

You probably learned how to calculate percentages in fourth grade and will be tempted to skip the next few paragraphs. Fair enough. But first do one simple exercise for me. Assume that a department store is selling a dress for \$100. The assistant manager marks down all merchandise by 25 percent. But then that assistant manager is fired for hanging out in a bar with Bill Gates,* and the new assistant manager raises all prices by 25 percent. What is the final price of the dress? If you said (or thought) \$100, then you had better not skip any paragraphs.

The final price of the dress is actually \$93.75. This is not merely a fun parlor trick that will win you applause and adulation at cocktail parties. Percentages are useful—but also potentially confusing or even deceptive. The formula for calculating a percentage difference (or change) is the following: $(\text{new figure} - \text{original figure}) / \text{original figure}$. The numerator (the part on the top of the fraction) gives us the size of the change in absolute terms; the denominator (the bottom of the fraction) is what puts this change in context by comparing it with our starting point. At first, this seems straightforward, as when the assistant store manager cuts the price of the \$100 dress by 25 percent. Twenty-five percent of the original \$100 price is \$25; that's the discount, which takes the price down to \$75. You can plug the numbers into the formula above and do some simple manipulation to get to the same place: $(\$100 - \$75) / \$100 = .25$, or 25 percent.

The dress is selling for \$75 when the new assistant manager demands that the price be raised 25 percent. That's where many of the people reading this paragraph probably made a mistake. The 25 percent markup is calculated as a percentage of the dress's new reduced price, which is \$75. The increase will be $.25(\$75)$, or \$18.75, which is how the final price ends up at \$93.75 (and not \$100). The point is that a percentage change always gives the value of some figure *relative to something else*. Therefore, we had better understand what that something else is.

We evaluate the academic performance of high school and college students by means of a grade point average, or GPA. A letter grade is assigned a point value; typically an A is worth 4 points, a B is worth 3, a C is worth 2, and so on. By graduation, when high school students are applying to college and college students are looking for jobs, the grade point average is a handy tool for assessing their academic potential. Someone who has a 3.7 GPA is clearly a stronger student than someone at the same school with a 2.5 GPA. That makes it a nice descriptive statistic. It's easy to calculate, it's easy to understand, and it's easy to compare across students.

But it's not perfect. The GPA does not reflect the difficulty of the courses that different students may have taken. How can we compare a student with a 3.4 GPA in classes that appear to be relatively nonchallenging and a student with a 2.9 GPA who has taken calculus, physics, and other tough subjects? I went to a high school that attempted to solve this problem by giving extra weight to difficult classes, so that an A in an "honors" class was worth five points instead of the usual four. This caused its own problems. My mother was quick to recognize the distortion caused by this GPA "fix." For a student taking a lot of honors classes (me), any A in a nonhonors course, such as gym or health education, would actually pull my GPA down, even though it is impossible to do better than an A in those classes. As a result, my parents forbade me to take driver's education in high school, lest even a perfect performance diminish my chances of getting into a competitive college and going on to write popular books. Instead, they paid to send me to a private driving school, at nights over the summer.

The curious thing is that the same people who are perfectly comfortable discussing statistics in the context of sports or the weather or grades will seize up with anxiety when a researcher starts to explain something like the Gini index, which is a standard tool in economics for measuring income inequality. I'll explain what the Gini index is in a moment, but for now *the most important thing to recognize is that the Gini index is just like the passer rating*. It's a handy tool for collapsing complex information into a single number. As such, it has the strengths of most descriptive statistics, namely that it provides an easy way to compare the income distribution in two countries, or in a single country at different points in time.

The Gini index measures how evenly wealth (or income) is shared within a country on a scale from zero to one. The statistic can be calculated for wealth or for annual income, and it can be calculated at the individual level or at the household level. (All of these statistics will be highly correlated but not identical.) The Gini index, like the passer rating, has no intrinsic meaning; it's a tool for comparison. A country in which every household had identical wealth would have a Gini index of zero. By contrast, a country in which a single household held the country's entire wealth would have a Gini index of one. As you can probably surmise, the closer a country is to one, the more unequal its distribution of wealth. The United States has a Gini index of .45, according to the Central Intelligence Agency (a great collector of statistics, by the way).¹ So what?

Once that number is put into context, it can tell us a lot. For example, Sweden has a Gini index of .23. Canada's is .32. China's is .42. Brazil's is .54. South Africa's is .65.* As we look across those numbers, we get a sense of where the United States falls relative to the rest of the world when it comes to income inequality. We can also compare different points in time. The Gini index for the United States was .41 in 1997 and grew to .45 over the next decade. (The most recent CIA data are for 2007.) This tells us in an objective way that while the United States grew richer over that period of time, the distribution of wealth grew more unequal. Again, we can compare the changes in the Gini index across countries over roughly the same time period. Inequality in Canada was basically unchanged over the same stretch. Sweden has had significant economic growth over the past two decades, but the Gini index in Sweden actually fell from .25 in 1992 to .23 in 2005, meaning that Sweden grew richer *and* more equal over that period.

Is the Gini index the perfect measure of inequality? Absolutely not—just as the passer rating is not a perfect measure of quarterback performance. But it certainly gives us some valuable information on a socially significant phenomenon in a convenient format.

How many homeless people live on the streets of Chicago? How often do married people have sex? These may seem like wildly different kinds of questions; in fact, they both can be answered (not perfectly) by the use of basic statistical tools. One key function of statistics is to use the data we have to make informed conjectures about larger questions for which we do not have full information. In short, we can use data from the “known world” to make informed inferences about the “unknown world.”

Let's begin with the homeless question. It is expensive and logistically difficult to count the homeless population in a large metropolitan area. Yet it is important to have a numerical estimate of this population for purposes of providing social services, earning eligibility for state and federal revenues, and gaining congressional representation. One important statistical practice is sampling, which is the process of gathering data for a small area, say, a handful of census tracts, and then using those data to make an informed judgment, or inference, about the homeless population for the city as a whole. Sampling requires far less resources than trying to count an entire population; done properly, it can be every bit as accurate.

A political poll is one form of sampling. A research organization will attempt to contact a sample of households that are broadly representative of the larger population and ask them their views about a particular issue or candidate. This is obviously much cheaper and faster than trying to contact every household in an entire state or country. The polling and research firm Gallup reckons that a methodologically sound poll of 1,000 households will produce roughly the same results as a poll that attempted to contact every household in America.

That's how we figured out how often Americans are having sex, with whom, and what kind. In the mid-1990s, the National Opinion Research Center at the University of Chicago carried out a remarkably ambitious study of American sexual behavior. The results were based on detailed surveys conducted in person with a large, representative sample of American adults. If you read on, Chapter 10 will tell you what they learned. *How many other statistics books can promise you that?*

Suppose you are at work, idly surfing the Web when you stumble across a riveting day-by-day account of Kim Kardashian's failed seventy-two-day marriage to professional basketball player Kris Humphries. You have finished reading about day seven of the marriage when your boss shows up with two enormous files of data. One file has warranty claim information for each of the 57,334 laser printers that your firm sold last year. (For each printer sold, the file documents the number of quality problems that were reported during the warranty period.) The other file has the same information for each of the 994,773 laser printers that your chief

competitor sold during the same stretch. Your boss wants to know how your firm's printers compare in terms of quality with the competition.

Fortunately the computer you've been using to read about the Kardashian marriage has a basics statistics package, but where do you begin? Your instincts are probably correct: The first descriptive task is often to find some measure of the “middle” of a set of data, or what statisticians might describe as its “central tendency.” What is the typical quality experience for your printers compared with those of the competition? The most basic measure of the “middle” of a distribution is the mean, or average. In this case, we want to know the average number of quality problems per printer sold for your firm and for your competitor. You would simply tally the total number of quality problems reported for all printers during the warranty period and then divide by the total number of printers sold. (Remember, the same printer can have multiple problems while under warranty.) You would do that for each firm, creating an important descriptive statistic: the average number of quality problems per printer sold.

Suppose it turns out that your competitor's printers have an average of 2.8 quality-related problems per printer during the warranty period compared with your firm's average of 9.1 reported defects. That was easy. You've just taken information on a million printers sold by two different companies and distilled it to the essence of the problem: your printers break a lot. Clearly it's time to send a short e-mail to your boss quantifying this quality gap and then get back to day eight of Kim Kardashian's marriage.

Or maybe not. I was deliberately vague earlier when I referred to the “middle” of a distribution. The mean, or average, turns out to have some problems in that regard, namely, that it is prone to distortion by “outliers,” which are observations that lie farther from the center. To get your mind around this concept, imagine that ten guys are sitting on bar stools in a middle-class drinking establishment in Seattle; each of these guys earns \$35,000 a year, which makes the mean annual income for the group \$35,000. Bill Gates walks into the bar with a talking parrot perched on his shoulder. (The parrot has nothing to do with the example, but it kind of spices things up.) Let's assume for the sake of the example that Bill Gates has an annual income of \$1 billion. When Bill sits down on the eleventh bar stool, the mean annual income for the bar patrons rises to about \$91 million. Obviously none of the original ten drinkers is any richer (though it might be reasonable to expect Bill Gates to buy a round or two). If I were to describe the patrons of this bar as having an average annual income of \$91 million, the statement would be both statistically correct and grossly misleading. This isn't a bar where multimillionaires hang out; it's a bar where a bunch of guys with relatively low incomes happen to be sitting next to Bill Gates and his talking parrot. The sensitivity of the mean to outliers is why we should not gauge the economic health of the American middle class by looking at per capita income. Because there has been explosive growth in incomes at the top end of the distribution—CEOs, hedge fund managers, and athletes like Derek Jeter—the average income in the United States could be heavily skewed by the megarich, making it look a lot like the bar stools with Bill Gates at the end.