

Calculate Frequencies in Google Sheets

Authored by
Mohammed looti

November 7, 2025

RECOMMENDED CITATION

Mohammed looti (2025). *Calculate Frequencies in Google Sheets*. PSYCHOLOGICAL STATISTICS. Retrieved from <https://statistics.arabpsychology.com/?p=12086>

Calculating [frequencies](#) is a fundamental task in [data analysis](#), serving as the foundation for understanding how often specific values or categories manifest within a given [dataset](#). Determining these distributions is the first step toward drawing meaningful conclusions from raw data.

For professionals and students utilizing [Google Sheets](#), this statistical operation is streamlined and made highly accessible through the dedicated [FREQUENCY\(\) function](#). This comprehensive tutorial will guide you through the precise steps required to compute both absolute counts and [relative frequencies](#), culminating in effective techniques for data visualization.

Understanding the FREQUENCY() Function Syntax

The [FREQUENCY\(\) function](#) is specifically engineered to determine the distribution of numerical data across specified intervals or 'bins'. A crucial distinction of this function is that it must be executed as an [array formula](#) in [Google Sheets](#), meaning it returns an array of results simultaneously, rather than a single value. Therefore, adequate space must be allocated in the destination column before entry.

The function employs a straightforward structure, requiring two primary arguments to define the calculation scope, ensuring precision when grouping data:

FREQUENCY(data, classes)

A clear understanding of these arguments ensures accurate computation of [frequencies](#):

data: This argument specifies the array or range containing the raw, numerical values you intend to analyze. This range is the source [dataset](#) that the function will process.

classes: This argument is the array or range defining the upper boundaries of the bins (or intervals) used for grouping the data. When calculating the frequency of distinct, individual values, the classes list should simply consist of those unique values themselves, acting as specific counting thresholds.

The following practical examples will systematically illustrate how to apply this versatile function within a typical data analysis workflow, translating theoretical syntax into actionable results in the spreadsheet environment.

Step 1: Preparing and Examining the Sample Dataset

To demonstrate the calculation process effectively, we will utilize a hypothetical [dataset](#), perhaps representing test scores or daily transaction volumes. Our sample consists of 15 numerical entries, organized vertically in a single column, starting at cell A2. This initial preparation step is crucial, as

the quality and structure of the input data directly impact the [FREQUENCY\(\) function](#) results.

For this tutorial, the raw data occupies the range A2:A16 within the [Google Sheets](#) environment, as depicted in the image below. Our primary objective is to calculate the absolute [frequency](#)--the raw count--of every unique score present in this collection of data points, providing a clear numerical summary of the data distribution.

<i>fx</i>				
	A	B	C	D
1	Data values			
2	12			
3	12			
4	13			
5	13			
6	13			
7	14			
8	14			
9	15			
10	15			
11	15			
12	15			
13	16			
14	17			
15	17			
16	17			
17				
18				
19				
20				
21				
22				

Step 2: Defining Data Classes (Unique Values)

Before executing the frequency calculation, we must precisely define the data classes, which are the unique values or intervals we wish to count against. Manually extracting every distinct value from a large dataset is inefficient and highly prone to error. Fortunately, [Google Sheets](#) offers powerful array functions that automate this preparation step, ensuring accuracy and saving significant time.

We can generate the required list of unique, sorted classes efficiently by combining the **UNIQUE()** and **SORT()** functions. In cell B2, enter the following array formula. This formula intelligently extracts all distinct numerical values from the source range (A2:A16) and arranges them in

ascending order, creating the necessary structure for the frequency bins:

=SORT(UNIQUE(A2:A16))

Upon execution, column B will be populated automatically with the unique values (12 through 17). These values now serve as the structured classes (bins) that the [FREQUENCY\(\) function](#) will use for grouping and counting, providing the framework for calculating absolute frequencies in the subsequent step.

<i>fx</i>				
	A	B	C	D
1	Data values	Classes		
2	12	12		
3	12	13		
4	13	14		
5	13	15		
6	13	16		
7	14	17		
8	14			
9	15			
10	15			
11	15			
12	15			
13	16			
14	17			
15	17			
16	17			
17				
18				
19				
20				
21				
22				
23				

Step 3: Calculating Absolute Frequencies

With the data range (A2:A16) and the defined classes (B2:B7) now ready, we can apply the central formula. As previously mentioned, the [FREQUENCY\(\) function](#) is an array function, meaning it will spill results across multiple cells equal in number to the classes defined. It is imperative that the destination range (starting at C2) is completely clear to prevent a #REF! error.

In cell C2, input the following formula. This command instructs [Google Sheets](#) to analyze the data in A2:A16 against the classes defined in B2:B7, efficiently calculating the distribution:

=FREQUENCY(A2:A16, B2:B7)

The resulting table, displayed in column C, provides the absolute [frequencies](#), which are the raw counts of how many times each unique score appeared in the original dataset:

	A	B	C	D
1	Data values	Classes	Frequency	
2	12	12	2	
3	12	13	3	
4	13	14	2	
5	13	15	4	
6	13	16	1	
7	14	17	3	
8	14		0	
9	15			
10	15			
11	15			
12	15			
13	16			
14	17			
15	17			
16	17			
17				
18				
19				
20				
21				

These calculated counts allow for immediate interpretation of the distribution of values, highlighting the central tendency and spread of the data:

The value **15** occurs most frequently, appearing **4** times.

The values **12** and **14** appear **2** times each.

The least frequent value, **16**, appears only **1** time.

Calculating Relative Frequencies for Statistical Comparison

While absolute counts are essential, calculating [relative frequencies](#) offers a more standardized

and statistically comparable measure. Relative frequency expresses each count as a proportion of the total [dataset](#) size, which is critical for comparing distributions across different sample sizes or contexts.

The calculation requires dividing each absolute frequency (Column C) by the total number of data points in the sample. We determine the total count using the standard **COUNT() function** applied to our original range (A2:A16). To ensure the denominator remains fixed when the formula is copied down the column, we must utilize absolute references (e.g., \$A\$2:\$A\$16).

Input the following formula into cell D2:

=C2/COUNT(\$A\$2:\$A\$16)

This initial calculation determines the [relative frequency](#) for the first class (value 12):

<i>fx</i>					
	A	B	C	D	
1	Data values	Classes	Frequency	Relative Frequency	
2	12	12	2	0.133	
3	12	13	3		
4	13	14	2		
5	13	15	4		
6	13	16	1		
7	14	17	3		
8	14		0		
9	15				
10	15				
11	15				
12	15				
13	16				
14	17				
15	17				
16	17				
17					
18					
19					
20					
21					
22					

To calculate the remaining proportions, simply use the fill handle (the small square at the bottom

right of cell D2) to drag the formula down to D7. This action yields the complete set of [relative frequencies](#), which are commonly formatted as percentages for improved statistical readability:

	A	B	C	D	
1	Data values	Classes	Frequency	Relative Frequency	
2	12	12	2	0.133	
3	12	13	3	0.200	
4	13	14	2	0.133	
5	13	15	4	0.267	
6	13	16	1	0.067	
7	14	17	3	0.200	
8	14		0	0	
9	15				
10	15				
11	15				
12	15				
13	16				
14	17				
15	17				
16	17				
17					
18					
19					
20					
21					

When interpreting these results, recall that [relative frequency](#) indicates the proportion of the whole. For instance, the value 13 accounts for 0.200 (or 20%) of all observed scores. A crucial validation check is that the sum of all calculated relative frequencies must precisely equal **1.0** (or 100%), confirming that the entire sample space has been accurately encompassed in the analysis.

Visualizing Frequency Distribution with a Histogram

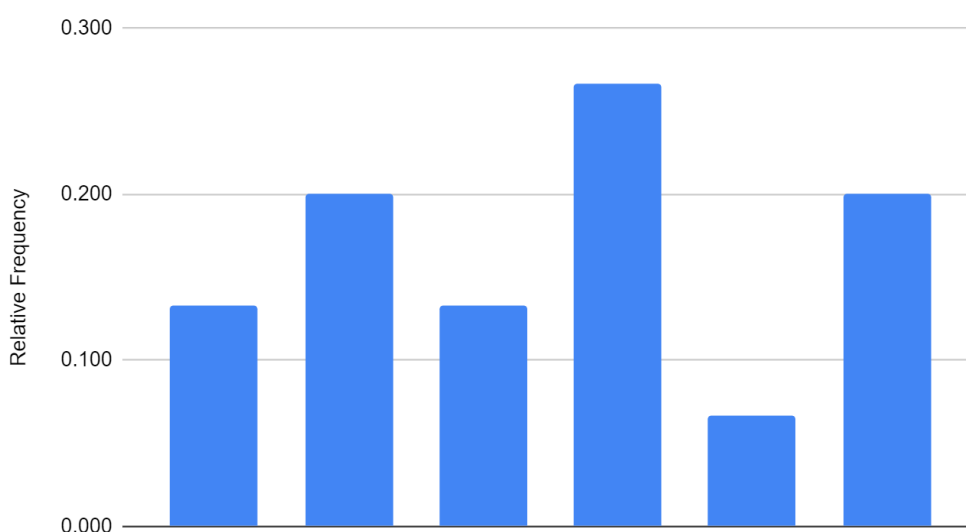
After successfully calculating both absolute and relative frequencies, the next step in effective data analysis is visualization. A [histogram](#) is the statistically appropriate and highly effective chart type for graphically representing the distribution of numerical [frequencies](#) within a [dataset](#).

To begin creating the visualization in [Google Sheets](#), highlight the range containing the calculated relative frequencies (D2:D7). This specific range contains the values that define the height of each bar in the resulting chart, representing the proportional occurrence of each score.

	A	B	C	D	
1	Data values	Classes	Frequency	Relative Frequency	
2	12	12	2	0.133	
3	12	13	3	0.200	
4	13	14	2	0.133	
5	13	15	4	0.267	
6	13	16	1	0.067	
7	14	17	3	0.200	
8	14		0	0	
9	15				
10	15				
11	15				
12	15				
13	16				
14	17				
15	17				
16	17				
17					
18					
19					
20					

Navigate to the main menu, click **Insert**, and select **Chart**. [Google Sheets](#)' smart charting tool will often suggest a column chart or a basic [histogram](#) based on the input data selection. If the default chart is not ideal, use the Chart Editor to manually select the appropriate column chart type.

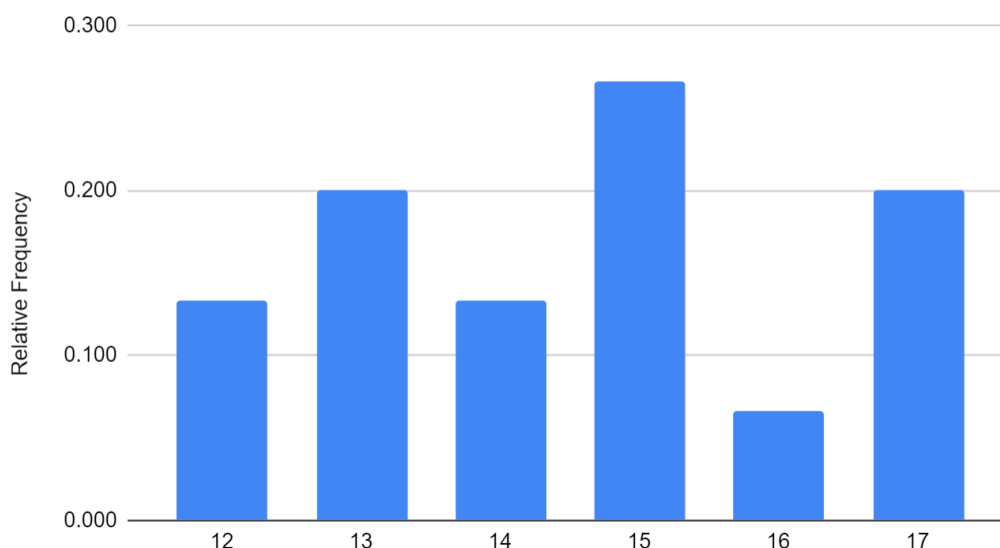
Relative Frequency



For the visualization to be fully interpretable, the x-axis must accurately reflect the unique values

(classes) we defined in Step 2. Using the **Chart Editor** panel, locate the X-axis configuration setting. Change the default range to link directly to our classes, specifying the range **B2:B7** as the horizontal axis labels. This customization transforms the generic column visualization into a precise and easily readable distribution chart:

Relative Frequency



The resulting [histogram](#) provides an immediate and impactful visual summary of the data distribution, allowing analysts to quickly identify skewness, locate the mode, and understand which specific values dominate the original dataset.