

The Logic of Sampling

CHAPTER OVERVIEW

Now you'll see how social scientists can select a few people for study—and discover things that apply to hundreds of millions of people not studied.



Introduction

A Brief History of Sampling

President Alf Landon

President Thomas E. Dewey

Two Types of Sampling Methods

Nonprobability Sampling

Reliance on Available Subjects

Purposive or Judgmental Sampling

Snowball Sampling

Quota Sampling

Selecting Informants

The Theory and Logic of Probability Sampling

—Conscious and Subconscious Sampling Bias

Representativeness and Probability of Selection

Random Selection

Probability Theory, Sampling Distributions, and Estimates of Sampling Error

Populations and Sampling Frames

Review of Populations

and Sampling Frames

Types of Sampling Designs

Simple Random Sampling

Systematic Sampling

Stratified Sampling

Implicit Stratification

in Systematic Sampling

Illustration: Sampling University Students

Multistage Cluster Sampling

Multistage Designs and Sampling Error

Stratification in Multistage Cluster Sampling

Probability Proportionate to Size (PPS) Sampling

Disproportionate Sampling and Weighting

Probability Sampling in Review

The Ethics of Sampling



aplia

Aplia for The Practice of Social Research

After reading, go to “Online Study Resources” at the end of this chapter for instructions on how to use Aplia’s homework and learning resources.

Introduction

One of the most visible uses of survey sampling lies in the political polling that is subsequently tested by election results. Whereas some people doubt the accuracy of sample surveys, others complain that political polls take all the suspense out of campaigns by foretelling the result.

Going into the 2008 presidential elections, pollsters were in agreement as to who would win, in contrast to their experiences in 2000 and 2004, which were closely contested races, Table 7-1 reports polls conducted during the few days preceding the election. Despite some variations, the overall picture they present is amazingly consistent and pretty well matches the election results.

Now, how many interviews do you suppose it took each of these pollsters to come within a couple of percentage points in estimating the behavior of more than 131 million voters? Often fewer than 2,000! In this chapter, we're going to find out how social researchers can pull off such wizardry.

For another powerful illustration of the potency of sampling, look at this graphic portrayal of then-President George W. Bush's approval ratings prior to and following the September 11, 2001, terrorist attack on the United States (see Figure 7-1). The data reported by several different polling agencies describe the same pattern.

Political polling, like other forms of social research, rests on observations. But neither pollsters nor other social researchers can observe everything that might be relevant to their interests. A critical part of social research, then, is deciding what to observe and what not. If you want to study voters, for example, which voters should you study?

The process of selecting observations is called *sampling*. Although sampling can mean any procedure for selecting units of observation—for example, interviewing every tenth passerby on a busy street—the key to generalizing from a sample to a larger population is probability sampling, which involves the important idea of random selection. Much of this chapter is devoted to the logic and skills of probability sampling. This topic is

TABLE 7-1
Election-Eve Polls Reporting Presidential Voting Plans, 2008

<i>Poll</i>	<i>Date Ended</i>	<i>Obama</i>	<i>McCain</i>
FOX	Nov 2	54	46
NBC/WSJ	Nov 2	54	46
Marist College	Nov 2	55	45
Harris Interactive	Nov 3	54	46
Reuters/C-SPAN/Zogby	Nov 3	56	44
ARG	Nov 3	54	46
Rasmussen	Nov 3	53	47
IBD/TIPP	Nov 3	54	46
DailyKos.com/Research 2000	Nov 3	53	47
GWU	Nov 3	53	47
Marist College	Nov 3	55	45
Actual vote	Nov 4	54	46

Source: Poll data are adapted from data presented at Pollster.com (<http://www.pollster.com/polls/us/08-us-pres-ge-mno.php>) on January 29, 2009. The official election results are from the Federal Election Commission (<http://www.fec.gov/pubrec/fe2008/2008presresults.pdf>) on the same date. For simplicity, since there were no undecideds in the official results and each of the third-party candidates received less than one percentage of the vote, I've apportioned the undecided and other votes according to the percentages, saying they were voting for Obama or McCain.

more rigorous and precise than some of the other topics in this book. Whereas social research as a whole is both art and science, sampling leans toward science. Although this subject is somewhat technical, the basic logic of sampling is not difficult to understand. In fact, the logical neatness of this topic can make it easier to comprehend than, say, conceptualization.

Although probability sampling is central to social research today, we'll take some time to examine a variety of nonprobability methods as well. These methods have their own logic and can provide useful samples for social inquiry.

Before we discuss the two major types of sampling, I'll introduce you to some basic ideas by way of a brief history of sampling. As you'll see, the pollsters who correctly predicted the election in 2008

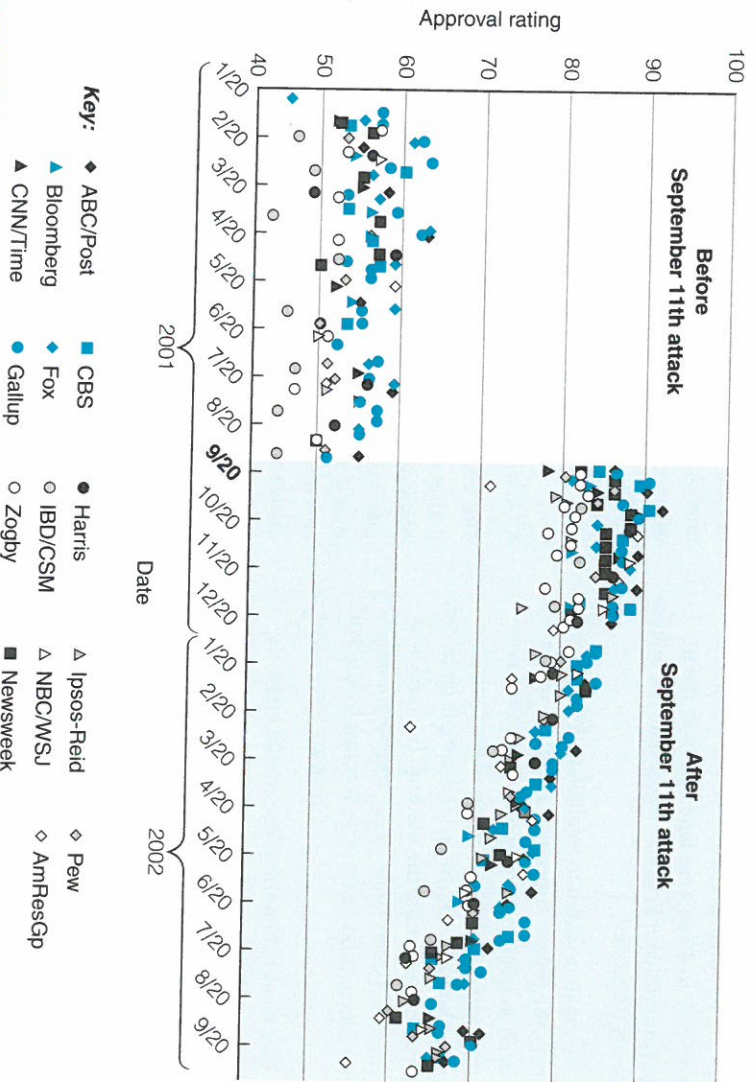


FIGURE 7-1 Bush Approval: Raw Poll Data. This graph demonstrates how independent polls produce the same picture of reality. This also shows the impact of a national crisis on the president's popularity: in this case, the September 11 terrorist attack and then-President George W. Bush's popularity. Source: Copyright © 2001, 2002 by dlineitka.com. (<http://www.pollkatz.homestead.com/files/MwHTML2.gif>). All rights reserved.

did so in part because researchers had learned to avoid some pitfalls that earlier pollsters had fallen into.

A Brief History of Sampling

Sampling in social research has developed hand in hand with political polling. This is the case, no doubt, because political polling is one of the few opportunities social researchers have to discover the accuracy of their estimates. On election day, they find out how well or how poorly they did.

President Alf Landon

President Alf Landon? Who's he? Did you sleep through an entire presidency in your U.S. history class? No—but Alf Landon would have been

president if a famous poll conducted by the *Literary Digest* had proved to be accurate. The *Literary Digest* was a popular newsmagazine published between 1890 and 1938. In 1916, *Digest* editors mailed postcards to people in six states, asking them whom they were planning to vote for in the presidential campaign between Woodrow Wilson and Charles Evans Hughes. Names were selected for the poll from telephone directories and automobile registration lists. Based on the postcards sent back, the *Digest* correctly predicted that Wilson would be elected. In the elections that followed, the *Literary Digest* expanded the size of its poll and made correct predictions in 1920, 1924, 1928, and 1932.

In 1936, the *Digest* conducted its most ambitious poll: Ten million ballots were sent to people listed in telephone directories and on lists of automobile owners. Over 2 million people responded,

giving the Republican contender, Alf Landon, a stunning 57 to 43 percent landslide over the incumbent, President Franklin Roosevelt. The editors modestly cautioned,

We make no claim to infallibility. We did not coin the phrase “uncanny accuracy” which has been so freely applied to our Polls. We know only too well the limitations of every straw vote, however enormous the sample gathered, however scientific the method. It would be a miracle if every State of the forty-eight behaved on Election Day exactly as forecast by the Poll.

(*Literary Digest 1936a: 6*)

Two weeks later, the *Digest* editors knew the limitations of straw polls even better: The voters gave Roosevelt a second term in office by the largest landslide in history, with 61 percent of the vote. Landon won only 8 electoral votes to Roosevelt’s 523.

The editors were puzzled by their unfortunate turn of luck. A part of the problem surely lay in the 22 percent return rate garnered by the poll. The editors asked,

Why did only one in five voters in Chicago to whom the *Digest* sent ballots take the trouble to reply? And why was there a preponderance of Republicans in the one-fifth that did reply? . . . We were getting better cooperation in what we have always regarded as a public service from Republicans than we were getting from Democrats. Do Republicans live nearer to mail-boxes? Do Democrats generally disapprove of straw polls?

(*Literary Digest 1936b: 7*)

Actually, there was a better explanation—what is technically called the *sampling frame* used by the *Digest*. In this case, the sampling frame consisted of telephone subscribers and automobile owners. In the context of 1936, this design selected a disproportionately wealthy sample of the voting population, especially coming on the tail end of the worst economic depression in the nation’s history. The sample effectively excluded poor people, and the poor voted predominantly for Roosevelt’s New Deal recovery program. The *Digest*’s poll may

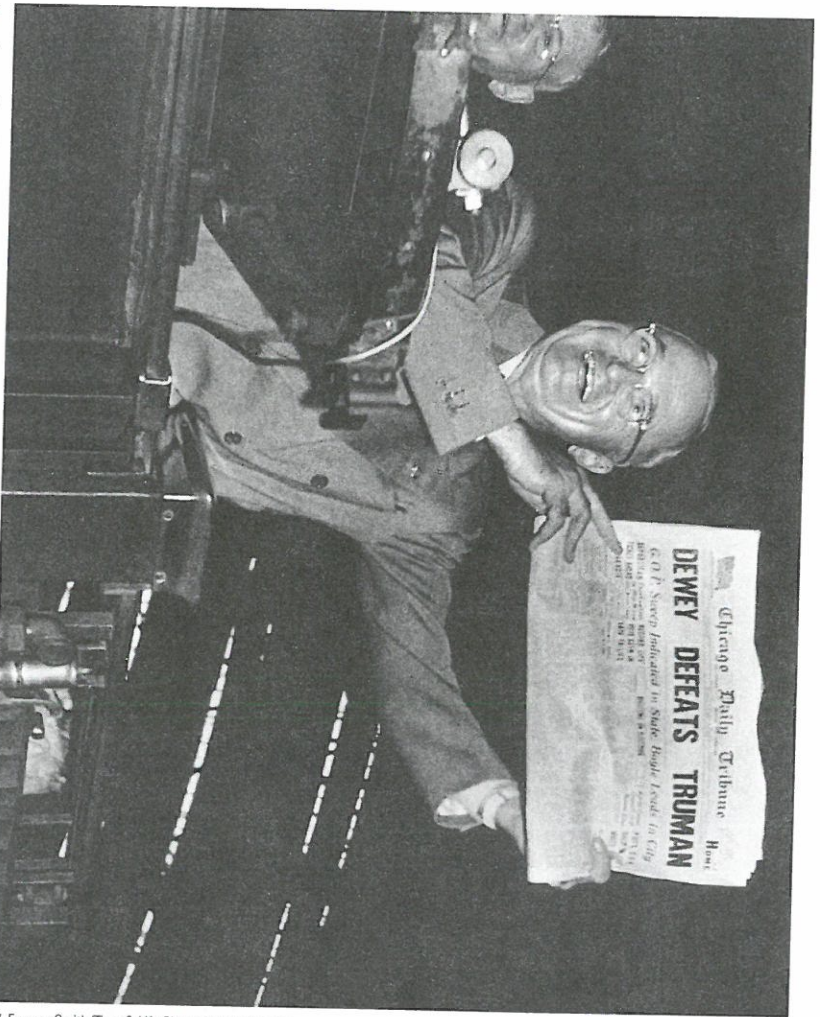
or may not have correctly represented the voting intentions of telephone subscribers and automobile owners. Unfortunately for the editors, it decidedly did not represent the voting intentions of the population as a whole.

President Thomas E. Dewey

The 1936 election also saw the emergence of a young pollster whose name would become synonymous with public opinion. In contrast to the *Literary Digest*, George Gallup correctly predicted that Roosevelt would beat Landon. Gallup’s success in 1936 hinged on his use of something called *quota sampling*, which we’ll look at more closely later in the chapter. For now, it’s enough to know that quota sampling is based on a knowledge of the characteristics of the population being sampled: what proportion are men, what proportion are women, what proportions are of various incomes, ages, and so on. Quota sampling selects people to match a set of these characteristics: the right number of poor, white, rural men; the right number of rich, African American, urban women; and so on. The quotas are based on those variables most relevant to the study. In the case of Gallup’s poll, the sample selection was based on levels of income; the selection procedure ensured the right proportion of respondents at each income level.

Gallup and his American Institute of Public Opinion used quota sampling to good effect in 1936, 1940, and 1944—correctly picking the presidential winner each of those years. Then, in 1948, Gallup and most political pollsters suffered the embarrassment of picking Governor Thomas Dewey of New York over the incumbent, President Harry Truman. The pollsters’ embarrassing miscue continued right up to election night. A famous photograph shows a jubilant Truman—whose followers’ battle cry was “Give ‘em hell, Harry!”—holding aloft a newspaper with the banner headline “Dewey Defeats Truman.”

Several factors accounted for the pollsters’ failure in 1948. First, most pollsters stopped polling in early October despite a steady trend toward Truman during the campaign. In addition, many voters were undecided throughout the campaign,



W. Eugene Smith/Time & Life Pictures/Getty Images

Based on early political polls that showed Dewey leading Truman, the *Chicago Tribune* sought to scoop the competition with this unfortunate headline.

and these went disproportionately for Truman when they stepped into the voting booth.

More important, Gallup's failure rested on the unrepresentativeness of his samples. Quota sampling—which had been effective in earlier years—was Gallup's undoing in 1948. This technique requires that the researcher know something about the total population (of voters in this instance). For national political polls, such information came primarily from census data. By 1948, however, World War II had produced a massive movement from the country to cities, radically changing the character of the U.S. population from what the 1940 census showed, and Gallup relied on 1940 census data. City dwellers, moreover, tended to vote Democratic; hence, the overrepresentation of rural voters in his poll had the effect of underestimating the number of Democratic votes.

Two Types of Sampling Methods

By 1948, some academic researchers had already been experimenting with a form of sampling based on probability theory. This technique involves the selection of a “random sample” from a list containing the names of everyone in the population being sampled. By and large, the probability-sampling methods used in 1948 were far more accurate than quota-sampling techniques.

Today, probability sampling remains the primary method of selecting large, representative samples for social research, including national political polls. At the same time, probability sampling can be impossible or inappropriate in many research situations. Accordingly, before turning to the logic and techniques of probability sampling, we'll first take a look at techniques for

nonprobability sampling and how they're used in social research.

Nonprobability Sampling

Social research is often conducted in situations that do not permit the kinds of probability samples used in large-scale social surveys. Suppose you wanted to study homelessness: There is no list of all homeless individuals, nor are you likely to create such a list. Moreover, as you'll see, there are times when probability sampling wouldn't be appropriate even if it were possible. Many such situations call for **nonprobability sampling**.

In this section, we'll examine four types of nonprobability sampling: reliance on available subjects, purposive (judgmental) sampling, snowball sampling, and quota sampling. We'll conclude with a brief discussion of techniques for obtaining information about social groups through the use of informants.

Reliance on Available Subjects

Relying on available subjects, such as stopping people at a street corner or some other location, is sometimes called "convenience" or "haphazard" sampling. This is a common method for journalists in their "person-on-the-street" interviews, but it is an extremely risky sampling method for social research. Clearly, this method does not permit any control over the representativeness of a sample. It's justified only if the researcher wants to study the characteristics of people passing the sampling point at specified times or if less-risky sampling methods

are not feasible. Even when this method is justified on grounds of feasibility, researchers must exercise great caution in generalizing from their data. Also, they should alert readers to the risks associated with this method.

University researchers frequently conduct surveys among the students enrolled in large lecture classes. The ease and frugality of such a method explains its popularity, but it seldom produces data of any general value. It may be useful for pretesting a questionnaire, but such a sampling method should not be used for a study purportedly describing students as a whole.

Consider this report on the sampling design in an examination of knowledge and opinions about nutrition and cancer among medical students and family physicians:

The fourth-year medical students of the University of Minnesota Medical School in Minneapolis comprised the student population in this study. The physician population consisted of all physicians attending a "Family Practice Review and Update" course sponsored by the University of Minnesota Department of Continuing Medical Education.

(Cooper-Stephenson and Theologides 1981: 472)

After all is said and done, what will the results of this study represent? The data do not provide a meaningful comparison of medical students and family physicians in the United States or even in Minnesota. Who were the physicians who attended the course? We can guess that they were probably more concerned about their continuing education than other physicians were, but we can't say for sure. Although such studies can provide useful insights, we must take care not to overgeneralize from them.

Purposive or Judgmental Sampling

Sometimes it's appropriate to select a sample on the basis of knowledge of a population, its elements, and the purpose of the study. This type of sampling is called **purposive** or **judgmental sampling**. In the initial design of a questionnaire, for example,

nonprobability sampling Any technique in which samples are selected in some way not suggested by probability theory. Examples include reliance on available subjects as well as purposive (judgmental), quota, and snowball sampling.

purposive (judgmental) sampling A type of nonprobability sampling in which the units to be observed are selected on the basis of the researcher's judgment about which ones will be the most useful or representative.

you might wish to select the widest variety of respondents to test the broad applicability of questions. Although the study findings would not represent any meaningful population, the test run might effectively uncover any peculiar defects in your questionnaire. This situation would be considered a pretest, however, rather than a final study.

In some instances, you may wish to study a small subset of a larger population in which many members of the subset are easily identified, but the enumeration of them all would be nearly impossible. For example, you might want to study the leadership of a student protest movement; many of the leaders are easily visible, but it would not be feasible to define and sample all the leaders. In studying all or a sample of the most visible leaders, you may collect data sufficient for your purposes.

Or let's say you want to compare left-wing and right-wing students. Because you may not be able to enumerate and sample from all such students, you might decide to sample the memberships of left- and right-leaning groups, such as the Green Party and the Tea Party. Although such a sample design would not provide a good description of either left-wing or right-wing students as a whole, it might suffice for general comparative purposes.

Field researchers are often particularly interested in studying deviant cases—cases that don't fit into fairly regular patterns of attitudes and behaviors—in order to improve their understanding of the more-regular pattern. For example, you might gain important insights into the nature of school spirit, as exhibited at a pep rally, by interviewing people who did not appear to be caught up in the emotions of the crowd or by interviewing students who did not attend the rally at all. Selecting deviant cases for study is another example of purposive study.

In qualitative research projects, the sampling of subjects may evolve as the structure of the situation being studied becomes clearer and certain types of subjects seem more central to understanding than others do. Let's say you're conducting an interview study among the members of a radical political group on campus. You may initially focus on friendship networks as a vehicle for the spread of group membership and participation. In the

course of your analysis of the earlier interviews, you may find several references to interactions with faculty members in one of the social science departments. As a consequence, you may expand your sample to include faculty in that department and other students that they interact with. This is called “theoretical sampling,” since the evolving theoretical understanding of the subject directs the sampling in certain directions.

Snowball Sampling

Another nonprobability sampling technique, which some consider to be a form of accidental sampling, is called **snowball sampling**. This procedure is appropriate when the members of a special population are difficult to locate, such as homeless individuals, migrant workers, or undocumented immigrants. In snowball sampling, the researcher collects data on the few members of the target population he or she can locate, then asks those individuals to provide the information needed to locate other members of that population whom they happen to know. “Snowball” refers to the process of accumulation as each located subject suggests other subjects. Because this procedure also results in samples with questionable representativeness, it's used primarily for exploratory purposes.

Suppose you wish to learn a community organization's pattern of recruitment over time. You might begin by interviewing fairly recent recruits, asking them who introduced them to the group. You might then interview the people named, asking them who introduced *them* to the group. You might then interview those people named, asking, in part, who introduced *them*. Or, in studying a loosely structured political group, you might ask one of the participants who he or she believes to be the most influential members of the group. You might interview those people and, in the course of the interviews, ask who *they* believe to be the most influential. In each of these examples, your sample

snowball sampling A nonprobability sampling method, often employed in field research, whereby each person interviewed may be asked to suggest additional people for interviewing.

would “snowball” as each of your interviewees suggested other people to interview.

Examples of this technique in social science research abound. Karen Farquharson (2005) provides a detailed discussion of how she used snowball sampling to discover a network of tobacco policy makers in Australia: both those at the core of the network and those on the periphery. Kath Browne (2005) used snowballing through social networks to develop a sample of non-heterosexual women in a small town in the United Kingdom. She reports that her own membership in such networks greatly facilitated this type of sampling, and that potential subjects in the study were more likely to trust her than to trust heterosexual researchers.

In more-general, theoretical terms, Chaim Noy argues that the process of selecting a snowball sample reveals important aspects of the populations being sampled, uncovering “the dynamics of natural and organic social networks” (2008: 329). Do the people you interview know others like themselves? Are they willing to identify those people to researchers? Thus, snowball sampling can be more than a simple technique for finding people to study. It can be a revealing part of the inquiry.

Quota Sampling

Quota sampling is the method that helped George Gallup avoid disaster in 1936—and set up the disaster of 1948. Like probability sampling, quota sampling addresses the issue of representativeness, although the two methods approach the issue quite differently.

Quota sampling begins with a matrix, or table, describing the characteristics of the target population. Depending on your research purposes, you may need to know what proportion of the population is male and what proportion female, as well as knowing what proportions of each sex fall into

various age categories, educational levels, ethnic groups, and so forth. In establishing a national quota sample, you might need to know what proportion of the national population is urban, eastern, male, under 25, white, working class, and the like, and all the possible combinations of these attributes. Once you’ve created such a matrix and assigned a relative proportion to each cell in the matrix, you proceed to collect data from people having all the characteristics of a given cell. You then assign to all the people in a given cell a weight appropriate to their portion of the total population. When all the sample elements are so weighted, the overall data should provide a reasonable representation of the total population.

Although quota sampling resembles probability sampling, it has several inherent problems. First, the *quota frame* (the proportions that different cells represent) must be accurate, and it’s often difficult to get up-to-date information for this purpose. The Gallup failure to predict Truman as the presidential victor in 1948 was due partly to this problem. Second, the selection of sample elements within a given cell may be biased even though its proportion of the population is accurately estimated. Instructed to interview five people who meet a given, complex set of characteristics, an interviewer may still avoid people living at the top of seven-story walk-ups, having particularly run-down homes, or owning vicious dogs.

In recent years, attempts have been made to combine probability- and quota-sampling methods, but the effectiveness of this effort remains to be seen. At present, you would be advised to treat quota sampling warily if your purpose is statistical description.

At the same time, the logic of quota sampling can sometimes be applied usefully to a field research project. In the study of a formal group, for example, you might wish to interview both leaders and nonleaders. In studying a student political organization, you might want to interview radical, moderate, and conservative members of that group. You may be able to achieve sufficient representativeness in such cases by using quota sampling to ensure that you interview both men and women, both younger and older people, and so forth.

quota sampling A type of nonprobability sampling in which units are selected into a sample on the basis of prespecified characteristics, so that the total sample will have the same distribution of characteristics assumed to exist in the population being studied.

Selecting Informants

When field research involves the researcher's attempt to understand some social setting—a juvenile gang or local neighborhood, for example—much of that understanding will come from a collaboration with some members of the group being studied.

Whereas social researchers speak of *respondents* as people who provide information about themselves, allowing the researcher to construct a composite picture of the group those respondents represent, an **informant** is a member of the group who can talk directly about the group *per se*.

Especially important to anthropologists, informants are important to other social researchers as well. If you wanted to learn about informal social networks in a local public-housing project, for example, you would do well to locate individuals who could understand what you were looking for and help you find it.

When Jeffrey Johnson (1990) set out to study a salmon-fishing community in North Carolina, he used several criteria to evaluate potential informants. Did their positions allow them to interact regularly with other members of the camp, for example, or were they isolated? (In this case, he found that the carpenter had a wider range of interactions than the boat captain did.) Was their information about the camp pretty much limited to their specific jobs, or did it cover many aspects of the operation? These and other criteria helped determine how useful the potential informants might be.

Usually, you'll want to select informants somewhat typical of the groups you're studying. Otherwise, their observations and opinions may be misleading. Interviewing only physicians will not give you a well-rounded view of how a community medical clinic is working, for example. Along the same lines, an anthropologist who interviews only men in a society where women are sheltered from outsiders will get a biased view. Similarly, although informants fluent in English are convenient for English-speaking researchers from the United States, they do not typify the members of many societies nor even many subgroups within English-speaking countries.

Simply because they're the ones willing to work with outside investigators, informants will

almost always be somewhat "marginal" or atypical within their group. Sometimes this is obvious. Other times, however, you'll learn about their marginality only in the course of your research.

In Jeffrey Johnson's study, a county agent identified one fisherman who seemed squarely in the mainstream of the community. Moreover, he was cooperative and helpful to Johnson's research. The more Johnson worked with the fisherman, however, the more he found the man to be a marginal member of the fishing community.

First, he was a Yankee in a southern town. Second, he had a pension from the Navy [so he was not seen as a "serious fisherman" by others in the community]. . . . Third, he was a major Republican activist in a mostly Democratic village. Finally, he kept his boat in an isolated anchorage, far from the community harbor.

(1990: 56)

Informants' marginality may not only bias the view you get, but their marginal status may also limit their access (and hence yours) to the different sections of the community you wish to study.

These comments should give you some sense of the concerns involved in nonprobability sampling, typically used in qualitative research projects. I conclude with the following injunction:

Your overall goal is to collect the *richest possible data*. By rich data, we mean a wide and diverse range of information collected over a relatively prolonged period of time in a persistent and systematic manner. Ideally, such data enable you to grasp the meanings associated with the actions of those you are studying and to understand the contexts in which those actions are embedded.

(Lofland et al. 2006: 15)

In other words, nonprobability sampling does have its uses, particularly in qualitative research projects. But researchers must take care to

informant Someone who is well versed in the social phenomenon that you wish to study and who is willing to tell you what he or she knows about it. Not to be confused with a *respondent*.

acknowledge the limitations of nonprobability sampling, especially regarding accurate and precise representations of populations. This point will become clearer as we discuss the logic and techniques of probability sampling.

The Theory and Logic of Probability Sampling

However appropriate to some research purposes, nonprobability sampling methods cannot guarantee that the sample we observed is representative of the whole population. When researchers want precise, statistical descriptions of large populations—for example, the percentage of the population who is unemployed, plan to vote for Candidate X, or feel a rape victim should have the right to an abortion—they turn to **probability sampling**. All large-scale surveys use probability-sampling methods.

Although the application of probability sampling involves some sophisticated use of statistics, the basic logic of probability sampling is not difficult to understand. If all members of a population were identical in all respects—all demographic characteristics, attitudes, experiences, behaviors, and so on—there would be no need for careful sampling procedures. In this extreme case of perfect homogeneity, in fact, any single case would suffice as a sample to study characteristics of the whole population.

In fact, of course, the human beings who compose any real population are quite heterogeneous, varying in many ways. Figure 7-2 offers a simplified illustration of a heterogeneous population: The 100 members of this small population differ by sex and race. We'll use this hypothetical micropopulation to illustrate various aspects of probability sampling.

The fundamental idea behind probability sampling is this: To provide useful descriptions of the total population, a sample of individuals from a population must contain essentially the same

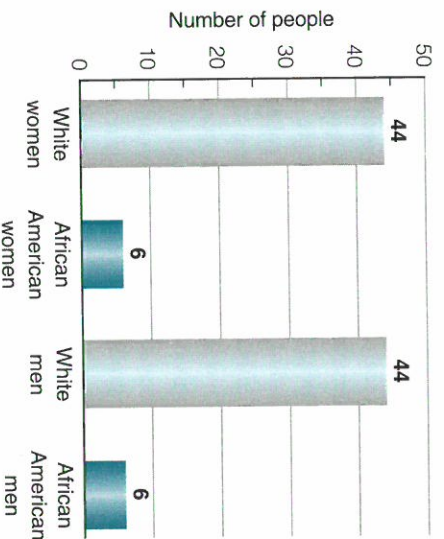


FIGURE 7-2
A Population of 100 Folks. Typically, sampling aims to reflect the characteristics and dynamics of large populations. For the purpose of some simple illustrations, let's assume our total population only has 100 members.

variations that exist in the population. This isn't as simple as it might seem, however. Let's take a minute to look at some of the ways researchers might go astray. Then, we'll see how probability sampling provides an efficient method for selecting a sample that should adequately reflect variations that exist in the population.

Conscious and Subconscious Sampling Bias

At first glance, it may look as though sampling is pretty straightforward. To select a sample of 100 university students, you might simply interview the first 100 students you find walking around campus. This kind of sampling method is often used by untrained researchers, but it runs a high risk of introducing biases into the samples.

In connection with sampling, *bias* simply means that those selected are not typical or representative of the larger populations they have been chosen from. This kind of bias does not have to be intentional. In fact, it is virtually inevitable when you pick people by the seat of your pants.

Figure 7-3 illustrates what can happen when researchers simply select people who are convenient for study. Although women are only 50 percent of our micropopulation, the people closest to

probability sampling The general term for samples selected in accord with probability theory, typically involving some random-selection mechanism. Specific types of probability sampling include EPSEM, PPS, simple random sampling, and systematic sampling.

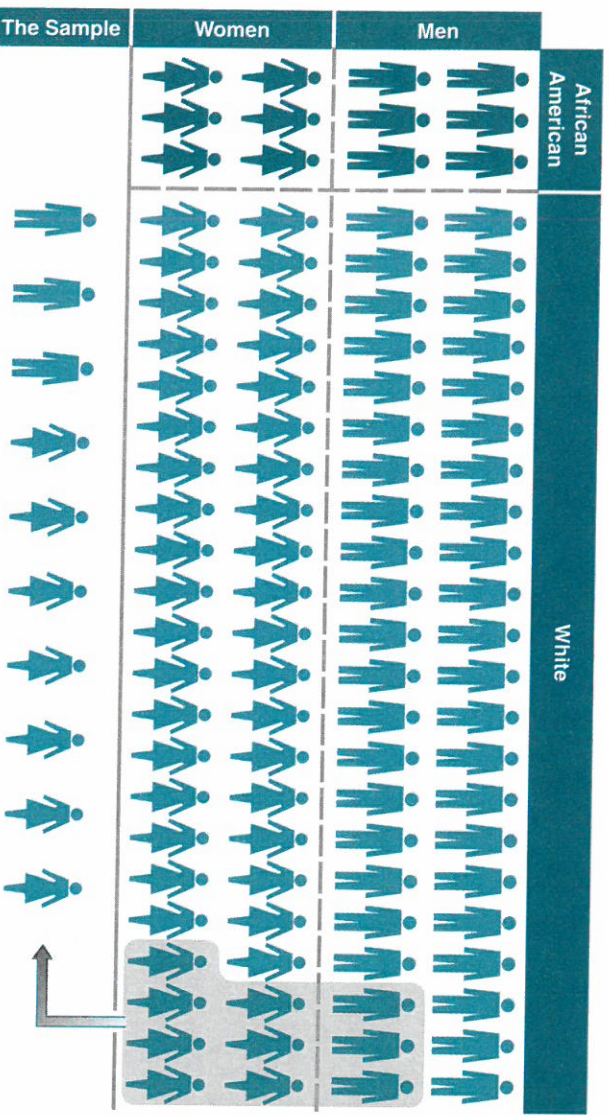


FIGURE 7-3

A Sample of Convenience: Easy, but Not Representative. Simply selecting and observing those people who are most readily at hand is the simplest method, perhaps, but it's unlikely to provide a sample that accurately reflects the total population.

the researcher (in the lower right corner) happen to be 70 percent women, and although the population is 12 percent African American, none was selected into the sample.

Beyond the risks inherent in simply studying people who are convenient, other problems can arise. To begin with, the researcher's personal leanings may affect the sample to the point where it does not truly represent the student population.

Suppose you're a little intimidated by students who look particularly "cool," feeling they might ridicule your research effort. You might consciously or subconsciously avoid interviewing such people. Or, you might feel that the attitudes of "super-straight-looking" students would be irrelevant to your research purposes and so avoid interviewing them.

Even if you sought to interview a "balanced" group of students, you wouldn't know the exact proportions of different types of students making up such a balance, and you wouldn't always be able to identify the different types just by watching them walk by.

Even if you made a conscientious effort to interview, say, every tenth student entering the

university library, you could not be sure of a representative sample, because different types of students visit the library with different frequencies. Your sample would overrepresent students who visit the library more often than others do.

The possibilities for inadvertent sampling bias are endless and not always obvious. Fortunately, many techniques can help us avoid bias.

Representativeness and Probability of Selection

Although the term **representativeness** has no precise, scientific meaning, it carries a

representativeness That quality of a sample of having the same distribution of characteristics as the population from which it was selected. By implication, descriptions and explanations derived from an analysis of the sample may be assumed to represent similar ones in the population. Representativeness is enhanced by probability sampling and provides for generalizability and the use of inferential statistics.

commonsense meaning that makes it useful here. For our purpose, a sample is representative of the population from which it is selected if the aggregate characteristics of the sample closely approximate those same aggregate characteristics in the population. If, for example, the population contains 50 percent women, then a sample must contain “close to” 50 percent women to be representative. Later, we’ll discuss “how close” in detail.

Note that samples need not be representative in all respects; representativeness is limited to those characteristics that are relevant to the substantive interests of the study. However, you may not know in advance which characteristics are relevant.

A basic principle of probability sampling is that a sample will be representative of the population from which it is selected if all members of the population have an equal chance of being selected in the sample. (We’ll see shortly that the size of the sample selected also affects the degree of representativeness.) Samples that have this quality are often labeled **EPSEM** samples (EPSEM stands for “equal probability of selection method”). Later, we’ll discuss variations of this principle, which forms the basis of probability sampling.

Moving beyond this basic principle, we must realize that samples—even carefully selected EPSEM samples—seldom if ever perfectly represent the populations from which they are drawn. Nevertheless, probability sampling offers two special advantages.

First, probability samples, although never perfectly representative, are typically more representative than other types of samples, because the biases previously discussed are avoided. In practice, a probability sample is more likely than

a nonprobability sample to be representative of the population from which it is drawn.

Second, and more important, probability theory permits us to estimate the accuracy or representativeness of the sample. Conceivably, an uninformed researcher might, through wholly haphazard means, select a sample that nearly perfectly represents the larger population. The odds are against doing so, however, and we would be unable to estimate the likelihood that he or she has achieved representativeness. The probability sampler, on the other hand, can provide an accurate estimate of success or failure. We’ll see exactly how this estimate can be achieved.

I’ve said that probability sampling ensures that samples are representative of the population we wish to study. As we’ll see in a moment, probability sampling rests on the use of a random-selection procedure. To develop this idea, though, we need to give more-precise meaning to two important terms: *element* and *population*.*

An **element** is that unit about which information is collected and that provides the basis of analysis. Typically, in survey research, elements are people or certain types of people. However, other kinds of units can constitute the elements for social research: Families, social clubs, or corporations might be the elements of a study. In a given study, elements are often the same as units of analysis, though the former are used in sample selection and the latter in data analysis.

Up to now we’ve used the term **population** to mean the group or collection that we’re interested in generalizing about. More formally, a *population* is the theoretically specified aggregation of study elements. Whereas the vague term *Americans* might be the target for a study, the delineation of the population would include the definition of the element *Americans* (for example, citizenship, residence) and the time referent for the study (Americans as of when?). Translating the abstract “adult New Yorkers” into a workable population would require a

EPSEM (equal probability of selection method) A sample design in which each member of a population has the same chance of being selected into the sample.

element That unit of which a population is composed and which is selected in a sample. Distinguished from *units of analysis*, which are used in data analysis.

population The theoretically specified aggregation of the elements in a study.

*I would like to acknowledge a debt to Leslie Kish and his excellent textbook *Survey Sampling*. Although I’ve modified some of the conventions used by Kish, his presentation is easily the most important source of this discussion.

specification of the age defining *adult* and the boundaries of New York. Specifying the term *college student* would include a consideration of full- and part-time students, degree candidates and non-degree candidates, undergraduate and graduate students, and so forth.

A **study population** is that aggregation of elements from which the sample is actually selected. As a practical matter, researchers are seldom in a position to guarantee that every element meeting the theoretical definitions laid down actually has a chance of being selected in the sample. Even where lists of elements exist for sampling purposes, the lists are usually somewhat incomplete. Some students are always inadvertently omitted from student rosters. Some telephone subscribers request that their names and numbers be unlisted.

Often, researchers decide to limit their study populations more severely than indicated in the preceding examples. National polling firms may limit their national samples to the 48 adjacent states, omitting Alaska and Hawaii for practical reasons. A researcher wishing to sample psychology professors may limit the study population to those in psychology departments, omitting those in other departments. Whenever the population under examination is altered in such fashions, you must make the revisions clear to your readers.

Random Selection

With these definitions in hand, we can define the ultimate purpose of sampling: to select a set of elements from a population in such a way that descriptions of those elements accurately portray the total population from which the elements are selected. Probability sampling enhances the likelihood of accomplishing this aim and also provides methods for estimating the degree of probable success.

Random selection is the key to this process. In **random selection**, each element has an equal chance of selection independent of any other event in the selection process. Flipping a coin is the most frequently cited example: Provided that the coin is perfect (that is, not biased in terms of coming up

heads or tails), the “selection” of a head or a tail is independent of previous selections of heads or tails. No matter how many heads turn up in a row, the chance that the next flip will produce “heads” is exactly 50–50. Rolling a perfect set of dice is another example.

Such images of random selection, although useful, seldom apply directly to sampling methods in social research. More typically, social researchers use tables of random numbers or computer programs that provide a random selection of sampling units. A **sampling unit** is that element or set of elements considered for selection in some stage of sampling. In Chapter 9, on survey research, we’ll see how computers are used to select random telephone numbers for interviewing, a technique called *random-digit dialing*.

The reasons for using random-selection methods are twofold. First, this procedure serves as a check on conscious or unconscious bias on the part of the researcher. The researcher who selects cases on an intuitive basis might very well select cases that would support his or her research expectations or hypotheses. Random selection erases this danger. More importantly, random selection offers access to the body of probability theory, which provides the basis for estimating the characteristics of the population as well as estimating the accuracy of samples. Let’s now examine probability theory in greater detail.

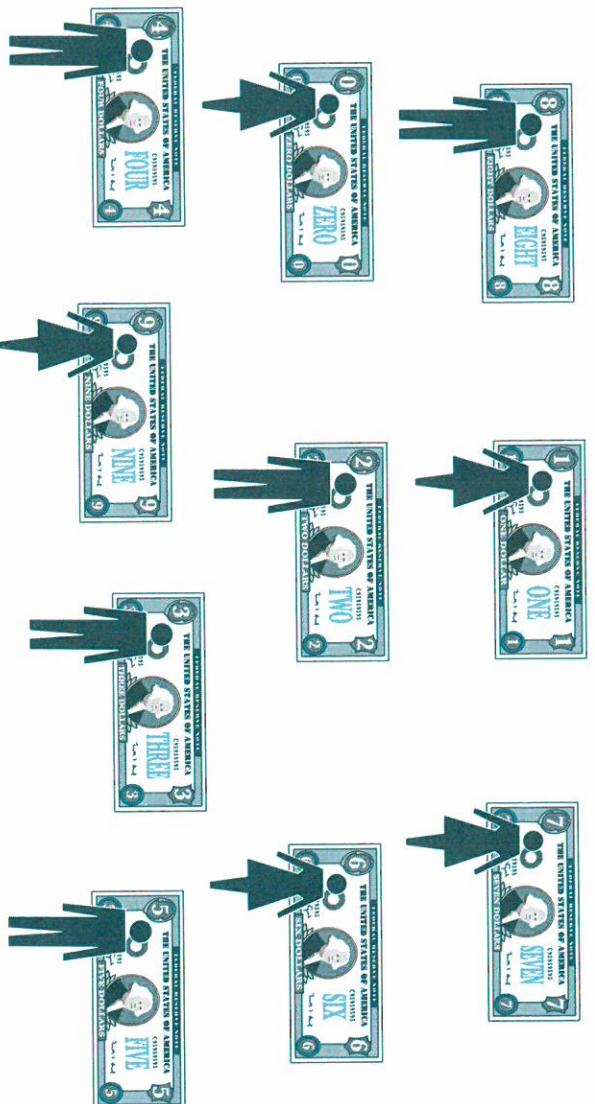
Probability Theory, Sampling Distributions, and Estimates of Sampling Error

Probability theory is a branch of mathematics that provides the tools researchers need to devise sampling techniques that produce representative

study population That aggregation of elements from which a sample is actually selected.

random selection A sampling method in which each element has an equal chance of selection independent of any other event in the selection process.

sampling unit That element or set of elements considered for selection in some stage of sampling.

**FIGURE 7-4**

A Population of 10 People with \$0–\$9. Let's simplify matters even more now by imagining a population of only 10 people with differing amounts of money in their pockets—ranging from \$0 to \$9.

samples and to analyze the results of their sampling statistically. More formally, probability theory provides the basis for estimating the parameters of a population. A **parameter** is the summary description of a given variable in a population. The mean income of all families in a city is a parameter; so is the age distribution of the city's population. When researchers generalize from a sample, they're using sample observations to estimate population parameters. Probability theory enables them to both make these estimates and arrive at a judgment of how likely the estimates will accurately represent the actual parameters in the population. For example, probability theory allows pollsters to infer from a sample of 2,000 voters how a population of 100 million voters is likely to vote—and to specify exactly what the probable margin of error of the estimates is.

Probability theory accomplishes these seemingly magical feats by way of the concept of sampling

distributions. A single sample selected from a population will give an estimate of the population parameter. Other samples would give the same or slightly different estimates. Probability theory tells us about the distribution of estimates that would be produced by a large number of such samples. To see how this works, we'll look at two examples of sampling distributions, beginning with a simple example in which our population consists of just ten cases, then moving on to a case of percentages that allows a clear illustration of probable margin of error.

The Sampling Distribution of Ten Cases

Suppose there are ten people in a group, and each has a certain amount of money in his or her pocket. To simplify, let's assume that one person has no money, another has one dollar, another has two dollars, and so forth up to the person with nine dollars. Figure 7-4 presents the population of ten people.*

parameter The summary description of a given variable in a population.

*I want to thank Hanan Selvin for suggesting this method of introducing probability sampling.

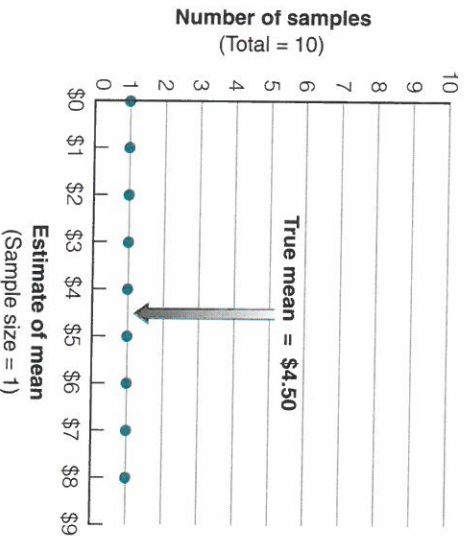


FIGURE 7-5

The Sampling Distribution of Samples of 1. In this simple example, the mean amount of money these people have is \$4.50 ($\$45/10$). If we picked 10 different samples of 1 person each, our “estimates” of the mean would range all across the board.

Our task is to determine the average amount of money one person has; specifically, the mean number of dollars. If you simply add up the money shown in Figure 7-4, you’ll find that the total is \$45, so the mean is \$4.50. Our purpose in the rest of this exercise is to estimate that mean without actually observing all ten individuals. We’ll do that by selecting random samples from the population and using the means of those samples to estimate the mean of the whole population.

To start, suppose we were to select—at random—a sample of only one person from the ten. Our ten possible samples thus consist of the ten cases shown in Figure 7-4.

The ten dots shown on the graph in Figure 7-5 represent these ten samples. Because we’re taking samples of only one, they also represent the “means” we would get as estimates of the population. The distribution of the dots on the graph is called the *sampling distribution*. Obviously, it wouldn’t be a very good idea to select a sample of only one, because the chances are great that we’ll miss the true mean of \$4.50 by quite a bit.

Now suppose we take a sample of two. As shown in Figure 7-6, increasing the sample size improves our estimations. There are now 45 possible

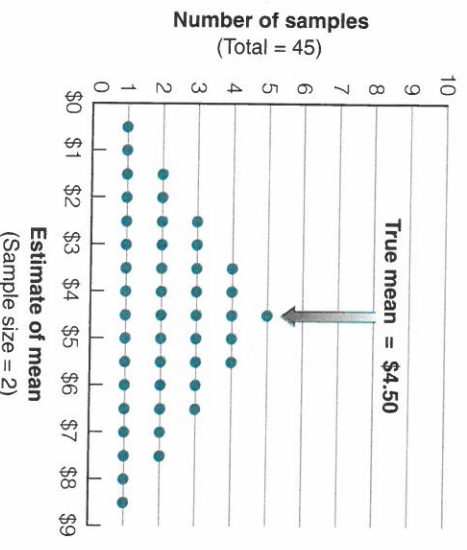


FIGURE 7-6

The Sampling Distribution of Samples of 2. By merely increasing our sample size to 2, we get possible samples that provide somewhat better estimates of the mean. We couldn’t get either \$0 or \$9, and the estimates are beginning to cluster around the true value of the mean: \$4.50.

samples: [\$0 \$1], [\$0 \$2], . . . , [\$7 \$8], [\$8 \$9]. Moreover, some of those samples produce the same means. For example, [\$0 \$6], [\$1 \$5], and [\$2 \$4] all produce means of \$3. In Figure 7-6, the three dots shown above the \$3 mean represent those three samples.

Moreover, the 45 samples are not evenly distributed, as they were when the sample size was only one. Rather, they’re somewhat clustered around the true value of \$4.50. Only two possible samples deviate by as much as \$4 from the true value ([\$0 \$1] and [\$8 \$9]), whereas five of the samples would give the true estimate of \$4.50; another eight samples miss the mark by only 50 cents (plus or minus).

Now suppose we select even larger samples. What do you think that will do to our estimates of the mean? Figure 7-7 presents the sampling distributions of samples of 3, 4, 5, and 6.

The progression of sampling distributions is clear. Every increase in sample size improves the distribution of estimates of the mean. The limiting case in this procedure, of course, is to select a sample of ten. There would be only one possible sample (everyone) and it would give us the true mean

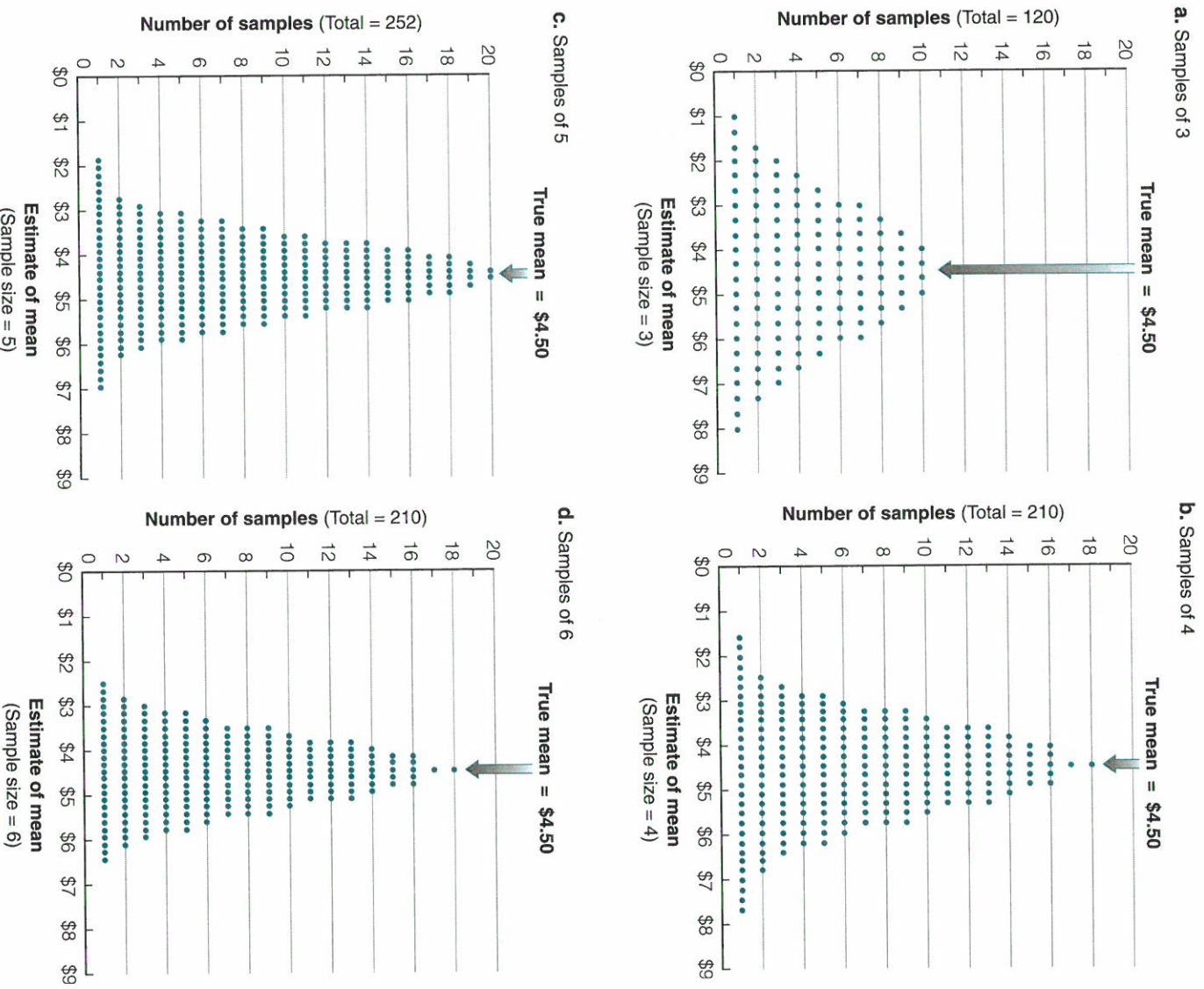
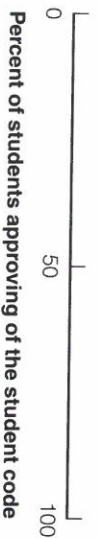


FIGURE 7-7

The Sampling Distributions of Samples of 3, 4, 5, and 6. As we increase the sample size, the possible samples cluster evermore tightly around the true value of the mean. The chance of extremely inaccurate estimates is reduced at the two ends of the distribution, and the percentage of the samples near the true value keeps increasing.

**FIGURE 7-8**

Range of Possible Sample Study Results. Shifting to a more realistic example, let's assume that we want to sample student attitudes concerning a proposed conduct code. Let's assume that 50 percent of the whole student body approves and 50 percent disapproves—though the researcher doesn't know that.

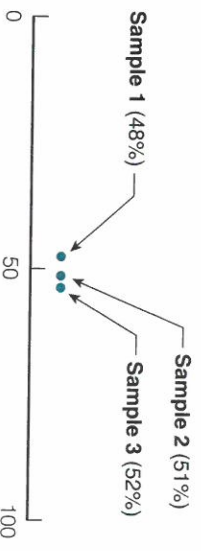
of \$4.50. As we'll see shortly, this principle applies to actual sampling of meaningful populations. The larger the sample selected, the more accurate it is as an estimation of the population from which it was drawn.

Sampling Distribution and Estimates of Sampling Error

Let's turn now to a more realistic sampling situation involving a much larger population and see how the notion of sampling distribution applies. Assume that we wish to study the student population of State University (SU) to determine the percentage of students who approve or disapprove of a student-conduct code proposed by the administration. The study population will be the aggregation of, say, 20,000 students contained in a student roster: the sampling frame. The elements will be the individual students at SU. We'll select a random sample of, say, 100 students for the purposes of estimating the entire student body. The variable under consideration will be *attitudes toward the code*, a binomial variable: *approve* and *disapprove*. (The logic of probability sampling applies to the examination of other types of variables, such as mean income, but the computations are somewhat more complicated. Consequently, this introduction focuses on binomials.)

The horizontal axis of Figure 7-8 presents all possible values of this parameter in the population—from 0 percent to 100 percent approval. The midpoint of the axis—50 percent—represents half the students approving of the code and the other half disapproving.

To choose our sample, we give each student on the student roster a number and select 100 random

**FIGURE 7-9**
Percent of students approving of the student code

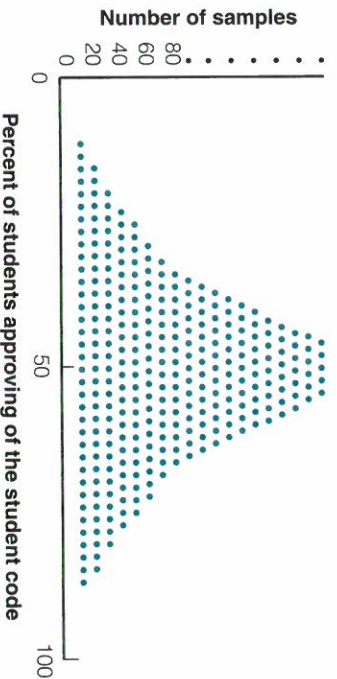
Results Produced by Three Hypothetical Studies. Assuming a large student body, let's suppose that we selected three different samples, each of substantial size. We would not necessarily expect those samples to perfectly reflect attitudes of the whole student body, but they should come reasonably close.

numbers from a table of random numbers. Then we interview the 100 students whose numbers have been selected and ask for their attitudes toward the student code: whether they approve or disapprove. Suppose this operation gives us 48 students who approve of the code and 52 who disapprove. This summary description of a variable in a sample is called a **statistic**. We present this statistic by placing a dot on the x axis at the point representing 48 percent.

Now let's suppose we select another sample of 100 students in exactly the same fashion and measure their approval or disapproval of the student code. Perhaps 51 students in the second sample approve of the code. We place another dot in the appropriate place on the x axis. Repeating this process once more, we may discover that 52 students in the third sample approve of the code.

Figure 7-9 presents the three different sample statistics representing the percentages of students in each of the three random samples who approved of the student code. The basic rule of random sampling is that such samples drawn from a population give estimates of the parameter that exists in the total population. Each of the random samples, then, gives us an estimate of the percentage of students in the total student body who approve of the student code. Unhappily, however, we have

statistic The summary description of a variable in a sample, used to estimate a population parameter.

**FIGURE 7-10**

The Sampling Distribution. If we were to select a large number of good samples, we would expect them to cluster around the true value (50 percent), but given enough such samples, a few would fall far from the mark.

selected three samples and now have three separate estimates.

To retrieve ourselves from this problem, let's draw more and more samples of 100 students each, question each of the samples concerning their approval or disapproval of the code, and plot the new sample statistics on our summary graph. In drawing many such samples, we discover that some of the new samples provide duplicate estimates, as in the illustration of ten cases. Figure 7-10 shows the sampling distribution of, say, hundreds of samples. This is often referred to as a *normal curve*.

Note that by increasing the number of samples selected and interviewed, we've also increased the range of estimates provided by the sampling operation. In one sense we've increased our dilemma in attempting to guess the parameter in the population. Probability theory, however, provides certain important rules regarding the sampling distribution presented in Figure 7-10.

First, if many independent random samples are selected from a population, the sample statistics provided by those samples will be distributed around the population parameter in a known way.

sampling error The degree of error to be expected by virtue of studying a sample instead of everyone. For probability sampling, the maximum error depends on three factors: the sample size, the diversity of the population, and the confidence level.

Thus, although Figure 7-10 shows a wide range of estimates, more of them are in the vicinity of 50 percent than elsewhere in the graph. Probability theory tells us, then, that the true value is in the vicinity of 50 percent.

Second, probability theory gives us a formula for estimating how closely the sample statistics are clustered around the true value. To put it another way, probability theory enables us to estimate the **sampling error**—the degree of error to be expected for a given sample design. This formula contains three factors: the parameter, the sample size, and the standard error (a measure of sampling error):

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

The symbols P and Q in the formula equal the population parameters for the binomial: If 60 percent of the student body approve of the code and 40 percent disapprove, P and Q are 60 percent and 40 percent, respectively, or 0.6 and 0.4. Note that $Q = 1 - P$ and $P = 1 - Q$. The symbol n equals the number of cases in each sample, and s is the standard error.

Let's assume that the population parameter in the student example is 50 percent approving of the code and 50 percent disapproving. Recall that we've been selecting samples of 100 cases each. When these numbers are put into the formula, we find that the standard error equals 0.05, or 5 percent.

In probability theory, the standard error is a valuable piece of information because it indicates the extent to which the sample estimates will be distributed around the population parameter. (If you're familiar with the standard deviation in statistics, you may recognize that the standard error, in this case, is the standard deviation of the sampling distribution.) Specifically, probability theory indicates that certain proportions of the sample estimates will fall within specified increments—each equal to one standard error—from the population parameter. Approximately 34 percent (0.3413) of the sample estimates will fall within one standard error increment above the population parameter, and another 34 percent will fall within one standard error below the parameter. In our example, the standard error increment is 5 percent, so we know that 34 percent of our samples will give estimates of student approval between 50 percent (the parameter) and 55 percent (one standard error above); another 34 percent of the samples will give estimates between 50 percent and 45 percent (one standard error below the parameter). Taken together, then, we know that roughly two-thirds (68 percent) of the samples will give estimates within 5 percent of the parameter.

Moreover, probability theory dictates that roughly 95 percent of the samples will fall within plus or minus two standard errors of the true value, and 99.9 percent of the samples will fall within plus or minus three standard errors. In our present example, then, we know that only one sample out of a thousand would give an estimate lower than 35 percent approval or higher than 65 percent.

The proportion of samples falling within one, two, or three standard errors of the parameter is constant for any random sampling procedure such as the one just described, providing that a large number of samples are selected. The size of the standard error in any given case, however, is a function of the population parameter and the sample size. If we return to the formula for a moment, we note that the standard error will increase as a function of an increase in the quantity P times Q . Note further that this quantity reaches its maximum in the situation of an even split in the population. If $P = 0.5$, $PQ = 0.25$; if $P = 0.6$, $PQ = 0.24$;

if $P = 0.8$, $PQ = 0.16$; if $P = 0.99$, $PQ = 0.0099$. By extension, if P is either 0.0 or 1.0 (either 0 percent or 100 percent approval of the student code), the standard error will be 0. If everyone in the population has the same attitude (no variation), then every sample will give exactly that estimate.

The standard error is also a function of the sample size—an inverse function. As the sample size increases, the standard error decreases. As the sample size increases, the several samples will be clustered nearer to the true value. Another general guideline is evident in the formula: Because of the square-root formula, the standard error is reduced by half if the sample size is quadrupled. In our present example, samples of 100 produce a standard error of 5 percent; to reduce the standard error to 2.5 percent, we must increase the sample size to 400.

All of this information is provided by established probability theory in reference to the selection of large numbers of random samples. (If you've taken a statistics course, you may know this as the Central Tendency Theorem.) If the population parameter is known and many random samples are selected, we can predict how many of the sample estimates will fall within specified intervals from the parameter.

Recognize that this discussion illustrates only the logic of probability sampling; it does not describe the way research is actually conducted. Usually, we don't know the parameter: The very reason we conduct a sample survey is to estimate that value. Moreover, we don't actually select large numbers of samples: We select only one sample. Nevertheless, the preceding discussion of probability theory provides the basis for inferences about the typical social research situation. Knowing what it would be like to select thousands of samples allows us to make assumptions about the one sample we do select and study.

Confidence Levels and Confidence Intervals

Whereas probability theory specifies that 68 percent of that fictitious large number of samples would produce estimates falling within one standard error of the parameter, we can turn the logic around and

infer that any single random sample estimate has a 68 percent chance of falling within that range. This observation leads us to the two key components of sampling error estimates: **confidence level** and **confidence interval**. We express the accuracy of our sample statistics in terms of a level of confidence that the statistics fall within a specified interval from the parameter. For example, we may say we are 95 percent confident that our sample statistics (for example, 50 percent favor the new student code) are within plus or minus 5 percentage points of the population parameter. As the confidence interval is expanded for a given statistic, our confidence increases. For example, we may say that we are 99.9 percent confident that our statistic falls within three standard errors of the true value. (Now perhaps you can appreciate the humorous quip of unknown origin: Statistics means never having to say you are certain.)

Although we may be confident (at some level) of being within a certain range of the parameter, we've already noted that we seldom know what the parameter is. To resolve this problem, we substitute our sample estimate for the parameter in the formula; that is, lacking the true value, we substitute the best available guess.

The result of these inferences and estimations is that we can estimate a population parameter and also the expected degree of error on the basis of one sample drawn from a population. Beginning with the question "What percentage of the student body approves of the student code?" you could select a random sample of 100 students and interview them. You might then report that your best estimate is that 50 percent of the student body approves of the code and that you are 95 percent confident that between 40 and 60 percent (plus or minus two standard errors) approve. The range

from 40 to 60 percent is the confidence interval. (At the 68 percent confidence level, the confidence interval would be 45–55 percent.)

The logic of confidence levels and confidence intervals also provides the basis for determining the appropriate sample size for a study. Once you've decided on the degree of sampling error you can tolerate, you'll be able to calculate the number of cases needed in your sample. Thus, for example, if you want to be 95 percent confident that your study findings are accurate within plus or minus 5 percentage points of the population parameters, you should select a sample of at least 400. (Appendix F is a convenient guide in this regard.)

This, then, is the basic logic of probability sampling. Random selection permits the researcher to link findings from a sample to the body of probability theory so as to estimate the accuracy of those findings. All statements of accuracy in sampling must specify both a confidence level and a confidence interval. The researcher must report that he or she is x percent confident that the population parameter is between two specific values. In this example, I've demonstrated the logic of sampling error using a variable analyzed in percentages. A different statistical procedure would be required to calculate the standard error for a mean, for example, but the overall logic is the same.

Notice that nowhere in this discussion of sample size and accuracy of estimates did we consider the size of the population being studied. This is because the population size is almost always irrelevant. A sample of 2,000 respondents drawn properly to represent Vermont voters will be no more accurate than a sample of 2,000 drawn properly to represent all voters in the United States, even though the Vermont sample would be a substantially larger proportion of that small state's voters than would the same number chosen to represent the nation's voters. The reason for this counterintuitive fact is that the equations for calculating sampling error all assume that the populations being sampled are infinitely large, so every sample would equal 0 percent of the whole.

Of course, this is not literally true in practice. However, a sample of 2,000 represents only 0.61 percent of the Vermonters who voted for president

confidence level The estimated probability that a population parameter lies within a given confidence interval. Thus, we might be 95 percent confident that between 35 and 45 percent of all voters favor Candidate A.

confidence interval The range of values within which a population parameter is estimated to lie.

in the 2008 election, and a sample of 2,000 U.S. voters represents a mere 0.0015 percent of the national electorate. Both of these proportions are sufficiently small as to approach the situation with infinitely large populations.

Unless a sample represents, say, 5 percent or more of the population it's drawn from, that proportion is irrelevant. In those rare cases of large proportions being selected, a "finite population correction" can be calculated to adjust the confidence intervals. The following formula calculates the proportion to be multiplied against the calculated error.

$$\text{finite population correction} = \sqrt{\frac{N - n}{N - 1}}$$

In the formula, N is the population size and n is the size of the sample. Notice that in the extreme case where you studied the whole population (hence $N = n$), the formula would yield zero as the finite population correction. Multiplying zero times the sampling error calculated by the earlier formula would give a final sampling error of zero, which would, of course, be precisely the case since you wouldn't have sampled at all.

Lest you weary of the statistical nature of this discussion, it is useful to realize what an amazing thing we have been examining. There is remarkable order within what might seem random and chaotic. One of the researchers to whom we owe this observation is Sir Francis Galton (1822–1911),

Order in Apparent Chaos—I know of scarcely anything so apt to impress the imagination as the wonderful form of cosmic order expressed by the "Law of Frequency of Error." The law would have been personified by the Greeks and deified, if they had known of it. It reigns with serenity and in complete self-effacement amidst the wildest confusion. The huger the mob, and the greater the apparent anarchy, the more perfect is its sway. It is the supreme law of Unreason (1889: 66).

Two cautions are in order before we conclude this discussion of the basic logic of probability sampling. First, the survey uses of probability theory as discussed here are technically not wholly justified.

The theory of sampling distribution makes assumptions that almost never apply in survey conditions. The exact proportion of samples contained within specified increments of standard errors, for example, mathematically assumes an infinitely large population, an infinite number of samples, and sampling with replacement—that is, every sampling unit selected is "thrown back into the pot" and could be selected again. Second, our discussion has greatly oversimplified the inferential jump from the distribution of several samples to the probable characteristics of one sample.

I offer these cautions to provide perspective on the uses of probability theory in sampling. Social researchers often appear to overestimate the precision of estimates produced by the use of probability theory. As I'll mention elsewhere in this chapter and throughout the book, variations in sampling techniques and nonsampling factors may further reduce the legitimacy of such estimates. For example, those selected in a sample who fail or refuse to participate detract further from the representativeness of the sample.

Nevertheless, the calculations discussed in this section can be extremely valuable to you in understanding and evaluating your data. Although the calculations do not provide as precise estimates as some researchers might assume, they can be quite valid for practical purposes. They are unquestionably more valid than less-rigorously derived estimates based on less-rigorous sampling methods. Most important, being familiar with the basic logic underlying the calculations can help you react sensibly both to your own data and to those reported by others.

Populations and Sampling Frames

The preceding section introduced the theoretical model for social research sampling. Although as students, research consumers, and researchers we need to understand that theory, it's no less important to appreciate the less-than-perfect conditions that exist in the field. In this section we'll look at one aspect of field conditions that requires a

compromise with idealized theoretical conditions and assumptions: the congruence of or disparity between populations of sampling frames.

Simply put, a **sampling frame** is the list or quasi list of elements from which a probability sample is selected. If a sample of students is selected from a student roster, the roster is the sampling frame. If the primary sampling unit for a complex population sample is the census block, the list of census blocks composes the sampling frame—in the form of a printed booklet or, better, some digital format permitting computer manipulation. Here are some reports of sampling frames appearing in research journals. In each example I've italicized the actual sampling frames.

We purchased a list of 50,000 Maryland residents who were registered to vote from Aristotle, which maintains a national database including 175 million registered voters. We refer to these residents as “registered voters” even though some of them have not actually gone to the polls in some time. The Aristotle database is compiled from state records, county boards of elections, state boards of registrars, etc.

(Tourangeau et al. 2010: 416)

Respondents were undergraduates enrolled in introductory psychology classes at Ohio State University in spring 2001.

(Chang and Krosnick, 2010: 135)

The data reported in this paper . . . were gathered from a probability sample of *adults aged 18 and over residing in households in the 48 contiguous United States*. Personal interviews with 1,914 respondents were conducted by the Survey Research Center of the University of Michigan during the fall of 1975.

(Jackman and Senter 1980: 345)

sampling frame That list or quasi list of units composing a population from which a sample is selected. If the sample is to be representative of the population, it is essential that the sampling frame include all (or nearly all) members of the population.

Properly drawn samples provide information appropriate for describing the population of elements composing the sampling frame—nothing more. I emphasize this point in view of the all-too-common tendency for researchers to select samples from a given sampling frame and then make assertions about a population similar to, but not identical to, the population defined by the sampling frame.

For example, take a look at this report, which discusses the drugs most frequently prescribed by U.S. physicians:

Information on prescription drug sales is not easy to obtain. But Rinaldo V. DeNuzzo, a professor of pharmacy at the Albany College of Pharmacy, Union University, Albany, NY, has been tracking prescription drug sales for 25 years by polling nearby drugstores. He publishes the results in an industry trade magazine, *MM@M*.

DeNuzzo's latest survey, covering 1980, is based on reports from 66 pharmacies in 48 communities in New York and New Jersey. Unless there is something peculiar about that part of the country, his findings can be taken as representative of what happens across the country.

(Moskowitz 1981: 33)

What is striking in the excerpt is the casual comment about whether there is anything peculiar about New York and New Jersey. There is. The lifestyle in these two states hardly typifies the other 48. We cannot assume that residents in these large, urbanized, eastern seaboard states necessarily have the same drug-use patterns that residents of Mississippi, Nebraska, or Vermont do.

Does the survey even represent prescription patterns in New York and New Jersey? To determine that, we would have to know something about the way the 48 communities and the 66 pharmacies were selected. We should be wary in this regard, in view of the reference to “polling nearby drugstores.” As we'll see, there are several methods for selecting samples that ensure representativeness, and unless they're used, we shouldn't generalize from the study findings.

A sampling frame, then, covers the population we wish to study. In the simplest sample design, the sampling frame is a list of the elements composing the study population. In practice, though, existing sampling frames often define the study population rather than the other way around. That is, we often begin with a population in mind for our study; then we search for possible sampling frames. Having examined and evaluated the frames available for our use, we decide which frame presents a study population most appropriate to our needs.

Studies of organizations are often the simplest from a sampling standpoint because organizations typically have membership lists. In such cases, the list of members constitutes an excellent sampling frame. If a random sample is selected from a membership list, the data collected from that sample may be taken as representative of all members—if all members are included in the list.

Populations that can be sampled from good organizational lists include elementary school, high school, and university students and faculty; church members; factory workers; fraternity or sorority members; members of social, service, or political clubs; and members of professional associations.

The preceding comments apply primarily to local organizations. Often, statewide or national organizations do not have a single membership list. There is, for example, no single list of Episcopalian church members. However, a slightly more complex sample design could take advantage of local church membership lists by first sampling churches and then subsampling the membership lists of those churches selected. (More about that later.)

Other lists of individuals may be especially relevant to the research needs of a particular study. Government agencies maintain lists of registered voters, for example, and some political pollsters use registration-based sampling (RBS), using those lists. In some cases, there may be delays in keeping such files up-to-date, and a person who is registered to vote may not actually do so in the election of interest.

Other lists that may be available contain the names of automobile owners, welfare recipients,

taxpayers, business permit holders, licensed professionals, and so forth. Although it may be difficult to gain access to some of these lists, they provide excellent sampling frames for specialized research purposes.

Of course, the sampling elements in a study need not be individuals. Social researchers might use lists of universities, businesses, cities, academic journals, newspapers, unions, political clubs, professional associations, and so forth.

Telephone directories were once used for “quick-and-dirty” public opinion polls. They’re easy and inexpensive to use—no doubt the reason for their popularity. And, if you want to make assertions about telephone subscribers, the directory is a fairly good sampling frame. (Realize, of course, that a given directory will not include new subscribers or those who have requested unlisted numbers.

Sampling is further complicated by the directories’ inclusion of nonresidential listings.) Unfortunately, telephone directories are all too often used as a listing of a city’s population or of its voters. Of the many defects in this reasoning, the chief one involves a bias, as we have seen. Poor people are less likely to have telephones; rich people may have more than one line. A telephone directory sample, therefore, is likely to have a middle- or upper-class bias. As we’ll see a little later, the telephone directory may produce an age bias, since many young people have only cell phones.

The class bias inherent in telephone directory samples is often hidden. Pre-election polls conducted in this fashion are sometimes quite accurate, perhaps because of the class bias evident in voting itself. Poor people are less likely to vote. Frequently, then, these two biases nearly coincide, so that the results of a telephone poll may come very close to the final election outcome. Unhappily, you never know for sure until after the election. And sometimes, as in the case of the 1936 *Literary Digest* poll, you may discover that the voters have not acted according to the expected class biases. The ultimate disadvantage of this method, then, is the researcher’s inability to estimate the degree of error to be expected in the sample findings.

In Chapter 9 we’ll return to the matter of sampling telephones, in connection with survey

research. We'll examine random-digit dialing, which was developed to resolve some of the problems just discussed, and we'll see that the growth in popularity of cell phones has further complicated matters.

Street directories and tax maps are sometimes used for easy samples of households, but they may present incompleteness and bias. For example, in strictly zoned urban regions, illegal housing units are unlikely to appear on official records. As a result, such units could not be selected, and sample findings could not be representative of those units, which are often poorer and more crowded than the average.

The preceding comments apply to the United States but not to all countries. In Japan, for example, the government maintains quite accurate population registration lists. Moreover, citizens are required by law to keep their information up-to-date, such as changes in residence or births and deaths in the household. As a consequence, you can select simple random samples of the population more easily in Japan than in the United States. Such a registration list in the United States would conflict directly with this country's norms regarding individual privacy.

In recent years, American researchers have begun experimenting with address files maintained by the U.S. Postal Service, such as the Special Delivery Sequence File. As problems have increasingly arisen with regard to the sampling of telephone numbers (discussed further in Chapter 9), address-based sampling (ABS) for use in mail surveys has been improving (Link et al. 2008).

Review of Populations and Sampling Frames

Because social research literature gives surprisingly little attention to the issues of populations and sampling frames, I've devoted special attention to them. Here is a summary of the main guidelines to remember:

1. Findings based on a sample can be taken as representing only the aggregation of elements that compose the sampling frame.

2. Often, sampling frames do not truly include all the elements their names might imply. Omissions are almost inevitable. Thus, a first concern of the researcher must be to assess the extent of the omissions and to correct them if possible. (Of course, the researcher may feel that he or she can safely ignore a small number of omissions that cannot easily be corrected.)
3. To be generalized even to the population composing the sampling frame, all elements must have equal representation in the frame. Typically, each element should appear only once. Elements that appear more than once will have a greater probability of selection, and the sample will, overall, overrepresent those elements.

Other, more practical matters relating to populations and sampling frames will be treated elsewhere in this book. For example, the form of the sampling frame—such as a list in a publication, a 3-by-5 card file, CD-ROM, or USB storage drive—can affect how easy it is to use. And ease of use may often take priority over scientific considerations: An “easier” list may be chosen over a “harder” one, even though the latter is more appropriate to the target population. We should not take a dogmatic position in this regard, but every researcher should carefully weigh the relative advantages and disadvantages of such alternatives.

Types of Sampling Designs

Up to this point, we've focused on simple random sampling. Indeed, the body of statistics typically used by social researchers assumes such a sample. As you'll see shortly, however, you have several options in choosing your sampling method, and you'll seldom if ever choose simple random sampling. There are two reasons for this. First, with all but the simplest sampling frame, simple random sampling is not feasible. Second, and probably surprisingly, simple random sampling may not be the most accurate method available. Let's turn now to a discussion of simple random sampling and the other options available.

Simple Random Sampling

As noted, **simple random sampling (SRS)** is the basic sampling method assumed in the statistical computations of social research. Because the mathematics of random sampling are especially complex, we'll detour around them in favor of describing the ways of employing this method in the field.

Once a sampling frame has been properly established, to use simple random sampling the researcher assigns a single number to each element in the list, not skipping any number in the process. A table of random numbers (Appendix C) is then used to select elements for the sample. See the Tips and Tools feature, "Using a Table of Random Numbers" for more about this process.

If your sampling frame is in a machine-readable form, such as CD-ROM or USB storage drive, a computer can automatically select a simple random sample. (In effect, the computer program numbers the elements in the sampling frame, generates its own series of random numbers, and prints out the list of elements selected.)

Figure 7-11 offers a graphic illustration of simple random sampling. Note that the members of our hypothetical micropopulation have been numbered from 1 to 100. Moving to Appendix C, we decide to use the last two digits of the first column and to begin with the third number from the top. This yields person number 30 as the first one selected into the sample. Number 67 is next, and so forth. (Person 100 would have been selected if "00" had come up in the list.)

Systematic Sampling

Simple random sampling is seldom used in practice. As you'll see, it's not usually the most efficient method, and it can be laborious if done manually. Typically, simple random sampling requires a list of elements. When such a list is available, researchers usually employ systematic sampling instead.

In **systematic sampling**, every k th element in the total list is chosen (systematically) for inclusion in the sample. If the list contained 10,000 elements and you wanted a sample of 1,000, you would

select every tenth element for your sample. To ensure against any possible human bias in using this method, you should select the first element at random. Thus, in the preceding example, you would begin by selecting a random number between one and ten. The element having that number is included in the sample, plus every tenth element following it. This method is technically referred to as a *systematic sample with a random start*. Two terms are frequently used in connection with systematic sampling. The **sampling interval** is the standard distance between elements selected in the sample; ten in the preceding sample. The **sampling ratio** is the proportion of elements in the population that are selected: 1/10 in the example.

$$\begin{aligned}\text{sampling interval} &= \frac{\text{population size}}{\text{sample size}} \\ \text{sampling ratio} &= \frac{\text{sample size}}{\text{population size}}\end{aligned}$$

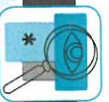
In practice, systematic sampling is virtually identical to simple random sampling. If the list of elements is indeed randomized before sampling, one might argue that a systematic sample drawn from that list is in fact a simple random sample. By now, debates over the relative merits of simple random sampling and systematic sampling have

simple random sampling (SRS) A type of probability sampling in which the units composing a population are assigned numbers. A set of random numbers is then generated, and the units having those numbers are included in the sample.

systematic sampling A type of probability sampling in which every k th unit in a list is selected for inclusion in the sample—for example, every 25th student in the college directory of students. You compute k by dividing the size of the population by the desired sample size; k is called the sampling interval. Within certain constraints, systematic sampling is a functional equivalent of simple random sampling and is usually easier to do. Typically, the first unit is selected at random.

sampling interval The standard distance between elements selected from a population for a sample.

sampling ratio The proportion of elements in the population that are selected to be in a sample.



Tips and Tools

Using a Table of Random Numbers

In social research, it's often appropriate to select a set of random numbers from a table such as the one in Appendix C. Here's how to do that.

Suppose you want to select a simple random sample of 100 people (or other units) out of a population totaling 980.

1. To begin, number the members of the population: in this case, from 1 to 980. Now the task is to select 100 random numbers. Once you've done that, your sample will consist of the people having the numbers you've selected. (*Note:* It's not essential to actually number them, as long as you're sure of the total. If you have them in a list, for example, you can always count through the list after you've selected the numbers.)
 2. The next step is to determine the number of digits you'll need in the random numbers you select. In our example, there are 980 members of the population, so you'll need three-digit numbers to give everyone a chance of selection. (If there were 11,825 members of the population, you'd need to select five-digit numbers.) Thus, we want to select 100 random numbers in the range from 001 to 980.
 3. Now turn to the first page of Appendix C. Notice there are several rows and columns of five-digit numbers, and are two pages, with the columns continuing from the first page to the second. The table represents a series of random numbers in the range from 00001 to 99999. To use the table for your hypothetical sample, you have to answer these questions:
 - a. How will you create three-digit numbers out of five-digit numbers?
 - b. What pattern will you follow in moving through the table to select your numbers?
 - c. Where will you start?
- Each of these questions has several satisfactory answers. The key is to create a plan and follow it. Here's an example.
4. To create three-digit numbers from five-digit numbers, let's agree to select five-digit numbers from the table but consider only the left-most three digits in each case. If we picked the first number on the first page—10480—we'd consider only the 104. (We could agree to take the digits farthest to the right, 480, or the middle three digits, 048, and any of these plans would work.) They key is to make a plan and stick with it. For convenience, let's use the left-most three digits.
 5. We can also choose to progress through the tables any way we want: down the columns, up them, across to the right or to the left, or diagonally. Again, any of these plans will work just fine as long as we stick to it. For convenience, let's agree to move down the columns. When we get to the bottom of one column, we'll go to the top of the next.
 6. Now, where do we start? You can close your eyes and stick a pencil into the table and start wherever the pencil point lands. (I know it doesn't sound scientific, but it works.) Or, if you're afraid you'll hurt the book or miss it altogether, close your eyes and make up a column number and a row number. ("I'll pick the number in the fifth row of column 2.") Start with that number.
 7. Let's suppose we decide to start with the fifth number in column 2. If you look on the first page of Appendix C, you'll see that the starting number is 39975. We've selected 399 as our first random number, and we have 99 more to go. Moving down the second column, we select 069, 729, 919, 143, 368, 695, 409, 939, and so forth. At the bottom of column 2 (on the second page of this table), we select number 017 and continue to the top of column 3: 015, 255, and so on.
 8. See how easy it is? But trouble lies ahead. When we reach column 5, we're speeding along, selecting 816, 309, 763, 078, 061, 277, 988 . . . Wait a minute! There are only 980 students in the senior class. How can we pick number 988? The solution is simple: Ignore it. Any time you come across a number that lies outside your range, skip it and continue on your way: 188, 174, and so forth. The same solution applies if the same number comes up more than once. If you select 399 again, for example, just ignore it the second time.
 9. That's it. You keep up the procedure until you've selected 100 random numbers. (Returning to your list, your sample consists of person number 399, person number 69, person number 729, and so forth.

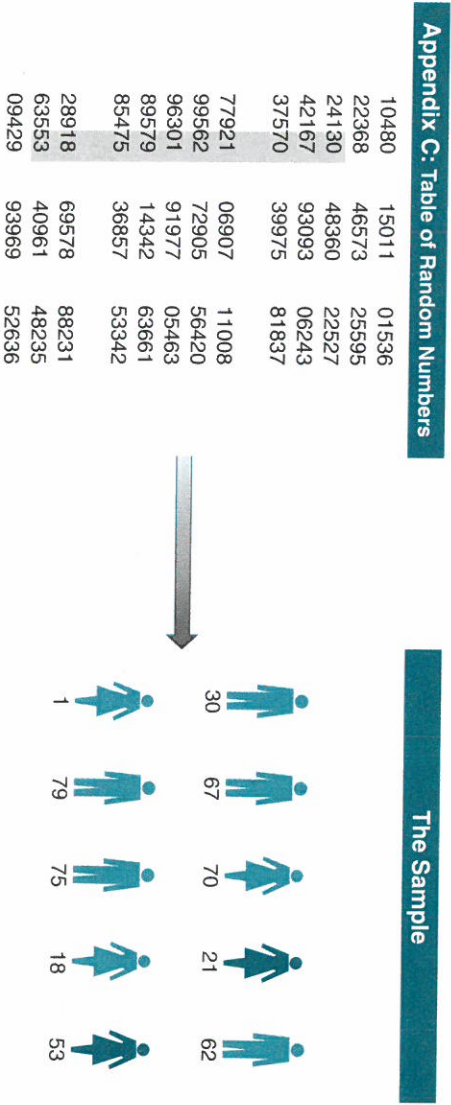
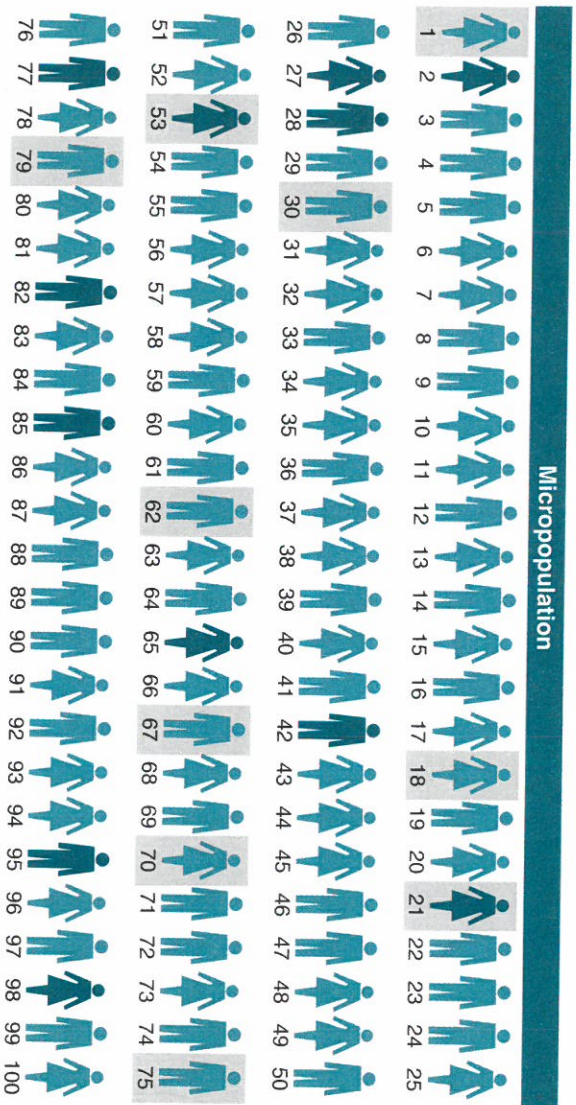


FIGURE 7-11

A Simple Random Sample. Having numbered everyone in the population, we can use a table of random numbers to select a representative sample from the overall population. Anyone whose number is chosen from the table is in the sample.

been resolved largely in favor of the latter, simpler method. Empirically, the results are virtually identical. And, as you'll see in a later section, systematic sampling, in some instances, is slightly more accurate than simple random sampling.

There is one danger involved in systematic sampling. The arrangement of elements in the list can make systematic sampling unwise. Such an arrangement is usually called periodicity. If the list

of elements is arranged in a cyclical pattern that coincides with the sampling interval, a grossly biased sample might be drawn. Here are two examples that illustrate this danger.

In a classic study of soldiers during World War II, the researchers selected a systematic sample from unit rosters. Every tenth soldier on the roster was selected for the study. The rosters, however, were arranged in a table of organizations: sergeants

first, then corporals and privates, squad by squad. Each squad had ten members. As a result, every tenth person on the roster was a squad sergeant. The systematic sample selected contained only sergeants. It could, of course, have been the case that no sergeants were selected for the same reason.

As another example, suppose we select a sample of apartments in an apartment building. If the sample is drawn from a list of apartments arranged in numerical order (for example, 101, 102, 103, 104, 201, 202, and so on), there is a danger of the sampling interval coinciding with the number of apartments on a floor or some multiple thereof. Then the samples might include only northwest-corner apartments or only apartments near the elevator. If these types of apartments have some other particular characteristic in common (for example, higher rent), the sample will be biased. The same danger would appear in a systematic sample of houses in a subdivision arranged with the same number of houses on a block.

In considering a systematic sample from a list, then, you should carefully examine the nature of that list. If the elements are arranged in any particular order, you should figure out whether that order will bias the sample to be selected, then you should take steps to counteract any possible bias (for example, take a simple random sample from cyclical portions).

Usually, however, systematic sampling is superior to simple random sampling, in convenience if nothing else. Problems in the ordering of elements in the sampling frame can usually be remedied quite easily.

Stratified Sampling

So far we've discussed two methods of sample selection from a list: random and systematic.

stratification The grouping of the units composing a population into homogeneous groups (or strata) before sampling. This procedure, which may be used in conjunction with simple random, systematic, or cluster sampling, improves the representativeness of a sample, at least in terms of the stratification variables.

Stratification is not an alternative to these methods; rather, it represents a possible modification of their use.

Simple random sampling and systematic sampling both ensure a degree of representativeness and permit an estimate of the error present. Stratified sampling is a method for obtaining a greater degree of representativeness by decreasing the probable sampling error. To understand this method, we must return briefly to the basic theory of sampling distribution.

Recall that sampling error is reduced by two factors in the sample design. First, a large sample produces a smaller sampling error than a small sample does. Second, a homogeneous population produces samples with smaller sampling errors than a heterogeneous population does. If 99 percent of the population agrees with a certain statement, it's extremely unlikely that any probability sample will greatly misrepresent the extent of agreement. If the population is split 50–50 on the statement, then the sampling error will be much greater.

Stratified sampling is based on this second factor in sampling theory. Rather than selecting a sample from the total population at large, the researcher ensures that appropriate numbers of elements are drawn from homogeneous subsets of that population. To get a stratified sample of university students, for example, you would first organize your population by college class and then draw appropriate numbers of freshmen, sophomores, juniors, and seniors. In a nonstratified sample, representation by class would be subjected to the same sampling error as other variables would. In a sample stratified by class, the sampling error on this variable is reduced to zero.

More-complex stratification methods are also possible. In addition to stratifying by class, you might also stratify by sex, by GPA, and so forth. In this fashion you might be able to ensure that your sample would contain the proper numbers of male sophomores with a 3.5 average, of female sophomores with a 4.0 average, and so forth.

The ultimate function of stratification, then, is to organize the population into homogeneous subsets (with heterogeneity between subsets)

and to select the appropriate number of elements from each. To the extent that the subsets are homogeneous on the stratification variables, they may be homogeneous on other variables as well. Because age is related to college class, a sample stratified by class will be more representative in terms of age as well, compared with an unstratified sample. Because occupational aspirations still seem to be related to sex, a sample stratified by sex will be more representative in terms of occupational aspirations.

The choice of stratification variables typically depends on what variables are available. Sex can often be determined in a list of names. University lists are typically arranged by class. Lists of faculty members may indicate their departmental affiliation. Government agency files may be arranged by geographic region. Voter registration lists are arranged according to precinct.

In selecting stratification variables from among those available, however, you should be concerned primarily with those that are presumably related to variables you want to represent accurately. Because sex is related to many variables and is often available for stratification, it's often used. Education is related to many variables, but it's often not available for stratification. Geographic location within a city, state, or nation is related to many things.

Within a city, stratification by geographic location usually increases representativeness in social class, ethnic group, and so forth. Within a nation, it increases representativeness in a broad range of attitudes as well as in social class and ethnicity.

When you're working with a simple list of all elements in the population, two methods of stratification predominate. In one method, you sort the population elements into discrete groups based on whatever stratification variables are being used. On the basis of the relative proportion of the population represented by a given group, you select—randomly or systematically—several elements from that group constituting the same proportion of your desired sample size. For example, if sophomore men with a 4.0 average compose 1 percent of the student population and you desire a sample of 1,000 students, you would select 10 sophomore men with a 4.0 average.

The other method is to group students as described and then put those groups together in a continuous list, beginning with all freshmen men with a 4.0 average and ending with all senior women with a 1.0 or below. You would then select a systematic sample, with a random start, from the entire list. Given the arrangement of the list, a systematic sample would select proper numbers (within an error range of 1 or 2) from each subgroup. (*Note:* A simple random sample drawn from such a composite list would cancel out the stratification.)

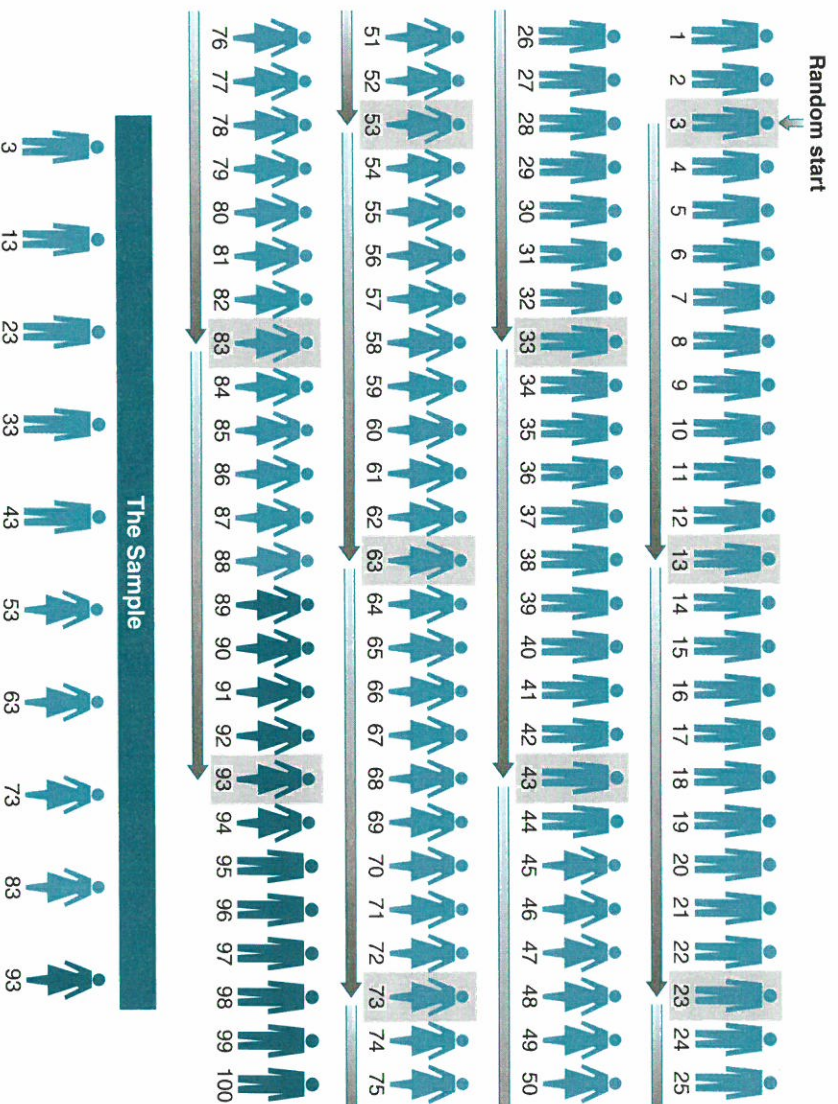
Figure 7-12 offers a graphic illustration of stratified, systematic sampling. As you can see, we lined up our micropopulation according to sex and race. Then, beginning with a random start of “3,” we've taken every tenth person thereafter: 3, 13, 23, . . . , 93.

Stratified sampling ensures the proper representation of the stratification variables; this, in turn, enhances the representation of other variables related to them. Taken as a whole, then, a stratified sample is more likely than a simple random sample to be more representative on several variables. Although the simple random sample is still regarded as somewhat sacred, it should now be clear that you can often do better.

Implicit Stratification in Systematic Sampling

I mentioned that systematic sampling can, under certain conditions, be more accurate than simple random sampling. This is the case whenever the arrangement of the list creates an implicit stratification. As already noted, if a list of university students is arranged by class, then a systematic sample provides a stratification by class where a simple random sample would not.

In a study of students at the University of Hawaii, after stratification by school class, the students were arranged by their student identification numbers. These numbers, however, were their social security numbers. The first three digits of the social security number indicate the state in which the number was issued. As a result, within a class, students were arranged by the state in which they

**FIGURE 7-12**

A Stratified, systematic Sample with a Random Start. A stratified, systematic sample involves two stages. First the members of the population are gathered into homogeneous strata; this simple example merely uses sex and race as a stratification variable, but more could be used. Then every k th (in this case, every tenth) person in the stratified arrangement is selected into the sample.

were issued a social security number, providing a rough stratification by geographic origin.

An ordered list of elements, therefore, may be more useful to you than an unordered, randomized list. I've stressed this point in view of the unfortunate belief that lists should be randomized before systematic sampling. Only if the arrangement presents the problems discussed earlier should the list be rearranged.

Illustration: Sampling University Students

Let's put these principles into practice by looking at an actual sampling design used to select a

sample of university students. The purpose of the study was to survey, with a mail-out questionnaire, a representative cross section of students attending the main campus of the University of Hawaii. The following sections describe the steps and decisions involved in selecting that sample.

Study Population and Sampling Frame

The obvious sampling frame available for use in this sample selection was the computerized file maintained by the university administration. The file contained students' names, local and permanent addresses, and social security numbers, as well as a variety of other information such as field of study, class, age, and sex.

The computer database, however, contained entries on all people who could, by any conceivable definition, be called students, many of whom seemed inappropriate for the purposes of the study. As a result, researchers needed to define the study population in a somewhat more restricted fashion. The final definition included those 15,225 day-program degree candidates who were registered for the fall semester on the Manoa campus of the university, including all colleges and departments, both undergraduate and graduate students, and both U.S. and foreign students. The computer program used for sampling then limited consideration to students fitting this definition.

Stratification

The sampling program also permitted stratification of students before sample selection. The researchers decided that stratification by college class would be sufficient, although the students might have been further stratified within class, if desired, by sex, college, major, and so forth.

Sample Selection

Once the students had been arranged by class, a systematic sample was selected across the entire rearranged list. The sample size for the study was initially set at 1,100. To achieve this sample, the sampling program was set for a 1/14 sampling ratio. The program generated a random number between 1 and 14; the student having that number and every 14th student thereafter was selected in the sample.

Once the sample had been selected, the computer was instructed to print each student's name and mailing address on self-adhesive mailing labels. These labels were then simply transferred to envelopes for mailing the questionnaires.

Sample Modification

This initial design of the sample had to be modified. Before the mailing of questionnaires, the researchers discovered that, because of unexpected expenses in the production of the questionnaires, they couldn't cover the costs of mailing to all 1,100 students. As a result, one-third of the mailing labels

were systematically selected (with a random start) for exclusion from the sample. The final sample for the study was thereby reduced to 733 students.

I mention this modification in order to illustrate the frequent need to alter a study plan in midstream. Because the excluded students were systematically omitted from the initial systematic sample, the remaining 733 students could still be taken as reasonably representing the study population. The reduction in sample size did, of course, increase the range of sampling error.

Multistage Cluster Sampling

The preceding sections have dealt with reasonably simple procedures for sampling from lists of elements. Such a situation is ideal. Unfortunately, however, much interesting social research requires the selection of samples from populations that cannot easily be listed for sampling purposes: the population of a city, state, or nation; all university students in the United States; and so forth. In such cases, the sample design must be much more complex. Such a design typically involves the initial sampling of groups of elements—*clusters*—followed by the selection of elements within each of the selected clusters.

Cluster sampling may be used when it's either impossible or impractical to compile an exhaustive list of the elements composing the target population, such as all church members in the United States. Often, however, the population elements are already grouped into subpopulations, and a list of those subpopulations either exists or can be created practically. For example, church members in the United States belong to discrete churches, which are either listed or could be. Following a cluster sample format, then, researchers

cluster sampling A multistage sampling in which natural groups (clusters) are sampled initially, with the members of each selected group being subsampled afterward. For example, you might select a sample of U.S. colleges and universities from a directory, get lists of the students at all the selected schools, then draw samples of students from each.



Research in Real Life

Sampling Iran

Whereas most of the examples given in this textbook are taken from its country of origin, the United States, the basic methods of sampling would apply in other national settings as well. At the same time, researchers may need to make modifications appropriate to local conditions. In selecting a national sample of Iran, for example, Hamid Abdollahyan and Taghi Azadarmaki (2000: 21) from the University of Tehran began by stratifying the nation on the basis of cultural differences, dividing the country into nine cultural zones as follows:

1. Tehran
2. Central region including Isfahan, Arak, Qum, Yazd, and Kerman
3. The southern provinces including Hormozgan, Khuzistan, Bushahr, and Fars
4. The marginal western region including Lorestan, Chahmahal and Bakhtiari, Kogiluyeh and Eelam
5. The western provinces including western and eastern Azarbaijan, Zanjan, Gilazvin, and Ardebil

could sample the list of churches in some manner (for example, a stratified, systematic sample). Next, they would obtain lists of members from each of the selected churches. Each of the lists would then be sampled, to provide samples of church members for study.

Another typical situation concerns sampling among population areas such as a city. Although there is no single list of a city's population, citizens reside on discrete city blocks or census blocks. Researchers can, therefore, select a sample of blocks initially, create a list of people living on each of the selected blocks, and take a subsample of the people on each block.

In a more complex design, researchers might sample blocks, list the households on each selected block, sample the households, list the people residing in each household, and, finally, sample the people within each selected household. This multistage sample design leads ultimately to a selection of a sample of individuals but does not require the initial listing of all individuals in the city's population.

6. The eastern provinces including Khuzasan and Semnan
7. The northern provinces including Gilan, Mazandran, and Golestan
8. System
9. Kurdistan

Within each of these cultural areas, the researchers selected samples of census blocks and, on each selected block, a sample of households. Their sample design made provisions for getting the proper numbers of men and women as respondents within households and provisions for replacing those households where no one was at home.

Though the United States and Iran are politically and culturally quite different, the sampling methods appropriate for selecting a representative sample of populations are the same. Later in this chapter, when you review a detailed description of sampling the household population of Oakland, California, you will find it strikingly similar to the methods used in Iran by Abdollahyan and Azadarmaki.

Source: Hamid Abdollahyan and Taghi Azadarmaki. 2000. "Sampling Design in a Survey Research: The Sampling Practice in Iran." Paper presented to the meetings of the American Sociological Association, August 12–16, Washington, DC.

Multistage cluster sampling, then, involves the repetition of two basic steps: listing and sampling. The list of primary sampling units (churches, blocks) is compiled and, perhaps, stratified for sampling. Then a sample of those units is selected. The selected primary sampling units are then listed and perhaps stratified. The list of secondary sampling units is then sampled, and so forth.

The listing of households on even the selected blocks is, of course, a labor-intensive and costly activity—one of the elements making face-to-face, household surveys quite expensive. Vincent Iannacchione, Jennifer Staab, and David Redden (2003) report some initial success using postal mailing lists for this purpose. Although the lists are not perfect, they may be close enough to warrant the significant savings in cost.

Multistage cluster sampling makes possible those studies that would otherwise be impossible. Specific research circumstances often call for special designs, as the feature Research in Real Life: "Sampling Iran" demonstrates.

Multistage Designs and Sampling Error

Although cluster sampling is highly efficient, the price of that efficiency is a less-accurate sample. A simple random sample drawn from a population list is subject to a single sampling error, but a two-stage cluster sample is subject to two sampling errors. First, the initial sample of clusters will represent the population of clusters only within a range of sampling error. Second, the sample of elements selected within a given cluster will represent all the elements in that cluster only within a range of sampling error. Thus, for example, a researcher runs a certain risk of selecting a sample of disproportionately wealthy city blocks, plus a sample of disproportionately wealthy households within those blocks. The best solution to this problem lies in the number of clusters selected initially and the number of elements within each cluster.

Typically, researchers are restricted to a total sample size; for example, you may be limited to conducting 2,000 interviews in a city. Given this broad limitation, however, you have several options in designing your cluster sample. At the extremes you could choose one cluster and select 2,000 elements within that cluster, or you could select 2,000 clusters with one element selected within each. Of course, neither approach is advisable, but a broad range of choices lies between them. Fortunately, the logic of sampling distributions provides a general guideline for this task.

Recall that sampling error is reduced by two factors: an increase in the sample size and increased homogeneity of the elements being sampled. These factors operate at each level of a multistage sample design. A sample of clusters will best represent all clusters if a large number are selected and if all clusters are very much alike. A sample of elements will best represent all elements in a given cluster if a large number are selected from the cluster and if all the elements in the cluster are very much alike.

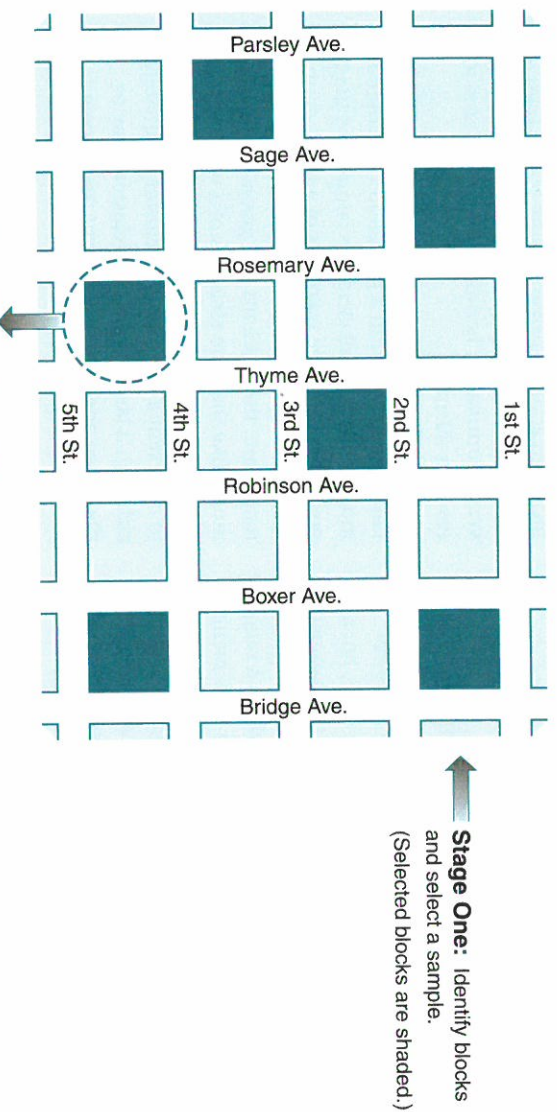
With a given total sample size, however, if the number of clusters is increased, the number of elements within a cluster must be decreased. In this respect, the representativeness of the clusters is

increased at the expense of more poorly representing the elements composing each cluster, or vice versa. Fortunately, homogeneity can be used to ease this dilemma.

Typically, the elements composing a given natural cluster within a population are more homogeneous than all elements composing the total population are. The members of a given church are more alike than all members of the denomination are; the residents of a given city block are more alike than the residents of a whole city are. As a result, relatively few elements may be needed to represent a given natural cluster adequately, although a larger number of clusters may be needed to adequately represent the diversity found among the clusters. This fact is most clearly seen in the extreme case of very different clusters composed of identical elements within each. In such a situation, a large number of clusters would adequately represent all its members. Although this extreme situation never exists in reality, it's closer to the truth in most cases than its opposite: identical clusters composed of grossly divergent elements.

The general guideline for cluster design, then, is to maximize the number of clusters selected while decreasing the number of elements within each cluster. However, this scientific guideline must be balanced against an administrative constraint. The efficiency of cluster sampling is based on the ability to minimize the listing of population elements. By initially selecting clusters, you need only list the elements composing the selected clusters, not all elements in the entire population. Increasing the number of clusters, however, goes directly against this efficiency factor. A small number of clusters may be listed more quickly and more cheaply than a large number. (Remember that all the elements in a selected cluster must be listed even if only a few are to be chosen in the sample.)

The final sample design will reflect these two constraints. In effect, you'll probably select as many clusters as you can afford. Lest this issue be left too open-ended at this point, here's one general guideline. Population researchers conventionally aim at the selection of 5 households per census block. If a total of 2,000 households are to be interviewed,



1. 491 Rosemary Ave.
2. 487 Rosemary Ave.
3. 473 Rosemary Ave.
4. 455 Rosemary Ave.
5. 437 Rosemary Ave.
6. 423 Rosemary Ave.
7. 411 Rosemary Ave.
8. 403 Rosemary Ave.
9. 1101 4th St.
10. 1123 4th St.
11. 1137 4th St.
12. 1157 4th St.
13. 1169 4th St.
14. 1187 4th St.
15. 402 Thyme Ave.
16. 408 Thyme Ave.
17. 424 Thyme Ave.
18. 446 Thyme Ave.
19. 458 Thyme Ave.
20. 480 Thyme Ave.
21. 498 Thyme Ave.
22. 1186 5th St.
23. 1174 5th St.
24. 1160 5th St.
25. 1140 5th St.
26. 1122 5th St.
27. 1118 5th St.
28. 1116 5th St.
29. 1104 5th St.
30. 1102 5th St.

FIGURE 7-13

Multistage Cluster Sampling. In multistage cluster sampling, we begin by selecting a sample of the clusters (in this case, city blocks). Then, we make a list of the elements (households, in this case) and select a sample of elements from each of the selected clusters.

you would aim at 400 blocks with 5 household interviews on each. Figure 7-13 presents a graphic overview of this process.

Before we turn to other, more-detailed procedures available to cluster sampling, let me reiterate that this method almost inevitably involves a loss of accuracy. The manner in which this appears, however, is somewhat complex. First, as noted earlier,

a multistage sample design is subject to a sampling error at each stage. Because the sample size is necessarily smaller at each stage than the total sample size, the sampling error at each stage will be greater than would be the case for a single-stage random sample of elements. Second, sampling error is estimated on the basis of observed variance among the sample elements. When those elements are drawn

from among relatively homogeneous clusters, the estimated sampling error will be too optimistic and must be corrected in the light of the cluster sample design.

Stratification in Multistage Cluster Sampling

Thus far, we've looked at cluster sampling as though a simple random sample were selected at each stage of the design. In fact, stratification techniques can be used to refine and improve the sample being selected.

The basic options here are essentially the same as those in single-stage sampling from a list. In selecting a national sample of churches, for example, you might initially stratify your list of churches by denomination, geographic region, size, rural or urban location, and perhaps by some measure of social class.

Once the primary sampling units (churches, blocks) have been grouped according to the relevant, available stratification variables, either simple random or systematic-sampling techniques can be used to select the sample. You might select a specified number of units from each group, or stratum, or you might arrange the stratified clusters in a continuous list and systematically sample that list.

To the extent that clusters are combined into homogeneous strata, the sampling error at this stage will be reduced. The primary goal of stratification, as before, is homogeneity.

There's no reason why stratification couldn't take place at each level of sampling. The elements listed within a selected cluster might be stratified before the next stage of sampling. Typically, however, this is not done. (Recall the assumption of relative homogeneity within clusters.)

Probability Proportionate to Size (PPS) Sampling

This section introduces you to a more sophisticated form of cluster sampling, one that is used in many large-scale survey-sampling projects. In the preceding discussion, I talked about selecting a random or systematic sample of clusters and then a random or systematic sample of elements within each cluster

selected. Notice that this produces an overall sampling scheme in which every element in the whole population has the same probability of selection.

Let's say we're selecting households within a city. If there are 1,000 city blocks and we initially select a sample of 100, that means that each block has a $100/1,000$ or 0.1 chance of being selected. If we next select 1 household in 10 from those residing on the selected blocks, each household has a 0.1 chance of selection within its block. To calculate the overall probability of a household being selected, we simply multiply the probabilities at the individual steps in sampling. That is, each household has a $1/10$ chance of its block being selected and a $1/10$ chance of that specific household being selected if the block is one of those chosen. Each household, in this case, has a $1/10 \times 1/10 = 1/100$ chance of selection overall. Because each household would have the same chance of selection, the sample so selected should be representative of all households in the city.

There are dangers in this procedure, however. In particular, the variation in the size of blocks (measured in numbers of households) presents a problem. Let's suppose that half the city's population resides in 10 densely packed blocks filled with high-rise apartment buildings, and suppose that the rest of the population lives in single-family dwellings spread out over the remaining 900 blocks. When we first select our sample of $1/10$ of the blocks, it's quite possible that we'll miss all of the 10 densely packed high-rise blocks. No matter what happens in the second stage of sampling, our final sample of households will be grossly unrepresentative of the city, comprising only single-family dwellings.

Whenever the clusters sampled are of greatly differing sizes, it's appropriate to use a modified sampling design called **PPS (probability proportionate to size)**. This design guards against

PPS (probability proportionate to size) This refers to a type of multistage cluster sample in which clusters are selected, not with equal probabilities (see *EPSEM*) but with probabilities proportionate to their sizes—as measured by the number of units to be subsampled.

the problem I've just described and still produces a final sample in which each element has the same chance of selection.

As the name suggests, each cluster is given a chance of selection proportionate to its size. Thus, a city block with 200 households has twice the chance of selection as one with only 100 households. Within each cluster, however, a fixed number of elements is selected, say, 5 households per block. Notice how this procedure results in each household having the same probability of selection overall.

Let's look at households of two different city blocks. Block A has 100 households; Block B has only 10. In PPS sampling, we would give Block A ten times as good a chance of being selected as Block B. So if, in the overall sample design, Block A has a 1/20 chance of being selected, that means Block B would only have a 1/200 chance. Notice that this means that all the households on Block A would have a 1/20 chance of having their block selected; Block B households have only a 1/200 chance.

If Block A is selected and we're taking 5 households from each selected block, then the households on Block A have a 5/100 chance of being selected into the block's sample. Because we can multiply probabilities in a case like this, we see that every household on Block A has an overall chance of selection equal to $1/20 \times 5/100 = 5/2,000 = 1/400$.

If Block B happens to be selected, on the other hand, its households stand a much better chance of being among the 5 chosen there: 5/10. When this is combined with their relatively poorer chance of having their block selected in the first place, however, they end up with the same chance of selection as those on Block A: $1/200 \times 5/10 = 5/2,000 = 1/400$.

weighting Assigning different weights to cases that were selected into a sample with different probabilities of selection. In the simplest scenario, each case is given a weight equal to the inverse of its probability of selection. When all cases have the same chance of selection, no weighting is necessary.

Further refinements to this design make it a very efficient and effective method for selecting large cluster samples. For now, however, it's enough to understand the basic logic involved.

Disproportionate Sampling and Weighting

Ultimately, a probability sample is representative of a population if all elements in the population have an equal chance of selection in that sample. Thus, in each of the preceding discussions, we've noted that the various sampling procedures result in an equal chance of selection—even though the ultimate selection probability is the product of several partial probabilities.

More generally, however, a probability sample is one in which each population element has a known nonzero probability of selection—even though different elements may have different probabilities. If controlled probability sampling procedures have been used, any such sample may be representative of the population from which it is drawn if each sample element is assigned a weight equal to the inverse of its probability of selection. Thus, where all sample elements have had the same chance of selection, each is given the same weight: 1. This is called a *self-weighting* sample.

Sometimes it's appropriate to give some cases more weight than others, a process called **weighting**. Disproportionate sampling and weighting come into play in two basic ways. First, you may sample subpopulations disproportionately to ensure sufficient numbers of cases from each for analysis. For example, a given city may have a suburban area containing one-fourth of its total population. Yet you might be especially interested in a detailed analysis of households in that area and may feel that one-fourth of this total sample size would be too few. As a result, you might decide to select the same number of households from the suburban area as from the remainder of the city. Households in the suburban area, then, are given a disproportionately better chance of selection than those located elsewhere in the city are.

As long as you analyze the two area samples separately or comparatively, you need not worry about the differential sampling. If you want to combine the two samples to create a composite picture of the entire city, however, you must take the disproportionate sampling into account. If n is the number of households selected from each area, then the households in the suburban area had a chance of selection equal to n divided by one-fourth of the total city population. Because the total city population and the sample size are the same for both areas, the suburban-area households should be given a weight of $1/4 n$, and the remaining households should be given a weight of $3/4 n$. This weighting procedure could be simplified by merely giving a weight of 3 to each of the households selected outside the suburban area.

Here's an example of the problems that can be created when disproportionate sampling is not accompanied by a weighting scheme. When the *Harvard Business Review* decided to survey its subscribers on the issue of sexual harassment at work, it seemed appropriate to oversample women because female subscribers were vastly outnumbered by male subscribers. Here's how G. C. Collins and Timothy Blodgett explained the matter:

We also skewed the sample another way: to ensure a representative response from women, we mailed a questionnaire to virtually every female subscriber, for a male/female ratio of 68% to 32%. This bias resulted in a response of 52% male and 44% female (and 4% who gave no indication of gender)—compared to HBR's U.S. subscriber proportion of 93% male and 7% female.

(1981: 78)

Notice a couple of things in this excerpt. First, it would be nice to know a little more about what “virtually every female” means. Evidently, the authors of the study didn't send questionnaires to all female subscribers, but there's no indication of who was omitted and why. Second, they didn't use the term *representative* with its normal social science usage. What they mean, of course, is that they wanted to get a substantial or “large enough”

response from women, and oversampling is a perfectly acceptable way of accomplishing that.

By sampling more women than a straightforward probability sample would have produced, the authors were able to “select” enough women (812) to compare with the men (960). Thus, when they report, for example, that 32 percent of the women and 66 percent of the men agree that “the amount of sexual harassment at work is greatly exaggerated,” we know that the female response is based on a substantial number of cases. That's good. There are problems, however.

To begin with, subscriber surveys are always problematic. In this case, the best the researchers can hope to talk about is “what subscribers to *Harvard Business Review* think.” In a loose way, it might make sense to think of that population as representing the more sophisticated portion of corporate management. Unfortunately, the overall response rate was 25 percent. Although that's quite good for subscriber surveys, it's a low response rate in terms of generalizing from probability samples.

Beyond that, however, the disproportionate sample design creates another problem. When the authors state that 73 percent of respondents favor company policies against harassment (Collins and Blodgett 1981: 78), that figure is undoubtedly too high, because the sample contains a disproportionately high percentage of women—who are more likely than men to favor such policies. And, when the researchers report that top managers are more likely to feel that claims of sexual harassment are exaggerated than are middle- and lower-level managers (1981: 81), that finding is also suspect. As the researchers report, women are disproportionately represented in lower management. That alone might account for the apparent differences among levels of management. In short, the failure to take account of the oversampling of women confounds all survey results that don't separate the findings by sex. The solution to this problem would have been to weight the responses by sex, as described earlier in this section.

In recent election campaign polling, survey weighting has become a controversial topic, as some polling agencies weight their results on the basis of party affiliation and other variables, whereas others

do not. Weighing in this instance involves assumptions regarding the differential participation of Republicans and Democrats in opinion polls and on election day—plus a determination of how many Republicans and Democrats there are. This is likely to be a topic of debate among pollsters and politicians in the years to come. Alan Reifman has created a website devoted to a discussion of this topic (link to it on your Sociology CourseMate at www.cengagebrain.com).

Probability Sampling in Review

Much of this chapter has been devoted to the key sampling method used in controlled survey research: probability sampling. In each of the variations examined, we've seen that elements are chosen for study from a population on a basis of random selection with known nonzero probabilities.

Depending on the field situation, probability sampling can be either very simple or extremely difficult, time-consuming, and expensive. Whatever the situation, however, it remains the most effective method for the selection of study elements. There are two reasons for this.

First, probability sampling avoids researchers' conscious or unconscious biases in element selection. If all elements in the population have an equal (or unequal and subsequently weighted) chance of selection, there is an excellent chance that the sample so selected will closely represent the population of all elements.

Second, probability sampling permits estimates of sampling error. Although no probability sample will be perfectly representative in all respects, controlled selection methods permit the researcher to estimate the degree of expected error.

In this lengthy chapter, we've taken on a basic issue in much social research: selecting observations that will tell us something more general than the specifics we've actually observed. This issue confronts field researchers, who face more action and more actors than they can observe and record fully, as well as political pollsters who want to predict an election but can't interview all voters. As

we proceed through the book, we'll see in greater detail how social researchers have found ways to deal with this issue.

The Ethics of Sampling

The key purpose of the sampling techniques discussed in this chapter is to allow researchers to make relatively few observations but gain an accurate picture of a much larger population. In the case of quantitative studies using probability sampling, the result should be a statistical profile, based on the sample, that closely mirrors the profile that would have been gained from observing the whole population. In addition to using legitimate sampling techniques, researchers should be careful to point out the possibility of errors: sampling error, flaws in the sampling frame, nonresponse error, or anything else that might make the results misleading.

Sometimes, more typically in qualitative studies, the purpose of sampling may be to tap into the breadth of variation within a population rather than to focus on the "average" or "typical" member of that population. While this is a legitimate and valuable approach, it poses the risk that readers may mistake the display of differences to reflect the distribution of characteristics in the population. In such a case, the researcher should make sure that the reader is not misled.

Main Points

Introduction

- Social researchers must select observations that will allow them to generalize to people and events not observed. Often this involves sampling a selection of people to observe.
- Understanding the logic of sampling is essential to doing social research.

A Brief History of Sampling

- Sometimes you can and should select probability samples using precise statistical techniques, but other times nonprobability techniques are more appropriate.

Nonprobability Sampling

- Nonprobability sampling techniques include relying on available subjects, purposive or judgmental sampling, snowball sampling, and quota sampling. In addition, researchers studying a social group may make use of informants. Each of these techniques has its uses, but none of them ensures that the resulting sample will be representative of the population being sampled.

The Theory and Logic of Probability Sampling

- Probability-sampling methods provide an excellent way of selecting representative samples from large, known populations. These methods counter the problems of conscious and unconscious sampling bias by giving each element in the population a known (nonzero) probability of selection.
 - Random selection is often a key element in probability sampling.
 - The most carefully selected sample will never provide a perfect representation of the population from which it was selected. There will always be some degree of sampling error.
 - By predicting the distribution of samples with respect to the target parameter, probability-sampling methods make it possible to estimate the amount of sampling error expected in a given sample.
 - The expected error in a sample is expressed in terms of confidence levels and confidence intervals.
- ### Populations and Sampling Frames
- A sampling frame is a list or quasi list of the members of a population. It is the resource used in the selection of a sample. A sample's representativeness depends directly on the extent to which a sampling frame contains all the members of the total population that the sample is intended to represent.

Types of Sampling Designs

- Several sampling designs are available to researchers.
- Simple random sampling is logically the most fundamental technique in probability sampling, but it is seldom used in practice.
- Systematic sampling involves the selection of every k th member from a sampling frame. This method is more practical than simple random sampling; with a few exceptions, it is functionally equivalent.

- Stratification, the process of grouping the members of a population into relatively homogeneous strata before sampling, improves the representativeness of a sample by reducing the degree of sampling error.

Multistage Cluster Sampling

- Multistage cluster sampling is a relatively complex sampling technique that frequently is used when a list of all the members of a population does not exist. Typically, researchers must balance the number of clusters and the size of each cluster to achieve a given sample size. Stratification can be used to reduce the sampling error involved in multistage cluster sampling.
- Probability proportionate to size (PPS) is a special, efficient method for multistage cluster sampling.
- If the members of a population have unequal probabilities of selection into the sample, researchers must assign weights to the different observations made, in order to provide a representative picture of the total population. The weight assigned to a particular sample member should be the inverse of its probability of selection.

Probability Sampling in Review

- Probability sampling remains the most effective method for the selection of study elements for two reasons: it avoids researcher bias in element selection and it permits estimates of sampling error.

The Ethics of Sampling

- Because probability sampling always carries a risk of error, the researcher must inform readers of any errors that might make results misleading.
- Sometimes, nonprobability sampling methods are used to obtain the breadth of variations in a population. In this case, the researcher must ensure that readers do not confuse variations with what's typical in the population.

KEY TERMS

The following terms are defined in context in the chapter and at the bottom of the page where the term is introduced, as well as in the comprehensive glossary at the back of the book.

cluster sampling	element
confidence interval	EPSEM
confidence level	informant

224 ■ Chapter 7: The Logic of Sampling

nonprobability sampling	sampling frame
parameter	sampling interval
population	sampling ratio
PPS	sampling unit
probability sampling	simple random sampling
purposive (judgmental) sampling	snowball sampling
quota sampling	statistic
random selection	stratification
representativeness	study population
sampling error	systematic sampling
	weighting

PROPOSING SOCIAL RESEARCH: SAMPLING

In this portion of the proposal, you'll describe how you'll select from among all the possible observations you might make. Depending on the data-collection method you plan to employ, either probability or nonprobability sampling may be more appropriate to your study. Similarly, this aspect of your proposal may involve the sampling of subjects or informants, or it could involve the sampling of corporations, cities, books, and so forth.

Your proposal, then, must specify what units you'll be sampling among, the data you'll use (such as a sampling frame) for purposes of your sample selection, and the actual sampling methods you'll use.

REVIEW QUESTIONS AND EXERCISES

1. Review the discussion of the 1948 Gallup Poll that predicted that Thomas Dewey would defeat Harry Truman for president. What are some ways Gallup could have modified his quota sample design to avoid the error?
2. Using Appendix C of this book, select a simple random sample of 10 numbers in the range of 1 to 9,876. What is each step in the process?
3. What are the steps involved in selecting a multi-stage cluster sample of students taking first-year English in U.S. colleges and universities?
4. In Chapter 9 we'll discuss surveys conducted on the Internet. Can you anticipate possible problems

concerning sampling frames, representativeness, and the like? Do you see any solutions?

5. Using InfoTrac College Edition on your Sociology CourseMate at www.cengagebrain.com, locate studies using (1) a quota sample, (2) a multistage cluster sample, and (3) a systematic sample. Write a brief description of each study.

SPSS EXERCISES

See the booklet that accompanies your text for exercises using SPSS (Statistical Package for the Social Sciences). There are exercises offered for each chapter, and you'll also find a detailed primer on using SPSS.

Online Study Resources

Access the resources your instructor has assigned. For this book, you can access:



**CourseMate for The
Practice of Social Research**

Login to CengageBrain.com to access chapter-specific learning tools including *Learning Objectives*, *Practice Quizzes*, *Videos*, *Internet Exercises*, *Flash Cards*, *Glossaries*, *Web Links*, and more from your Sociology CourseMate.



If your professor has assigned Aplia homework:

1. Sign into your account.
 2. After you complete each page of questions, click "Grade It Now" to see detailed explanations of every answer.
 3. Click "Try Another Version" for an opportunity to improve your score.
- Visit www.cengagebrain.com to access your account and purchase materials.