

Inference for a Difference in Proportions

Nate Wells

Math 141, 4/8/22

Outline

In this lecture, we will. . .

Outline

In this lecture, we will. . .

- Calculate confidence intervals and perform hypothesis tests for proportions using the theory-based method
- Investigate the theoretical distribution for differences in proportions
- Calculate confidence intervals and conduct hypothesis tests for differences in proportions

Taste Test

- On Wednesday, Math 141 students participated in an experiment to determine whether the typical Reed student can distinguish between two different flavors of carbonated water.

Taste Test

- On Wednesday, Math 141 students participated in an experiment to determine whether the typical Reed student can distinguish between two different flavors of carbonated water.
 - Each student was provided 3 cups; 2 of the cups had the same flavor, and the other cup had a different flavor. Students were asked to identify the cup that was different.
- Let p denote the true proportion of the population who can correctly identify the cup that is different.

Taste Test

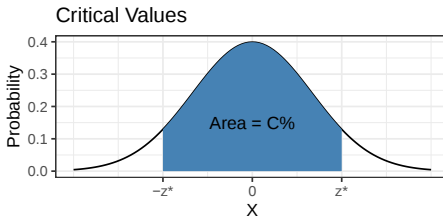
- On Wednesday, Math 141 students participated in an experiment to determine whether the typical Reed student can distinguish between two different flavors of carbonated water.
 - Each student was provided 3 cups; 2 of the cups had the same flavor, and the other cup had a different flavor. Students were asked to identify the cup that was different.
- Let p denote the true proportion of the population who can correctly identify the cup that is different.
- Suppose the two flavors of carbonated water are distinguishable. Why is it still plausible $p < 1$?

Critical Values

- The **critical value** z^* for a $C\%$ confidence interval is the value so that $C\%$ of area is between $-z^*$ and z^* in the standard Normal distribution

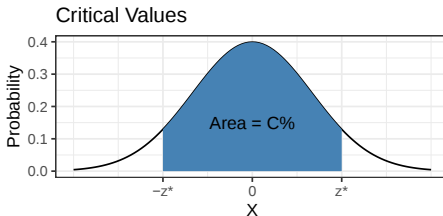
Critical Values

- The **critical value** z^* for a $C\%$ confidence interval is the value so that $C\%$ of area is between $-z^*$ and z^* in the standard Normal distribution



Critical Values

- The **critical value** z^* for a $C\%$ confidence interval is the value so that $C\%$ of area is between $-z^*$ and z^* in the standard Normal distribution



- For Normal distributions, approximately 95% of observations are within 2 standard deviations of the mean.
 - So the critical value for 95% confidence is approximately
$$z^* = 2 \quad (\text{exact value is } z^* = 1.96)$$

Confidence Intervals

When a sample statistic is approximately Normally distribution, the $C\%$ confidence interval is

$$\text{statistic} \pm z^* \cdot SE$$

where z^* is the critical value for $C\%$ confidence and SE is the standard error for the statistic.

Confidence Intervals

When a sample statistic is approximately Normally distribution, the $C\%$ confidence interval is

$$\text{statistic} \pm z^* \cdot SE$$

where z^* is the critical value for $C\%$ confidence and SE is the standard error for the statistic.

- The standard error for a sample proportion $\hat{p}$ is $SE = \sqrt{\frac{p(1-p)}{n}}$. Since we don't know p , we estimate it in the SE formula with $\hat{p}$.

Confidence Intervals

When a sample statistic is approximately Normally distributed, the $C\%$ confidence interval is

$$\text{statistic} \pm z^* \cdot SE$$

where z^* is the critical value for $C\%$ confidence and SE is the standard error for the statistic.

- The standard error for a sample proportion $\hat{p}$ is $SE = \sqrt{\frac{p(1-p)}{n}}$. Since we don't know p , we estimate it in the SE formula with $\hat{p}$.

Theorem

Suppose an SRS of size n is collected from a population with parameter p . If n is large enough so that both $n\hat{p}$ and $n(1 - \hat{p})$ are at least 10, then the confidence interval for p is

$$\hat{p} \pm z^* \sqrt{\frac{\hat{p}(1 - \hat{p})}{n}}$$

Taste Test Continued

- Suppose we are interested in estimating the value of p , the proportion of the population who will correctly identify the different cup.

Taste Test Continued

- Suppose we are interested in estimating the value of p , the proportion of the population who will correctly identify the different cup.
 - Create a 90% confidence interval for this parameter.

Taste Test Continued

- Suppose we are interested in estimating the value of p , the proportion of the population who will correctly identify the different cup.
 - Create a 90% confidence interval for this parameter.
- As before, our sample statistic is $\hat{p} = \frac{29}{59}$.

Taste Test Continued

- Suppose we are interested in estimating the value of p , the proportion of the population who will correctly identify the different cup.
 - Create a 90% confidence interval for this parameter.
- As before, our sample statistic is $\hat{p} = \frac{29}{59}$.
- The critical value for a 90% confidence interval is the number z^* so that 90% area is between $-z^*$ and z^* . It is the 0.95 **quantile**

Taste Test Continued

- Suppose we are interested in estimating the value of p , the proportion of the population who will correctly identify the different cup.
 - Create a 90% confidence interval for this parameter.
- As before, our sample statistic is $\hat{p} = \frac{29}{59}$.
- The critical value for a 90% confidence interval is the number z^* so that 90% area is between $-z^*$ and z^* . It is the 0.95 **quantile**

```
qnorm(p = .95, mean = 0, sd = 1)
```

```
## [1] 1.644854
```

Taste Test Continued

- Suppose we are interested in estimating the value of p , the proportion of the population who will correctly identify the different cup.
 - Create a 90% confidence interval for this parameter.
- As before, our sample statistic is $\hat{p} = \frac{29}{59}$.
- The critical value for a 90% confidence interval is the number z^* so that 90% area is between $-z^*$ and z^* . It is the 0.95 **quantile**

```
qnorm(p = .95, mean = 0, sd = 1)
```

```
## [1] 1.644854
```

- The standard error for $\hat{p}$ is

$$SE(\hat{p}) \approx \sqrt{\frac{\hat{p}(1 - \hat{p})}{n}} = \sqrt{\frac{0.49(1 - 0.49)}{59}} = 0.065$$

An Example

- The theory-based confidence interval takes the form

$$\hat{p} \pm z^* \cdot SE$$

An Example

- The theory-based confidence interval takes the form

$$\hat{p} \pm z^* \cdot SE$$

- In this case,

$$0.49 \pm 1.64 \cdot 0.065 \quad \text{or} \quad 0.49 \pm 0.1066$$

An Example

- The theory-based confidence interval takes the form

$$\hat{p} \pm z^* \cdot SE$$

- In this case,

$$0.49 \pm 1.64 \cdot 0.065 \quad \text{or} \quad 0.49 \pm 0.1066$$

- That is, a plausible range of values for p is 0.38 to 0.60, with confidence 90%.

An Example

- The theory-based confidence interval takes the form

$$\hat{p} \pm z^* \cdot SE$$

- In this case,

$$0.49 \pm 1.64 \cdot 0.065 \quad \text{or} \quad 0.49 \pm 0.1066$$

- That is, a plausible range of values for p is 0.38 to 0.60, with confidence 90%.
- How does this compare to the bootstrap method?

An Example

- The theory-based confidence interval takes the form

$$\hat{p} \pm z^* \cdot SE$$

- In this case,

$$0.49 \pm 1.64 \cdot 0.065 \quad \text{or} \quad 0.49 \pm 0.1066$$

- That is, a plausible range of values for p is 0.38 to 0.60, with confidence 90%.
- How does this compare to the bootstrap method?

```
set.seed(84)
lacroix %>% specify(response = correct, success = "yes") %>%
  generate(reps=5000, type = "bootstrap") %>%
  calculate(stat = "prop") %>%
  get_ci(level = .9, type = "percentile")
```

```
## # A tibble: 1 x 2
##   lower_ci upper_ci
##   <dbl>    <dbl>
## 1    0.390    0.593
```

Section 2

Difference in Proportions

Difference in Proportions

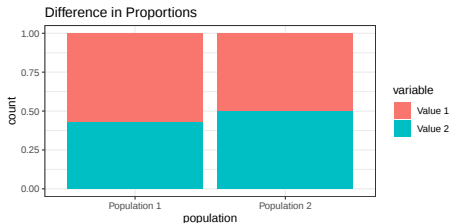
- Suppose we have two populations and wish to compare the proportions p_1 and p_2 of the level of a categorical variable in each population.

Difference in Proportions

- Suppose we have two populations and wish to compare the proportions p_1 and p_2 of the level of a categorical variable in each population.
- That is, we want to know the value of the difference $p_1 - p_2$ in proportion.

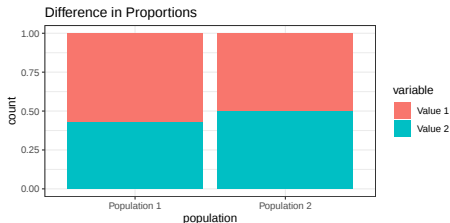
Difference in Proportions

- Suppose we have two populations and wish to compare the proportions p_1 and p_2 of the level of a categorical variable in each population.
- That is, we want to know the value of the difference $p_1 - p_2$ in proportion.



Difference in Proportions

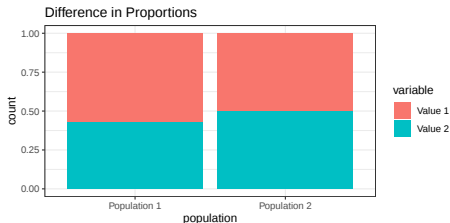
- Suppose we have two populations and wish to compare the proportions p_1 and p_2 of the level of a categorical variable in each population.
- That is, we want to know the value of the difference $p_1 - p_2$ in proportion.



- A reasonable point estimate for $p_1 - p_2$ is the difference in sample proportions $\hat{p}_1 - \hat{p}_2$ for a sample taken from the 1st and 2nd populations.

Difference in Proportions

- Suppose we have two populations and wish to compare the proportions p_1 and p_2 of the level of a categorical variable in each population.
- That is, we want to know the value of the difference $p_1 - p_2$ in proportion.



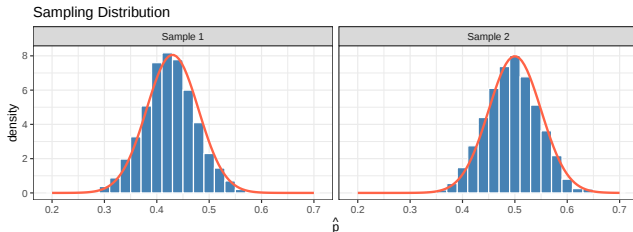
- A reasonable point estimate for $p_1 - p_2$ is the difference in sample proportions $\hat{p}_1 - \hat{p}_2$ for a sample taken from the 1st and 2nd populations.
- As long as we can verify that the statistic $\hat{p}_1 - \hat{p}_2$ has an approximately Normal distribution, we can use the same techniques we used for single sample proportions.

Distribution for $\hat{p}_1 - \hat{p}_2$

- We know that individually, both $\hat{p}_1$ and $\hat{p}_2$ are approximately Normal:

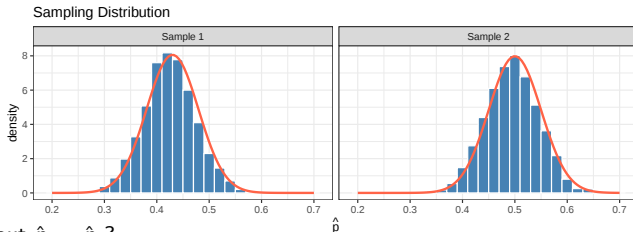
Distribution for $\hat{p}_1 - \hat{p}_2$

- We know that individually, both $\hat{p}_1$ and $\hat{p}_2$ are approximately Normal:



Distribution for $\hat{p}_1 - \hat{p}_2$

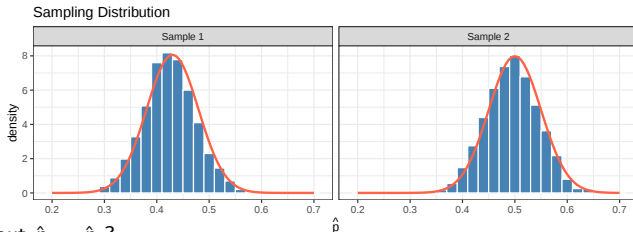
- We know that individually, both $\hat{p}_1$ and $\hat{p}_2$ are approximately Normal:



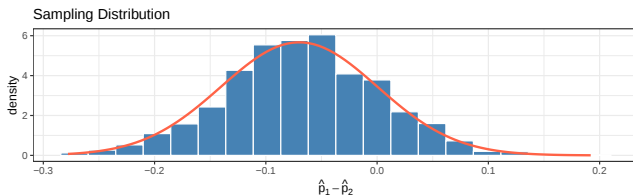
- What about $\hat{p}_1 - \hat{p}_2$?

Distribution for $\hat{p}_1 - \hat{p}_2$

- We know that individually, both $\hat{p}_1$ and $\hat{p}_2$ are approximately Normal:

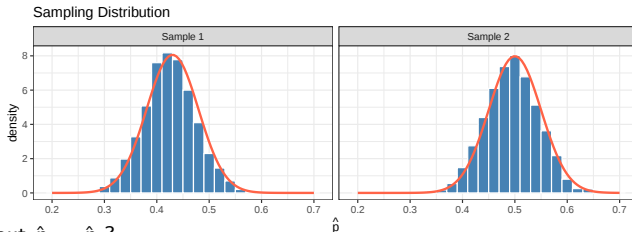


- What about $\hat{p}_1 - \hat{p}_2$?

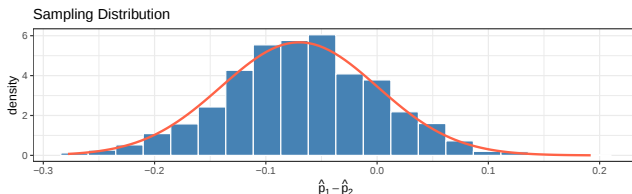


Distribution for $\hat{p}_1 - \hat{p}_2$

- We know that individually, both $\hat{p}_1$ and $\hat{p}_2$ are approximately Normal:



- What about $\hat{p}_1 - \hat{p}_2$?



- The sum or difference of **independent** Normal variables will also be Normal, with variance equal to the sum of individual variances.

Conditions for Theory-based Normal Approximation

Theorem

The difference $\hat{p}_1 - \hat{p}_2$ is approximately Normal when

- ① Each sample proportion is approximately normal (≥ 10 success/failure)
- ② The two samples are independent of each other

In this case, the standard error of the difference in sample proportions is

$$SE_{\hat{p}_1 - \hat{p}_2} = \sqrt{SE_{\hat{p}_1}^2 + SE_{\hat{p}_2}^2} = \sqrt{\frac{\hat{p}_1(1 - \hat{p}_1)}{n_1} + \frac{\hat{p}_2(1 - \hat{p}_2)}{n_2}}$$

Conditions for Theory-based Normal Approximation

Theorem

The difference $\hat{p}_1 - \hat{p}_2$ is approximately Normal when

- ① Each sample proportion is approximately normal (≥ 10 success/failure)
- ② The two samples are independent of each other

In this case, the standard error of the difference in sample proportions is

$$SE_{\hat{p}_1 - \hat{p}_2} = \sqrt{SE_{\hat{p}_1}^2 + SE_{\hat{p}_2}^2} = \sqrt{\frac{\hat{p}_1(1 - \hat{p}_1)}{n_1} + \frac{\hat{p}_2(1 - \hat{p}_2)}{n_2}}$$

- Importantly, we know the distribution is Normal and we have the standard error

Conditions for Theory-based Normal Approximation

Theorem

The difference $\hat{p}_1 - \hat{p}_2$ is approximately Normal when

- ① Each sample proportion is approximately normal (≥ 10 success/failure)
- ② The two samples are independent of each other

In this case, the standard error of the difference in sample proportions is

$$SE_{\hat{p}_1 - \hat{p}_2} = \sqrt{SE_{\hat{p}_1}^2 + SE_{\hat{p}_2}^2} = \sqrt{\frac{\hat{p}_1(1 - \hat{p}_1)}{n_1} + \frac{\hat{p}_2(1 - \hat{p}_2)}{n_2}}$$

- Importantly, we know the distribution is Normal and we have the standard error
 - We can use `qnorm` to find critical values for confidence intervals and `pnorm` to compute P-values for hypothesis tests

Partisanship

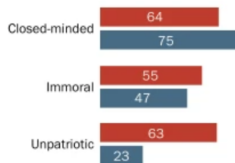
U.S. POLITICS | OCTOBER 10, 2019

Partisan Antipathy: More Intense, More Personal

The share of Republicans who give Democrats a "cold" rating on a 0-100 thermometer has risen 14 percentage points since 2016. Similarly, 57% of Democrats give Republicans a very cold rating, up from 2016.

% who say members of the **other** party are a lot/somewhat more ____ compared to other Americans

- Republicans say Democrats are more ...
- Democrats say Republicans are more ...



Partisanship

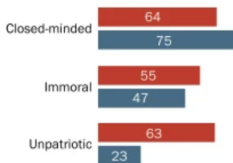
U.S. POLITICS | OCTOBER 10, 2019

Partisan Antipathy: More Intense, More Personal

The share of Republicans who give Democrats a "cold" rating on a 0-100 thermometer has risen 14 percentage points since 2016. Similarly, 57% of Democrats give Republicans a very cold rating, up from 2016.

% who say members of the *other* party are a lot/somewhat more ____ compared to other Americans

■ Republicans say Democrats are more ...
■ Democrats say Republicans are more ...



- Was there really a difference in the proportion of Democrats that view Republicans as close-minded compared to Republicans that view Democrats the same? Or is the difference just due to random sampling?

Confidence Intervals

Let's use the Normal approximation.

Confidence Intervals

Let's use the Normal approximation.

Elsewhere in the study, we find the number of Republicans and Democrats surveyed were 4948 and 4947, respectively.

Confidence Intervals

Let's use the Normal approximation.

Elsewhere in the study, we find the number of Republicans and Democrats surveyed were 4948 and 4947, respectively.

- Our standard error is therefore 0.009

```
SE<-sqrt(p_hat_r*(1-p_hat_r)/n_r + p_hat_d*(1-p_hat_d)/n_d )  
SE
```

```
## [1] 0.00919054
```

Confidence Intervals

Let's use the Normal approximation.

Elsewhere in the study, we find the number of Republicans and Democrats surveyed were 4948 and 4947, respectively.

- Our standard error is therefore 0.009

```
SE<-sqrt(p_hat_r*(1-p_hat_r)/n_r + p_hat_d*(1-p_hat_d)/n_d )  
SE
```

```
## [1] 0.00919054
```

- At a 95% confidence level, the critical value is $z^* = 1.96$

```
z<-qnorm(.975)  
z
```

```
## [1] 1.959964
```

Confidence Intervals II

- Assembling these pieces, the confidence interval for $p_r - p_d$ is

$$(\hat{p}_r - \hat{p}_d) \pm z^* \cdot SE$$

```
ci_low<-p_hat_r - p_hat_d - z*SE  
ci_high<-p_hat_r - p_hat_d + z*SE  
c(ci_low, ci_high)
```

```
## [1] -0.12801313 -0.09198687
```

Confidence Intervals II

- Assembling these pieces, the confidence interval for $p_r - p_d$ is

$$(\hat{p}_r - \hat{p}_d) \pm z^* \cdot SE$$

```
ci_low<-p_hat_r - p_hat_d - z*SE
ci_high<-p_hat_r - p_hat_d + z*SE
c(ci_low, ci_high)
```

```
## [1] -0.12801313 -0.09198687
```

- Note that both endpoints of the interval are less than 0, suggesting that the true difference in proportions between Republicans and Democrats is negative

Confidence Intervals II

- Assembling these pieces, the confidence interval for $p_r - p_d$ is

$$(\hat{p}_r - \hat{p}_d) \pm z^* \cdot SE$$

```
ci_low<-p_hat_r - p_hat_d - z*SE
ci_high<-p_hat_r - p_hat_d + z*SE
c(ci_low, ci_high)
```

```
## [1] -0.12801313 -0.09198687
```

- Note that both endpoints of the interval are less than 0, suggesting that the true difference in proportions between Republicans and Democrats is negative
 - i.e. a greater proportion of Democrats hold the view that Republicans are closed-minded compared to the converse

Confidence Interval via `infer`

Alternatively, we can use `infer` to compute confidence intervals.

Confidence Interval via `infer`

Alternatively, we can use `infer` to compute confidence intervals.

- We'll use the `pew` data set.

Confidence Interval via infer

Alternatively, we can use `infer` to compute confidence intervals.

- We'll use the `pew` data set.

```
pew %>% group_by(party,close_minded) %>%  
  summarize(N = n()) %>%  
  mutate(prop = N / sum(N))
```

```
## # A tibble: 4 x 4  
## # Groups:   party [2]  
##   party      close_minded      N  prop  
##   <chr>      <chr>      <int> <dbl>  
## 1 Democrat    no        1237 0.250  
## 2 Democrat    yes        3710 0.750  
## 3 Republican no        1781 0.360  
## 4 Republican yes        3167 0.640
```

Confidence Interval via infer II

```
boot<-pew %>%  
  specify(close_minded ~ party, success = "yes" ) %>%  
  generate(reps = 1000, type = "bootstrap" ) %>%  
  calculate( "diff in props", order = c("Republican", "Democrat") )
```

Confidence Interval via infer II

```
boot<-pew %>%
  specify(close_minded ~ party, success = "yes" ) %>%
  generate(reps = 1000, type = "bootstrap" ) %>%
  calculate( "diff in props", order = c("Republican", "Democrat") )

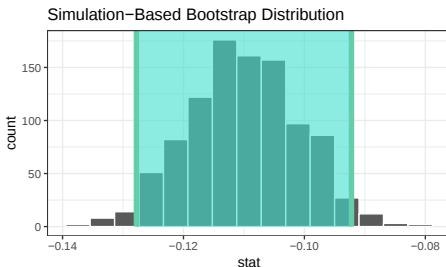
interval <-boot %>% get_confidence_interval(level = .95, type = "se",
  point_estimate = p_hat_r - p_hat_d)
interval

## # A tibble: 1 x 2
##   lower_ci upper_ci
##   <dbl>    <dbl>
## 1   -0.128  -0.0922
```

Confidence Interval via infer II

```
boot<-pew %>%  
  specify(close_minded ~ party, success = "yes" ) %>%  
  generate(reps = 1000, type = "bootstrap" ) %>%  
  calculate( "diff in props", order = c("Republican", "Democrat") )  
  
interval <-boot %>% get_confidence_interval(level = .95, type = "se",  
  point_estimate = p_hat_r - p_hat_d)  
interval
```

```
## # A tibble: 1 x 2  
##   lower_ci upper_ci  
##   <dbl>   <dbl>  
## 1   -0.128  -0.0922
```



Pooled sample for Hypothesis Tests

- Suppose we are interested in testing the following hypotheses

$$H_0 : p_1 = p_2 \quad H_a : p_1 \neq p_2$$

Pooled sample for Hypothesis Tests

- Suppose we are interested in testing the following hypotheses

$$H_0 : p_1 = p_2 \quad H_a : p_1 \neq p_2$$

- **If the null hypothesis is true**, collecting a sample of sizes n_1 and n_2 from each population is the same as collecting a single sample of size $n_1 + n_2$.

Pooled sample for Hypothesis Tests

- Suppose we are interested in testing the following hypotheses

$$H_0 : p_1 = p_2 \quad H_a : p_1 \neq p_2$$

- **If the null hypothesis is true**, collecting a sample of sizes n_1 and n_2 from each population is the same as collecting a single sample of size $n_1 + n_2$.
 - So we may instead consider the pooled proportion $\hat{p}$ given by

$$\hat{p} = \frac{\text{overall successes}}{\text{overall sample size}} = \frac{n_1 \hat{p}_1 + n_2 \hat{p}_2}{n_1 + n_2}$$

Pooled sample for Hypothesis Tests

- Suppose we are interested in testing the following hypotheses

$$H_0 : p_1 = p_2 \quad H_a : p_1 \neq p_2$$

- **If the null hypothesis is true**, collecting a sample of sizes n_1 and n_2 from each population is the same as collecting a single sample of size $n_1 + n_2$.
 - So we may instead consider the pooled proportion $\hat{p}$ given by

$$\hat{p} = \frac{\text{overall successes}}{\text{overall sample size}} = \frac{n_1 \hat{p}_1 + n_2 \hat{p}_2}{n_1 + n_2}$$

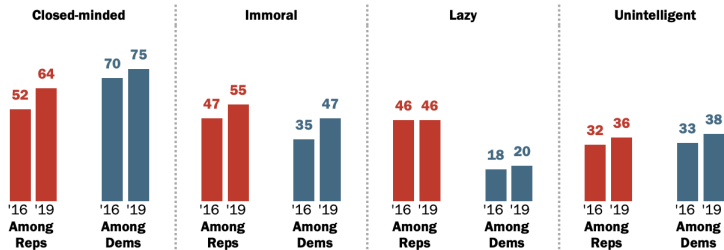
- This gives a standard error for the null distribution of

$$SE = \sqrt{\frac{\hat{p}(1 - \hat{p})}{n_1} + \frac{\hat{p}(1 - \hat{p})}{n_2}}$$

Partisanship over Time

Increasing shares of partisans see members of the other party as 'closed-minded' and 'immoral'

% who say members of the other party are a lot/somewhat more ____ compared to other Americans



Note: Partisans do not include leaners.

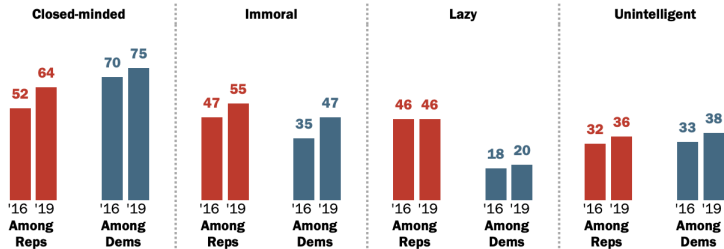
Source: Survey of U.S. adults conducted Sept. 3-15, 2019.

PEW RESEARCH CENTER

Partisanship over Time

Increasing shares of partisans see members of the other party as 'closed-minded' and 'immoral'

% who say members of the other party are a lot/somewhat more ____ compared to other Americans



Note: Partisans do not include leaners.

Source: Survey of U.S. adults conducted Sept. 3-15, 2019.

PEW RESEARCH CENTER

- Was there really a change in the proportion of Democrats that view Republicans as close-minded between 2016 and 2019?

Hypothesis Tests

We test

$$H_0 : p_{16} = p_{19} \qquad H_a : p_{16} \neq p_{19}$$

Hypothesis Tests

We test

$$H_0 : p_{16} = p_{19} \quad H_a : p_{16} \neq p_{19}$$

- Let's use the Normal approximation. In 2016, the number of participants was 4948 and in 2019, the number was 2947. This gives a pooled proportion of $\hat{p} = 0.725$

Hypothesis Tests

We test

$$H_0 : p_{16} = p_{19} \quad H_a : p_{16} \neq p_{19}$$

- Let's use the Normal approximation. In 2016, the number of participants was 4948 and in 2019, the number was 2947. This gives a pooled proportion of $\hat{p} = 0.725$

```
n_16<-4948  
n_19<-2947
```

```
p_hat_16<-.7  
p_hat_19<-.75
```

```
p_hat<-(p_hat_16*n_16 + p_hat_19*n_19)/(n_16 + n_19)
```

```
p_hat
```

```
## [1] 0.7249975
```

Hypothesis Tests

We test

$$H_0 : p_{16} = p_{19} \quad H_a : p_{16} \neq p_{19}$$

- Let's use the Normal approximation. In 2016, the number of participants was 4948 and in 2019, the number was 2947. This gives a pooled proportion of $\hat{p} = 0.725$

```
n_16<-4948  
n_19<-2947
```

```
p_hat_16<-.7  
p_hat_19<-.75
```

```
p_hat<-(p_hat_16*n_16 + p_hat_19*n_19)/(n_16 + n_19)  
p_hat
```

```
## [1] 0.7249975
```

- The standard error for the null distribution is 0.009

```
SE <- sqrt( p_hat*(1- p_hat)/n_16 + p_hat*(1- p_hat)/n_19 )  
SE
```

```
## [1] 0.008977568
```

Hypothesis Tests II

- Our test statistic is

$$z = \frac{\hat{p}_{16} - \hat{p}_{19}}{SE} = -5.57$$

```
z <- (p_hat_16 - p_hat_19)/SE  
z
```

```
## [1] -5.569437
```

Hypothesis Tests II

- Our test statistic is

$$z = \frac{\hat{p}_{16} - \hat{p}_{19}}{SE} = -5.57$$

```
z <- (p_hat_16 - p_hat_19)/SE  
z
```

```
## [1] -5.569437
```

- The P-value for this statistic is 0.00000002

```
P_value<-2*pnorm(z,0,1)  
P_value
```

```
## [1] 2.555634e-08
```

Hypothesis Tests II

- Our test statistic is

$$z = \frac{\hat{p}_{16} - \hat{p}_{19}}{SE} = -5.57$$

```
z <- (p_hat_16 - p_hat_19)/SE  
z
```

```
## [1] -5.569437
```

- The P-value for this statistic is 0.00000002

```
P_value<-2*pnorm(z,0,1)  
P_value
```

```
## [1] 2.555634e-08
```

- The test is significant at $\alpha = 0.01$ and we reject the null hypothesis.

Hypothesis Tests II

- Our test statistic is

$$z = \frac{\hat{p}_{16} - \hat{p}_{19}}{SE} = -5.57$$

```
z <- (p_hat_16 - p_hat_19)/SE  
z
```

```
## [1] -5.569437
```

- The P-value for this statistic is 0.00000002

```
P_value<-2*pnorm(z,0,1)  
P_value
```

```
## [1] 2.555634e-08
```

- The test is significant at $\alpha = 0.01$ and we reject the null hypothesis.
 - It is unlikely that the observed difference in proportions is due to chance, if the populations truly had the same proportion.

Hypothesis Test via `infer`

Let's now use the `pew2` data

Hypothesis Test via infer

Let's now use the pew2 data

```
pew2 %>% group_by(year,close_minded) %>%  
  summarize(N = n()) %>%  
  mutate(prop = N / sum(N))
```

```
## # A tibble: 4 x 4  
## # Groups:   year [2]  
##   year close_minded      N prop  
##   <chr> <chr>      <int> <dbl>  
## 1 2016 no          1484 0.300  
## 2 2016 yes         3464 0.700  
## 3 2019 no          1237 0.250  
## 4 2019 yes         3710 0.750
```

Hypothesis Tests via infer II

```
nullldist<-pew2 %>%  
  specify(close_minded ~ year, success = "yes" ) %>%  
  hypothesize(null = "independence") %>%  
  generate(reps = 1000, type = "permute" ) %>%  
  calculate( "diff in props", order = c("2016", "2019") )
```

Hypothesis Tests via infer II

```
nullldist<-pew2 %>%  
  specify(close_minded ~ year, success = "yes" ) %>%  
  hypothesize(null = "independence") %>%  
  generate(reps = 1000, type = "permute" ) %>%  
  calculate( "diff in props", order = c("2016", "2019") )
```

```
p_value <-nullldist %>% get_p_value(obs_stat = (p_hat_16 - p_hat_19),  
  direction = "both")
```

```
p_value
```

```
## # A tibble: 1 x 1  
##   p_value  
##   <dbl>  
## 1      0
```

Hypothesis Tests via infer II

```
nullldist<-pew2 %>%  
  specify(close_minded ~ year, success = "yes" ) %>%  
  hypothesize(null = "independence") %>%  
  generate(reps = 1000, type = "permute" ) %>%  
  calculate( "diff in props", order = c("2016", "2019") )
```

```
p_value <-nullldist %>% get_p_value(obs_stat = (p_hat_16 - p_hat_19),  
  direction = "both")
```

```
p_value
```

```
## # A tibble: 1 x 1  
##   p_value  
##   <dbl>  
## 1      0
```

