

Statistical Power

Overview

1. The Central Limit Theorem revisited
2. From Null worlds to True effect worlds
3. Statistical power
4. A power simulation

The Central Limit Theorem revisited

Are action movies better than comedies?

In the previous session on hypothesis testing, we invented a null world:

We simulated a population of 1'000'000 movies with no difference.

```
1 set.seed(1234) # For reproducibility
2
3 imaginary_movies_null <- tibble(
4   movie_id = 1:1000000,
5   rating = sample(seq(1, 10, by = 0.1), size = 1000000, replace = TRUE),
6   genre = sample(c("Comedy", "Action"), size = 1000000, replace = TRUE)
7 )
```

Sampling distribution

We randomly drew 1,000 samples from this population and calculated the difference between action movies and comedies for each.

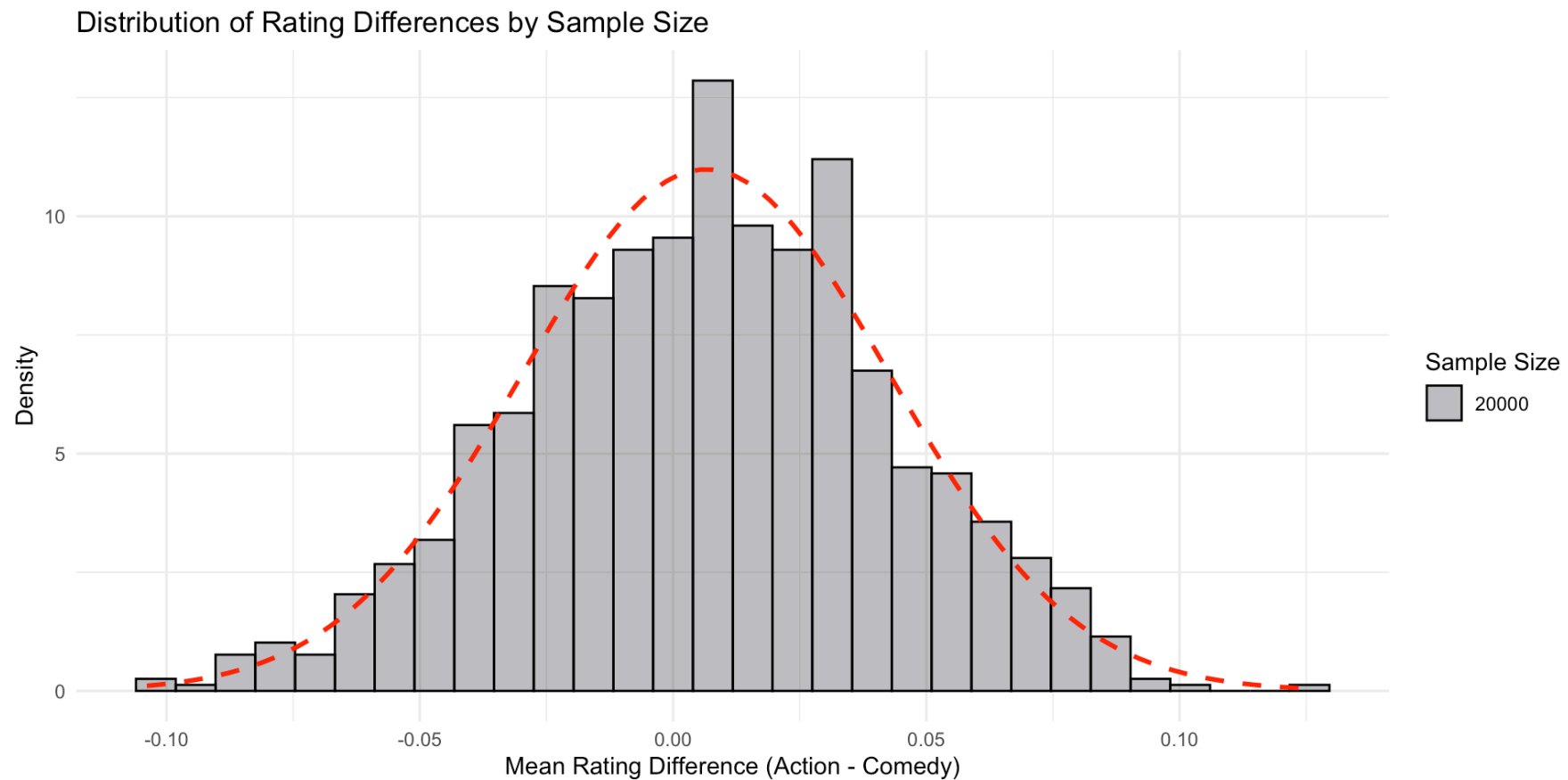
We called the distribution of the differences from the different samples the **sampling distribution**

Our sample size was always the same: 20,000.

```
1 n_simulations <- 1000
2 differences <- c() # make an empty vector
3 sample_size <- 20000
4
5 for (i in 1:n_simulations) {
6   # draw a sample of 20'000 films
7   imaginary_sample <- imaginary_movies |>
8     sample_n(sample_size)
9   # compute rating difference in the sample
10  estimate <- imaginary_sample |>
11    group_by(genre) |>
12    summarize(avg_rating = mean(rating)) |>
13    summarise(diff = avg_rating[genre == "Action"] - avg_rating[genre == "Comedy"]) %>%
14    pull(diff)
15
16  differences[i] <- estimate
17 }
```

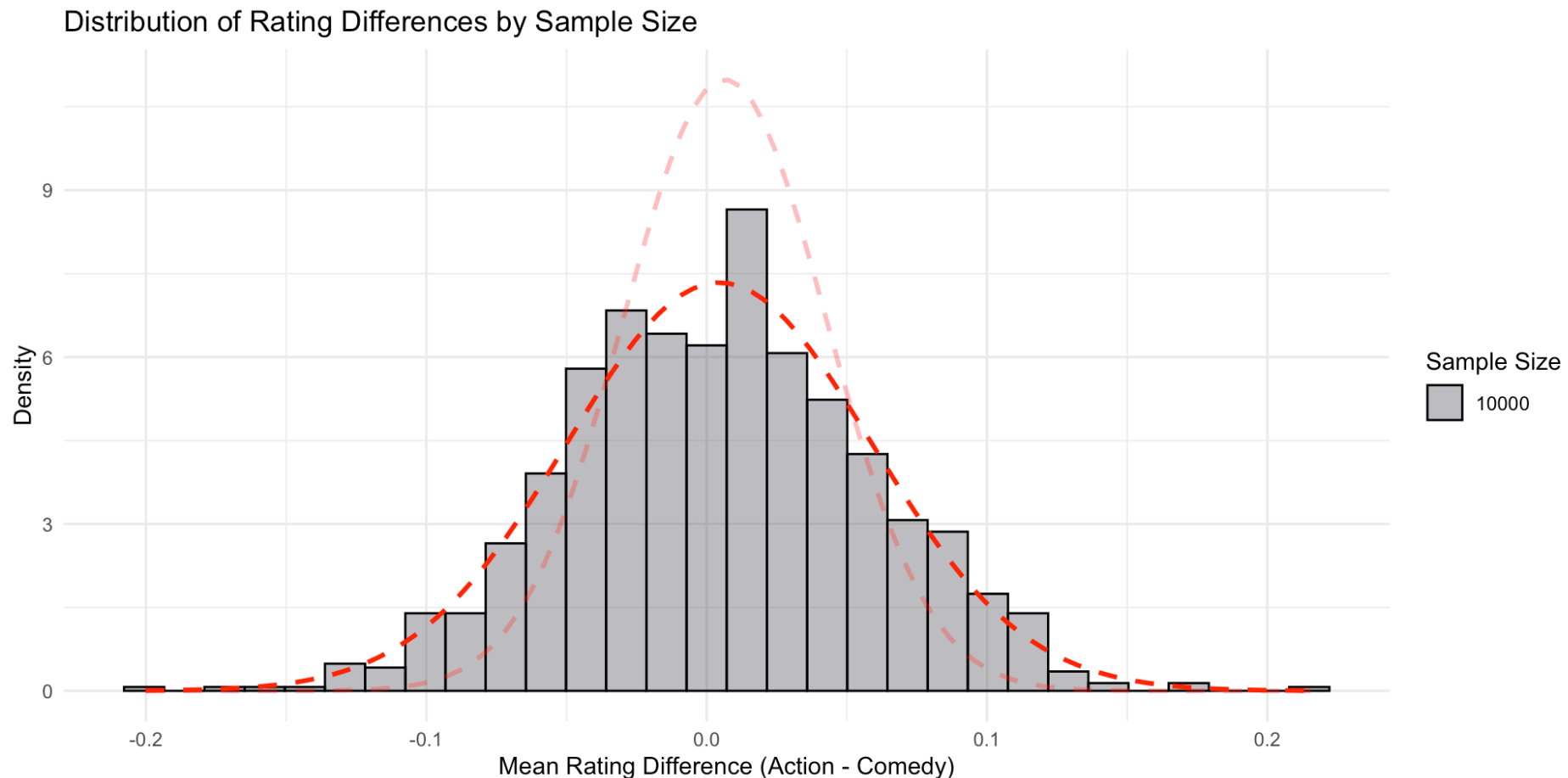
The Central Limit Theorem (*part I*)

The sampling distribution approximates the shape of a normal distribution



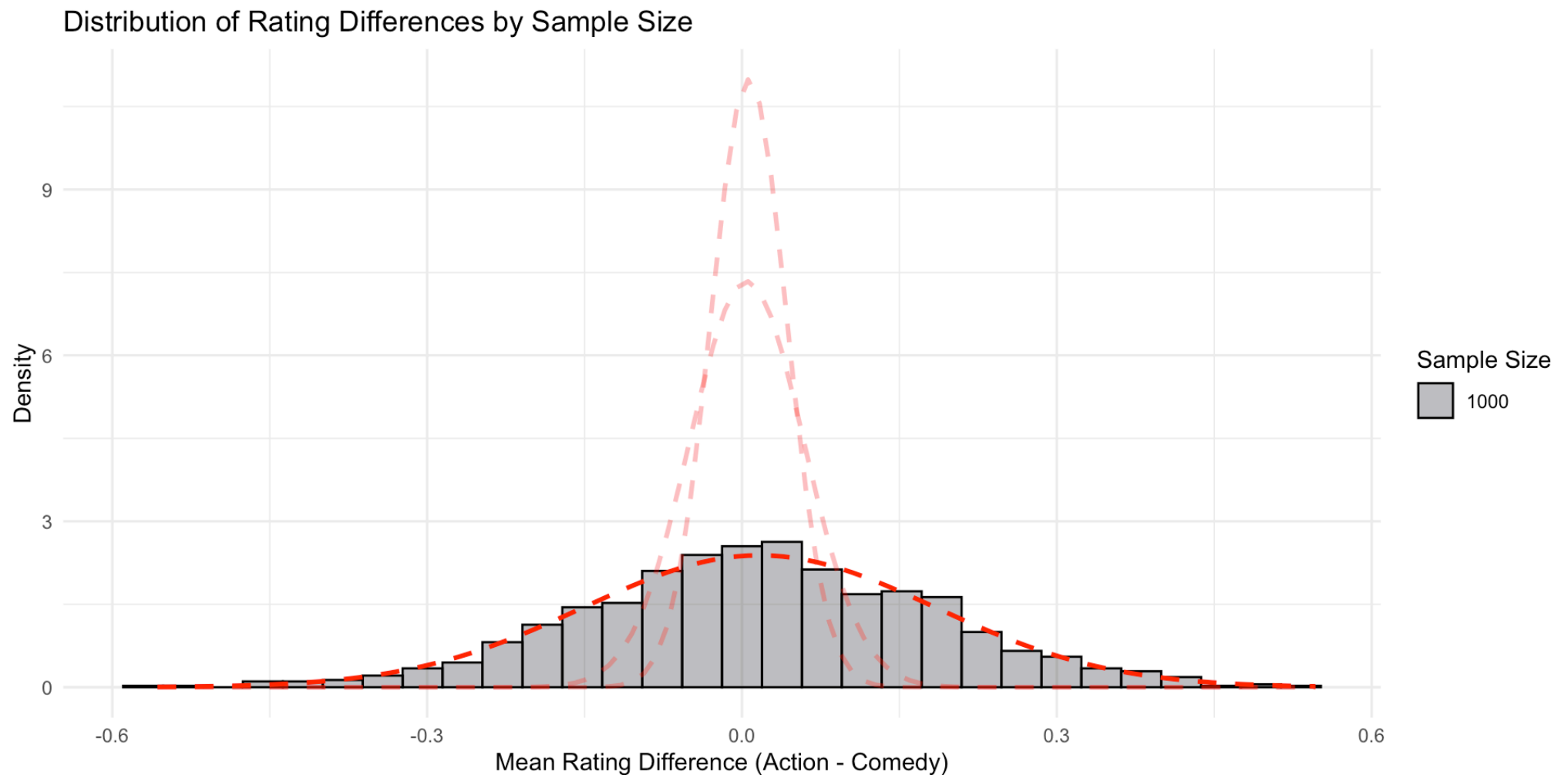
The Central Limit Theorem *(part II)*

This is what the sampling distribution looks like with samples of size 10,000



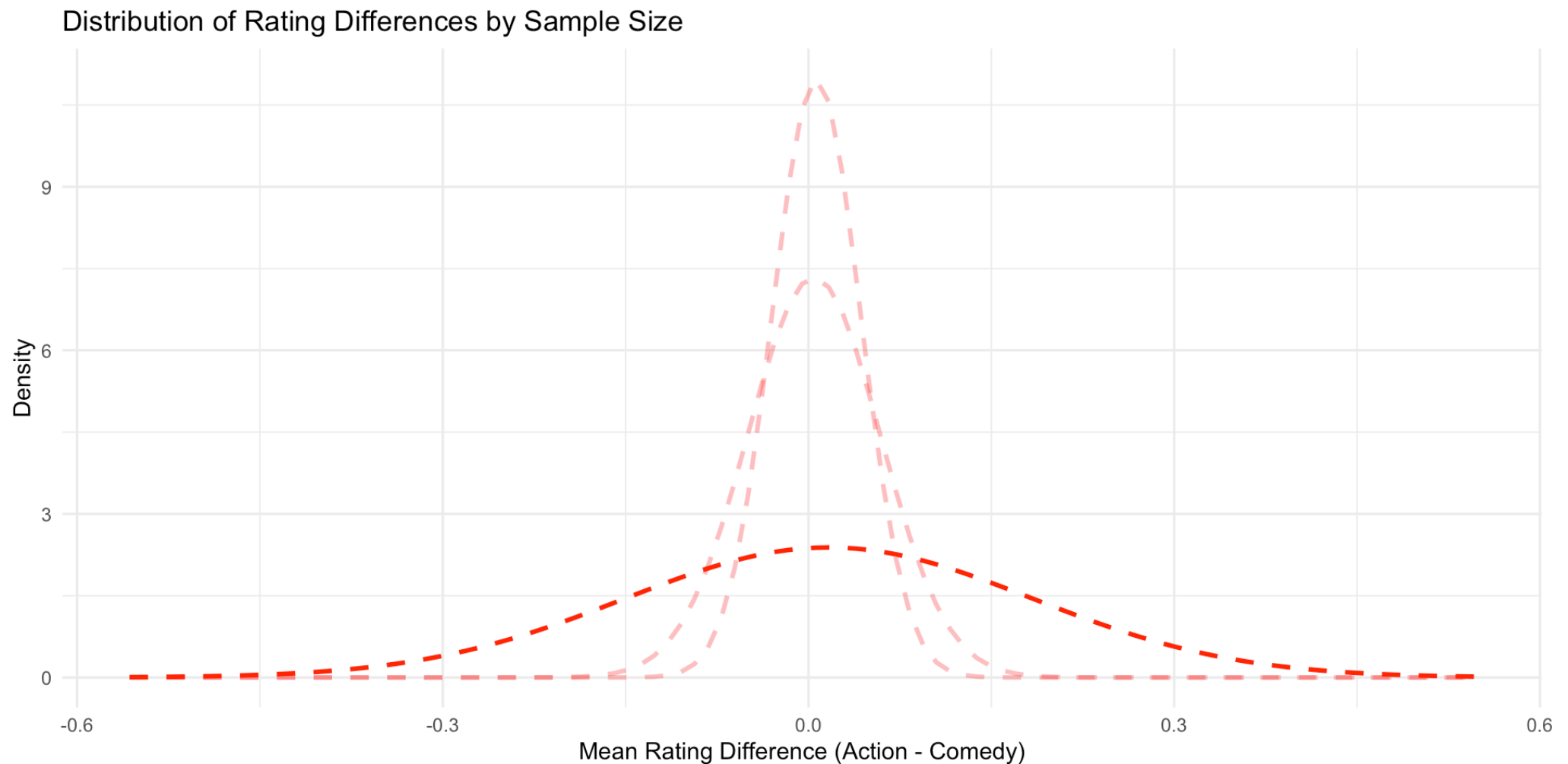
The Central Limit Theorem (*part II*)

And with samples of size 1,000



The Central Limit Theorem (*part II*)

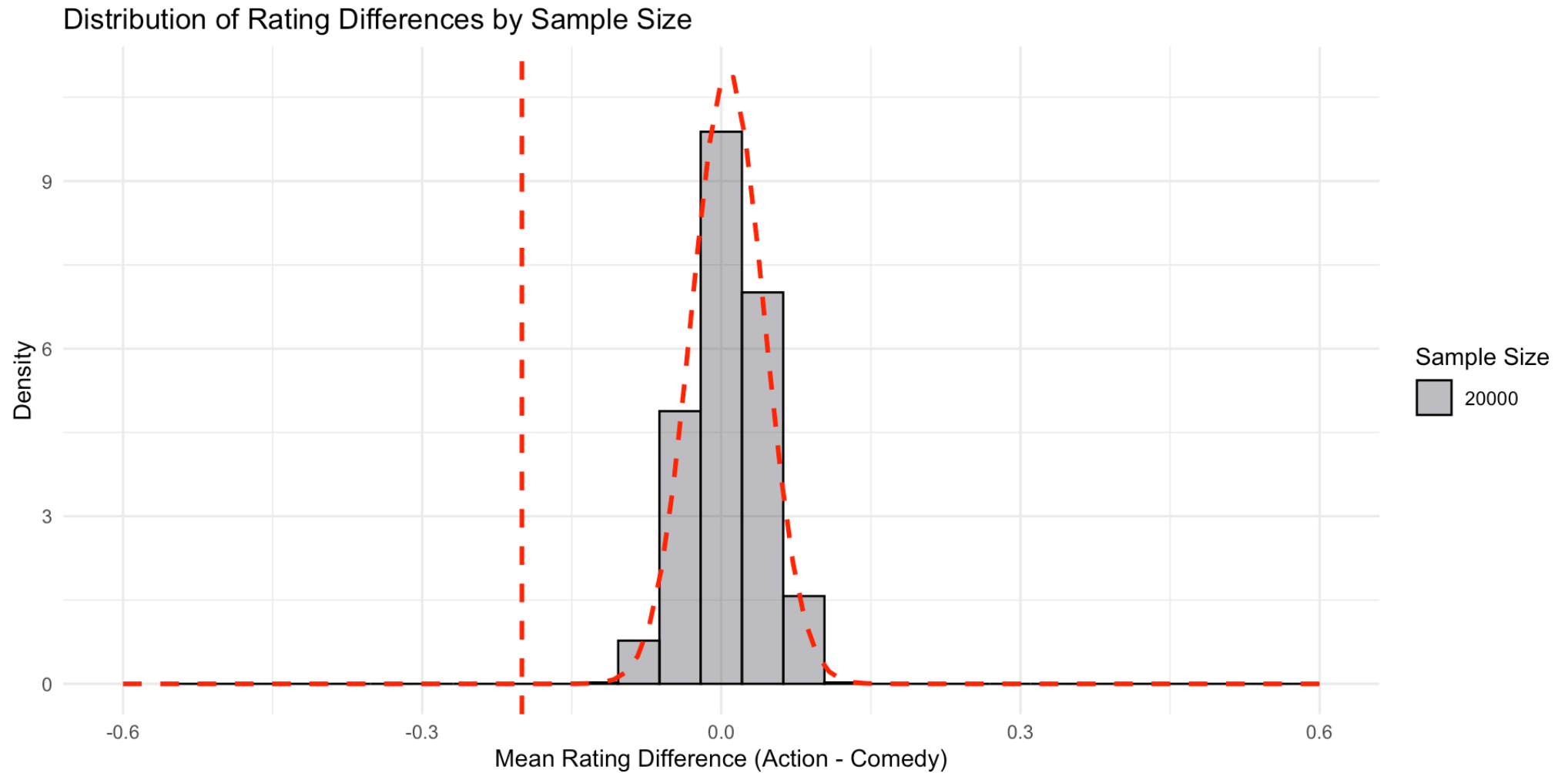
The smaller the sample, the larger the standard deviation of the sampling distribution



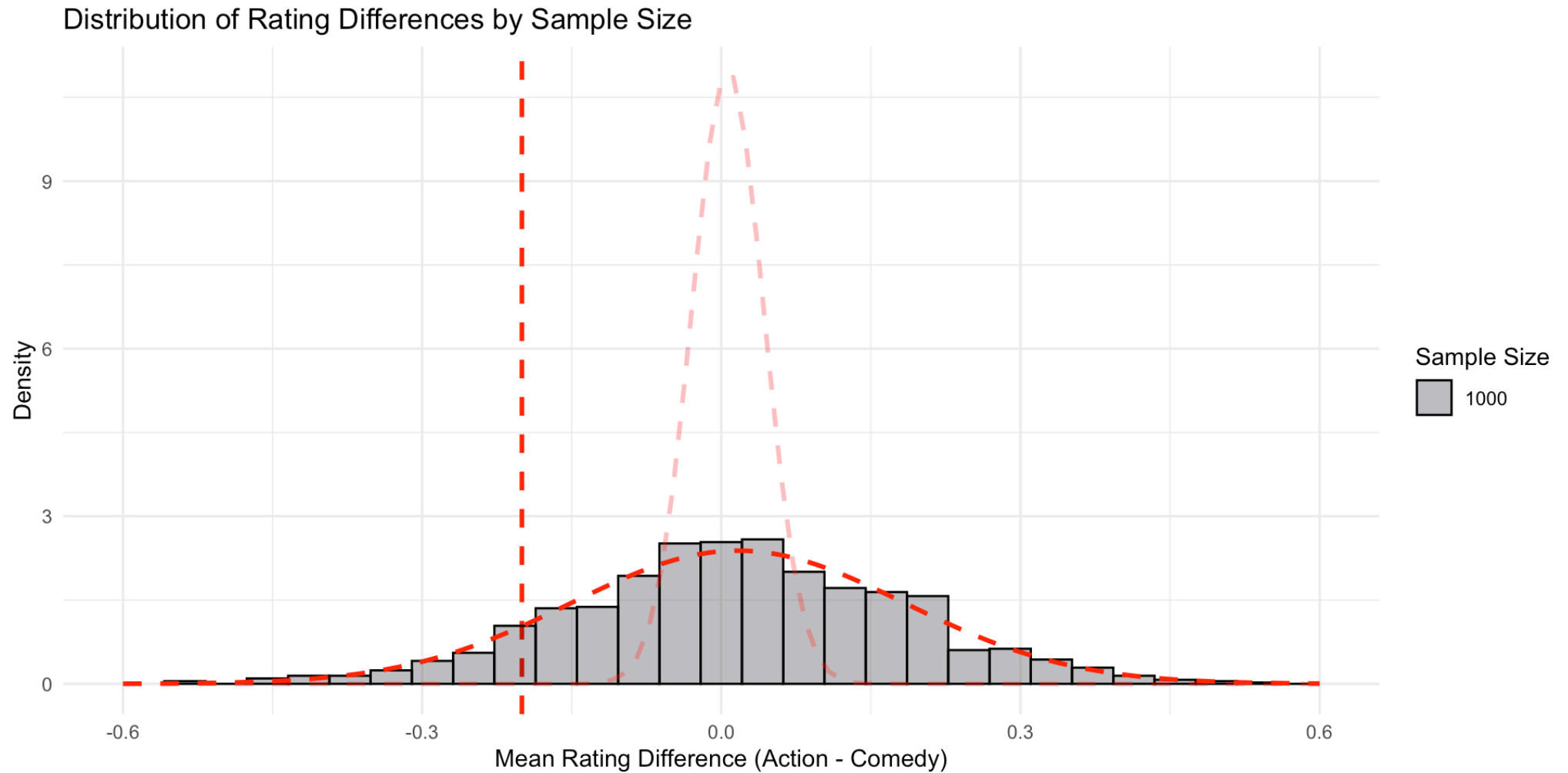
How is this relevant for hypothesis testing?

Imagine we find an effect of -0.2 in our sample

In a sample based on 20'000 movies, that seems reasonably unlikely in a Null world



But the same effect in a sample of 1000 seems not so unlikely...



The Central Limit Theorem

(part I)

- The sampling distribution approximates the shape of a normal distribution

(part II)

- The smaller the sample, the larger the standard deviation of the sampling distribution

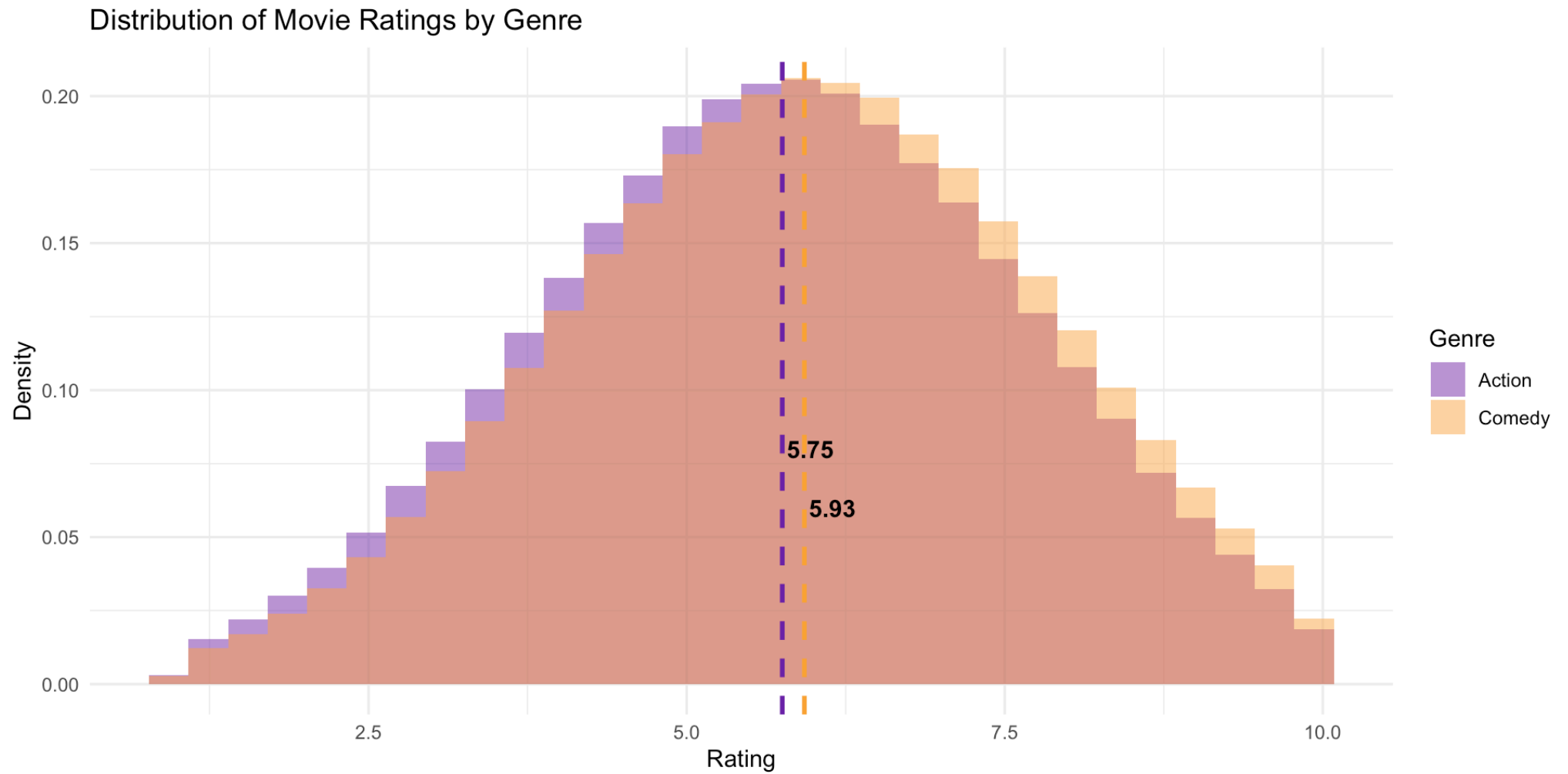
From Null worlds to True effect worlds

For example, let's imagine a world where the true difference between action and comedy movies is -0.2

```
1 # Generate Comedy and Action movie ratings using truncated normal distributions
2 imaginary_movies_true <- tibble(
3   movie_id = 1:1000000,
4   genre = sample(c("Comedy", "Action"), size = 1000000, replace = TRUE),
5   rating = ifelse(
6     genre == "Comedy",
7     rtruncnorm(1000000, a = 1, b = 10, mean = 6.0, sd = 2),
8     rtruncnorm(1000000, a = 1, b = 10, mean = 5.8, sd = 2)
9   )
10 )
```

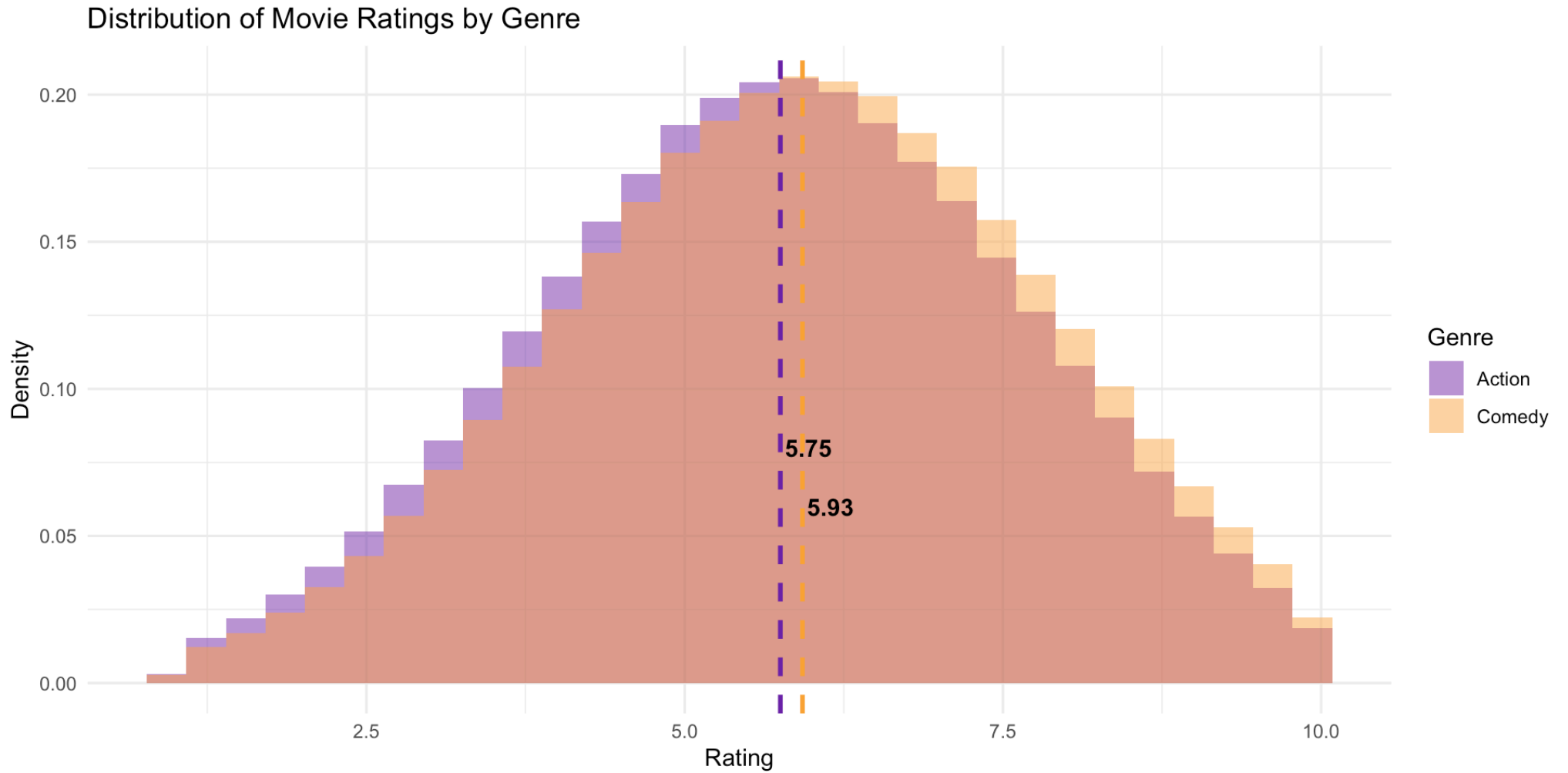
Let's plot the population data

► Code



🎉 It worked, we see our imagined effect of ~ -0.2

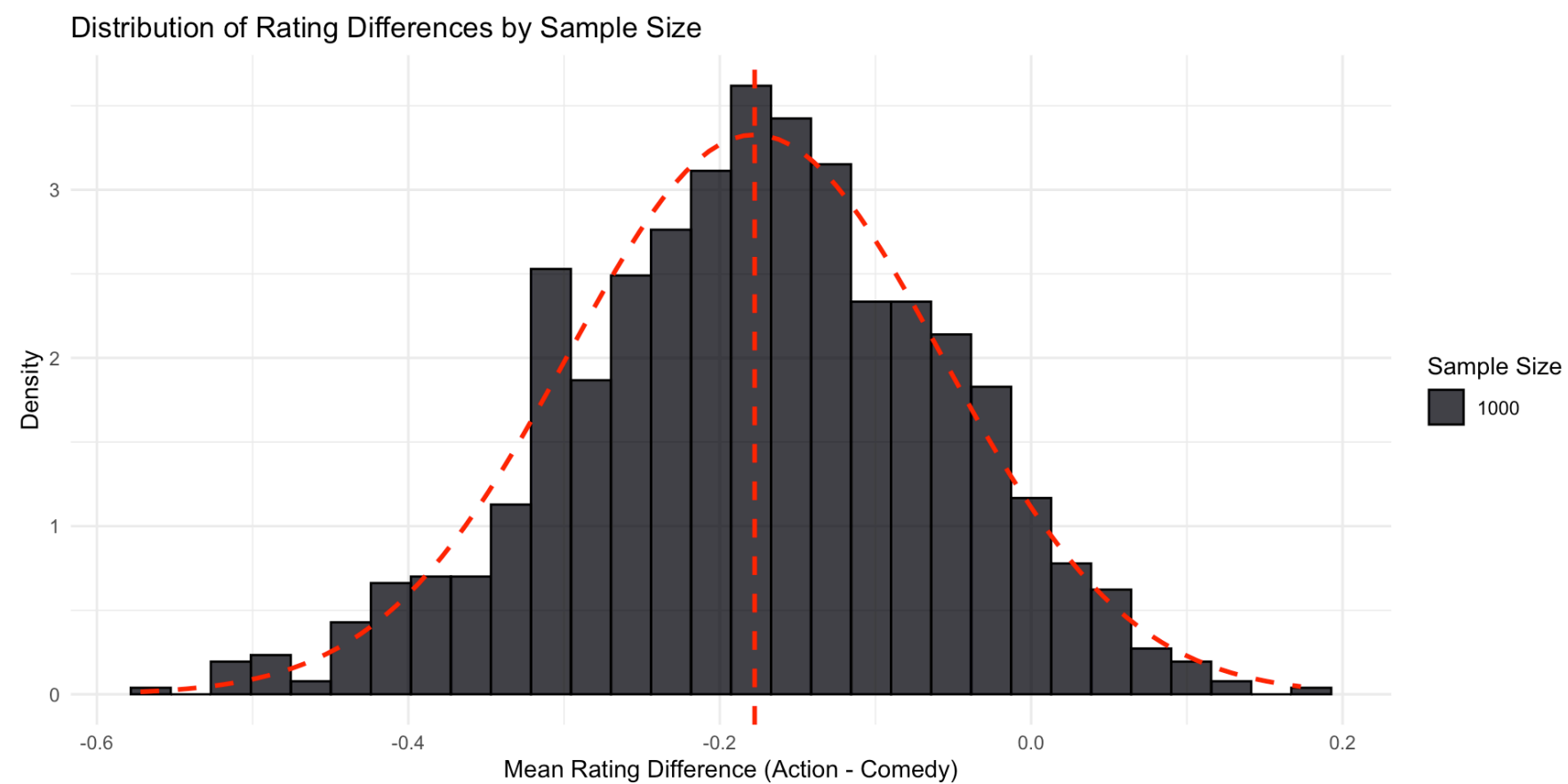
► Code



Imagine we look at a sample of 1000 movie ratings.

```
1 n_simulations <- 1000
2 differences <- c() # make an empty vector
3 sample_size <- 1000
4
5 for (i in 1:n_simulations) {
6   # draw a sample of 20'000 films
7   imaginary_sample <- imaginary_movies |>
8     sample_n(sample_size)
9   # compute rating difference in the sample
10  estimate <- imaginary_sample |>
11    group_by(genre) |>
12    summarize(avg_rating = mean(rating)) |>
13    summarise(diff = avg_rating[genre == "Action"] - avg_rating[genre == "Comedy"]) %>%
14    pull(diff)
15
16  differences[i] <- estimate
17 }
```

This is what our sampling distribution would look like



Now imagine we do a hypothesis test.

We can simulate both a null world and a true effect world (so far, nothing new)

Null world

```
1 # Null world population
2 imaginary_movies_null <- tibble(
3   movie_id = 1:1000000,
4   rating = sample(seq(1, 10, by = 0.1), size = 100
5   genre = sample(c("Comedy", "Action"), size = 100
6 )
```

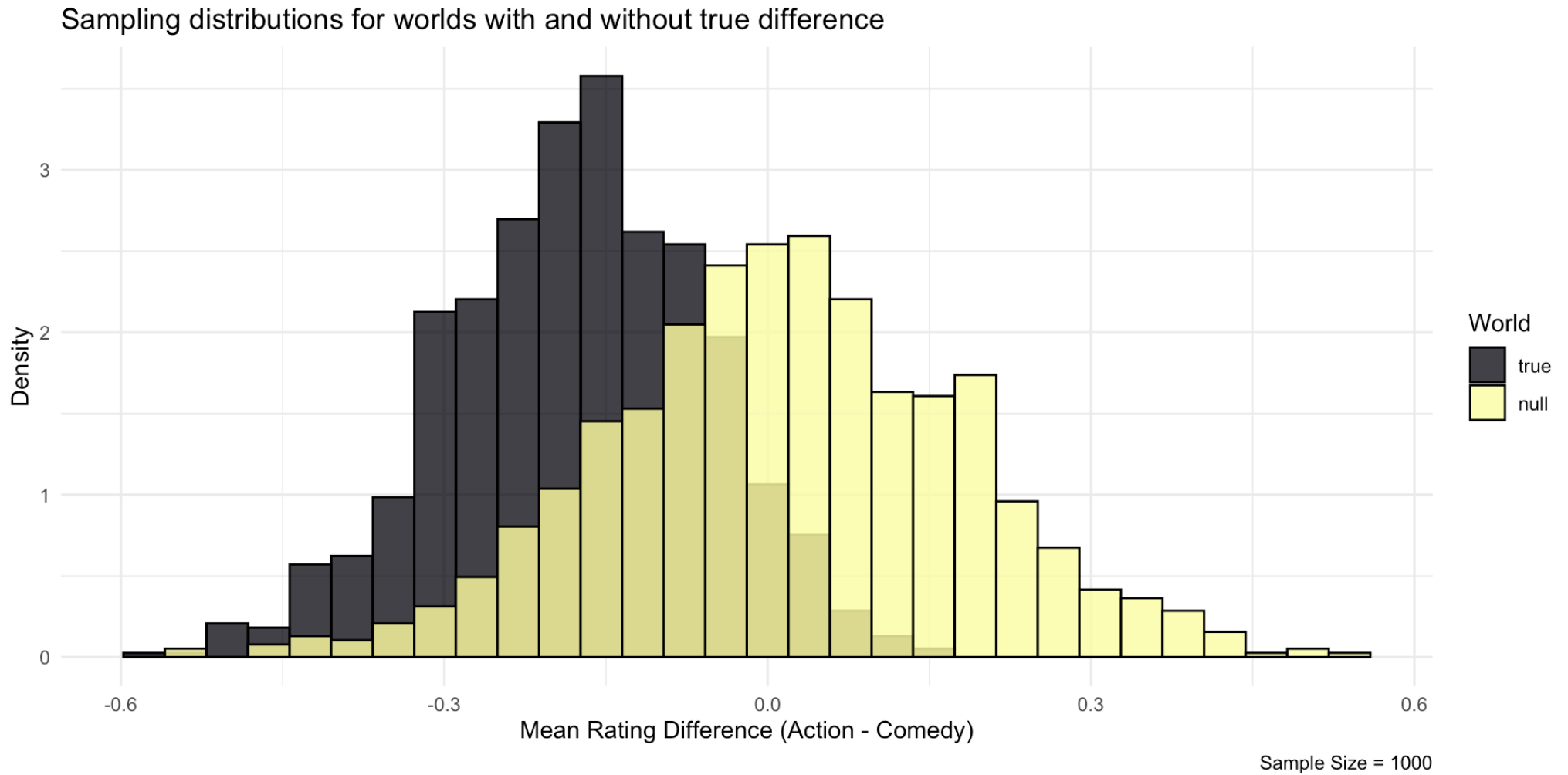
True effect world

```
1 # True effect world population
2 imaginary_movies_true <- tibble(
3   movie_id = 1:1000000,
4   genre = sample(c("Comedy", "Action"), size = 1000000, repla
5   rating = ifelse(
6     genre == "Comedy",
7     rtruncnorm(1000000, a = 1, b = 10, mean = 6.0, sd = 2),
8     rtruncnorm(1000000, a = 1, b = 10, mean = 5.8, sd = 2)
9   )
10 )
```

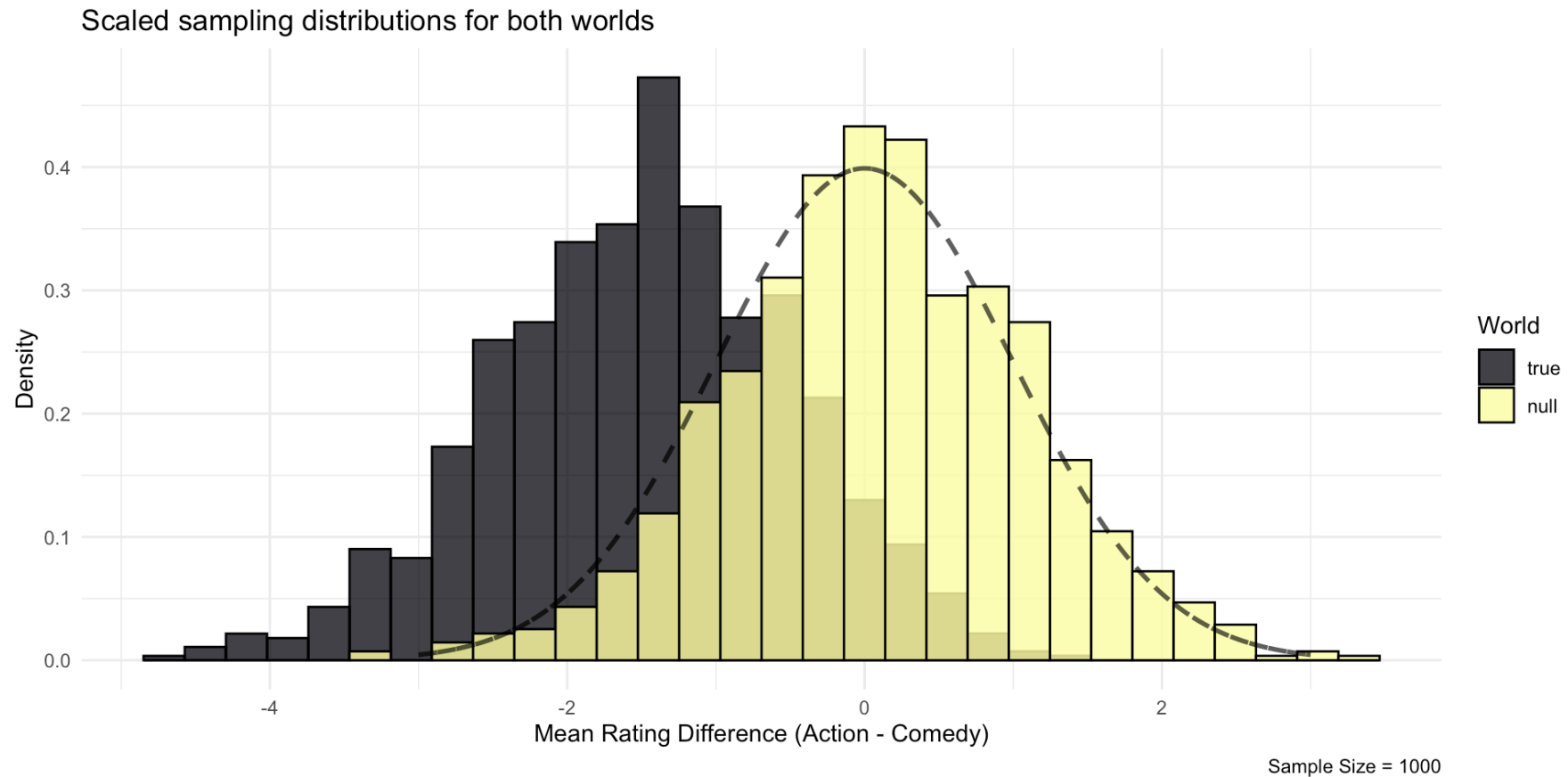
And make a sampling distribution with sample size 1000 for both worlds (also nothing new)

```
1 n_simulations <- 1000
2 differences <- c() # make an empty vector
3 sample_size <- 1000
4
5 data <- imaginary_movies_true # replace with imaginary_movies_null
6
7 for (i in 1:n_simulations) {
8   # draw a sample of films
9   imaginary_sample <- data |>
10     sample_n(sample_size)
11   # compute rating difference in the sample
12   estimate <- imaginary_sample |>
13     group_by(genre) |>
14     summarize(avg_rating = mean(rating)) |>
15     summarise(diff = avg_rating[genre == "Action"] - avg_rating[genre == "Comedy"]) %>%
16     pull(diff)
17
18   differences[i] <- estimate
19 }
```

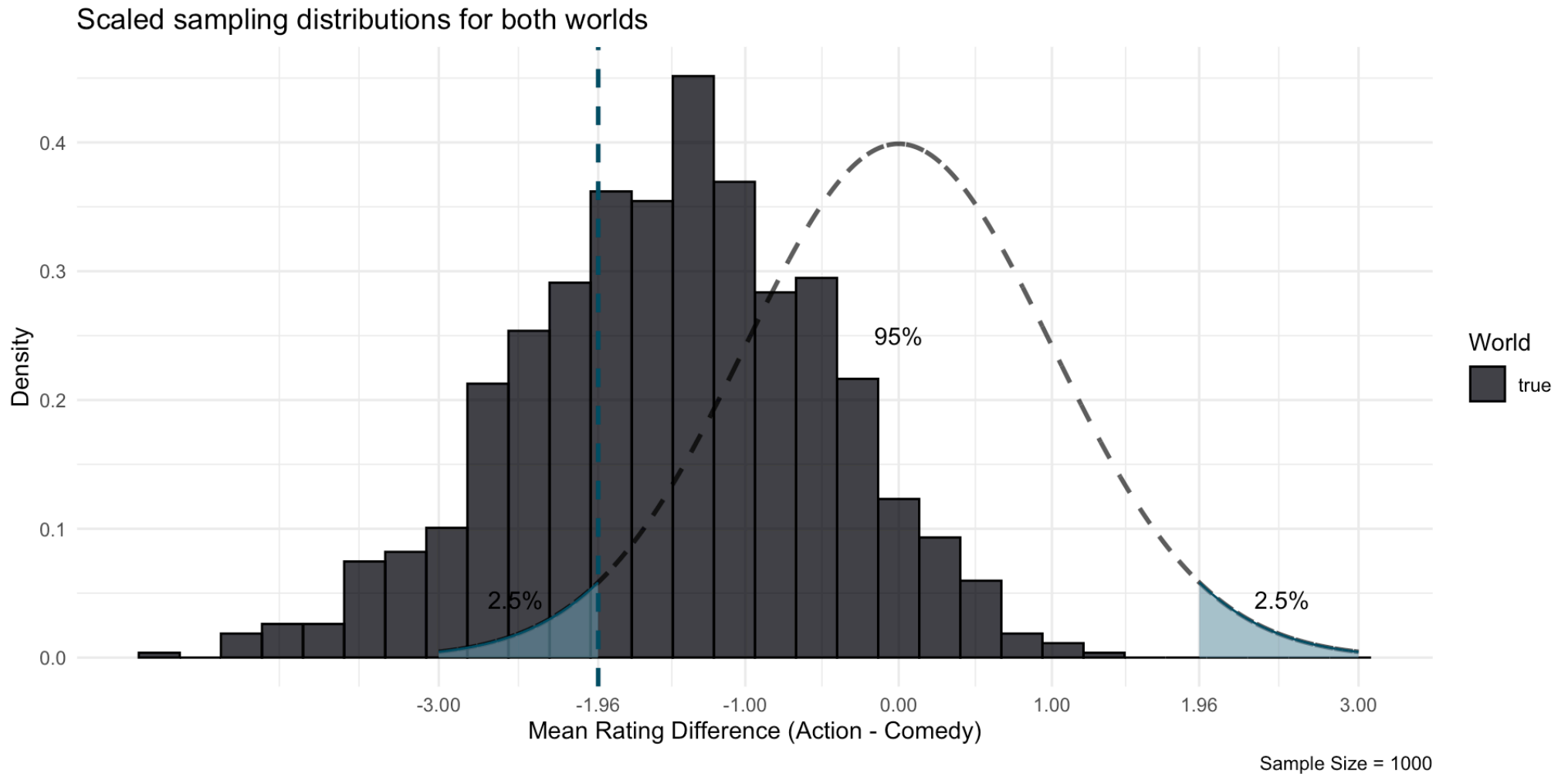
We can plot both simulated worlds together.



Let's bring both worlds on the scale of a standard normal distribution, dividing their respective standard deviation

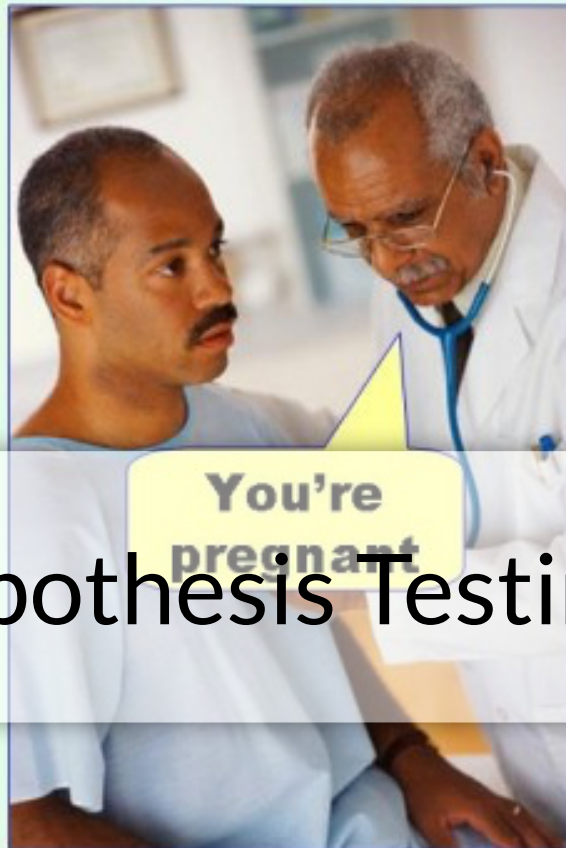


With a sample size of 1000, in many cases, we would say: “This could have occurred in a Null World”



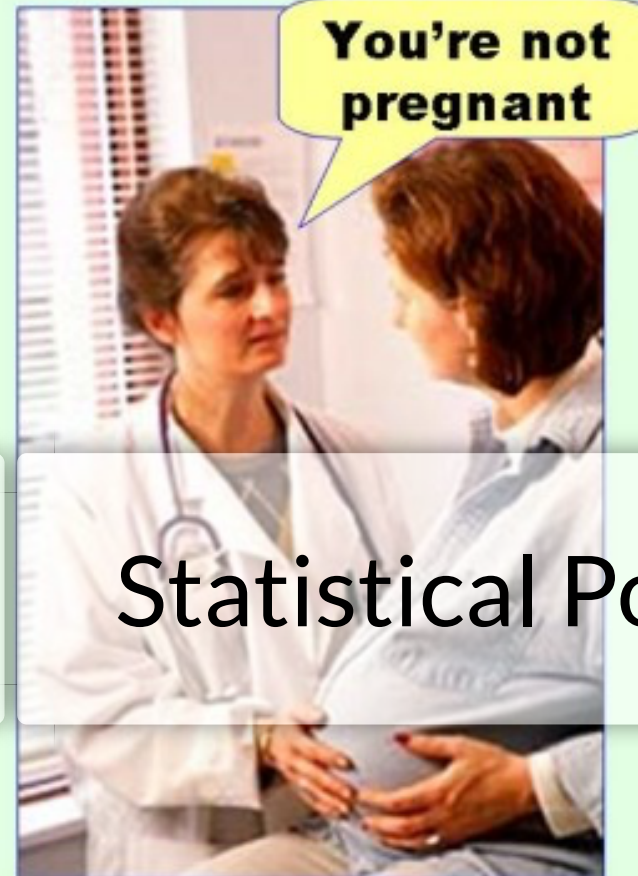
Two errors

Type I error
(false positive)



Hypothesis Testing

Type II error
(false negative)



Statistical Power

Two errors

	Hypothesis Testing	Power Analysis
Aim	Rule out that we observe something just by chance.	Ensure that we would find an effect.
Error question	“What are the chances that we find an effect at least this large in our sample, given that there is no effect in the population?”	“What are the chances that we do not find a statistically significant effect in our sample, although there is a certain effect in the population?”
Typical threshold for acceptable error	$\alpha = 5 \%$	$\beta = 20 \%$

Statistical Power

Statistical Power

Statistical power is the probability of detecting an effect with a hypothesis test, *given a certain effect size*

(Or $1 - \beta$)

Your turn #1: Calculating power

1. Create your true effect world: Simulate a population with a true difference of -0.5 between action and comedy movies.
2. Get your sampling distribution: Simulate 1000 random samples of size $n = 1000$ and store the results.
3. Plot your Sampling distribution (use `data.frame()` to turn your vector into a data frame that can be read by `ggplot`)
4. Prepare for hypothesis testing: bring the results on scale of the standard normal distribution (divide by the standard deviation of your distribution)
5. Check the transformation: plot the new, standardized values
6. Calculate your power: Check how many (standardized) differences are below the 5% threshold, i.e. smaller than or equal to -1.96 (hint: use `mutate()` in combination with `ifelse` to create a new variable `significant` that takes the values of `TRUE` or `FALSE`. Then use `summarize()` and `sum()` to calculate the share). Are you above the 80% power threshold?

1. Create your true effect world: Simulate a population with a true difference of -0.5 between action and comedy movies.

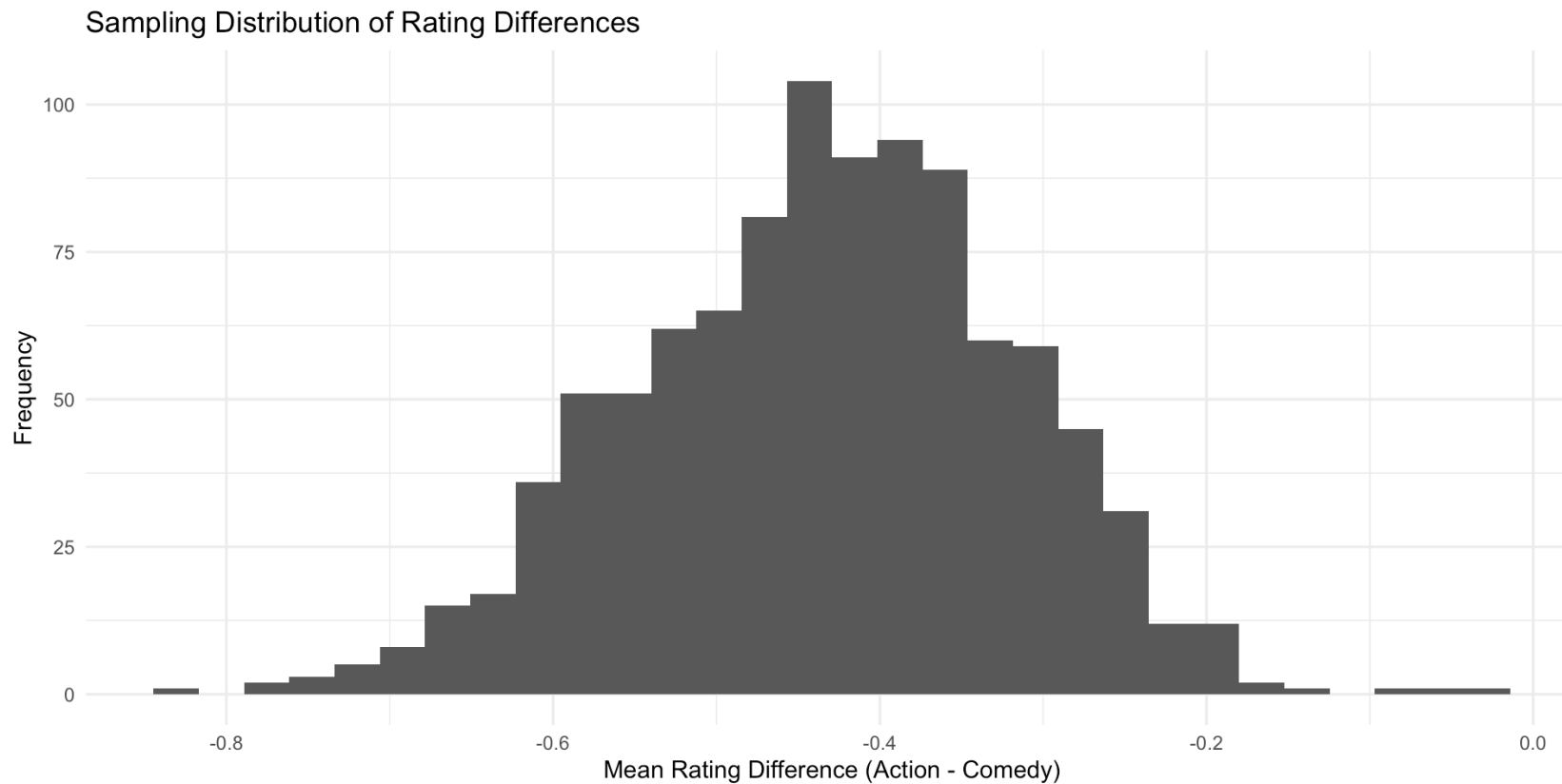
```
1 # True effect world population
2 imaginary_movies_true <- tibble(
3   movie_id = 1:1000000,
4   genre = sample(c("Comedy", "Action"), size = 1000000, replace = TRUE),
5   rating = ifelse(
6     genre == "Comedy",
7     rtruncnorm(1000000, a = 1, b = 10, mean = 6.0, sd = 2),
8     rtruncnorm(1000000, a = 1, b = 10, mean = 5.5, sd = 2)
9   )
10 )
```

2. Get your sampling distribution: Simulate 1000 random samples of size $n = 1000$ and store the results.

```
1 n_simulations <- 1000
2 differences <- c() # make an empty vector
3 sample_size <- 1000
4
5 data <- imaginary_movies_true # replace with imaginary_movies_null
6
7 for (i in 1:n_simulations) {
8   # draw a sample of films
9   imaginary_sample <- data |>
10     sample_n(sample_size)
11   # compute rating difference in the sample
12   estimate <- imaginary_sample |>
13     group_by(genre) |>
14     summarize(avg_rating = mean(rating)) |>
15     summarise(diff = avg_rating[genre == "Action"] - avg_rating[genre == "Comedy"]) %>%
16     pull(diff)
17
18   differences[i] <- estimate
19 }
```

3. Plot your Sampling distribution (use `data.frame()` to turn your vector into a data frame that can be read by `ggplot`)

```
1 ggplot(data.frame(differences), aes(x = differences)) +  
2   geom_histogram() +  
3   labs(title = "Sampling Distribution of Rating Differences",  
4         x = "Mean Rating Difference (Action - Comedy)",  
5         y = "Frequency") +  
6   theme_minimal()
```

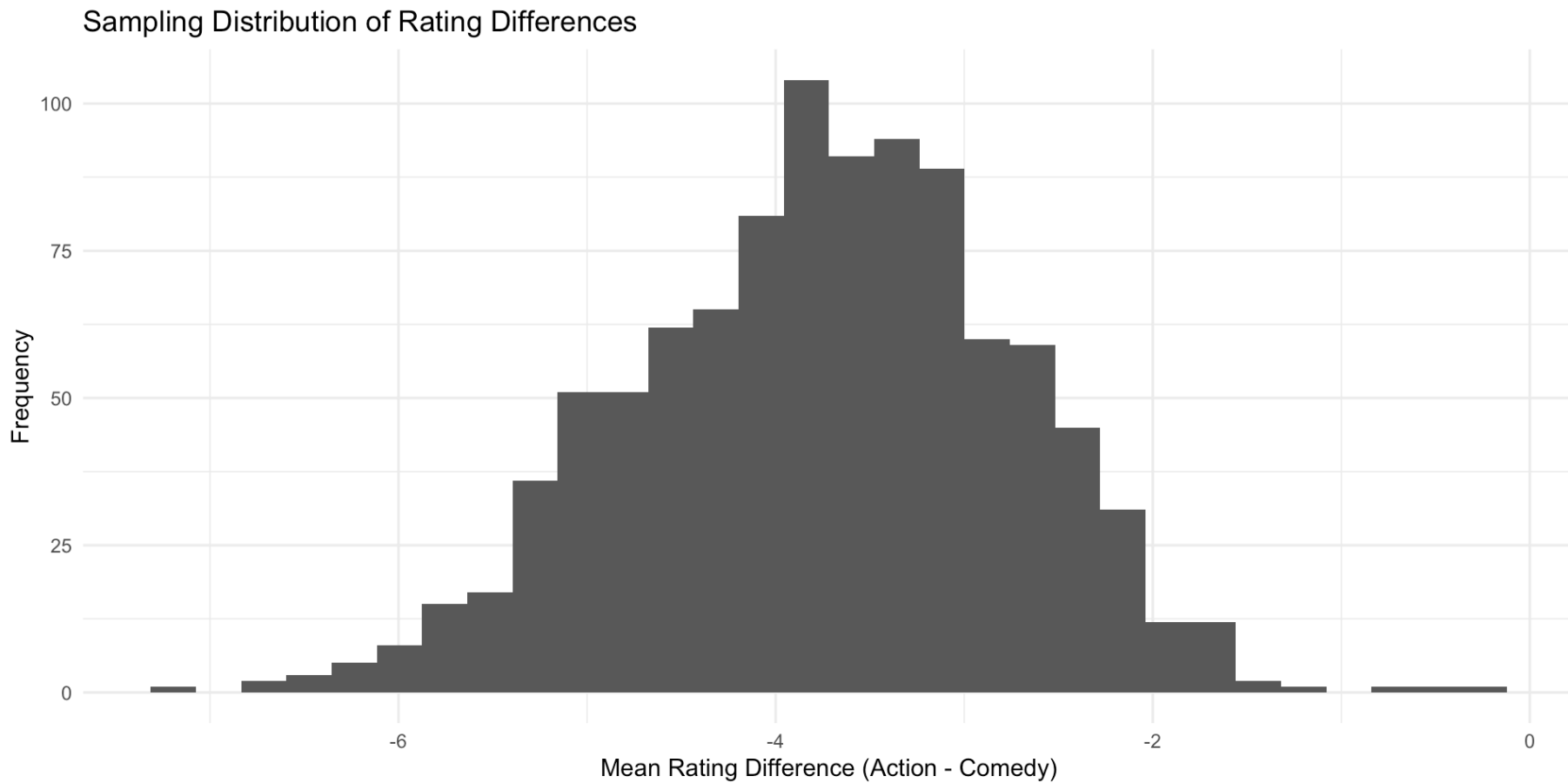


4. Prepare for hypothesis testing: bring the results on scale of the standard normal distribution (divide by the standard deviation of your distribution)

```
1 differences_standardized <- differences/sd(differences)
```

5. Check the transformation: plot the new, standardized values

```
1 ggplot(data.frame(differences_standardized), aes(x = differences_standardized)) +  
2   geom_histogram() +  
3   labs(title = "Sampling Distribution of Rating Differences",  
4         x = "Mean Rating Difference (Action - Comedy)",  
5         y = "Frequency") +  
6   theme_minimal()
```



6. Calculate your power: Check how many (standardized) differences are below the 5% threshold, i.e. smaller than or equal to -1.96 (hint: use `mutate()` in combination with `ifelse` to create a new variable `significant` that takes the values of `TRUE` or `FALSE`. Then use `summarize()` and `sum()` to calculate the share). Are you above the 80% power threshold?

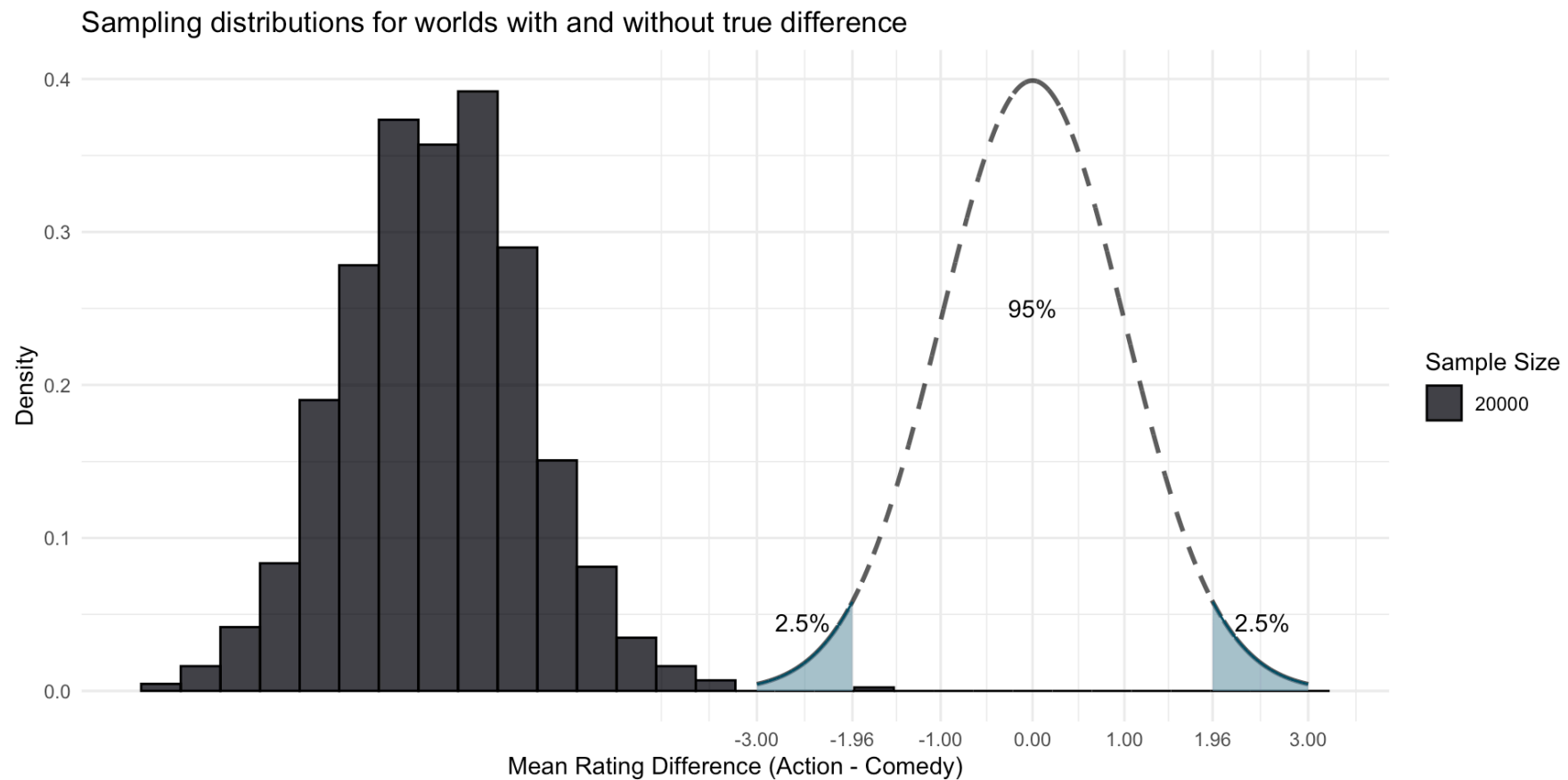
```
1 data.frame(differences_standardized) |>
2   mutate(significant = ifelse(differences_standardized <= -1.96, TRUE, FALSE)) |>
3   summarize(sum_significant = sum(significant),
4             # you can also calculate the share directly
5             share_significant = sum(significant) / n()
6   )
```

```
sum_significant share_significant
1              975              0.975
```

A power simulation

The central limit theorem (again)

The larger the sample size, the more statistical power



An example of a power analysis

For an example, head over to the [guide on power analysis](#) on the course website.

That's it for today :)