

Technology Companion

Doing Statistics on the TI-84

A hands-on companion to
You're Smarter Than You Think: Statistics

Safaa Dabagh

Keyed chapter-by-chapter to the main book

*"The calculator does the arithmetic. Your job is to ask the right question
and read the answer. You can absolutely do that."*

Open Educational Resource. Free to Use, Share, and Adapt (CC BY 4.0)

2026

Before You Press a Single Button

If the sight of all those little buttons makes your shoulders climb toward your ears, breathe. The TI-84 is not a puzzle you have to solve. It is a very good calculator that does one thing at a time. You choose a menu, you type a few numbers, you press **ENTER**, and it answers. That is the whole loop.

Here is the promise that matters most: **you cannot break it**. There is no wrong button that ruins anything. If a screen looks strange, press **CLEAR** to wipe the current line, or press **2nd** then **MODE** (which is labeled **QUIT**) to jump all the way back to the home screen. Those two keys are your escape hatch, always. Lost? **CLEAR** and **2nd QUIT**. You are never stuck.

This companion follows the main book chapter for chapter. When you finish Chapter 4 in the book, open Chapter 4 here and watch the same ideas happen on the calculator. You never have to guess which menu goes with which idea.

One more promise: every set of keystrokes in this book is short, and every one is followed by what the answer looks like on the screen and what it means. You are never left staring at the display wondering what you are seeing.

Meet Your Calculator

You will use a handful of keys again and again. Learn these six and you can do the whole course. Everything else is a detour you rarely take.

- **STAT** opens the statistics menus. Its **EDIT** tab is where you type your data into lists; its **CALC** tab computes summaries; its **TESTS** tab runs every hypothesis test and confidence interval.
- **2nd then VARS** is labeled **DISTR** (distributions). This is your table in the back of the book, replaced: `normalcdf`, `invNorm`, `binompdf`, and friends live here.
- **STAT then TESTS** holds the tests and intervals: t-tests, proportion tests, ANOVA, chi-square, regression tests. One menu, the whole back half of the course.
- **2nd then Y=** is labeled **STAT PLOT**. This turns on the picture: histogram, boxplot, or scatterplot of your lists.
- **ZOOM** then choice **9:ZoomStat** auto-fits any statistics plot to the screen so you never have to set the window by hand.
- **MATH then PROB** (press the right arrow to the **PROB** tab) holds the randomizers: `randInt`, `randIntNoRep`, and `rand` for drawing samples.

Typing data into a list (do this constantly)

Press **STAT**, then choose 1:Edit. You will see columns named **L1**, **L2**, **L3**. Put your cursor in **L1**, type a number, press **ENTER**, type the next, press **ENTER**, and so on down the column. A second variable goes in **L2** the same way. To wipe a list, arrow up onto its name (the **L1** at the very top), press **CLEAR**, then **ENTER**. Nothing is ever stuck; you can always clear and start the column again.

TI-84 (press these keys)

```
STAT > 1:Edit  
put cursor in L1, type: 2 ENTER 2 ENTER
```

That is the single move you will repeat in almost every chapter. Once it feels routine, the whole calculator feels routine.

Getting Data Into the Calculator

Be honest with yourself about one thing up front: **the TI-84 cannot open a file**. It has no way to read a .csv spreadsheet. You get your numbers into it by *typing them into lists*. That sounds like a chore, but the data sets in this book are small on purpose, and every chapter tells you the exact numbers to type.

1. Numbers go into a list. Press **STAT**, choose 1:Edit, and type a quantitative column straight down **L1** (and a second column down **L2** when you have one).

TI-84 (press these keys)

```
STAT > 1:Edit
```

```
L1: 14 15 13 15 16 12 14 15 (each followed by ENTER)
```

2. Two variables go into two lists side by side. When you have paired numbers, like hours and grades, put one in **L1** and its partner in **L2**, row for row.

3. Words become numbers. The calculator only stores numbers, so a categorical variable (Home, Library, Cafe) gets *coded*: Home = 1, Library = 2, Cafe = 3, for example. For a two-way table you usually just type the *counts*, not the raw words. Each chapter shows you how.

That is the whole toolkit. The numbers you need are printed in every “What you’ll see” box and gathered in the appendix. Now let’s walk the book.

Chapter 1

What's the Story in the Data?

■ WHAT THIS CHAPTER DOES ON THE TI-84

Type a column of data into a list, and understand why a word variable has to be turned into a number before the calculator can hold it.

In Chapter 1 you learned the difference between a population and a sample, and between categorical and quantitative variables. On the TI-84 that difference is physical: numbers go into a list as they are, but words cannot. A category has to be coded as a number first.

TI-84 (press these keys)

```
STAT > 1:Edit  
L1: 14 15 13 15 16 12 14 15 14 13 15 16 14 15 13 14  
(press ENTER after each value)  
  
then to code study_spot in L2:  
Home = 1, Library = 2, Cafe = 3  
L2: 1 2 3 1 1 2 ...
```

What you'll see (on the screen)

```
L1(1)=14  L1(2)=15  ...  L1(16)=14  
L2 holds the codes 1, 2, 3 for the categories
```

■ READ THE OUTPUT

L1 now holds the quiz scores, real numbers the calculator can average. L2 holds codes standing in for words, because the TI cannot store the word "Home." That is the "words vs. numbers" idea from the chapter, made concrete: quantitative data lives in a list as itself, and categorical data has to be translated into numbers first. When you later count categories, you count the codes.

■ YOUR TURN

Type five made-up quiz scores into **L1**. Now arrow up onto the name **L1**, press **CLEAR**, then **ENTER**, and watch the column empty. You just learned the undo you will use most.

Chapter 2

Where Did This Data Come From?

■ WHAT THIS CHAPTER DOES ON THE TI-84

Draw a *simple random sample* of 5 students from 40, and see how seeding the randomizer makes the draw repeatable.

Chapter 2 is about how good data is collected. The gold standard is the simple random sample: everyone has an equal chance. The TI can draw one for you from the **MATH PROB** menu.

TI-84 (press these keys)

```
first, seed for repeatability:  
5 > STO> then MATH > PROB > 1:rand, then ENTER  
  
then draw 5 different students from 1 to 40:  
MATH > PROB > 8:randIntNoRep(1, 40)  
read the FIRST 5 numbers it lists
```

What you'll see (on the screen)

```
{9 37 1 23 6 ...}  
the first five: 9, 37, 1, 23, 6
```

■ READ THE OUTPUT

`randIntNoRep(1,40)` shuffles the numbers 1 through 40 with no repeats, and you take the first five as your sample of students. Seeding first (a number, then `STO`, then `rand`) is a promise: “start the randomness from the same spot.” Anyone who seeds with the same number gets the same shuffle. That is how a random study is made checkable. Change the seed, get a different sample.

■ YOUR TURN

Seed with 5 again and re-run `randIntNoRep(1,40)`. Same first five? Now skip the seed and run it twice. What changes each time?

Chapter 3

Show Me a Picture

■ WHAT THIS CHAPTER DOES ON THE TI-84

Make the core graphs from the chapter: a histogram and a boxplot of a quantitative list, drawn straight from your data.

The TI draws pictures from your lists. Type the numbers once, then let **STAT PLOT** and **ZoomStat** do the rest. Bar charts for categories are drawn from the coded counts; histograms and boxplots come straight from the numbers.

TI-84 (press these keys)

put quiz scores in L1 first (STAT > 1:Edit)

histogram:

2nd > STAT PLOT > 1:Plot1 > On

Type = histogram, Xlist = L1

then ZOOM > 9:ZoomStat

boxplot (shows the middle, spread, outliers):

2nd > STAT PLOT > 1:Plot1 > On

Type = modified **boxplot**, Xlist = L1

then ZOOM > 9:ZoomStat

What you'll see (on the screen)

a histogram with bars piling up around 14 to 15

a boxplot with its box from about 13.75 to 15

▪ READ THE OUTPUT

Turning Plot1 **On** tells the calculator to draw, and Type chooses which picture. ZoomStat sizes the window so the whole graph fits, so you never fight the window settings. The histogram bins the scores into touching bars, and you see the pile near 14 to 15, exactly like the book. The modified boxplot shows the middle 50% as a box and flags any outliers with dots.

▪ YOUR TURN

With the histogram on screen, press **WINDOW** and change **Xscl** (the bar width) from 1 to 2. Fewer, wider bars now. Does the pile still sit near 14 to 15?

Chapter 4

The Typical and the Spread

■ WHAT THIS CHAPTER DOES ON THE TI-84

Compute the center (mean, median) and the spread (standard deviation, IQR) with one command, 1-Var Stats.

Chapter 4's formulas are a single menu choice on the TI. You do not compute them by hand here; you check your by-hand work against the calculator.

TI-84 (press these keys)

put quiz scores in L1 first (STAT > 1:Edit)

STAT > CALC > 1:1-Var Stats

List: L1, then Calculate

scroll down with the arrow to see **all** the numbers

What you'll see (on the screen)

x-bar = 14	(the mean)
Sx = 1.15	(sample standard deviation)
n = 16	
med = 14	(the median)
Q1 = 13.75	
Q3 = 15	so IQR = Q3 - Q1 = 1.25

▪ READ THE OUTPUT

1-Var Stats is your whole five-number summary plus the mean and standard deviation, on one scrolling screen. $\bar{x}$ is the mean (14) and med is the median (14); when they match this closely, the scores are fairly symmetric. The calculator does not print IQR directly, so you subtract: Q3 minus Q1 is $15 - 13.75 = 1.25$. Read Sx, not the other one below it: Sx is the *sample* standard deviation, which is what you almost always want.

▪ YOUR TURN

Add one unusually high score, 40, to the end of **L1** and run 1-Var Stats again. Which moved more, the mean or the median? That is why we report the median for skewed data.

Chapter 5

What Are the Chances?

■ WHAT THIS CHAPTER DOES ON THE TI-84

Store a two-way table as a matrix, and read probabilities and conditional probabilities off the counts.

Chapter 5's two-way tables are where probability gets concrete. On the TI you hold the table in a matrix, then reason from the counts by hand, which keeps the ideas visible.

TI-84 (press these keys)

```
2nd > MATRX (the x-inverse key) > EDIT > 1:[A]
set size to 2 x 2, then type the counts:
[A] = [[30, 20],
       [10, 40]]
(row 1 = caffeine Yes; row 2 = caffeine No)
(col 1 = slept well Yes; col 2 = No)
```

What you'll see (on the screen)

```
[A] =
[ 30  20 ]   row total 50 (had caffeine)
[ 10  40 ]   row total 50 (no caffeine)
grand total = 100
```

■ READ THE OUTPUT

The whole 100 people sit in [A]. A plain probability like “fraction who had caffeine and slept well” is one cell over the grand total: $30/100 = 0.30$. A *conditional* probability, “given they had caffeine, what fraction slept well,” divides *within the row*: $30/50 = 0.60$. That within-a-row division is the trickiest idea in the chapter, and here you can literally point at the row you are conditioning on.

■ YOUR TURN

Using $[A]$, find $P(\text{slept well})$ for everyone: add the top of each column, then divide by 100. Now find $P(\text{slept well given no caffeine})$ by dividing within the second row. Are they the same? What does it mean if they are not?

Chapter 6

Predictable Randomness

■ WHAT THIS CHAPTER DOES ON THE TI-84

Get exact binomial and normal probabilities from the **DISTR** menu, no table in the back of the book.

Chapter 6's distribution tables are replaced by four commands under **2nd DISTR**. For the normal curve, `normalcdf` gives "probability between two values," and `invNorm` runs it backward from a percentile to a score.

TI-84 (press these keys)

BINOMIAL: 10 quizzes, each correct with probability 0.7

2nd > DISTR > `binompdf`(10, 0.7, 8) P(exactly 8)

2nd > DISTR > `binomcdf`(10, 0.7, 8) P(8 or fewer)

NORMAL: exam scores, mean 70, sd 8

2nd > DISTR > `normalcdf`(-1E99, 80, 70, 8) P(below 80)

2nd > DISTR > `invNorm`(0.90, 70, 8) the 90th percentile

(type 1E99 as: 1, then 2nd EE, then 99)

What you'll see (on the screen)

`binompdf`(10,0.7,8) = 0.233

`binomcdf`(10,0.7,8) = 0.851

`normalcdf`(-1E99,80,70,8) = 0.894

`invNorm`(0.90,70,8) = 80.25

▪ READ THE OUTPUT

`normalcdf(-1E99, 80, 70, 8) = 0.894` says about 89% of scores fall below 80; the `-1E99` is just “negative infinity,” a floor so low it means “from the very bottom.” `invNorm(0.90, 70, 8) = 80.25` runs it backward: to sit in the top 10%, you need about 80.25. `normalcdf` takes scores and returns a probability; `invNorm` takes a probability and returns a score. They are opposites.

▪ YOUR TURN

Find the probability a score is *above* 85 with `normalcdf(85, 1E99, 70, 8)`. Here `1E99` is positive infinity, a ceiling so high it means “all the way to the top.”

Chapter 7

From Sample to Truth

■ WHAT THIS CHAPTER DOES ON THE TI-84

Understand the Central Limit Theorem through the standard-error formula, and use `normalcdf` with that standard error to answer questions about a sample mean.

Chapter 7 claims something wild: even from a lopsided population, the averages of samples pile up into a bell. The TI cannot run a thousand simulations comfortably, so we lean on the honest shortcut the theorem gives us: a formula for how spread-out those sample means are.

TI-84 (press these keys)

the spread of `sample` means is the STANDARD ERROR:

$SE = Sx / \text{sqrt}(n)$

`example:` $Sx = 1.15, n = 16$

$1.15 / \text{sqrt}(16) = 0.2875$

then treat the `sample mean` as normal with that SE.

$P(\text{sample mean below } 14.6), \text{ center } 14, SE \text{ } 0.2875:$

`2nd > DISTR > normalcdf(-1E99, 14.6, 14, 0.2875)`

What you'll see (on the screen)

$SE = 1.15 / 4 = 0.2875$

$\text{normalcdf}(-1E99, 14.6, 14, 0.2875) = 0.982$

▪ READ THE OUTPUT

The Central Limit Theorem promises that sample means follow a bell shape centered at the true mean, with spread equal to the *standard error*, $Sx/\sqrt{n}$. That is why bigger samples give tighter, more trustworthy averages: dividing by a larger $\sqrt{n}$ shrinks the SE. Once you have the SE, a question about a sample mean is just a `normalcdf` with that SE in the last slot. The TI is limited for big simulations, so we let the formula carry the idea, and it carries it exactly.

▪ YOUR TURN

Recompute the standard error with $n = 4$ instead of 16 (use $Sx = 1.15$). Bigger or smaller than 0.2875? Smaller samples give a larger standard error, exactly what the chapter warned.

Chapter 8

How Sure Are We?

■ WHAT THIS CHAPTER DOES ON THE TI-84

Build a confidence interval for a mean and for a proportion, and read it the way the chapter taught.

A confidence interval is a range of believable values for the truth. The TI hands you the whole interval from the **TESTS** menu.

TI-84 (press these keys)

```
MEAN: 95% interval for the quiz scores (in L1)
STAT > TESTS > 8:TInterval
Inpt: Data, List: L1, C-Level: .95, then Calculate

PROPORTION: 95% interval, 18 "yes" out of 40
STAT > TESTS > A:1-PropZInt
x: 18, n: 40, C-Level: .95, then Calculate
```

What you'll see (on the screen)

```
TInterval:      (13.4, 14.6)
1-PropZInt:     (0.30, 0.61)
```

■ READ THE OUTPUT

The mean interval, about 13.4 to 14.6, says: we are 95% confident the true average score sits in there. The proportion interval, about 0.30 to 0.61, does the same for the true “yes” rate. Want 90% instead of 95%? Change C-Level to .90 and watch the interval shrink. The calculator asks for the confidence level right on the input screen, so the trade-off is always in front of you.

■ YOUR TURN

Re-run the TInterval with C-Level: .99. Wider or narrower than 95%? Why does being *more confident* cost you a *wider* net?

Chapter 9

Testing What We Believe

■ WHAT THIS CHAPTER DOES ON THE TI-84

Run a one-sample hypothesis test for a mean and for a proportion, and find the p-value.

Chapter 9 is the logic of the p-value. The TI runs the test and prints the p-value plainly; your job is to decide what it means.

TI-84 (press these keys)

MEAN: H0 says the true mean is 13 (scores in L1)

STAT > TESTS > 2:T-Test

Inpt: Data, mu0: 13, List: L1, then Calculate

PROPORTION: H0 says the true "yes" rate is 0.5

STAT > TESTS > 5:1-PropZTest

p0: .5, x: 18, n: 40, then Calculate

What you'll see (on the screen)

T-Test:

t = 3.46

p = 0.003

x-bar = 14

▪ READ THE OUTPUT

Find the p value in the output. Here it is 0.003, smaller than 0.05, so we reject the claim that the mean is 13; the evidence says the true mean is higher. Same move for the proportion test: read the p -value, compare it to your significance level, decide. The calculator never decides for you, and that is the whole point of the chapter. One honest note: the TI-84 has *no* built-in test for a single variance, so the chapter's variance test is one you set up by formula rather than by menu.

▪ CHECK YOUR VOICE

When the output fills the screen, the anxious voice says “I don’t know what any of this means.” You do. You are looking for one number, the p value, and asking one question: smaller than 0.05 or not? Everything else on that screen is supporting detail.

Chapter 10

Comparing Two Groups

■ WHAT THIS CHAPTER DOES ON THE TI-84

Compare two group means (independent and paired) and two proportions.

Chapter 10 asks: is the difference between two groups real, or just noise? One menu choice per situation, and one careful decision about which one fits.

TI-84 (press these keys)

INDEPENDENT groups: A in L1, B in L2

STAT > TESTS > 4:2-SampTTest

List1: L1, List2: L2, then Calculate

PAIRED (same people twice): before in L1, after in L2

first **make** the differences:

L2 - L1 > STO> L3 (on the list-name line at top of L3)

then STAT > TESTS > 2:T-Test on L3, μ_0 : 0

TWO PROPORTIONS: 18/40 vs 25/45

STAT > TESTS > 6:2-PropZTest

x1: 18, n1: 40, x2: 25, n2: 45, then Calculate

What you'll see (on the screen)

2-SampTTest:

means 78.2 and 66.5

p = 0.036

▪ READ THE OUTPUT

The two group means are about 78 and 67, and the p-value 0.036 is below 0.05, so the gap is bigger than noise would easily explain. Here is the decision the chapter cares about: 2-SampTTest is for two *separate* groups. For the *same* people measured twice, that test is wrong; instead you build the differences in L3 ($L2 - L1$) and run a plain T-Test against 0. Choosing the right design is the chapter's key move, and on the TI it is the difference between one menu path and the other.

▪ YOUR TURN

Which design, independent or paired, matches a weight-loss study that weighs the *same* people in January and June? Build the differences in **L3** and run T-Test on them.

Chapter 11

When Three or More Groups Disagree

■ WHAT THIS CHAPTER DOES ON THE TI-84

Run a one-way ANOVA to ask whether three or more group means differ.

When there are three or more groups, you do not run many t-tests. You run one ANOVA. Chapter 11's whole method is one command on the TI.

TI-84 (press these keys)

put the three groups in three lists:

L1: 72 75 68 L2: 80 83 78 L3: 65 70 60

STAT > TESTS > F:ANOVA(

then type: ANOVA(L1, L2, L3)

press ENTER

What you'll see (on the screen)

ANOVA:

F = 15.44

p = 0.004

■ READ THE OUTPUT

Read the p value, 0.004, below 0.05, so at least one group's mean really differs from the others. But ANOVA does not say *which* one. Here is an honest limit: the TI-84 has **no** built-in Tukey follow-up. When ANOVA is significant, you follow up by comparing the groups two at a time, or by eye from their means. That two-step logic, "is there a difference, then where," is the heart of the chapter, and on the TI the second step is on you.

■ YOUR TURN

Look at the three group means (run 1-Var Stats on each list). Which two groups sit farthest apart? That pair is your best guess for where the real difference lives.

Chapter 12

Counts, Categories, and Surprises

■ WHAT THIS CHAPTER DOES ON THE TI-84

Run both chi-square tests: goodness-of-fit (one variable) and independence (two variables).

Chapter 12 is about counts that surprise us. Are the categories as evenly split as expected? Are two categorical variables related? Two flavors of one test, both on the **TESTS** menu.

TI-84 (press these keys)

GOODNESS OF FIT: a die rolled 60 times, **is** it fair?

observed counts in L1: 8 12 9 11 7 13

expected counts in L2: 10 10 10 10 10 10

STAT > TESTS > **D**:X2GOF-Test

Observed: L1, Expected: L2, **df**: 5, then Calculate

INDEPENDENCE: caffeine and sleep related?

put the 2x2 counts in **matrix** [A] (from Chapter 5)

STAT > TESTS > **C**:X2-Test

Observed: [A], Expected: [B], then Calculate

([B] **is filled in for** you)

What you'll see (on the screen)

X2GOF-Test:

X-squared = 2.8

p = 0.73

▪ READ THE OUTPUT

For the die, the p-value 0.73 is large (well above 0.05), so we do *not* have evidence the die is unfair; the ups and downs are ordinary luck. For the caffeine-and-sleep table, the X2-Test on matrix [A] would give a small p-value if the two variables were related, and a large one if they act independently. The calculator even writes the expected counts into [B] for you. Same idea, two questions, exactly as the chapter frames them.

▪ YOUR TURN

Change the die counts in **L1** to something lopsided, like 2 3 4 20 6 25, and re-run X2GOF-Test. Does the p-value drop below 0.05 now? What would you conclude about that die?

Chapter 13

Drawing the Line

■ WHAT THIS CHAPTER DOES ON THE TI-84

Measure correlation, fit a regression line, read its slope and R^2 , and predict.

The last chapter draws the line through a scatterplot and asks what it means. The TI gives you the correlation, the line, and a prediction, all from two lists.

TI-84 (press these keys)

hours in L1: 1 2 3 4 5 6

grade in L2: 52 58 63 70 74 80

if **r** is not shown, turn it **on** ONCE:

2nd > CATALOG > DiagnosticOn, then ENTER

fit the line:

STAT > CALC > 8:LinReg(a+bx)

Xlist: L1, Ylist: L2, then Calculate

scatterplot:

2nd > STAT PLOT > 1:Plot1 > On, Type = scatter

then ZOOM > 9:ZoomStat

What you'll see (on the screen)

LinReg(a+bx):

a = 46.47 (intercept)

b = 5.69 (slope)

r = 0.9985

r-squared = 0.997

▪ READ THE OUTPUT

$r = 0.9985$ is a very strong positive relationship (if r is missing, run `DiagnosticOn` once and refit). The slope $b = 5.69$ means each extra study hour adds about 5.7 points. r -squared $= 0.997$ says hours explain about 99.7% of the variation in grade (unusually clean, because this is teaching data). To predict at 3.5 hours, put the numbers into the equation by hand: $46.47 + 5.69 \times 3.5 \approx 66.4$. For a formal test of the slope, use `STAT > TESTS > LinRegTTest`. Slope, strength, prediction: the whole chapter, on two lists.

▪ CHECK YOUR VOICE

You just fit a regression model, the thing that sounded impossible on day one. Notice that. The keystrokes got shorter as the ideas got bigger, because the calculator carries the arithmetic. You carried the understanding the whole way.

The Datasets, Chapter by Chapter

Every example in this book comes with a ready-made data file, so you always have the exact numbers to type. All nine files live in a folder called `data`, sitting right next to this PDF. They are plain `.csv` files, which means they open in **StatCrunch**, **jamovi**, Excel, and Google Sheets. The TI-84 is the one tool in the set that *cannot* open them: it has no file reader. So on the calculator, you open the `.csv` on a computer (or read the number in this book), and type the columns into lists by hand.

How to load a dataset onto the TI-84 (do this once per file)

There is no import button; you type. Press **STAT**, choose 1:Edit, and enter each column into a list:

- **Numeric columns** go straight into **L1** (and **L2** for a second number), one value per row.
- **Categorical columns** get *coded* as numbers (Home = 1, Library = 2, Cafe = 3), or, for a two-way table, you type the *counts* instead of the raw categories.

TI-84 (press these keys)

```
STAT > 1:Edit
```

```
L1: type the quiz_score column, one value per ENTER
```

```
L2: type a second numeric column, or coded categories
```

It is more typing than the other tools, but the data sets in this book are small on purpose, and every “What you’ll see” box gives you the exact numbers.

The nine datasets

File	Chapters	What's inside (columns)
class.csv	1, 3, 4, 8, 9	16 students: student, study_spot (Home/Library/Cafe), quiz_score
population40.csv	2	40 students to sample from: student_id, gpa
caffeine_sleep.csv	5, 12	100 people: caffeine (Yes/No), slept_well (Yes/No)
survey40.csv	8, 9	40 people: person, response (Yes/No; 18 said Yes)
two_groups.csv	10	Two independent groups: id, group (A/B), score
before_after.csv	10	Same people twice: student, before, after
three_groups.csv	11	Three groups for ANOVA: id, group (A/B/C), score
dice60.csv	12	A die rolled 60 times: face (1 to 6), count
study_hours.csv	13	hours studied and the grade earned (6 students)

Two chapters need no data typed in. Chapter 6 (Predictable Randomness) uses the calculator's built-in DISTR functions (`normalcdf`, `binompdf`, and `friends`), and Chapter 7 (From Sample to Truth) works from the standard-error *formula* rather than a data file. For those two, the keystrokes make the numbers; there is nothing to type into a list beyond the values already inside the commands.

Typing each one is the same move. For the two-way table in Chapter 5 and 12, type the four counts into matrix [A] instead of a list. For the two-group test in Chapter 10, put group A in L1 and group B in L2, then point the test at both lists. One set of numbers, entered by hand, every chapter.

Quick Reference: The Whole Book on the TI-84

One keystroke path per idea

Idea (chapter)	TI-84
Type data into a list (1)	STAT > 1:Edit
Random sample (2)	MATH > PROB > randIntNoRep(1,40)
Histogram / boxplot (3)	2nd STAT PLOT > On, then ZOOM > 9
Center & spread (4)	STAT > CALC > 1:1-Var Stats
Two-way table (5)	2nd MATRX > EDIT [A], divide by hand
Normal / binomial (6)	2nd DISTR: normalcdf, invNorm, binompdf
Standard error (7)	Sx / sqrt(n), then normalcdf
Confidence interval (8)	STAT > TESTS > 8:TInterval / A:1-PropZInt
One-sample test (9)	STAT > TESTS > 2:T-Test / 5:1-PropZTest
Two groups (10)	STAT > TESTS > 4:2-SampTTest / 6:2-PropZTest
ANOVA (11)	STAT > TESTS > F:ANOVA(L1,L2,L3)
Chi-square (12)	STAT > TESTS > D:X2GOF-Test / C:X2-Test
Regression (13)	STAT > CALC > 8:LinReg(a+bx)

You did statistics, and you did it on a calculator that fits in your bag. That is not a small thing. Keep this page next to you; it is the whole course in thirteen keystroke paths.