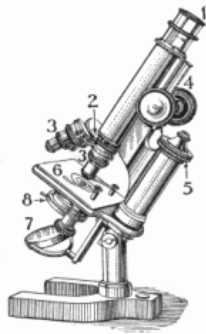


MODULE 4

APPENDIX

Reliability of Results



Module 4 Appendix: Reliability of Results

This section returns to the assessment of how many people fear that they will be victims of a crime. Using NSDstat, we can construct a *frequency table* of the responses to the question 'How worried are you about your home being broken into?' The results are:

v5 How worried about your home broken into?			
Value labels	Code	Number	% valid
Very worried	1	3675	18.9
Fairly worried	2	7361	37.9
Not very worried	3	6737	34.7
Not at all worried	4	1628	8.4
TOTAL		19401	100
<i>Included 19401 cases from a total of 19411</i>			

For example, 3675 people reported that they were 'very worried' when asked: 'How worried are you about your home being broken into?' but this reflects the sample size which is specific to this particular survey. Consequently we *standardize* the count by calculating the *percentage* that this category is of all the people who were asked 'How worried are you about your home being broken into?'. In the 2000 British Crime Survey, 19,401 people responded to the question¹. About 19% of all the people asked if they were worried that their house would be broken into who replied reported that they were 'Very Worried'². Because the total all percentages is 100%, by converting raw counts into percentages, we can see immediately the *relative* importance (or, conversely, the *relative unimportance* of a particular category). In this case, although 19% of the people asked the question were 'Very Worried' that their house might be broken into, 81% were not 'Very Worried'. The commentator (you, in this case) has to interpret findings like these. However, before doing so, the commentator has to be sure about who is being described.

The sampling problem

In the last section, we reported that 3675 of the 19,401 people who were asked how worried they were that their house would be burgled reported that they were 'Very Worried'. Who were these 19,401 people? Why should we be interested in them?

They were chosen by the researchers who conducted the British Crime Survey to *represent* the range of opinion about, and experience of, crime in the entire population of England and Wales in 2000. As the population of England and Wales in 2000 numbered over 35 million people, the ability to *sample* their views and experience of crime by speaking to about 20,000 people is highly beneficial both in terms of money and time. However gaining these benefits introduces new costs.

Representative samples must be selected methodically following acceptable procedures. The selection of people for the sample *must* be random. Even given this procedure, the researcher must accept that different samples of the same population asked exactly the

same questions will provide different estimates of the population's views. Because the population has only one view (even though the individuals who make up the population may have different views), when several samples give different estimates of that view, they are in error. The size of error can be assessed using the [Central Limit Theorem](#).

The Central Limit Theorem, which underpins our use of samples to estimate population characteristics ('parameters') states:

'The [means of random samples](#) drawn from large [populations](#) will be [normally distributed](#) around the original population mean with standard deviation inversely proportional to the *square root* of sample size and directly proportional to the population [standard deviation](#).'

If the Central Limit Theorem makes us doubt whether we know that 19% of the British population are very worried that their house would be burgled³, the Central Limit Theorem also allows the social researcher to benefit from the economies derived from [sampling](#) by describing how the sample estimates will be distributed around the population [parameter](#), the value that we are really interested in assessing. The [normal distribution](#) is a bell-shaped curve which is symmetrical around the real population value (i.e. 'parameter'). That is, there are as many underestimates as there are overestimates but the bell-shape implies that most of the estimates will be close to the real population value. In fact 66% of all the samples will estimate a value which is within 1 standard unit of the population parameter; 95% of all the samples will estimate a value which is within 2 standard units of the population parameter; and 99% of all the samples will estimate a value which is within 3 standard units of the population parameter. Of course, this raises the question of what a standard unit is. Fortunately, the Central Limit Theorem tells us how to calculate the size of a standard unit by relating the breadth of the distribution to how [varied](#) the population is on the characteristic that we are trying to assess and to how big the sample is.

Being Confident in the Face of Error

The Central Limit Theorem leads us to doubt whether we have precisely estimated a population parameter with a statistic drawn from one sample, but the Theorem also gives us the basis for estimating what the possible range of values might be. NSDstat's *Error Margin Calculator* is designed to calculate *confidence intervals* for percentages.

To run the *Error Margin Calculator*:

→Start→NSDStat→Tools→Error margin . Because in our example, roughly 20% of our sample reported that they were 'very worried' about their homes being burgled, 80% chose a different response. This is a measure of how varied the sample was on this characteristic. To enter this, **→Percentage↓80-20**. Sample size is the other factor associated with standard units by the Central Limit Theorem — it is automatically entered by NSDstat. The last piece of information that is required describes *your* attitude to risk — or to being wrong. To enter this,

→Confidence interval↓{95%| 99%| 99.9%}→Display.

As an experiment, try each Confidence interval — what happens?

What you are being told is that 95% of the samples of size 19411 would return a value of $18.9\% \pm 0.56\%$ or, in other words, you can be 95% confident that the real population value ([parameter](#)) is between 18.34% and 19.46%. However you can be 99% confident that the population parameter lies between 18.16% and 19.64%. In other words, it seems that the trade-off is between confidence and precision. What happens if you insist on being 99.9% confident?

Of course, sample sizes of close to 20,000 are very expensive. What would be the confidence interval if the sample size were halved? Quartered? A tenth? To find out,

→N→10000→Confidence interval↓{95% | 99% | 99.9%}→Display and

→N→5000→Confidence interval↓{95% | 99% | 99.9%}→Display and

→N→1900→Confidence interval↓{95% | 99% | 99.9%}→Display.

Notes

¹In fact, 19411 people were asked the question but only 19401 gave valid information to the interviewer. The 'missing' 10 may have said that they did not know how much they feared having things stolen from their house or they may have refused to answer or they may have been homeless or any number of other reasons.

²Percentages are calculated using the formula

((number in a category ÷ total providing valid information) × 100). Applying the formula to our example, $((3675 \div 19401) \times 100) = ((0.19) \times 100) = 19\%$

³We know that if we drew another sample of 19,411 people from the population resident in Great Britain in 2000 and asked them if they feared being burgled, it is unlikely that exactly 18.9% of that sample would report that they were very worried.