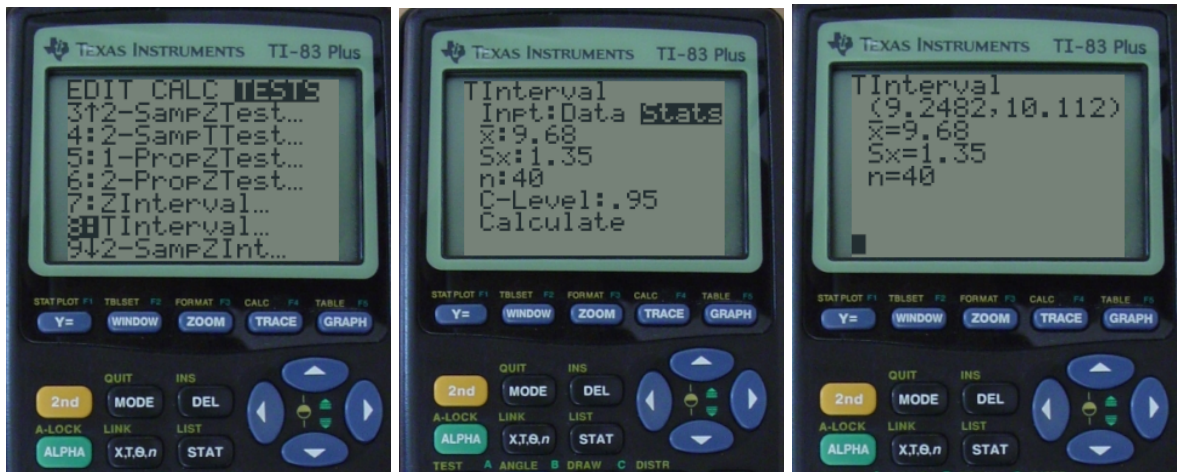


Tue, Apr 7 - Slide 279-280

- Imagine there is some **theoretical truth** represented by a “**population**” statistic (usually represented by Greek letters like μ or σ). Often there is no concrete “population”, it’s a conceptual “long-run” statistic.
- It is often impossible, impractical, or prohibitively expensive to know the precise truth, so we settle for **estimates: sample statistics** (e.g. $\bar{x}$ or s) that are based on relatively small samples.
- We **infer** that, as the sample size n gets larger:
 - **mean**: μ should be “around” $\bar{x}$
 - **standard deviation**: σ should be “close to” s
 - **proportion**: p should be “approximately” $\hat{p}$
 - **correlation**: ρ (rho) should be “near” r
- For example, suppose I ask a random sample of $n = 20$ C-N students when they got up this morning, and get $\bar{x} = 9.89$. Another professor might poll a different sample of $n = 20$ students and get an answer of $\bar{x} = 9.47$. Every sample will yield different sample statistics because of luck of the draw, and measurement noise.
- Maybe I combine the samples to get $n = 40$ and $\bar{x} = 9.68$. Then I can infer that the average waking time for all 2000 C-N students is $\mu \approx 9.68$. Maybe it’s really 9.83 or 9.56; we know there must be some wiggle room, but it’s probably fairly close to 9.68.
- In this example, the population size is 2000, and the **sample size** is only 40. It’s much easier to survey 40 students than all 2000. Practically speaking, we don’t need to know the exact μ for all 2000 students. A rough estimate is often good enough, but it’s good to know how far off we might be.
- The next topic is **confidence intervals**, where we’ll learn how much wiggle room to put around our estimates.

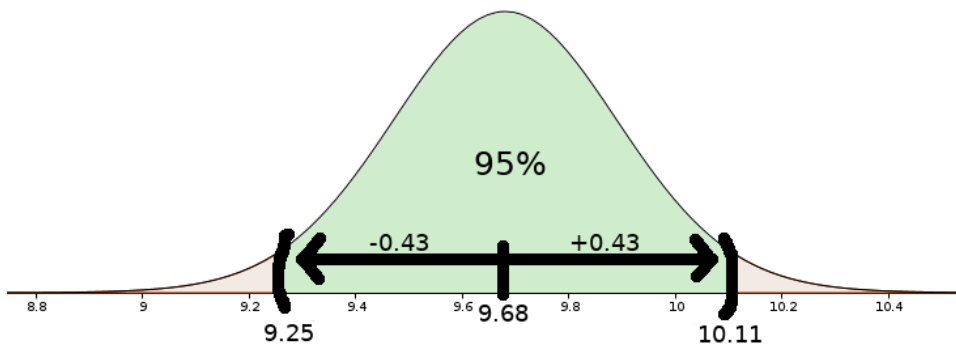
Tue, Apr 7 - Slide 285,288

- We will skip over some mathematical details, but suffice it to say that under certain assumptions, there is a bell curve distribution of the errors you make when estimating a statistic.
- Imagine going far enough out in the tails so that you are very likely to **cover** the truth with a **confidence interval** around your **point estimate**.
- Consider our example of $\bar{x} = 9.68$ average wake time for a sample of $n = 40$ students. We also need to know the sample standard deviation; suppose we had the data in Excel and got $s = 1.35$. Let’s find a 95% confidence interval (CI) around 9.68. This is a **window** which is 95% sure of containing the true mean μ for all 2000 students.
- In your TI calculator, hit STAT, go over to TESTS, and scroll down to **tInterval**.



You have the “Stats”, so enter them and set the “C-level” to .95.

- It should give you the CI: (9.25, 10.11). Sketch that interval on a number line, and notice that 9.68 would be the midpoint, and the endpoints are .43 in each direction.



- It is common to write the CI as $9.68 \pm .43$. Here the **sample mean** $\bar{x} = 9.68$ is the **point estimate**, and 0.43 is the **margin of error** (MOE).
- A good rule of thumb is that the 95% MOE for μ is about double the standard error:

$$\approx 2 \frac{1.35}{\sqrt{40}} \approx .43$$

- We could say that there is a 95% chance that the average wake time for all 2000 C-N students was between 9.25 and 10.11 (that is 9:15 and 10:07 AM).
- We are implicitly assuming that there was no bias either in our choice of $n = 40$ students or the answers they gave. Statistics can be worse than useless if you have bad sample data, or if mathematical assumptions aren't valid.

Tue, Apr 7 - Slide 290

Sketch the confidence interval on a number line to visualize the margin of error.

1. 74 ± 3 is the same as (71, 77)
2. (12.4, 13.1) is the same as 12.75 ± 0.35

Tue, Apr 7 - Slide 292

1. In this example, you have the raw data, not the stats. So first enter the data in L1, then select DATA when you do the **tInterval**. Also, set the confidence level to .90.



We could write the 90% CI as $5.55 \pm .78$.

2. Suppose you aren't comfortable with such a wide CI. As a **thought experiment**, keep $\bar{x}$ and s the same, but for a bigger sample of $n = 14$.



So now the 90% CI is $5.55 \pm .55$, and even the low-end estimate is at least 5 oz/ton.

It's logical that bigger sample sizes shrink the confidence interval (lower the MOE).

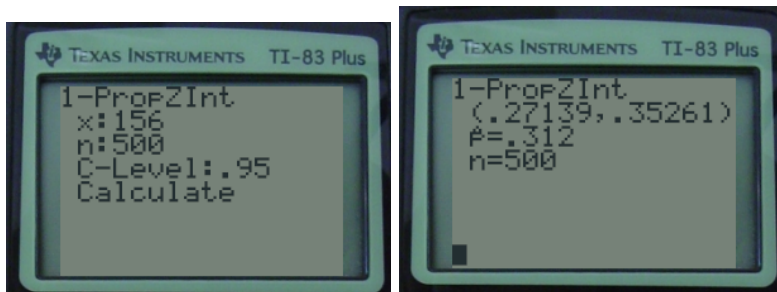
Tue, Apr 7 - Slide 293

We can also place confidence intervals around a sample proportion.

- Suppose we flipped $n = 500$ Hershey kisses, and 156 of them landed on the base. The **point estimate** is our **sample proportion**:

$$\hat{p} = \frac{156}{500} = 0.312$$

- Go to STAT-TESTS-(scroll down)-**1propZint**, and enter $x = 156$ and $n = 500$, and get a 95% CI.



We could write the CI $(.271, .353) = .312 \pm .041$.

- The results of an experiment like this might be reported as 31.2% with a margin of error of 4.1 percentage points. This essentially means there is a 95% chance that the long-run proportion would settle down within 4.1 percentage points of 31.2%.

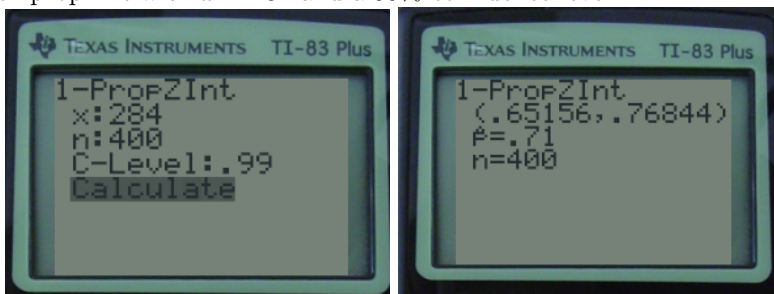
Tue, Apr 7 - Slide 295

Every CI has a **confidence level** attached to it. 95% is very common, but you might see 90% or 99%, etc. The complement to this number is called the **significance level** α .

confidence level	significance level
.95	$\alpha = .05$
.90	$\alpha = .10$
.99	$\alpha = .01$
.999	$\alpha = .001$

Tue, Apr 7 - Slide 297

Let's do $n = 400$ and $\alpha = .01$. She made 284 shots (71% of 400).
Use 1propZint with $x = 284$ and a 99% confidence level.



Assuming her skill is constant, we estimate with 99% certainty that her long-run free throw percentage would be between 65.16% and 76.84%. Or we could write that as $71 \pm 5.84\%$.

Tue, Apr 7 - Slide 298

If you are interested in **totals**, multiply your estimates by the population size.

1. $\hat{p} = \frac{243}{18000} = .0135$, so 1.35% of those surveyed watch the show.
Use 1propZint, and multiply by 320 million.
2. Those surveyed filled an average of 12.2 prescriptions last year.
Since this is a mean, use tInterval with the given stats, and then multiply by 320 million.

Tue, Apr 7 - Slide 299

NOTE: imagine this is testing for COVID-19 instead of vampires.

1. 1propZint with $x = 35$ and $n = 500$ gives $(.0476, .0924) = 7.0 \pm 2.24\%$ of the population is infected.
2. "victims per vampire" indicates we are talking about a mean.
The sample mean was 203 victims for 35 vampires, or $\bar{x} = \frac{203}{35} = 5.8$ victims per vampire.
tInterval with $n = 35$, $\bar{x} = 5.8$, and $s = 2$ gives $(5.11, 6.49) = 5.80 \pm .69$.

We didn't test everyone, but accounting for sample size and variation, we estimate that $7.0 \pm 2.24\%$ of the population are vampires, and that the town's vampires averaged 5.8 ± 0.69 victims last year.