

En términos de las opciones que revisamos en el capítulo, es muy probable que su investigación fuera exploratoria. Quizá tiene intereses tanto descriptivos como explicativos: ¿qué porcentaje de estudiantes apoyan el derecho de las mujeres a abortar y cuáles son las causas de que algunos estén en favor y otros se opongan? Las unidades de análisis son individuos: los estudiantes. Tal vez decida que un estudio transversal sería adecuado para sus propósitos. Digamos que usted estaría contento con saber un poco sobre el estado actual de las cosas. Aunque no tendría pruebas de los procesos que estén en curso, sería capaz de aproximar algunos análisis longitudinales.

Preparación

La parte superior de la figura 4.2 contiene varias actividades posibles. Al llevar adelante su interés en las actitudes de los estudiantes sobre el derecho de abortar, sin duda querrá leer sobre el tema. Si usted tiene la corazonada de que las actitudes se relacionan de alguna manera con la especialidad universitaria, podría averiguar lo que otros investigadores han escrito al respecto. El apéndice A lo ayudará a consultar la biblioteca de su escuela. Además, es probable que quiera hablar con algunos de los que apoyan el derecho de abortar y con algunos que se oponen. También querrá asistir a las reuniones de grupos relacionados con el tema. Todas estas actividades lo prepararán para tomar las decisiones sobre el diseño de la investigación que estamos a punto de examinar. A medida que revise la bibliografía sobre el derecho de abortar, observe las decisiones que han tomado otros investigadores sobre el diseño de su proyecto, preguntándose siempre si las mismas decisiones cumplirían su propósito.

Por cierto, ¿cuál es su propósito? Es importante que lo aclare antes de diseñar su estudio. ¿Planea escribir un ensayo basado en su investigación para cumplir con un requisito del curso o como una

tesis honorífica? ¿Su objetivo es conseguir información para respaldar sus argumentos en favor o en contra del derecho de abortar? ¿Quiere escribir un artículo para el periódico de la universidad o para una publicación académica?

Habitualmente, su propósito al emprender una investigación se puede expresar como un informe. El apéndice C le ayudará a organizar los informes de investigación; le recomiendo que la primera etapa del diseño de su proyecto consista en el borrador de su informe. En concreto, debe tener en claro las afirmaciones que desea sostener cuando termine la investigación. Estos son ejemplos de tales afirmaciones: "Los estudiantes suelen mencionar el derecho de abortar en el contexto de la discusión de temas sociales que les preocupan en lo personal". "X por ciento de los estudiantes de la universidad estatal está en favor del derecho de las mujeres a elegir el aborto." "Los ingenieros se inclinan (más/menos) que los sociólogos a apoyar el derecho de abortar."

Aunque su informe final no se parezca mucho a su imagen inicial, este ejercicio le dará material con el cual probar la pertinencia de los diseños de investigación.

Conceptuación

Con frecuencia hablamos en forma bastante casual de conceptos de las ciencias sociales como prejuicios, enajenación, religiosidad y liberalismo, pero es necesario aclarar lo que entendemos por esos conceptos para llegar a conclusiones significativas sobre ellos. En el capítulo 5 examinaremos a fondo este proceso de **conceptuación**. Por ahora, veamos en qué puede consistir en el caso de nuestro ejemplo hipotético.

Si usted va a estudiar lo que opinan los estudiantes universitarios sobre el aborto y por qué, lo primero que debe especificar es lo que entiende por "derecho de abortar". En concreto, deberá prestar atención a las condiciones en las que la gente aprobaría o desaprobaba el aborto; por ejemplo, cuando la vida de la mujer está en peligro, en el caso de violación o incesto o simplemente porque la mujer así lo desea. Descubrirá que el apoyo general al aborto varía de acuerdo con las circunstancias.

Desde luego, necesitará especificar todos los conceptos que planea estudiar. Si quiere estudiar el posible efecto de las especializaciones universita-

rias, tendrá que decidir si quiere considerar sólo las especialidades oficiales o también las intenciones de los estudiantes. ¿Qué va a hacer con los que no tienen especialidad?

Si realiza una encuesta o un experimento, deberá especificar de antemano esos conceptos. Si planea una investigación menos estructurada, como las entrevistas abiertas, una parte importante de su estudio consistirá en descubrir las dimensiones, aspectos o matices de los conceptos. Así, será capaz de revelar e informar aspectos de la vida social que no son accesibles con el uso más casual o menos riguroso del lenguaje.

Elección del método de investigación

Como veremos en la parte 3, el científico social cuenta con una variedad de métodos de investigación. Cada método tiene sus ventajas y desventajas, y se aplican mejor al estudio de ciertos conceptos que a otros.

En términos de nuestro estudio hipotético de las actitudes acerca del derecho de abortar, una encuesta podría ser el método más adecuado: ya sea entrevistar a los estudiantes o pedirles que llenen un cuestionario. Como veremos en el capítulo 10, las encuestas se prestan particularmente bien al estudio de la opinión pública. Esto no quiere decir que no pueda hacer un buen uso de los otros métodos presentados en la parte 3. Por ejemplo, mediante el *análisis de contenidos* (véase el capítulo 12) puede examinar las cartas al editor y analizar las posturas que tienen sobre el aborto los redactores de las cartas. La *investigación de campo* (capítulo 11) ofrecerá una vía para comprender cómo se relacionan las personas en cuanto al tema del aborto, cómo lo discuten y cómo cambian de opinión. Cuando lea la parte 3 verá cómo aplicar otros métodos de investigación al estudio de su tema. Habitualmente, el mejor diseño de investigación utiliza más de un método para aprovechar las ventajas de cada uno.

Operacionalización

Después de especificar los conceptos por estudiar y de elegir el método de investigación, debemos elegir nuestras técnicas u operaciones de medición (véase el capítulo 6). En algunos casos, esto requiere delinear concretamente las técnicas, como

la redacción de las preguntas de un cuestionario. En cualquier caso, debemos decidir cómo recopilar los datos deseados: observación directa, revisión de documentos oficiales, un cuestionario u otra técnica.

Si usted decidió estudiar con una encuesta las actitudes acerca del derecho de abortar, puede operacionalizar su principal variable preguntando a los entrevistados si aprobarían el derecho de las mujeres a abortar en varias condiciones que usted conceptuó: en el caso de violación o incesto, si su vida estuviera amenazada por el embarazo, etc. Le pediría a los entrevistados que aprobaran o desaprobaran por separado cada situación.

Población y muestreo

Además de perfeccionar conceptos y mediciones, debe decidir *quién* o *qué* estudiar. La *población* de un estudio es aquel grupo (por lo regular, de personas) del que queremos obtener conclusiones. Ahora bien, casi nunca podemos estudiar a todos los miembros de la población que nos interesa ni podemos hacer todas las observaciones posibles. Por tanto, en cada caso escogemos una *muestra* de los datos que podríamos recopilar y estudiar. Desde luego, el muestreo de la información ocurre en la vida diaria y a menudo produce observaciones sesgadas (recuerde nuestro análisis de la "observación selectiva" en el capítulo 1). Los investigadores sociales son más cuidadosos al tomar muestras que van a observar.

En el capítulo 8 se describen los métodos para seleccionar muestras que reflejen adecuadamente la población total que nos interesa. Observe en la figura 4.2 que las decisiones sobre la población y el muestreo se relacionan con el método de investigación elegido. Mientras que las técnicas de muestreo probabilístico serían importantes en una encuesta a gran escala o un análisis de contenidos, un investigador de campo necesitaría sólo a los informantes que le den una imagen equilibrada de la situación que estudia, y un experimentador asignaría los sujetos a los grupos experimental y de control de forma tal que fueran comparables.

En nuestro estudio hipotético de las actitudes hacia el aborto, la población pertinente serían los estudiantes de su universidad. Sin embargo, como descubriremos en el capítulo 8, para seleccionar una muestra se requiere ser aún más específico.

¿Incluirá estudiantes de medio tiempo igual que de tiempo completo? ¿Sólo estudiantes de posgrado o todos? ¿Ciudadanos nacionales o también extranjeros? ¿Estudiantes de licenciatura, posgrado o ambos? Hay muchas preguntas de esta clase, y debe responderlas de acuerdo con los propósitos de su investigación. Si su propósito es predecir cómo votarían los estudiantes en un referéndum local sobre el aborto, sería mejor limitar su población a quienes tienen el derecho de votar y sea probable que voten.

Observaciones

Tras decidir qué estudiar en quiénes y con qué método, usted está listo para hacer observaciones, para recopilar datos empíricos. Los capítulos de la parte 3, que describen los diversos métodos de investigación, indican las técnicas de observación adecuadas para cada uno.

En el caso de la encuesta sobre el aborto, lo mejor sería imprimir cuestionarios y enviarlos por correo a la muestra elegida de los estudiantes, o bien podría organizar a un grupo de encuestadores para que realice el estudio por teléfono. En el capítulo 10 revisaremos las ventajas y desventajas relativas de éstas y otras posibilidades.

Procesamiento de datos

Según el método de investigación elegido, habrá amasado un volumen de información en una forma que tal vez no pueda interpretarse de inmediato. Si ha dedicado un mes a observar de primera mano a una pandilla callejera, ahora tendrá suficientes notas de campo para escribir un libro. En un estudio histórico de la diversidad étnica de su escuela, quizá recopiló grandes cantidades de documentos oficiales, entrevistas con directores y otros, etc. En el capítulo 14 se describen algunas formas en que se procesan o transforman los datos científicos sociales para su análisis cuantitativo o cualitativo.

En el caso de una encuesta, las observaciones "crudas" suelen estar en la forma de cuestionarios con recuadros marcados, repuestas escritas en espacios en blanco, etc. La fase de procesamiento de los datos de una encuesta comprende la clasificación (*codificación*) de las respuestas escritas y la transferencia de toda la información a una computadora.

Análisis

Finalmente, interpretamos los datos reunidos con el fin de llegar a conclusiones que reflejen los intereses, ideas y teorías que iniciaron la investigación. En los capítulos 11 y 15 a 17 describiremos algunas de las opciones disponibles para analizar sus datos. Observe que los resultados de sus análisis regresan a sus intereses, ideas y teorías originales. En la práctica, esta vuelta bien puede representar el inicio de otro ciclo de investigación.

En la encuesta de las actitudes de los estudiantes hacia el derecho de abortar, la fase de análisis tendría objetivos tanto descriptivos como explicativos. Podría comenzar por calcular los porcentajes de estudiantes que están en favor o se oponen al derecho de abortar. En conjunto, estos porcentajes brindarían una buena imagen de la opinión estudiantil sobre el asunto.

Más allá de la simple descripción, podría referir las opiniones de varios subconjuntos de los estudiantes: hombres frente a mujeres, alumnos de primero, segundo, tercer o cuarto año y de posgrado, estudiantes de ingeniería, de sociología, de literatura, etc. La descripción de subgrupos podría llevarlo a un análisis explicativo, como se expone en el capítulo 15.

Aplicación

La última etapa del proceso de investigación consiste en los usos del estudio que realizó y las conclusiones a las que llegó. Para empezar, es probable que quiera comunicar sus descubrimientos, para que los demás conozcan lo que aprendió. Tal vez sería apropiado preparar e incluso publicar un informe escrito. Quizá hará presentaciones orales, como textos pronunciados en reuniones profesionales y científicas. Tal vez a otros estudiantes les interese oír lo que aprendió de ellos.

Tal vez usted quiera ir más allá de sólo informar lo que aprendió y discutir las implicaciones de sus descubrimientos. ¿Indican éstos algo sobre las acciones que se pueden emprender para apoyar las medidas políticas? Esto sería de interés tanto para los defensores como para quienes se oponen al derecho de abortar.

Por último, debe considerar el planteamiento de su investigación en cuanto a nuevos estudios de su tema. ¿Qué errores habría que corregir en estudios

futuros? ¿Qué vías sugeridas por su estudio habría que ahondar en investigaciones posteriores?

Revisión

Como muestra este panorama, el diseño de investigación comprende un conjunto de decisiones sobre *qué tema estudiar, en qué población, con qué métodos de investigación y con qué objetivo*. Mientras que las secciones anteriores sobre los propósitos de la investigación, las unidades de análisis y los puntos de interés se centraban en ampliar sus perspectivas respecto de todos estos aspectos, el diseño de investigación es el proceso de concentrar, de estrechar su punto de vista a los propósitos del estudio.

Si usted realiza un proyecto de investigación para uno de sus cursos, quizá le adelantaron muchos aspectos del diseño. Si debe hacer un proyecto para un curso de métodos experimentales, le habrán especificado el método de investigación. Si el proyecto es para un curso de comportamiento electoral, el tema de la investigación se especificará de alguna manera. Como no está a mi alcance prever todas esas restricciones, los siguientes párrafos asumirán que no hay ninguna.

Al diseñar un proyecto de investigación, descubrirá que es útil comenzar evaluando tres cosas: sus intereses, sus capacidades y los recursos disponibles. Cada consideración le sugerirá muchos estudios posibles.

Simule el comienzo de un proyecto de investigación convencional: pregúntese qué le interesa conocer. De seguro tiene varias preguntas sobre el comportamiento y las actitudes sociales. ¿Por qué algunas personas tienen opiniones políticas liberales y otras conservadoras? ¿Por qué algunos son más religiosos que otros? ¿Por qué se une la gente a grupos militares? ¿Aún discriminan las universidades a los catedráticos pertenecientes a grupos minoritarios? ¿Por qué debe conservar una mujer una relación de maltrato? Deténgase un momento y medite en las preguntas que le interesan y preocupan.

Cuando tenga unas cuantas preguntas que a usted mismo le interese responder, reflexione en la información necesaria para contestarlas. ¿Qué unidades de análisis le darán la información más útil: estudiantes universitarios, empresas, votan-

tes, ciudades o qué? En sus reflexiones, esta pregunta será inseparable de la del tema de la investigación. Entonces, decida *qué aspectos* de las unidades de análisis ofrecerán la información necesaria para responder la pregunta de la investigación.

Ya que cuente con algunas ideas sobre la información que conviene a sus propósitos, pregúntese cómo hará para conseguirla. ¿Existe la probabilidad de que los datos relevantes ya se encuentren en alguna parte (digamos, en una publicación gubernamental) o deberá reunirlos usted mismo? Si cree que tendrá que reunirlos, ¿cómo va a hacerlo? ¿Necesita aplicar una encuesta a un gran número de personas o entrevistar a unas cuantas? ¿Puede saber lo que necesita asistiendo a las reuniones de ciertos grupos? ¿Puede recopilar los datos de libros de la biblioteca?

A medida que responda estas preguntas, se encontrará inmerso en el diseño de la investigación. Tenga siempre presente sus capacidades de investigador y los recursos disponibles. No diseñe un estudio perfecto que no pueda llevar a cabo. Si lo desea, ensaye un método de investigación que no haya usado antes para que aprenda mucho más, pero no se ponga en grandes desventajas.

Una vez que tenga una idea general de lo que quiere estudiar y la manera de hacerlo, repase cuidadosamente otras investigaciones en publicaciones y libros para saber cómo han abordado el tema otros investigadores y qué descubrieron. Un repaso de la bibliografía puede llevarlo a replantear su diseño de investigación: tal vez decida utilizar el método de algún investigador anterior o incluso repetir un estudio previo. La repetición independiente de los proyectos de investigación es un procedimiento habitual en las ciencias físicas, y es igualmente importante en las ciencias sociales, aunque estos científicos tienden a pasarlo por alto. O, si quiere, haga algo más que repetir y estudie algún aspecto del tema que en su opinión los investigadores anteriores descuidaron.

Veamos otro planteamiento que puede adoptar. Supongamos que el tema ha sido estudiado con métodos de investigación de campo. ¿Puede diseñar un experimento que ponga a prueba los descubrimientos de los investigadores anteriores? ¿O conoce estadísticas con las que pueda probar las conclusiones? ¿Arrojó alguna encuesta amplia resultados que le gustaría explorar con mayor detalle mediante observaciones en el sitio o entrevistas

rà una imagen completa del funcionamiento de una clínica local. En el mismo sentido, un antropólogo que sólo estudiara a los hombres en una sociedad donde las mujeres están ocultas a los extraños obtendría un cuadro deformado. Igualmente, si bien serían deseables en ciertos aspectos obvios los informantes occidentalizados con facilidad de expresión en español, no serían característicos de los miembros una sociedad aislada y analfabeta.

Simplemente porque son los que están dispuestos a trabajar con investigadores foráneos, los informantes serán casi siempre un poco "marginales" o atípicos dentro de su grupo. En ocasiones es evidente; sin embargo, en otras únicamente se descubre su marginalidad en el curso de la investigación.

En el estudio de Jeffrey Johnson de los pescadores en Carolina del Norte, el delegado del condado señaló a un pescador que parecía ir en contra de la comunidad. No obstante, cooperó y ayudó a la investigación. Ahora bien, entre más trabajaba Johnson con el pescador, más descubría que era un miembro marginal de la comunidad.

Primero, era un norteno en un pueblo del sur. Segundo, contaba con una pensión de la Marina [por lo que los demás no lo veían como un "pescador serio"]... Tercero, era un gran activista republicano en una población principalmente demócrata. Por último, guardaba su embarcación en un fondeadero alejado del puerto común.

(1990:56)

La marginalidad de los informantes no sólo sesga la imagen que uno se forma, sino que también puede limitar su propio acceso (y por ende el de uno) a los sectores de la comunidad que se desea estudiar.

Estos comentarios deben darle alguna idea de las cuestiones que pueden despertar alguna preocupación sobre el muestreo no probabilístico que se suele emplear en los proyectos de investigación cualitativa. Concluyo con la siguiente advertencia (Lofland y Lofland, 1995:16):

Su meta general es recopilar los *datos más ricos posibles*. En términos ideales, "datos ricos" significa un conjunto amplio y diverso de información reunida durante un periodo relativamente largo. También idealmente, esto se consigue mediante el contacto directo y una prolongada inmersión en alguna ubicación o circunstancia social. [cursivas en el original]

Cambiamos ahora nuestra atención para ver el muestreo en las encuestas de gran escala, destinadas a entregar descripciones precisas y estadísticas de poblaciones numerosas. A veces queremos conocer el porcentaje de la población desempleada, de la que piensa votar por el candidato X o de la que opina que las víctimas de una violación deben tener el derecho de abortar. Estas tareas se cumplen con la lógica y las técnicas del muestreo probabilístico.

La lógica del muestreo probabilístico

Si todos los miembros de una población fueran idénticos en todos los aspectos —todas las características demográficas, las actitudes, experiencias, conductas etc.—, no habría necesidad de procedimientos cuidadosos de muestreo. En tal circunstancia, cualquier muestra bastaría. De hecho, en este caso extremo de homogeneidad un sujeto sería suficiente como muestra para estudiar las características de toda la población.

Por supuesto, las personas que componen cualquier población real son muy heterogéneas, varían de muchas maneras. La figura 8.1 muestra una ilustración simplificada de una población heterogénea: el sexo y la raza difieren entre los 100 miembros de esta pequeña población. A lo largo del capítulo nos valdremos de esta micropoblación hipotética para ejemplificar varios aspectos del muestreo.

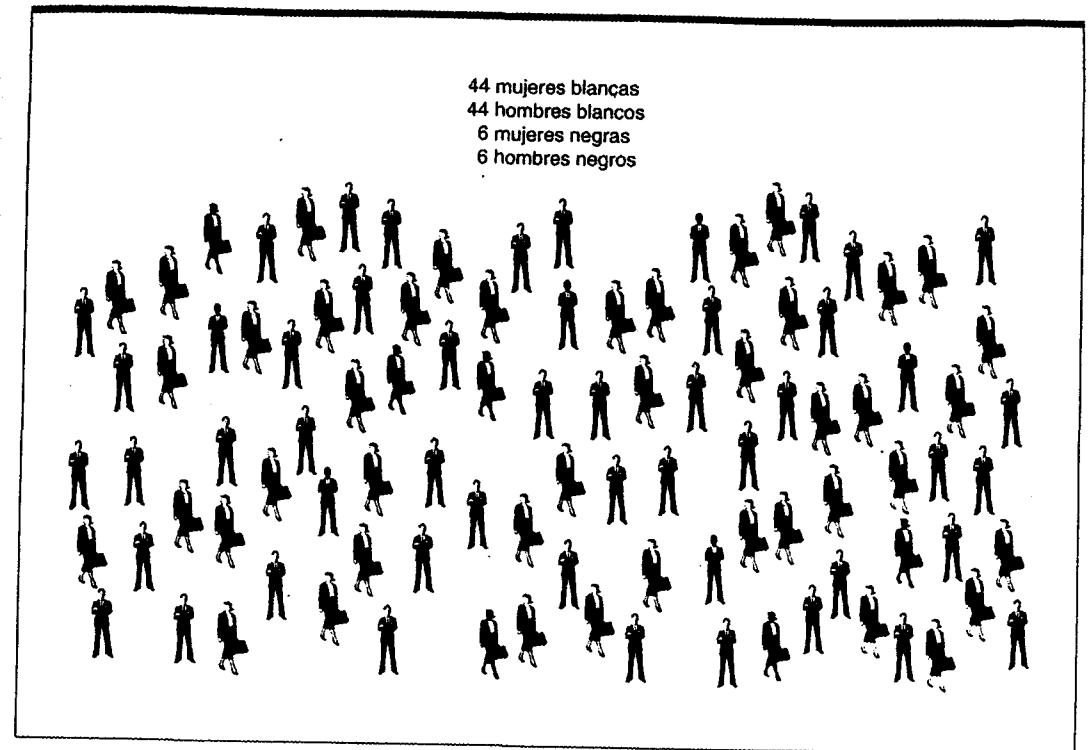
Para ofrecer descripciones útiles del total de la población, una muestra de sus individuos debe contener en esencia las mismas variaciones que aquélla, y lograr esto no es tan fácil como parece.

Dediquemos un momento a considerar algunas de las formas en que se extravían los investigadores. Después, veremos que el muestreo probabilístico brinda un método eficaz para elegir una muestra que refleje bien las variaciones de la población.

Sesgos conscientes e inconscientes en el muestreo

A primera vista, parecería que el muestreo es un asunto muy sencillo. Para elegir una muestra de un centenar de estudiantes universitarios uno podría ir al campus y entrevistar a los primeros 100 alumnos que pasearan por el lugar. Los investigadores sin capacitación emplean a menudo este método de muestreo, pero tiene graves problemas.

Figura 8.1
Población de 100 individuos



La figura 8.2 ilustra lo que ocurre cuando uno elige a las personas del estudio por conveniencia. Aunque las mujeres suman nada más 50 por ciento de nuestra micropoblación, en el grupo más cercano al investigador (esquina superior derecha) resulta que hay 70 por ciento de mujeres; además, aunque 12 por ciento de la población es de negros, no se eligió a ninguno para la muestra.

Hay otros problemas además de los riesgos inherentes de estudiar la muestra por conveniencia. Para empezar, sus propias inclinaciones personales pueden influir en ella al grado de que no represente verdaderamente a la población que investiga. Supongamos que en su estudio de los universitarios usted se siente intimidado por los alumnos de aspecto más "duro" y piensa que acaso ridiculicen su investigación; entonces, tal vez —de manera consciente o inconsciente— no los entreviste. O bien le parece que las actitudes de los estudiantes "pe-

ripuestos" serían irrelevantes para los propósitos de su investigación y tampoco los interroga.

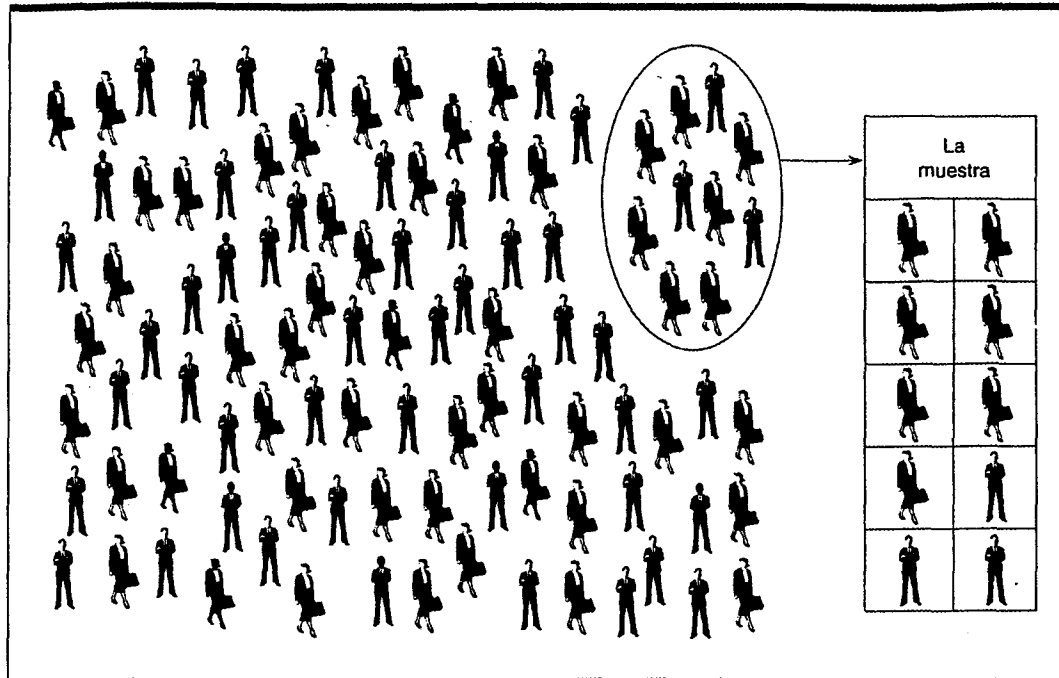
Aunque trate de entrevistar a un grupo "equilibrado" de estudiantes, no conocería las proporciones exactas de los tipos de alumnos que componen tal equilibrio ni sería siempre capaz de identificar los tipos con sólo verlos pasar.

Y aunque hiciera un esfuerzo consciente por entrevistar a cada décimo estudiante que entrara a la biblioteca, no estaría seguro de tener una muestra *representativa*, puesto que distintas clases de alumnos visitan la biblioteca con frecuencias diversas. En su muestra estarían sobrerrepresentados los que van más a menudo.

Cuando hablamos de un sesgo en el muestreo queremos decir, simplemente, que los elegidos no son "característicos" o "representativos" de las poblaciones mayores de donde se tomaron. Este tipo de sesgo es prácticamente inevitable cuando uno elige a la gente de manera aleatoria.

Figura 8.2

Muestra de conveniencia: fácil pero no representativa



Del mismo modo, no se puede confiar en que los "sondeos telefónicos de opinión pública"—en los que las estaciones de radio y televisión le piden al auditorio que marque cierto número para registrar su opinión— representen poblaciones generales. Para empezar, no todos los miembros de la población estarán al tanto del sondeo. Este problema también invalida los sondeos de revistas y periódicos que publican cupones para que los lectores los llenen y envíen por correo. Incluso entre quienes tienen noticia de tales sondeos, no todos expresarán su opinión, especialmente si hacerlo les costará una estampilla, un sobre o un cargo telefónico.

Irónicamente, Philip Perinelli (1986), jefe de personal del servicio DIAL-IT 900 de AT&T Communications (que ofrece a las empresas un sistema de sondeo telefónico), sin darse cuenta reconoció la incapacidad de estos sondeos para representar todas las opiniones por igual. Perinelli quiso responder a las críticas argumentando que "el cargo de 50 centavos de dólar asegura que respondan sólo los

interesados y también que nadie 'rellene' la urna". No podemos determinar la opinión pública general si consideramos "sólo a los interesados". Esto excluye a quienes no creen que valga la pena gastar los 50 centavos, así como a los que saben que tales sondeos no son válidos. Ambos grupos podrían tener una opinión e incluso votar el día de las elecciones. La afirmación de Perinelli de que el cargo de 50 centavos impedirá que se llene la urna significa exactamente que sólo los ricos pueden darse el lujo de practicar este ejercicio.

Las posibilidades de sesgos inadvertidos en el muestreo son infinitas y no siempre obvias. Por fortuna, hay técnicas que nos permiten evitarlas.

Representatividad y probabilidad de la selección

Aunque el término *representatividad* no tiene un significado científico preciso, tiene un significado de sentido común útil. Para nuestros propósitos,

una muestra será representativa de la población de la que fue tomada si la suma de sus características se aproxima al conjunto de características de la población (las muestras no tienen que ser representativas en todos los aspectos; la representatividad se limita a las características importantes para los intereses reales del estudio, aunque al principio uno no sepa cuáles son las que importan). Por ejemplo, si la población tiene 50 por ciento de mujeres, una muestra representativa contendría también "cerca de" 50 por ciento de mujeres. Posteriormente explicaremos en detalle "cuán cerca".

Un principio básico del muestreo probabilístico es que una muestra será representativa de la población de la que proviene si todos los miembros de esta última tienen la misma probabilidad de ser elegidos (dentro de poco veremos que el tamaño de la muestra influye en el grado de representatividad). Las muestras que poseen esta cualidad suelen llevar las iniciales MESEPI (método de selección de probabilidad igual). Más adelante expondremos las variaciones de este principio, que forma la base del muestreo probabilístico.

Más allá de este principio básico, debemos advertir que las muestras—incluso las muestras MESEPI elegidas con sumo cuidado— rara vez o nunca representan *perfectamente* a las poblaciones de las que se extraen. No obstante, el muestreo probabilístico tiene dos ventajas particulares.

Primera, las muestras probabilísticas, aunque no sean perfectas, en general son *más representativas* que otras debido a que evitan los sesgos que acabamos de explicar. En la práctica, es más probable que una muestra probabilística sea representativa de su población que una no probabilística.

Segunda, y más importante, la teoría de la probabilidad nos permite estimar la exactitud o representatividad de la muestra. Cabe pensar que un investigador mal informado pueda, por medios totalmente azarosos, elegir una muestra que represente casi a la perfección a la población mayor. Sin embargo, las probabilidades están en contra, y no seríamos capaces de estimar la verosimilitud de que el investigador haya alcanzado la representatividad. En cambio, quien tome una muestra probabilística puede hacer un estimado preciso de su éxito o fracaso.

Después de un breve glosario de la terminología del muestreo, examinaremos los medios del inves-

tigador para estimar la representatividad de su muestra probabilística.

Conceptos y terminología del muestreo

El siguiente examen de la teoría y la práctica del muestreo usa muchos términos técnicos cuya definición doy rápidamente aquí. En su mayor parte, empleo los términos utilizados en los libros de muestreo y estadística para que usted comprenda mejor esas otras fuentes.

Al presentar este glosario quisiera reconocer mi deuda con Leslie Kish y su excelente libro *Survey Sampling*. Aunque he modificado algunas de las convenciones que sigue Kish, su exposición es con mucho la referencia más importante de esta sección.

Elemento Un *elemento* es una unidad de la que se recopila información y que brinda la base para el análisis. Habitualmente, en las encuestas los elementos son las personas, o ciertas clases de personas; sin embargo, otras unidades también constituyen los elementos para la investigación social: familias, clubes sociales o empresas pueden ser los elementos de un estudio. (Nota: Elementos y unidades de análisis suelen ser los mismos en un estudio, aunque los primeros se refieran a la selección de la muestra y las segundas remitan al análisis de datos.)

Población Una *población* es la suma—especificada por una teoría— de los elementos de estudio. Mientras que el término vago *portugueses* puede ser el objeto de un estudio, la descripción de la población comprendería la definición del elemento *portugueses* (por ejemplo, ciudadanía, residencia) y el referente temporal (¿portugueses desde cuándo?). Para traducir el abstracto *meridenses adultos* en una población manejable se requeriría una especificación de la edad que definiera *adultos* y los límites de *Mérida*. Para especificar el término *estudiante universitario* se incluiría una consideración de los estudiantes de tiempo completo y medio tiempo, los candidatos a posgrados y los pasantes de licenciatura, los estudiantes de licenciatura y de posgrado, etcétera.

Aunque los investigadores deben comenzar con una especificación cuidadosa de su población, las licencias poéticas los autorizan a redactar sus in-

formes en términos de un universo hipotético. Para facilitar la presentación, aun los investigadores más concienzudos acostumbran hablar de, digamos, "nigerianos" y no de "ciudadanos residentes en la república de Nigeria desde el 12 de noviembre de 1999". La norma principal en este punto, como en los demás, es que uno no debe confundir ni engañar a los lectores.

Población de estudio Una *población de estudio* es la suma de los elementos de los que se eligió la muestra. Como asunto práctico, rara vez se encontrará en la posición de garantizar que todos los elementos que satisfagan las definiciones teóricas asentadas tengan realmente la probabilidad de ser elegidos para la muestra. Incluso aunque existan listas de elementos para propósitos de muestreo, éstas suelen estar incompletas de alguna manera. Siempre faltan, sin querer, algunos estudiantes en los listados. Algunos suscriptores telefónicos solicitan que sus nombres y números se guarden confidencialmente.

A menudo, los investigadores deciden limitar más sus poblaciones de estudio que lo que indican los ejemplos anteriores. Las empresas de sondeos pueden limitar sus muestras nacionales y excluir a las poblaciones más remotas por razones obvias. El investigador que desea tomar una muestra de profesores de psicología puede limitar su estudio a quienes ejercen en el departamento de la materia y omitir a los que trabajan en otros departamentos (en cierto sentido, diríamos que estos investigadores han redefinido sus universos y poblaciones, en cuyo caso deben aclarar estas revisiones a sus lectores).

Unidad de muestreo Una *unidad de muestreo* es aquel elemento o conjunto de elementos cuya elección se considera en alguna etapa del muestreo. En una muestra de una sola etapa, las unidades de muestreo son las mismas que los elementos y probablemente también sean las unidades de análisis. En cambio, en las muestras más complicadas se aplican varios niveles de unidades de muestreo. Por ejemplo, digamos que usted toma del censo una muestra de manzanas de una ciudad, luego una muestra de casas de las manzanas elegidas y por último una muestra de adultos de las casas seleccionadas. En cada etapa, las unidades de muestreo son las manzanas del censo, las casas y los adultos, pero sólo estos últimos son elementos. En

concreto, las frases *unidades de muestreo primario*, *unidades de muestreo secundario* y *unidades de muestreo final* designan las etapas sucesivas.

Marco de muestreo Un *marco de muestreo* es la lista concreta de unidades de muestreo de la que se elige la muestra, o una de sus etapas. En los diseños de muestreo de una sola etapa, el marco de muestreo no es más que la lista de la población de estudio, que ya definimos. Si se elige una muestra simple de estudiantes del listado, éste es el marco de muestreo. Si la unidad de muestreo primario para una muestra compleja de la población es la manzana del censo, la lista de manzanas del censo compone el marco de muestreo (en la forma de un folleto impreso, una cinta magnética de almacenamiento u otro registro computarizado).

Como dijimos, en el diseño de muestras de una sola etapa el marco de muestreo es la lista de elementos que componen la población de estudio. En la práctica, los marcos de muestreo suelen definir a la población de estudio, y no al contrario. A menudo comenzamos con la idea de una población para nuestro estudio y después buscamos los marcos de muestreo viables. Examinamos y evaluamos estos marcos y decidimos cuál presenta una población de estudio más adecuada a nuestras necesidades.

Aunque la relación entre poblaciones y marcos de muestreo es crucial, no ha recibido suficiente atención. En una sección posterior continuaremos el tema con mayores detalles.

Unidad de observación Una *unidad de observación*, o unidad de recopilación de datos, es un elemento o grupo de elementos del que se reúne la información. Aquí también, la unidad de análisis y la de observación suelen ser la misma —el individuo—, pero no tiene que serlo por fuerza. Así, un investigador podría investigar a los jefes de familias (las unidades de observación) para conseguir información sobre los miembros de la casa (las unidades de análisis).

Nuestra tarea se simplifica cuando las unidades de análisis y de observación son las mismas. Sin embargo, muchas veces no es posible o viable, y, en tales situaciones, necesitamos un poco de ingenio para recopilar los datos importantes para nuestras unidades de análisis sin observarlas directamente.

Variable Como ya hemos visto, una *variable* es un conjunto de atributos mutuamente excluyentes: género, edad, ocupación, etc. Es posible describir los elementos de una población por sus atributos en determinada *-variable*. La investigación social aspira a describir la distribución de atributos que componen una variable en una población. Así, un investigador puede describir la distribución de edades de una población examinando la frecuencia relativa de las distintas edades de los miembros.

Por definición, las variables deben *variar*; si todos los elementos de la población poseen el mismo atributo, éste es una *constante* suya, y no parte de una variable.

Parámetro Un *parámetro* es la descripción resumida de cierta variable en una población. Son parámetros el ingreso medio de todas las familias de una ciudad y la distribución de edades de sus habitantes. Una parte importante de la investigación social atañe a la estimación de los parámetros poblacionales a partir de observaciones de la muestra.

Estadísticos Los *estadísticos* son descripciones resumidas de cierta variable de la muestra. Así, el ingreso medio calculado de una muestra y la distribución de edades de ésta son estadísticos. Los estadísticos de las muestras sirven para hacer estimaciones de los parámetros de la población.

Error de muestreo Los métodos de muestreo probabilístico rara vez, si acaso, dan estadísticos exactamente iguales a los parámetros que estiman. Sin embargo, la teoría de la probabilidad nos permite estimar el grado de error esperado en determinado diseño de muestra. Más adelante profundizaremos en el *error de muestreo*.

Niveles e intervalos de confianza Los dos componentes claves de las estimaciones de los errores de muestreo son los *niveles de confianza* y los *intervalos de confianza*. Expresamos la exactitud de los estadísticos de nuestra muestra en términos de un nivel de confianza de que los valores caen dentro de un intervalo especificado del parámetro. Por ejemplo, podríamos decir que tenemos un 95 por ciento de confianza de que nuestros estadísticos (digamos, 50 por ciento en favor del candidato X) están dentro de más o menos cinco puntos porcentuales del parámetro de la población. A medida que el intervalo de confianza para un estadístico determinado se amplía, aumenta nuestra confianza y podríamos decir que tenemos un 99.9 por ciento de

confianza en que nuestro estadístico se encuentra a 7.5 puntos porcentuales del parámetro. En la siguiente sección explicaremos la forma de calcular los intervalos y niveles del muestreo, y aclararemos más estos dos conceptos.

Teoría del muestreo probabilístico y distribución muestral o de muestreo

Con las definiciones anteriores podemos ahora examinar la teoría básica del muestreo probabilístico en lo que atañe a la investigación social. También consideraremos la lógica de la distribución y el error de muestreo con respecto a una *variable binominal*, es decir, una variable compuesta por dos atributos.

Teoría del muestreo probabilístico

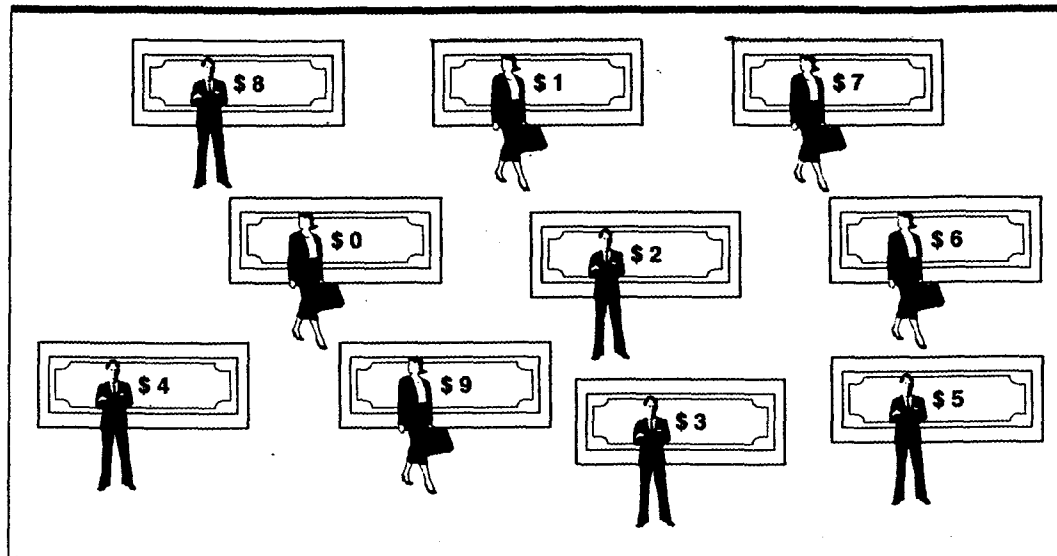
El propósito final del muestreo es elegir un conjunto de elementos de una población de modo tal que la descripción de dichos elementos (estadísticos) represente en forma precisa los parámetros de esa población total de la que fueron tomados. El muestreo probabilístico aumenta la probabilidad de alcanzar este objetivo y también proporciona los métodos para calcular el grado probable de éxito.

La *selección aleatoria* es la clave de este proceso. Aquí cada elemento tiene la misma probabilidad de selección independientemente de cualquier otro suceso en el proceso. El ejemplo más citado es el de arrojar una moneda perfecta: la "selección" de una cara o una cruz es independiente de las selecciones anteriores de una u otra. Otro ejemplo es el rodar de unos dados perfectos. Con todo, estas imágenes de selección aleatoria casi nunca se aplican en forma directa a los métodos de muestreo de la investigación social. Generalmente, el investigador social usa tablas de números aleatorios o programas de computadora que proporcionan una selección aleatoria de unidades de muestreo. En el capítulo 10, que trata de las encuestas, veremos que se emplean computadoras para seleccionar aleatoriamente números telefónicos y realizar entrevistas (el llamado *marcado de dígitos aleatorios*).

Las razones para aplicar los métodos de selección aleatoria —tablas de números aleatorios o programas de computadora— son dos. Primera, el

Figura 8.3

Población de 10 personas con cero a nueve dólares



procedimiento hace las veces de verificación para evitar los sesgos conscientes o inconscientes del investigador. El estudioso que escoge sus casos en forma intuitiva bien puede elegir los que sustenten sus expectativas o sus hipótesis. La selección aleatoria elimina este peligro. Segunda, y más importante, la selección aleatoria brinda acceso al cuerpo de la teoría de la probabilidad, que es la base para las estimaciones de los errores y los parámetros poblacionales. Examinaremos ahora con más detalle la teoría de la probabilidad.

La distribución muestral de 10 casos

Para hacer la introducción a las estadísticas del muestreo probabilístico, comencemos con un ejemplo sencillo de sólo 10 casos.* Supongamos que hay un grupo de 10 personas y que cada una tiene cierta suma en su bolsillo. Para simplificar, digamos que una persona no tiene dinero, otra tiene un dólar, otra dos, etc., hasta llegar a la décima, que tiene nueve dólares. La figura 8.3 presenta esta población de 10 personas.

*Quiero dar las gracias a Hanan Selvin por sugerirme este método de presentar el muestreo probabilístico.

Nuestra tarea es determinar la suma promedio de dinero que tiene una persona: en concreto, la cantidad media de dólares. Si sumamos las cantidades de la figura 8.3, el resultado es 45 dólares, así que la media es 4.5. Nuestro objetivo en el resto del ejercicio es estimar esta media sin observar realmente a los 10 individuos. Para ello tomamos de la población muestras aleatorias y estimamos con las medias de éstas la media de toda la población.

Para empezar, supongamos que vamos a elegir —de manera aleatoria— una muestra de una sola persona de las 10. Dependiendo de la persona que escojamos, estimaríamos la media del grupo en cualquier cantidad entre cero y nueve dólares. La figura 8.3 exhibe las 10 muestras posibles.

Los 10 puntos de la gráfica que se muestra en la figura 8.4 representan las 10 medias "muestrales" que obtendríamos como estimados de la población. El ordenamiento de los puntos se denomina *distribución muestral*. Como es obvio, no sería una muy buena idea elegir una muestra de sólo uno, puesto que tendríamos muchas probabilidades de errar la media verdadera, de 4.5.

¿Pero qué pasa si tomamos medias de dos individuos? Como se aprecia en la figura 8.5, aumen-

Figura 8.4

Distribución muestral de muestras de tamaño 1

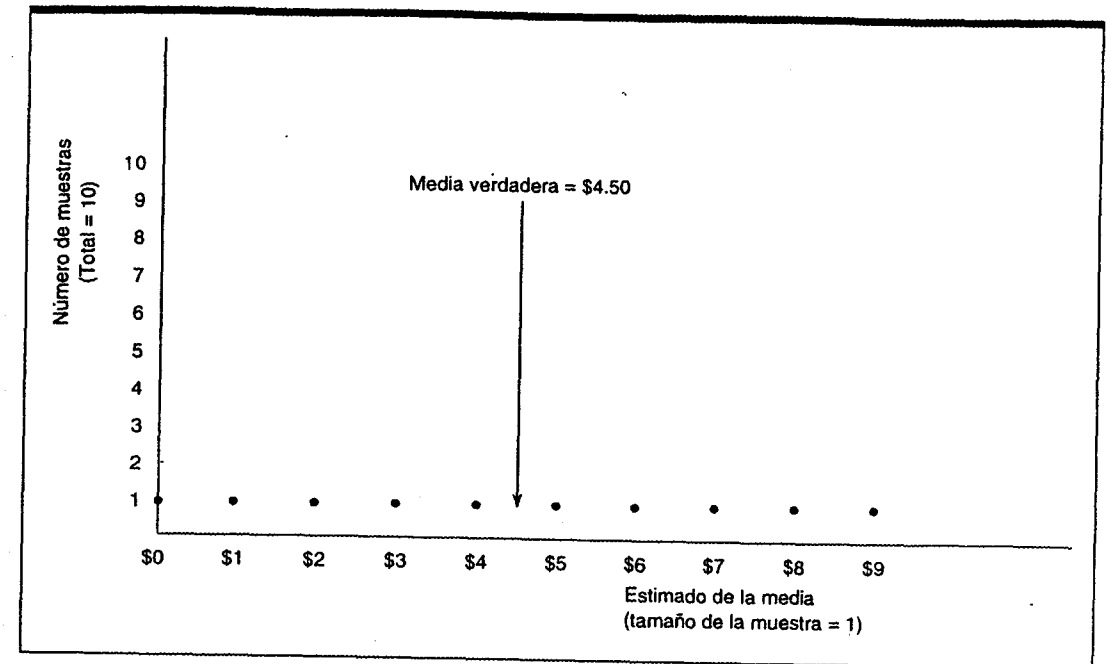


Figura 8.5

Distribución muestral de muestras de dos individuos

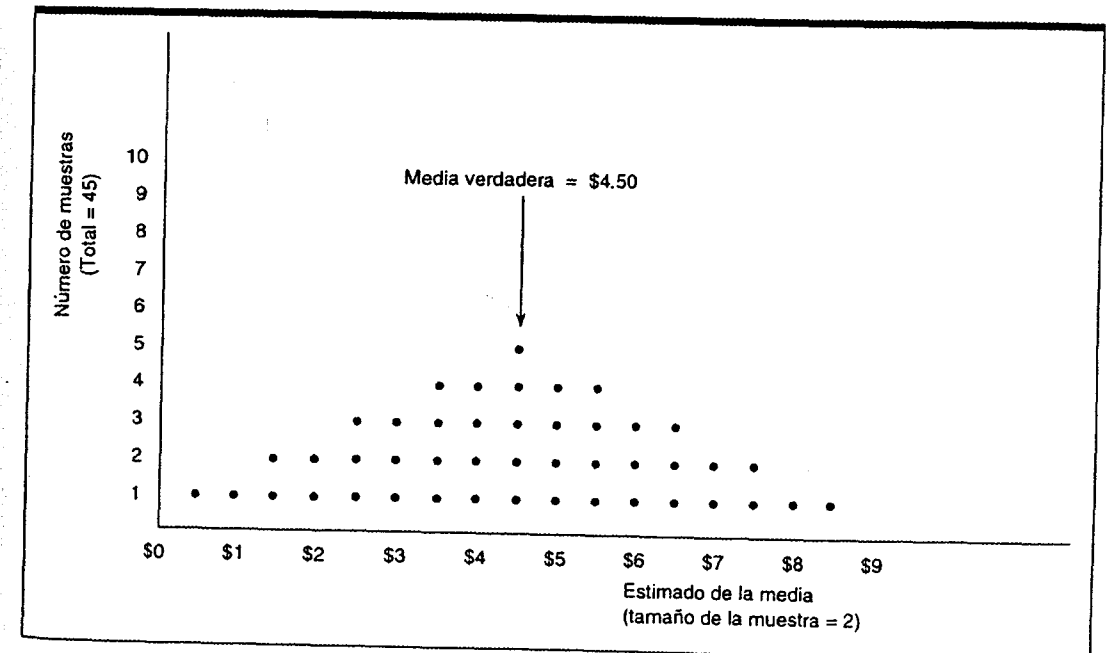


Figura 8.6

Distribución muestral de muestras de tamaño tres, cuatro, cinco y seis individuos

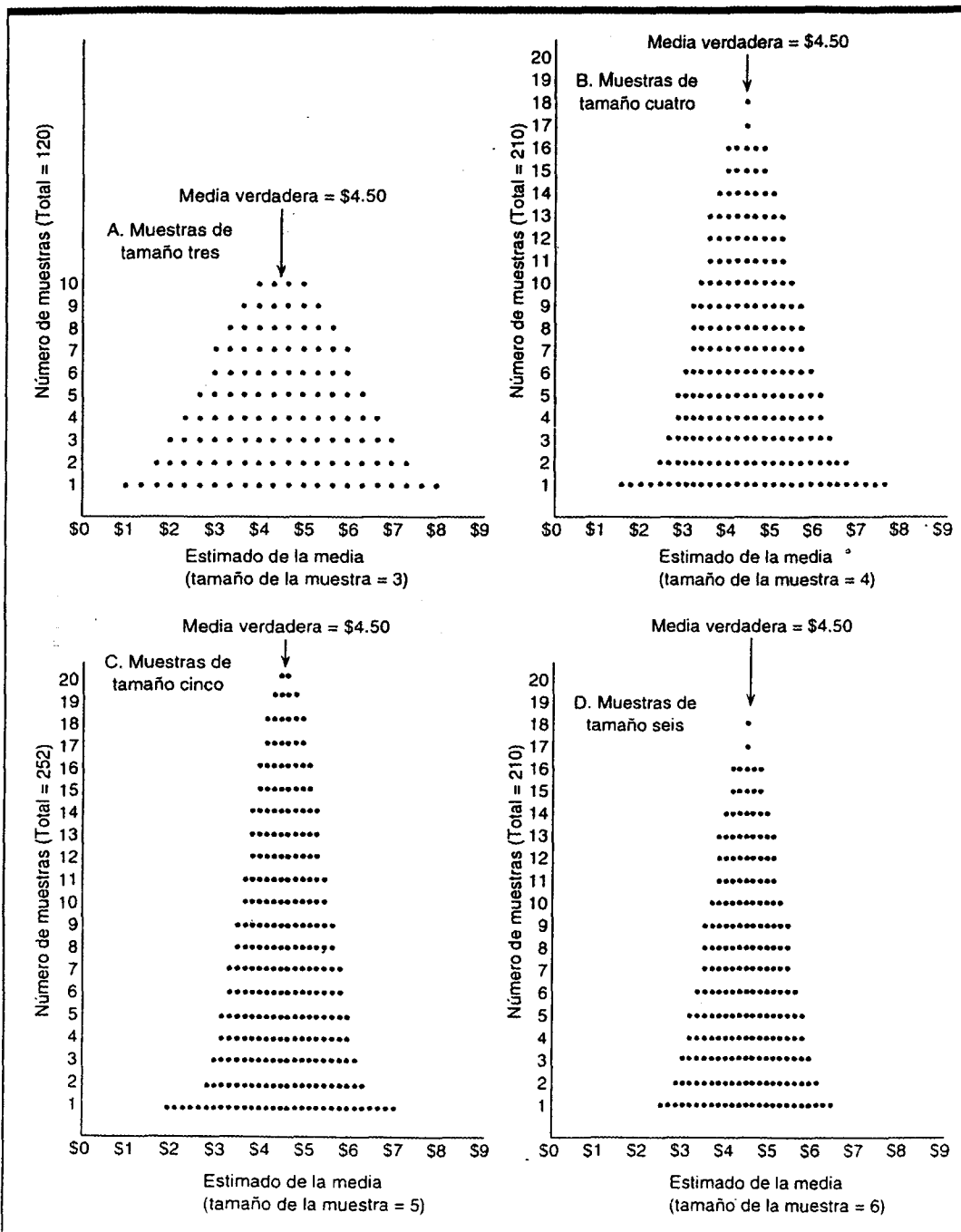
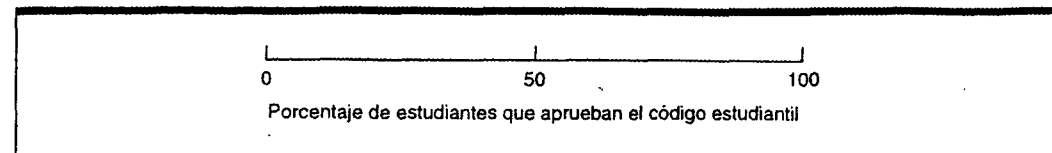


Figura 8.7

Intervalo de los posibles resultados del estudio de las muestras



tar el tamaño de la muestra mejora nuestras estimaciones. Ahora hay 45 muestras posibles: [0,1], [0,2], ..., [7,8], [8,9]. Más aún, algunas de estas muestras arrojan la misma media. Por ejemplo, [0,6], [1,5] y [2,4] dan una media de 3. En la figura 8.5, los tres puntos verticales sobre la media de tres dólares representan estas tres muestras.

También se observa que las 45 medias de las muestras no están distribuidas de manera uniforme, sino que se agrupan alrededor del valor verdadero de 4.5. Sólo dos muestras se desvían tanto como cuatro dólares del valor real [0,1] y [8,9], mientras que cinco muestras dan el estimado verdadero de 4.5; otras cinco muestras yerran la marca por apenas 50 centavos (de más o de menos).

Ahora supongamos que elegimos muestras mayores. ¿Qué cree que ocurrirá con nuestros estimados de la media? La figura 8.6 presenta las distribuciones de muestras de tres, cuatro, cinco y seis individuos.

La progresión de las distribuciones de muestreo es obvia. Cada vez que aumenta el tamaño de la muestra mejora la distribución de los estimados de la media. Desde luego, el caso límite del procedimiento es elegir una muestra de 10. Sólo habría una muestra posible (de todos los individuos) y nos daría la media verdadera, de 4.5.

Distribución de muestreo binomial

Vayamos ahora a una situación de muestreo más realista para ver en la práctica la noción de distribución de muestreo. Tomemos un ejemplo sencillo con una población mucho mayor que 10 individuos. Supongamos por ahora que deseamos investigar a la población de alumnos de la Universidad Estatal (UE) para determinar si aprueban o desaprueban un código de conducta estudiantil propuesto por la dirección. La población estudiantil será el conjunto de, digamos, los 20 000 individuos

del listado de alumnos: el marco de muestreo. Los elementos serán los estudiantes de la UE. La variable considerada la conformarán las actitudes hacia el código, es decir, una variable binomial: aprobación o desaprobación. Elegiremos una muestra aleatoria de, digamos, 100 estudiantes, con el objetivo de estimar la postura de todo el alumnado.

El eje horizontal de la figura 8.7 presenta todos los valores posibles de este parámetro, de cero a 100 por ciento de aprobación. El punto medio del eje —50 por ciento— representa la situación en que la mitad de los estudiantes aprueba el código y la otra mitad lo rechaza.

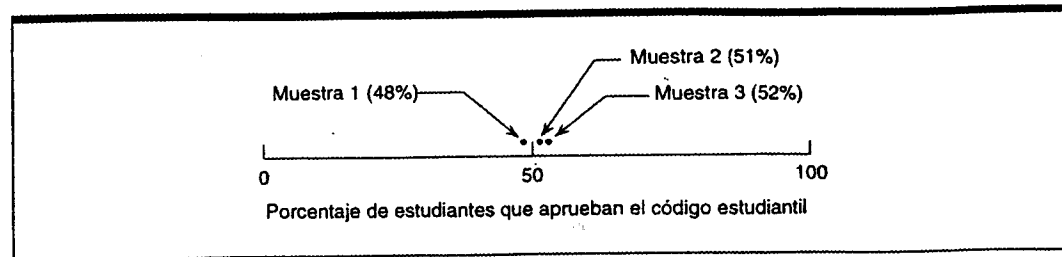
Para elegir nuestra muestra asignamos a cada estudiante del listado un número y seleccionamos 100 números aleatorios de una tabla. Entonces, entrevistamos a los estudiantes cuyos números elegimos y les preguntamos sus actitudes hacia el código: lo aprueban o lo desaprueban. Supongamos que esta operación nos da 48 estudiantes que aprueban el código y 52 que lo desaprueban. Asentamos este estadístico colocando un punto sobre el eje de las x en el sitio correspondiente al 48 por ciento.

Ahora supongamos que elegimos otra muestra de 100 estudiantes exactamente de la misma manera y medimos su aprobación o desaprobación del código. Aquí, quizá 51 estudiantes lo aprobaron, así que colocamos un punto en el lugar apropiado del eje de las x . Si repetimos este proceso una vez más, descubriremos, digamos, que 52 estudiantes de la tercera muestra aprobaron el código.

La figura 8.8 muestra los tres estadísticos que representan los porcentajes de los estudiantes de cada muestra aleatoria que aprobaron el código. La regla básica del muestreo aleatorio es que tales muestras dan estimados aproximados del parámetro que es propio de toda la población. Así, las muestras aleatorias dan un estimado del porcentaje de estudiantes del conjunto del alumnado que

Figura 8.8

Resultados de los tres estudios hipotéticos



aprueba el código. Pero por desgracia escogimos tres muestras y ahora tenemos tres estimados distintos.

Para salir del problema vamos a tomar más muestras de 100 estudiantes, a pedirle a cada uno su aprobación o desaprobación y a trazar los nuevos estadísticos en nuestra gráfica de resumen. Al tomar muchas muestras, descubrimos que algunas arrojan estimados duplicados, como en el ejemplo de los 10 casos. La figura 8.9 muestra la distribución de muestreo de, digamos, cientos de muestras. El resultado es la llamada *curva normal*.

Observe que, al incrementar el número de muestras elegidas y entrevistadas, también ampliamos el intervalo de estimados que arroja la operación de muestreo. En cierto sentido, aumentamos nuestro dilema de tratar de conjeturar el parámetro de la población. Sin embargo, la teoría de la probabilidad establece ciertas reglas importantes que atañen a la distribución de muestreo presentada en la figura 8.9.

Primera, si se toman de una población muchas muestras aleatorias independientes, las estadísticas que producen *estarán distribuidas alrededor del parámetro de la población* en una forma conocida. Así, aunque la figura 8.9 muestra un intervalo amplio de estimados, la mayoría se encuentra en las cercanías del 50 por ciento y no en otra parte de la gráfica. Por tanto, la teoría de la probabilidad nos dice que el valor verdadero se encuentra alrededor del 50 por ciento.

Segunda, la teoría de la probabilidad nos da una fórmula para calcular *cuán cerca* están los estadísticos del valor verdadero. La fórmula contiene tres factores: el parámetro, el tamaño de la muestra y el

error estándar (una medida del error de muestreo):

$$e = \sqrt{\frac{P \times Q}{n}}$$

Los símbolos P y Q de la fórmula corresponden a los parámetros de la población de la variable binomial: si 60 por ciento del alumnado aprobó el código estudiantil y 40 por ciento lo rechazó, P y Q son, respectivamente, 60 y 40 por ciento, o .6 y .4. Observe que $Q = 1 - P$ y que $P = 1 - Q$. El símbolo n es el número de casos en cada muestra y e es el error estándar.

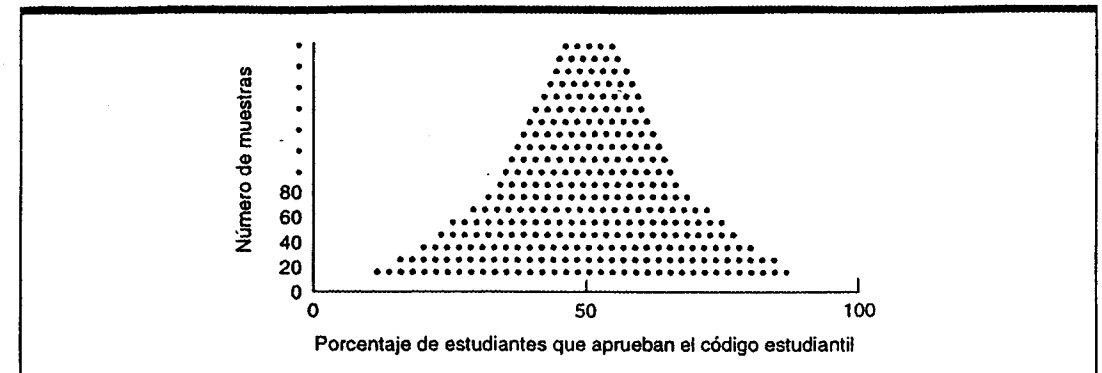
Supongamos que el parámetro de la población en el ejemplo de los estudiantes es de 50 por ciento de aprobación del código y 50 por ciento de rechazo. Recuerde que elegimos muestras de 100 casos cada una. Cuando estas cifras se sustituyen en la fórmula, descubrimos que el error estándar es de .05, o cinco por ciento.

En la teoría de la probabilidad, el error estándar es un dato valioso, pues indica el grado en que los estimados de la muestra se distribuirán alrededor del parámetro de la población. Si usted está familiarizado con la *desviación estándar* en estadística, habrá reconocido que el error estándar, en este caso, es la desviación estándar de la distribución de muestreo.

En concreto, la teoría de la probabilidad indica que cierta proporción de los estimados de las muestras se situarán dentro de incrementos específicos—cada uno igual a un error estándar—del parámetro de la población. Aproximadamente 34 por ciento (.3413) de los estimados caerán dentro de un error estándar por arriba del parámetro y

Figura 8.9

Distribución de muestreo



otro 34 por ciento caerán dentro de un error estándar por debajo. En nuestro ejemplo, el incremento del error estándar es de cinco por ciento, así que sabemos que 34 por ciento de nuestras muestras darán estimados de la aprobación estudiantil entre 50 por ciento (el parámetro) y 55 (un error estándar por arriba); otro 34 por ciento de las muestras dará estimados entre 50 y 45 por ciento (un error estándar por abajo del parámetro). Por tanto, tomados en conjunto, sabemos que alrededor de dos tercios (68 por ciento) de las muestras darán estimados dentro de más o menos cinco por ciento del parámetro.

Más aún, la teoría de la probabilidad dicta que aproximadamente 95 por ciento de las muestras caerán dentro de más o menos dos errores estándar del valor verdadero y que 99.9 por ciento de las muestras caerán dentro de más o menos tres errores estándar. Así, en nuestro ejemplo, sabemos que sólo una muestra de mil dará un estimado menor a 35 por ciento de apoyo o mayor a 65 por ciento.

La proporción de las muestras que quedan a uno, dos o tres errores estándar del parámetro es constante en cualquier procedimiento de muestreo aleatorio como el que acabamos de describir, siempre que se tome un gran número de muestras. En cualquier caso, el tamaño del error estándar está en función del parámetro de la población y del tamaño de la muestra. Si volvemos un momento a nuestra fórmula, advertimos que el error estándar aumentará en función del incremento en el producto de P por Q . Observe también que este producto

alcanza su máximo en la situación de una división uniforme en la población. Si $P = .5$, $PQ = .25$; si $P = .6$, $PQ = .24$; si $P = .8$, $PQ = .16$; si $P = .99$, $PQ = .0099$. Por extensión, si P es 0.0 o 1.0 (o bien cero o 100 por ciento de aprobación del código estudiantil), el error estándar será igual a cero. Si todos los miembros de la población tienen la misma actitud (ninguna variación), cualquier muestra dará exactamente el mismo estimado.

El error estándar también es una función del tamaño de la muestra, pero una función *inversa*: a medida que aumenta el tamaño de la muestra, disminuye el error estándar. Cuando se incrementa su tamaño, las muestras se agruparán más cerca del valor verdadero. De la fórmula se desprende otro lineamiento general: como ésta es una raíz cuadrada, el error estándar se reduce a la mitad si el tamaño de la muestra *se cuadruplica*. En nuestro ejemplo, si las muestras de 100 individuos producen un error estándar de cinco por ciento, para reducirlo a 2.5 por ciento debemos aumentar el tamaño de la muestra a 400.

Toda esta información proviene de la teoría de la probabilidad, que asume la selección de un gran número de muestras aleatorias (si alguna vez siguió un curso de estadística, sabrá que se trata del "teorema de la tendencia central"). Si conocemos el parámetro de la población y tomamos muchas muestras aleatorias, podemos predecir cuántas de las muestras caerán dentro de los intervalos especificados del parámetro. Tenga presente que esta exposición sólo ilustra la *lógica* del muestreo probabilístico y que no describe la forma en que se

realizan las investigaciones. Habitualmente no conocemos el parámetro: efectuamos una inspección de las muestras para estimar esa cifra. Más aún, en realidad no elegimos muchas muestras: tomamos apenas una. No obstante, este estudio de la teoría de la probabilidad sienta las bases para hacer inferencias sobre la situación de investigación social típica. Saber qué ocurriría si se eligiera miles de muestras nos permite hacer suposiciones sobre la muestra que escogemos y examinamos.

La teoría de la probabilidad especifica que 68 por ciento de ese gran número de muestras ficticias arrojará estimados que caerán dentro de un error estándar del parámetro, mas nosotros le damos la vuelta al razonamiento e inferimos que cualquier muestra aleatoria única tiene 68 por ciento de probabilidades de caer dentro de ese margen. En este sentido hablamos de *niveles de confianza*: tenemos 68 por ciento de confianza de que el estimado de nuestra muestra cae dentro de un error estándar del parámetro. O también decimos que tenemos 95 por ciento de confianza de que el estadístico muestral caerá dentro de dos errores estándar del parámetro, etc. Con fundamento, nuestra confianza aumenta a medida que se amplía el margen de error. Estamos casi seguros (99.9 por ciento) de que nos encontramos a tres errores estándar del valor verdadero.

Aunque tengamos la confianza (en cierto nivel) de estar dentro de determinado margen del parámetro, ya dijimos que rara vez conocemos este parámetro. Para resolver el problema sustituimos en la fórmula nuestro estimado del parámetro; es decir, a falta del valor verdadero, lo remplazamos con nuestra mejor conjetura.

El resultado de estas inferencias y estimaciones es que podemos hacer una estimación aproximada del parámetro de la población así como también el grado esperado de error a partir de una muestra tomada de esa población. Si usted parte de la pregunta "¿qué porcentaje del alumnado aprueba el código estudiantil?", podría elegir una muestra aleatoria de 100 estudiantes y entrevistarlos. Entonces informaría que su mejor estimado es que 50 por ciento del estudiantado aprueba el código y que tiene 95 por ciento de confianza que entre 40 y 60 por ciento (más o menos dos errores estándar) lo aprueban. El margen de 40 a 60 por ciento se denomina *intervalo de confianza* (en el nivel de confianza de 68 por ciento, el intervalo de confianza sería de 45 a 55 por ciento).

La lógica de los niveles y los intervalos de confianza también es la base para determinar el tamaño apropiado de la muestra que se quiere estudiar. Si usted ya decidió el error de muestreo que puede tolerar, será capaz de calcular el número de casos que requiere su muestra. Así, por ejemplo, si quiere tener 95 por ciento de confianza en que los resultados de su estudio serán exactos en el margen de más o menos cinco puntos porcentuales de los parámetros de la población, tomaría una muestra de por lo menos 400 (el apéndice H es una buena guía al respecto).

Ésta es, pues, la lógica básica del muestreo probabilístico. La selección aleatoria le permite al investigador vincular los resultados de una muestra a la teoría de la probabilidad para estimar su exactitud. Todas las afirmaciones de exactitud del muestreo deben especificar tanto un nivel como un intervalo de frecuencia. El investigador debe informar que tiene el x por ciento de confianza en que el parámetro de la población se encuentra entre dos valores específicos.

George Gallup (1984:7) explica de esta manera su error de muestreo en un artículo periodístico sobre un sondeo reciente sobre las actitudes de hijos y padres:

Los resultados de los adultos se basan en entrevistas personales con 1 520 mayores de 18 años realizadas entre el 26 y el 29 de octubre en más de 300 localidades de todo el país elegidas en forma científica. De los resultados que arrojan las muestras de este tamaño, podemos decir con 95 por ciento de confianza que el error atribuible al muestreo y otros efectos aleatorios estaría a tres puntos porcentuales en cada dirección.

O ver lo que el *New York Times* ("How the Poll Was Conducted", 1995:15) expresó acerca de un sondeo realizado sobre opiniones religiosas:

En teoría, en 19 casos de 20 los resultados basados en estas muestras diferirán en no más de tres puntos porcentuales en cualquier dirección de lo que se habría obtenido de recurrir a todos los estadounidenses adultos.

La próxima vez que lea esta clase de declaraciones en el periódico, tendrán más sentido para usted. Con todo, tenga cuidado, pues a veces se expresan tales afirmaciones sin que estén acreditadas; pero ahora usted es capaz de decidirlo.

En esta exposición sólo consideramos una clase de estadístico: los porcentajes arrojados por una **variable dicotómica** o **binominal**. Sin embargo, la misma lógica se aplicaría al examen de otros datos estadísticos, como el ingreso medio. Debido a que en tal caso los cálculos son algo más complicados, preferi considerar sólo las variables binominales en esta introducción.

Debe estar advertido de que, en términos técnicos, los usos de la teoría de la probabilidad en las encuestas que acabamos de explicar no están completamente justificados. La teoría de la distribución de muestreo tiene supuestos que casi nunca se dan en las condiciones de las encuestas. Por ejemplo, el número de muestras contenido en incrementos especificados de errores estándar supone una población infinitamente grande, un número infinito de muestras y un muestreo sin remplazos. Más aún, aquí hemos simplificado en exceso el salto de la inferencia de la distribución de varias muestras a las características de una sola.

Estas advertencias pretenden situarlo en perspectiva. Los investigadores suelen sobrestimar la precisión de los estimados de la teoría de la probabilidad en lo que atañe a las ciencias sociales. Como diremos en otras partes del capítulo y el libro, las variaciones de las técnicas de muestreo y factores ajenos pueden reducir aún más la legitimidad de las estimaciones. No obstante, los cálculos que estudiamos en esta sección serán muy provechosos cuando usted desee comprender y evaluar sus datos. Estos cálculos no dan las estimaciones precisas que algunos investigadores suponen, pero pueden ser válidos para los propósitos prácticos. Sin duda, son más válidos que las estimaciones menos meticulosas basadas en métodos de muestreo menos rigurosos.

Más importante, debe familiarizarse con la lógica básica de estos cálculos. Ya con esta información será capaz de reaccionar en forma razonable ante sus propios datos y los que informen otros.

Poblaciones y marcos de muestreo

La sección inmediata anterior se ocupó del modelo teórico para el muestreo en la investigación social. Si bien el lector, el estudiante y el investigador necesitan comprender la teoría, no es menos impor-

tante que aprecien las condiciones imperfectas que se dan en el campo. Esta sección se dedica al análisis de un aspecto de las condiciones del campo que exige un equilibrio entre las condiciones y las premisas teóricas. Vamos a considerar aquí la congruencia o la disparidad entre poblaciones de los marcos de muestreo.

Para decirlo de manera sencilla, un marco de muestreo es la lista o la cuasi lista de elementos de la que se elige una muestra probabilística. Los siguientes son informes de marcos de muestreo aparecidos en publicaciones científicas:

Obtuvimos los datos para esta investigación de una muestra aleatoria de padres con hijos en el tercer grado de escuelas públicas y parroquiales del condado de Yakima, Washington.

(PETERSEN Y MAYNARD, 1981:92)

La muestra del Tiempo 1 consistió en 160 nombres tomados al azar del directorio telefónico de Lubbock, Texas.

(TAN, 1980:242)

Reunimos los datos informados en este artículo [...] de una muestra probabilística de adultos de 18 años y mayores, residentes de los 48 estados contiguos de Estados Unidos. Durante el otoño de 1975, el Centro de Investigación y Encuestas de la Universidad de Michigan realizó entrevistas personales con 1914 individuos.

(JACKMAN Y SENTER, 1980:345)

Las muestras tomadas en la forma correcta proporcionan información adecuada para describir la población de elementos que componen el marco de muestreo, y nada más. Destaco el punto en vista de la tendencia tan común de los investigadores a elegir muestras de cierto marco y a hacer afirmaciones sobre poblaciones similares, pero no idénticas, a la población que definió ese marco de muestreo.

Por ejemplo, echemos un vistazo a este artículo que se ocupa de los fármacos que más recetan los médicos:

No es fácil conseguir información sobre las ventas de medicamentos, pero durante 25 años Rinaldo V. DeNuzzo, profesor de la materia en el Albany College de Farmacéutica, de la Universidad Unión, en Albany, N.Y., ha rastreado tales ventas investigando en las farmacias de las cercanías. Publicó los resultados en *MM&M*, una revista de comercio industrial.

La última encuesta de DeNuzzo, que abarcó el año de 1980, se basa en informes de 66 farmacias en 48 localidades de Nueva York y Nueva Jersey. A menos que haya algo peculiar en esa zona del país, se pueden tomar sus descubrimientos como representativos de lo que ocurre en toda la nación.

(Mosokowitz, 1981:33)

El principal hecho que debe haberlo sorprendido es el comentario casual sobre si habría algo peculiar en Nueva York y Nueva Jersey. Lo hay. El estilo de vida de estos dos estados no es el característico de los otros 48. No podemos suponer que los habitantes de estos estados costeros del este, grandes y urbanizados, tengan los mismos hábitos de uso de medicamentos que los habitantes de Mississippi, Nebraska o Vermont.

¿Pero acaso representa la encuesta siquiera los esquemas de prescripción de Nueva York y Nueva Jersey? Para determinarlo tendríamos que conocer un poco la forma en que se eligieron las 48 localidades y las 66 farmacias. Debemos ser precavidos al respecto, en vista de la referencia a una investigación "en las farmacias de las cercanías". Como veremos, hay varios métodos para elegir muestras que garantizan la representatividad y, a menos que los hayamos utilizado, no debemos generalizar a partir de los resultados del estudio.

Los estudios de empresas suelen ser los más fáciles desde el punto de vista del muestreo, porque tienen listas de sus miembros. En estos casos, tales listas constituyen excelentes marcos de muestreo. Si se toma de la lista una muestra aleatoria, los datos reunidos pueden considerarse representativos de todos los miembros (si todos los miembros aparecen en la lista).

Entre las poblaciones que cuentan con listas adecuadas, de las que se pueden tomar muestras, se encuentran estudiantes y maestros de escuelas elementales, secundarias y universidades, feligreses, obreros, miembros de fraternidades, miembros de clubes sociales, políticos y de servicios, y miembros de asociaciones de profesionistas.

Los comentarios anteriores atañen sobre todo a las organizaciones locales, pues con frecuencia las instituciones estatales o nacionales no tienen una lista única de miembros; por ejemplo, no hay una lista única de episcopistas. Sin embargo, un diseño de muestra ligeramente más complicado aprovecharía las listas de membresía de las iglesias locales: primero tomaría una muestra de las iglesias y luego una submuestra de las listas de los feligreses de las iglesias elegidas (veremos más de este tema más adelante).

Otras listas de individuos podrían tener una importancia particularmente especial para las necesidades de investigación de un estudio. Por ejemplo, las dependencias gubernamentales tienen padrones de votantes que usted podría emplear si desea realizar un sondeo antes de las elecciones o un examen detallado de la conducta electoral, pero debe estar seguro de que el padrón está actualizado. Hay también listas de propietarios de automóviles, beneficiarios de la asistencia social, contribuyentes, dueños de franquicias, profesionistas acreditados, etc. Aunque sea difícil conseguir algunas de estas listas, ofrecen excelentes marcos de muestreo para los propósitos especializados de una investigación.

Ya sabemos que los elementos de muestreo de los estudios no tienen que ser siempre individuos, así que podemos notar que también hay listas de otras clases de elementos: universidades, empresas de varios tipos, ciudades, publicaciones académicas, periódicos, sindicatos, clubes políticos, asociaciones profesionales, etcétera.

Los directorios telefónicos se emplean a menudo para sondeos de opinión "rápidos y descuidados". Nadie niega que son fáciles y baratos, ni hay dudas de las razones de su popularidad. Y si usted quiere hacer algunas afirmaciones sobre los suscriptores telefónicos, el directorio es un buen marco de muestreo (desde luego, recuerde que un directorio no incluirá a los nuevos suscriptores ni a los que tienen números privados, y que el muestreo se complicará aún más por la probable inclusión de teléfonos que no son residenciales). Por desgracia, se abusa de los directorios telefónicos como listados de los habitantes de una ciudad o de sus votantes. De los muchos defectos de este recurso, el principal es un sesgo de clase social. Es menos probable que los pobres tengan teléfono, y los ricos pueden tener más de una línea. Por tanto, las muestras tomadas de directorios telefónicos suelen tener un sesgo hacia las clases media y alta.

El sesgo de clase inherente a los directorios telefónicos a menudo está oculta. Los sondeos prelectorales que se realizan de esta manera son a veces muy precisos, quizá por el propio sesgo de clase que es evidente en el voto: las personas pobres votan menos. Entonces ocurre con frecuencia que ambos sesgos casi coincidan y que los resultados del sondeo telefónico se aproximen al desenlace de las elecciones. Por desgracia, uno nunca lo sabe con certeza hasta que terminan los conteos, y, a veces, como en el caso del sondeo de 1936 del Li-

terary Digest, uno descubre que los votantes no actuaron de acuerdo con los sesgos de clase que se anticipaban. Así, la principal desventaja del método es que el investigador carece de la capacidad de estimar el grado de error esperado en los resultados de la muestra. Se usan a menudo los directorios por calles y los mapas de contribuyentes para tomar muestras cómodas de viviendas, pero también pueden tener problemas de no compleción y posibles sesgos. Por ejemplo, dentro de las regiones urbanas zonificadas rigurosamente, es poco probable que aparezcan las unidades habitacionales ilegales en los registros oficiales. En consecuencia, no se elegirían estas unidades y los resultados de la muestra no serían representativas de ellas, que suelen ser más pobres y sobrepobladas que el promedio.

La mayoría de estos comentarios se aplican a Estados Unidos, pero en otros países la situación es muy diferente. Por ejemplo, en Japón el gobierno elabora listas muy precisas de la población. Además, se obliga por ley a los ciudadanos a mantenerlas actualizadas con los cambios de residencia, nacimientos y muertes en el hogar. Por eso es posible elegir con más facilidad muestras aleatorias simples de la población japonesa. Estas listas de registro se elaboran con criterios que chocan directamente con las normas de muchos países occidentales en cuanto a la privacidad de los individuos.

Clases de diseño de muestreo

Hasta este punto nos hemos concentrado en el muestreo aleatorio simple (MAS). En realidad, la generalidad de las estadísticas de que se valen los investigadores sociales supone tales muestras. Sin embargo, como veremos en breve, hay varias opciones para elegir el método de muestreo y usted rara vez escogerá el simple. Hay dos razones para ello. Primera, salvo con los marcos más sencillos, el muestreo aleatorio simple no es viable. Segunda —y quizá se sorprenda—, el muestreo aleatorio simple puede no ser el método más exacto disponible. Pasemos pues a explicar este método de muestreo y las demás opciones.

Muestreo aleatorio simple (MAS)

Como dijimos, el muestreo aleatorio simple es el método básico en los cálculos estadísticos de la investigación social. Debido a que las matemáticas

del muestreo aleatorio son especialmente complicadas, tomaremos un atajo en favor de la descripción de los modos de empleo de este método en el campo.

Ya que se ha establecido un marco de muestreo apropiado, para aplicar el muestreo aleatorio simple el investigador asigna un número único a cada elemento de la lista, sin saltar ningún número. En ese momento usa una tabla de números aleatorios (apéndice E) para elegir los elementos que compondrán la muestra. El recuadro titulado "Uso de las tablas de números aleatorios" explica el procedimiento.

Si el marco de muestreo se almacenó en una forma mecánica de lectura, como un disquete o una cinta magnética, la computadora puede elegir automáticamente una muestra aleatoria simple (en efecto, el programa de computadora numera los elementos del marco, genera su propia serie de números aleatorios e imprime la lista de los elementos que seleccionó).

La figura 8.10 ofrece una ilustración gráfica del muestreo aleatorio simple. Observe que todos los miembros de nuestra micropoblación hipotética están numerados del 1 al 100. Pasamos al apéndice E y decidimos tomar los últimos dos dígitos de la primera columna y empezar con el tercer número de arriba abajo. Esto nos señala a la persona número 30 como la primera elegida para la muestra, seguida por la número 67, etc. (la persona 100 habría sido elegida si en la lista hubiera aparecido el "00").

Muestreo sistemático

Rara vez se utiliza en la práctica el muestreo aleatorio simple. Como veremos, no suele ser el método más eficaz, y puede ser laborioso si se hace a mano. El MAS requiere una lista de elementos. Cuando se tiene una lista, los investigadores prefieren el muestreo sistemático en lugar del aleatorio simple.

En el muestreo sistemático, cada k -ésimo elemento de la lista se elige (sistemáticamente) para incluirlo en la muestra. Si la lista contiene 10 000 elementos y usted quiere una muestra de 1 000, escoge a cada décimo elemento. Para precaverse de cualquier posible sesgo humano en el uso de este método, elija al azar el primer elemento. Así, en el ejemplo anterior, comenzaría por elegir un número al azar entre uno y 10. El elemento que tiene ese número se incluye en la muestra y luego cada $dé-$