

Name: _____

TA: _____

Imagine you run an experiment testing the effect of weather on people's mood. You administer the PANAS (Positive and Negative Affect Scale) to one group of subjects on a rainy day and to another group on a sunny day and get the following data.

Rain: {35, 29, 72, 16, 58, 46, 47, 29, 85, 33}

Sun: {33, 32, 79, 20, 66, 47, 53, 39, 83, 38}

Now you'll do an independent-samples t-test to determine whether there's a reliable difference between the means of your samples. To be safe, do a two-tailed test.

1. Write the null hypothesis in words.

There is no effect of weather on a person's average mood.

2. Write the null hypothesis as an equation.

$$\mu_{\text{rain}} = \mu_{\text{sun}}$$

3. Write the alternative hypothesis in words.

A person's average mood differs on rainy and sunny days.

4. Write the alternative hypothesis as an equation.

$$\mu_{\text{rain}} \neq \mu_{\text{sun}}$$

5. Calculate the means of your samples.

$$M_{\text{rain}} = \frac{35 + 29 + 72 + 16 + 58 + 46 + 47 + 29 + 85 + 33}{10} = 45$$

$$M_{\text{sun}} = \frac{33 + 32 + 79 + 20 + 66 + 47 + 53 + 39 + 83 + 38}{10} = 49$$

6. Calculate the mean square based on both samples.

$$MS = \frac{\sum_{\text{rain group}} (X - M_{\text{rain}})^2 + \sum_{\text{sun group}} (X - M_{\text{sun}})^2}{n_{\text{rain}} + n_{\text{sun}} - 2}$$

$$= \frac{10^2 + 16^2 + 27^2 + 29^2 + 13^2 + 1^2 + 2^2 + 16^2 + 40^2 + 12^2 + 16^2 + 17^2 + 30^2 + 29^2 + 17^2 + 2^2 + 4^2 + 10^2 + 34^2 + 11^2}{18}$$

$$= 448.44$$

7. Calculate the standard error for the difference between sample means.

$$\sigma_{M_{\text{rain}} - M_{\text{sun}}} = \sqrt{MS \left(\frac{1}{n_{\text{rain}}} + \frac{1}{n_{\text{sun}}} \right)} = \sqrt{448.44 \left(\frac{1}{10} + \frac{1}{10} \right)} = 9.47$$

8. Calculate the t statistic for testing the null hypothesis you wrote above.

$$t = \frac{M_{\text{rain}} - M_{\text{sun}}}{\sigma_{M_{\text{rain}} - M_{\text{sun}}}} = \frac{45 - 49}{9.47} = -.42$$

9. The two-tailed critical value for this test is ± 2.10 . Which hypothesis do the data support?

Null hypothesis: There is no reliable difference between sample means.

Now imagine you had the same data, but the samples were based on the same set of subjects. That is, the first entries in both samples are from the same person, and so on.

10. What kind of test should you do now?

Paired-samples t-test

11. Calculate the difference scores and their mean.

$$X_{\text{diff}} = \{2, -3, -7, -4, -8, -1, -6, -10, 2, -5\}$$

$$M_{\text{diff}} = -40/10 = -4$$

12. Calculate the mean square (i.e., the sample variance) from your difference scores.

$$MS = s^2 = \frac{\sum (X_{\text{diff}} - M_{\text{diff}})^2}{n-1} = \frac{6^2 + 1^2 + 3^2 + 0^2 + 4^2 + 3^2 + 2^2 + 6^2 + 6^2 + 1^2}{9} = 16.44$$

13. Calculate the standard error of your answer in question 11.

$$\sigma_{M_{\text{diff}}} = \frac{s}{\sqrt{n}} = \sqrt{\frac{16.44}{10}} = 1.28$$

14. Calculate the new t statistic.

$$t = \frac{M_{\text{diff}}}{\sigma_{M_{\text{diff}}}} = \frac{-4}{1.28} = -3.12$$

15. The two-tailed critical value for this test is ± 2.26 . Which hypothesis do the data support now?

Alternative hypothesis: There is a reliable difference between the sample means.

16. Think about why the conclusion changed from question 9 to question 15. Notice the mean difference score from question 11 is the same as the difference of your means in question 5. So, the effect size is the same, but it was reliable in one case and not in the other. This is because the standard error changed, which in turn is because the mean square changed. Look at your mean squares from questions 6 and 12, and think about what those numbers mean. Write some thoughts here on why the mean square differs so much between tests, and why this difference explains the different conclusions between the two tests.

The mean square in each case is an estimate of the population variance. In the first case, we're talking about variance of individual PANAS scores. In the second case, we're talking about variances of **differences** in PANAS scores between rainy and sunny days. The first variance is greater because of variability across people. Notice that in the dataset, some people are really unhappy overall (e.g., the 4th subject) and others are really happy overall (e.g., the 9th subject). This variance across people doesn't contribute to the variance of differences scores, because each difference score is based on a single person. If that person is really happy on a sunny day, they're probably pretty happy on a rainy day.

Because difference scores are less variable than individual scores, the means of difference scores are more reliable than the means of individual scores. Therefore the difference between sample means is more reliable for the paired-samples test, even though the difference itself is the same.