

Prof. Green

Intro Stats // Fall 1998

Homework Assignment 3:
Politics Section Supplement

1. Using Minitab, calculate the mean for index of racial tolerance in this sample . How much does it differ from the value of 12.7 that is characteristic of youth samples of white respondents in the country as a whole? Calculate the approximate standard error of the sample mean. Construct a 95% confidence interval around the sample mean by adding and subtracting 1.96 times the standard error. Does 12.7 fall inside this interval?

T Confidence Intervals

Variable	N	Mean	StDev	SE Mean	95.0 % CI
ractoler	264	12.8598	1.5548	0.0957	(12.6714, 13.0483)

Answer: just barely.

2. Examine a crosstabulation of 'ractoler' by 'treatmnt' that percentages the table in a meaningful way. What do you infer from this table?

Rows: ractoler		Columns: treatmnt	
	0	1	All
7	0	1	1
	--	0.56	0.38
8	5	0	5
	5.81	--	1.89
9	3	8	11
	3.49	4.49	4.17
10	2	7	9
	2.33	3.93	3.41
11	8	8	16
	9.30	4.49	6.06
12	12	19	31
	13.95	10.67	11.74
13	22	41	63
	25.58	23.03	23.86
14	34	94	128
	39.53	52.81	48.48
All	86	178	264
	100.00	100.00	100.00

A smaller proportion of the treatment=0 group scores highly on racial tolerance. For example, only 39.5% of the treatment=0 group obtains a score of 14, as compared to 52.8% of the treatment=1 group.

3. Pull down the menu Stat > Basic Stats > 2 sample t > Samples in one column; insert ractoler in the Samples slot and treatmnt in the Subscript slot. Test whether the treatment (attending an integrated course) increases racial tolerance -- that is, the null hypothesis is "greater than." What is the mean level of tolerance in each group? Is this difference bigger than one would ordinarily attribute to random chance?

Two sample T for ractoler

treatmnt	N	Mean	StDev	SE Mean
0	86	12.57	1.73	0.19
1	178	13.00	1.45	0.11

95% CI for mu (0) - mu (1): (-0.86, -0.00)

T-Test mu (0) = mu (1) (vs <): T = -2.00 P = 0.024 DF = 144

The means are 12.57 and 13.00, for a difference of -.43. A difference of zero is what we would expect if both treatment and control groups had been drawn from the same population.

When we speak of an X% confidence interval, we are referring to a process or algorithm that captures the true difference in means X% of the time across hypothetical replications of the experiment.

The 95% confidence interval around this mean extends from -.86 to slightly less than 0, suggesting that the true difference in means may not be zero. Note that the level of confidence is rather arbitrary: a 99% interval would have included zero, while a 90% interval would not have. As it stands, zero is right at the border of the 95% interval.

For future reference, the "p-value" here (.024) represents the probability that one would observe a mean difference this small (-.43) if the true population difference had been zero.

4. Create a new column of data that contains only the respondents in the control group. We wish to simulate what would happen if respondents were drawn at random from this "population"; by picking a 'lucky' collection of control respondents, could one obtain a mean for 'ractoler' as high as what was actually observed in the treatment group? First, sort the data by 'treatment' and put the sorted data in a new pair of columns. Then copy and paste the 'ractoler' scores for the control group into a new column. Next, use the Calc > Random Data > Sample from Columns menu (with replacement) to generate samples with the same number of respondents (178) as in the treatment group. Calculate the mean for each of these fictitious samples (20 or so samples should give you a pretty good sense of the sampling variability). What percentage of these hypothetical samples had a mean 'ractoler' score as high as was observed in the treatment group?

Although this procedure is not a rigorous way of demonstrating that the two treatment groups differ by more than random chance, it is a nice way to get a feel for chance variation. As it happened, just one of my 20 random samples of size 178 had a mean of 13. But when I repeated the procedure again, none of my 20 samples had a mean this large.

5. *Briefly propose a nonexperimental design to test this substantive hypothesis. What are the strengths and weaknesses of experimentation as applied to this problem?*

One could survey students in Outward Bound courses and simply compare those who happened to attend integrated courses with those who attended racially homogeneous courses. This design is not as strong as the experiment described above, since we cannot be sure that students do not self-select which group to attend based on preexisting differences in racial attitudes.

1. *Construct a Minitab dataset of these results and produce a crosstabulation of availability by treatment/control condition. (I realize that you could do this problem set faster if the data were pre-packaged, but you will learn about how to code your own data if you do it once yourself; this lesson may prove quite valuable down the road.)*

Your dataset should have 180 rows (one for each phone call) and four columns: a 0/1 variable indicating whether the apartment was available, a 0/1 variable for gay self-identification, a 0/1 variable for sex, and a variable for city.

Percentage the table in a way helps answer the question: did gay self-identification increase the odds that the landlord said that the apartment was unavailable?

Control: sex = 0			
Rows: negative		Columns: treat	
	0	1	All
0	33 73.33	12 26.67	45 50.00
1	12 26.67	33 73.33	45 50.00
All	45 100.00	45 100.00	90 100.00

Among women, gay self-identification caused the probability of being turned away to climb from 26.7% to 73.3%.

Control: sex = 1			
Rows: negative		Columns: treat	
	0	1	All

0	30	10	40
	66.67	22.22	44.44
1	15	35	50
	33.33	77.78	55.56
All	45	45	90
	100.00	100.00	100.00

Among men, gay self-identification caused the probability of being turned away to climb from 33.3% to 77.8%.

Rows: negative		Columns: treat		
	0	1	All	
0	63	22	85	
	70.00	24.44	47.22	
1	27	68	95	
	30.00	75.56	52.78	
All	90	90	180	
	100.00	100.00	100.00	

Overall, gay self-identification caused the probability of being turned away to climb from 30% to 75.6%.

2. Produce a crosstabulation of availability by treatment by sex of caller. Percentage the results intelligently and interpret.

Rows: negative		Columns: sex		
	0	1	All	
0	45	40	85	
	50.00	44.44	47.22	
1	45	50	95	
	50.00	55.56	52.78	
All	90	90	180	
	100.00	100.00	100.00	

Men were somewhat more likely to receive negative replies. 55.6% of men were turned away, as compared to 50% of women.

3. Using Stat > Basic Stats > 2 Proportions, conduct a statistical test to assess whether the differences in availability rates differ between treatment and control conditions. Can the observed differences between conditions be attributed to random chance?

Two sample T for negative

treat	N	Mean	StDev	SE Mean
0	90	0.300	0.461	0.049
1	90	0.756	0.432	0.046

95% CI for $\mu(0) - \mu(1)$: (-0.587, -0.324)

T-Test $\mu(0) = \mu(1)$ (vs not =): T = -6.84 P = 0.0000 DF = 177

The observed difference in proportions (-.456) is our best guess of the effect of gay self-identification. The 95% confidence interval extends from -.587 to -.324, which does not include zero. It does not appear that these samples are drawn from the same population.