

Learn How to Perform a Chi-Square Goodness of Fit Test in SPSS

Authored by
Mohammed looti

November 8, 2025

RECOMMENDED CITATION

Mohammed looti (2025). *Learn How to Perform a Chi-Square Goodness of Fit Test in SPSS*. PSYCHOLOGICAL STATISTICS. Retrieved from <https://statistics.arabpsychology.com/?p=12884>

The [Chi-Square Goodness of Fit Test](#) is a fundamental statistical tool utilized to ascertain whether the observed frequency distribution of a single [categorical variable](#) significantly deviates from a hypothesized or expected distribution. In essence, this test determines if a sample taken from a population accurately reflects a theoretical probability distribution.

This comprehensive tutorial provides step-by-step instructions on how to effectively conduct the [Chi-Square Goodness of Fit Test](#) using the industry-standard statistical software, [SPSS](#) (Statistical Package for the Social Sciences). We will walk through the process from initial data input and preparation to the final interpretation of the results, ensuring clarity for researchers and analysts alike.

Understanding the Chi-Square Goodness of Fit Test

The primary goal of the Goodness of Fit Test is to compare the observed counts within various categories to the counts that would be expected if the data perfectly matched a specified distribution. This hypothesized distribution often assumes uniformity (meaning all categories are equally likely), but it can also be based on historical data or theoretical predictions. When applying this test, we are essentially testing two competing hypotheses.

The formal structure of the test relies on the establishment of the [null hypothesis](#) (H_0) and the alternative hypothesis (H_A). The [null hypothesis](#) posits that there is no significant difference between the observed distribution and the expected distribution; that is, the categorical variable follows the hypothesized distribution. Conversely, the alternative hypothesis states that the observed data does not fit the hypothesized distribution, suggesting a statistically significant difference between the observed and expected frequencies.

The calculation yields a Chi-Square test statistic, which measures the aggregated discrepancy between the observed and expected values. A large Chi-Square value indicates a significant deviation from the expected distribution, leading to the rejection of the [null hypothesis](#). To accurately gauge this deviation, the test requires the calculation of [degrees of freedom](#), which is determined by the number of categories minus one. This parameter is crucial for referencing the appropriate Chi-Square distribution table to find the corresponding p-value.

Case Study: Testing Customer Flow Uniformity

To illustrate the application of this test in [SPSS](#), consider a practical scenario. A small business owner operates a retail shop and hypothesizes that customer arrivals are uniformly distributed across the five standard weekdays. Specifically, the owner claims that an equal number of customers visit the shop on Monday, Tuesday, Wednesday, Thursday, and Friday. We want to test if this claim of uniform distribution holds true based on collected empirical data.

A researcher meticulously records the total number of customers who entered the shop during a specific week, generating the following observed frequency data for the [categorical variable](#) "Day":

Monday: 50 customers

Tuesday: 60 customers

Wednesday: 40 customers

Thursday: 47 customers

Friday: 53 customers

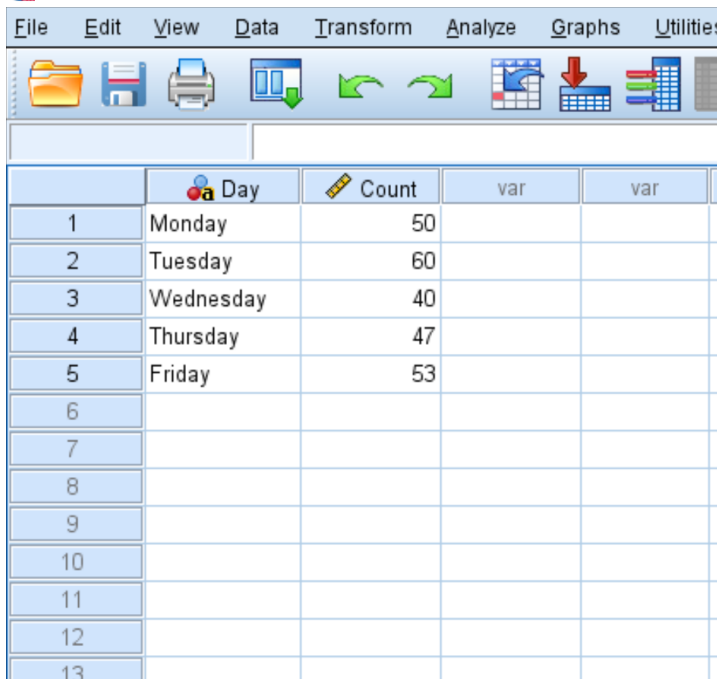
The total number of customers observed over the five days is $50 + 60 + 40 + 47 + 53 = 250$. If the owner's claim of equal distribution were perfectly true, we would expect $250 / 5 = 50$ customers on each of the five weekdays. The goal of the upcoming steps is to use the [Chi-Square Goodness of Fit Test](#) in [SPSS](#) to statistically determine if the observed differences (e.g., 60 on Tuesday vs. 50 expected) are simply due to random chance or if they represent a genuine non-uniform distribution.

Step 1 & 2: Data Entry and Weighting Cases in SPSS

The first critical stage involves setting up the data correctly within the [SPSS](#) environment. Unlike raw data where each row represents a single observation, frequency data requires a specific structure. We must create two variables: one representing the categories (e.g., "Day") and one representing the observed frequencies (e.g., "Count").

Step 1: Input the Data.

Enter the categorical variable (Day) and its corresponding observed frequency (Count) into the Data View tab of [SPSS](#). It is important that the "Day" variable is defined as a nominal or ordinal categorical type, and the "Count" variable is defined as numeric. The data should appear in the following required format:

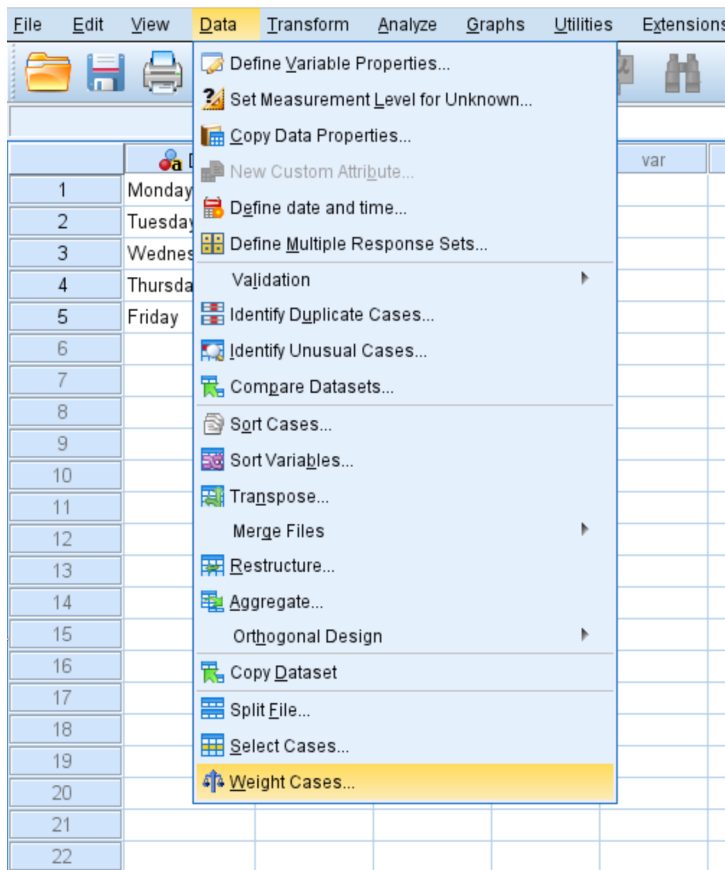


	Day	Count	var	var
1	Monday	50		
2	Tuesday	60		
3	Wednesday	40		
4	Thursday	47		
5	Friday	53		
6				
7				
8				
9				
10				
11				
12				
13				

Step 2: Use Weighted Cases.

Since we are inputting summarized frequency data rather than 250 individual customer entries, we must instruct [SPSS](#) to treat the values in the "Count" column as weights for the corresponding categories in the "Day" column. This step is crucial for the statistical calculation to accurately reflect the total number of observations.

To execute this weighting, navigate to the main menu and click the **Data** tab, followed by **Weight Cases**. This action opens a dialogue box necessary for assigning the weight variable. As shown in the navigation image below, this ensures the program correctly interprets the input as aggregated data.

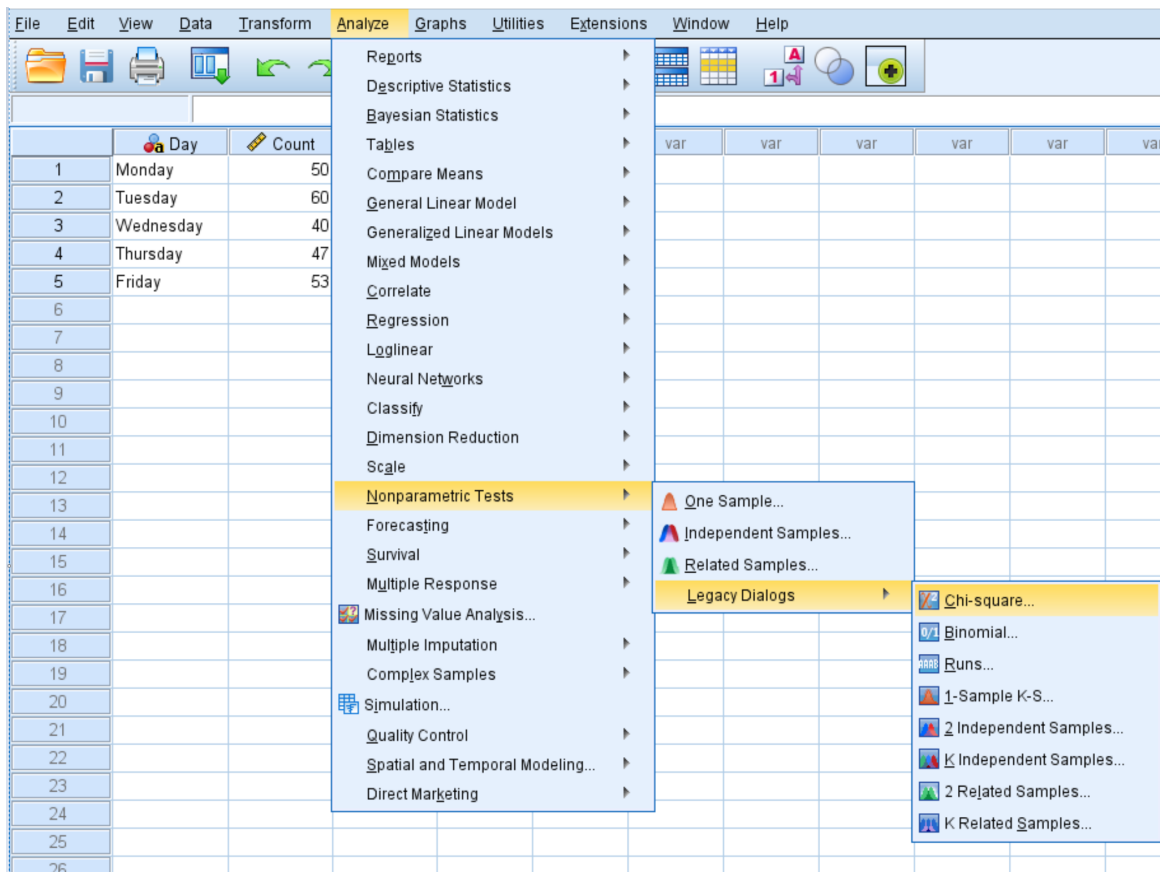


In the resulting Weight Cases dialogue box, select the option **Weight cases by**. Then, drag the variable **Count** into the box labeled Frequency Variable. Once the variable is placed, click **OK**. [SPSS](#) is now prepared to treat the "Day" variable based on the frequencies listed in the "Count" column for the subsequent Chi-Square test.

Step 3: Executing the Chi-Square Analysis

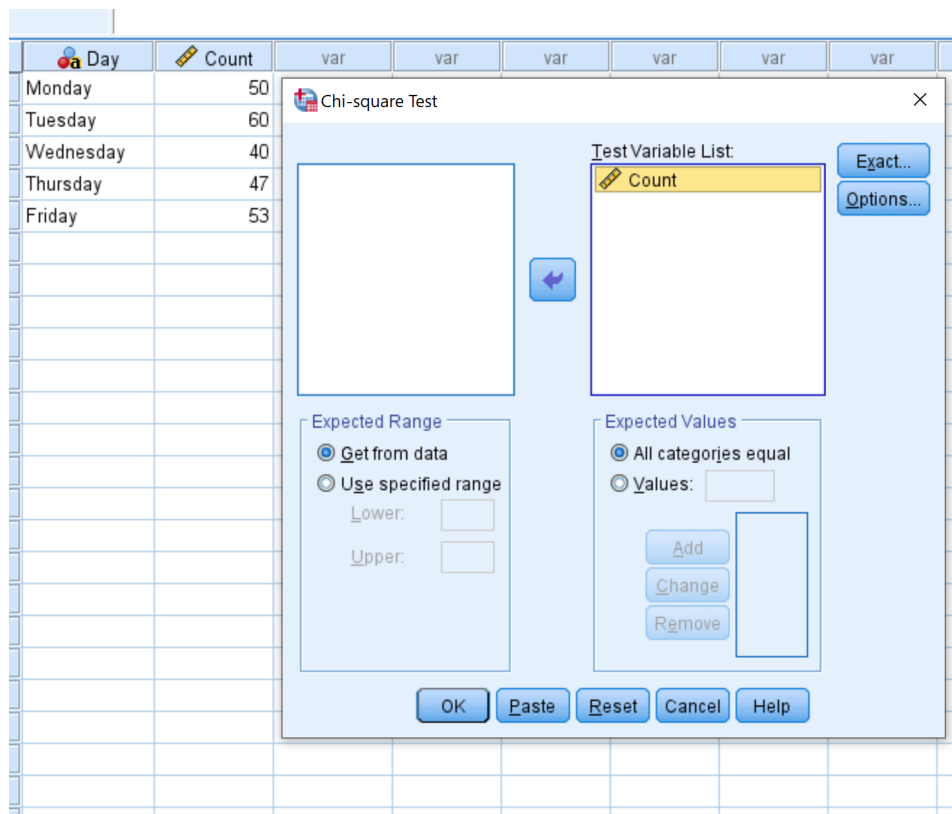
With the data properly weighted, we can now proceed to perform the [Chi-Square Goodness of Fit Test](#) itself. This test is located within the Nonparametric Tests suite in [SPSS](#), as it does not rely on assumptions about the population distribution shape (like normality), which are required for parametric tests.

Click on the **Analyze** tab in the main menu. Hover over **Nonparametric Tests**, then select **Legacy Dialogs**, and finally click **Chi-Square**. This sequence of commands initiates the setup for the Goodness of Fit Test. The navigation path is detailed in the image below, illustrating the precise menu selections needed.



In the Chi-Square Test dialogue box that appears, you must specify the variable being tested against the hypothesized distribution. Drag the variable **Day** (the [categorical variable](#)) into the Test Variable List box. Note that although we used "Count" to weight the cases, "Day" is the variable whose distribution we are assessing.

Under the Expected Values section, ensure that the label checked next to **All categories equal** remains selected. We choose this option because the shop owner's initial claim specified a uniform distribution, meaning we expect an equal proportion of customers on Monday, Tuesday, Wednesday, Thursday, and Friday. If the expected values were based on a different theoretical distribution (e.g., 20% Monday, 30% Tuesday), we would instead select the "Values" option and input those specific expected percentages or ratios. Since our assumption is uniform, we proceed with the default setting and click **OK** to generate the output.



Step 4: Interpreting the Statistical Output

Upon clicking **OK**, [SPSS](#) generates the results in the Output Viewer, presenting two main tables that are essential for decision-making regarding the [null hypothesis](#). The first table summarizes the observed versus expected frequencies, while the second table provides the core statistical measures necessary for inference.

The first output table, typically labeled "Day (Observed and Expected Frequencies)," provides a breakdown for each category. It lists the **Observed N** (the counts we collected), the **Expected N** (the count of 50 for each day, based on the uniform distribution claim), and the **Residual** (the difference between Observed N and Expected N). Analyzing the residuals visually gives an initial sense of deviation; for example, Tuesday has a residual of +10 (60 observed - 50 expected), indicating 10 more customers than expected.

➔ NPar Tests

Chi-Square Test

Frequencies

Count			
	Observed N	Expected N	Residual
40	40	50.0	-10.0
47	47	50.0	-3.0
50	50	50.0	.0
53	53	50.0	3.0
60	60	50.0	10.0
Total	250		

Test Statistics

Count	
Chi-Square	4.360 ^a
df	4
Asymp. Sig.	.359

a. 0 cells (0.0%) have expected frequencies less than 5. The minimum expected cell frequency is 50.0.

The second table, labeled "Test Statistics," contains the quantitative results of the [Chi-Square Goodness of Fit Test](#):

Chi-Square: This is the calculated test statistic, found to be 4.36. This value represents the total squared standardized difference between the observed and expected frequencies across all categories.

df: This stands for the [degrees of freedom](#). For the Goodness of Fit Test, it is calculated as the number of categories minus one. Since we have five days (categories), the calculation is $5 - 1 = 4$.

Asymp. Sig.: This is the asymptotic significance, or the p-value. This value corresponds to a Chi-Square statistic of 4.36 with 4 [degrees of freedom](#), and is found to be 0.359. The p-value is the probability of observing data as extreme as (or more extreme than) the sample data, assuming the [null hypothesis](#) is true.

Conclusion and Decision Making

The final step in the hypothesis testing process involves comparing the p-value (Asymp. Sig.) generated by [SPSS](#) to the predetermined significance level (α). Standard practice dictates a significance level of $\alpha = 0.05$. The decision rule is straightforward: if the p-value is less than α , we reject the [null hypothesis](#); otherwise, we fail to reject it.

In this case study, the calculated p-value is 0.359. Since $0.359 > 0.05$, we fail to reject the [null hypothesis](#) (H_0).

This failure to reject the [null hypothesis](#) leads to the conclusion that we do not have sufficient statistical evidence, at the 5% significance level, to claim that the true distribution of customer arrivals is significantly different from the uniform distribution claimed by the shop owner. While there were observed variations in daily customer counts, these deviations are likely due to random sampling variability and are not statistically compelling enough to refute the claim that the customer flow is equally distributed across the weekdays.

Therefore, the data is consistent with the shop owner's assertion of uniform customer traffic. The [Chi-Square Goodness of Fit Test](#) confirms that the observed frequencies do not significantly misfit the expected uniform distribution, providing valuable insight into the underlying dynamics of the shop's customer base.