

Learn Exploratory Data Analysis (EDA) Using Excel

Authored by
Mohammed looti

October 30, 2025

RECOMMENDED CITATION

Mohammed looti (2025). *Learn Exploratory Data Analysis (EDA) Using Excel*. PSYCHOLOGICAL STATISTICS. Retrieved from <https://statistics.arabpsychology.com/?p=5841>

In the vast and evolving landscape of [data science](#), the initial and most crucial phase of any successful project is [Exploratory Data Analysis](#) (EDA). EDA is not merely a preliminary check; it is a meticulous, investigative process that empowers analysts to immerse themselves fully in a [dataset](#). By systematically examining the data, we aim to uncover hidden patterns, identify critical anomalies, and establish foundational hypotheses, all before diving into complex statistical modeling or rigorous hypothesis testing. This early stage ensures that the analyst possesses a deep, intuitive understanding of the data's inherent structure and limitations.

This comprehensive discovery process typically relies on a triumvirate of core activities, which, when executed together, deliver a holistic picture of the data's characteristics, quality, and potential predictive power. Mastering these activities is essential for any data practitioner, regardless of the tools they use.

Summarizing Data: Applying fundamental [descriptive statistics](#) to quantify the central tendencies (like the mean), dispersion (variance or range), and the overall shape of the data distribution.

Visualizing Data: Utilizing various charts, plots, and graphs to transform raw numbers into immediate, actionable insights, making trends, relationships, and glaring outliers instantly visible.

Identifying Gaps: Thoroughly detecting, locating, and quantifying the prevalence of [missing values](#) (or NA values), which are almost inevitable in real-world data collection and pose a significant threat to analytical reliability if ignored.

Through these systematic actions, data professionals secure invaluable knowledge regarding how values are distributed, pinpointing potential data quality issues and isolating problematic entries that require special attention. This proactive foundational work ensures that the data is thoroughly understood and adequately prepared before advancing to more sophisticated stages, such as [predictive analytics](#) or fitting an advanced [statistical model](#). The following practical guide will walk you through a detailed, step-by-step example of how to execute an effective exploratory data analysis using the ubiquitous spreadsheet application, [Microsoft Excel](#).

Step 1: Establishing and Preparing Your Sample Dataset

To practically illustrate the core methodologies of [Exploratory Data Analysis](#), we will first construct a simple, yet realistic, sample [dataset](#) directly within Excel. This foundational step is critical, as the quality and structure of the initial data dictate the success of subsequent analytical steps. Our illustrative dataset focuses on performance metrics for ten distinct basketball players, providing a manageable scope for demonstrating various EDA techniques.

As shown in the image below, the dataset incorporates three essential numerical [variables](#) for each athlete: the total **Points** scored, the number of **Rebounds** collected, and the amount of **Assists** made. This combination of metrics allows us to explore diverse quantitative data points

and their underlying distributions effectively.

	A	B	C	D	E	F	G	H
1		Points	Rebounds	Assists				
2		21	8	0				
3		15	7	4				
4		15	9	NA				
5		30		5				
6		29	5	5				
7		26	7	6				
8		22	4	7				
9		10		12				
10		13	2	13				
11		16	1	11				
12								
13								
14								
15								
16								
17								
18								
19								
20								
21								
22								

An intentional feature of this sample dataset is the inclusion of incomplete entries, reflecting the challenging nature of real-world data. You will note that several cells corresponding to our numerical variables contain blank fields or designated **NA values**. These gaps, representing [missing values](#), can arise from various collection issues, such as logging errors or non-response. Acknowledging and accurately handling these omissions is paramount, as unaddressed missingness can introduce severe bias and compromise the validity of any statistical conclusions drawn later in the analysis process.

Step 2: Summarizing Data Characteristics with Descriptive Statistics

Once the data has been prepared, the immediate next phase of [Exploratory Data Analysis](#) is the calculation of summary metrics. This is accomplished by leveraging [descriptive statistics](#), which provide a quantitative blueprint of the central location, variability, and overall shape of each numerical [variable](#). Specifically, we will compute key indicators: the [mean](#) (average), the [median](#) (middle value), the first and third [quartiles](#), and the absolute [minimum and maximum values](#) for our 'Points', 'Rebounds', and 'Assists' columns.

These statistical summaries offer complementary perspectives essential for initial data understanding. The **mean** gives us a weighted average, useful for assessing typical performance, while the **median** offers a robust measure of centrality, being less sensitive to extreme outliers. The **quartiles** are crucial for defining the interquartile range (IQR), illustrating the concentration and spread of the central 50% of the data. Finally, the **minimum and maximum values** clearly define the bounds of our observations. The resulting summary statistics for our basketball data are presented in the figure below:

	A	B	C	D	E	F	G
1		Points	Rebounds	Assists			
2		21	8	0			
3		15	7	4			
4		15	9	NA			
5		30		5			
6		29	5	5			
7		26	7	6			
8		22	4	7			
9		10		12			
10		13	2	13			
11		16	1	11			
12							
13	Mean	19.7	5.375	7			
14	Median	18.5	6	6			
15	Q1	15	3.5	5			
16	Q3	25	7.25	11			
17	Min	10	1	0			
18	Max	30	9	13			
19							
20							
21							
22							

To calculate these metrics efficiently within Excel, we utilized several powerful built-in functions. The exact formulas applied to the 'Points' variable (located in column B) are detailed below. These formulas were subsequently replicated across the other variables using Excel's drag-and-fill functionality:

B13: [=AVERAGE\(B2:B11\)](#) - Calculates the arithmetic average of all numeric values within the defined range.

B14: [=MEDIAN\(B2:B11\)](#) - Identifies the exact middle value in the sorted dataset.

B15: [=QUARTILE\(B2:B11, 1\)](#) - Returns the value at the 25th percentile, marking the first quartile boundary.

B16: [=QUARTILE\(B2:B11, 3\)](#) - Returns the value at the 75th percentile, defining the third quartile boundary.

B17: [=MIN\(B2:B11\)](#) - Locates the lowest numerical observation within the selected data range.

B18: [=MAX\(B2:B11\)](#) - Locates the highest numerical observation within the selected data range.

A key advantage of using Excel's native statistical functions is their inherent capacity to manage incomplete data seamlessly. Every formula listed above is specifically engineered to automatically disregard blank cells or other non-numeric [missing values](#) during the calculation of the respective [descriptive statistic](#). By applying these formulas to the 'Points' column and then dragging them horizontally, we ensure comprehensive, consistent summarization across the 'Rebounds' (column C) and 'Assists' (column D) variables without manual recalculation.

Step 3: Visualizing Data Distributions for Deeper Insights

[Data visualization](#) constitutes an indispensable cornerstone of [Exploratory Data Analysis](#). It offers a uniquely powerful mechanism for perceiving patterns, detecting hidden trends, and identifying anomalies that often remain obscure when reviewing only numerical summaries. By converting quantitative data into graphical representations, we gain immediate insight into the underlying [distribution](#) of values and the relationships between various [variables](#). Excel provides accessible and robust charting tools to facilitate this essential process.

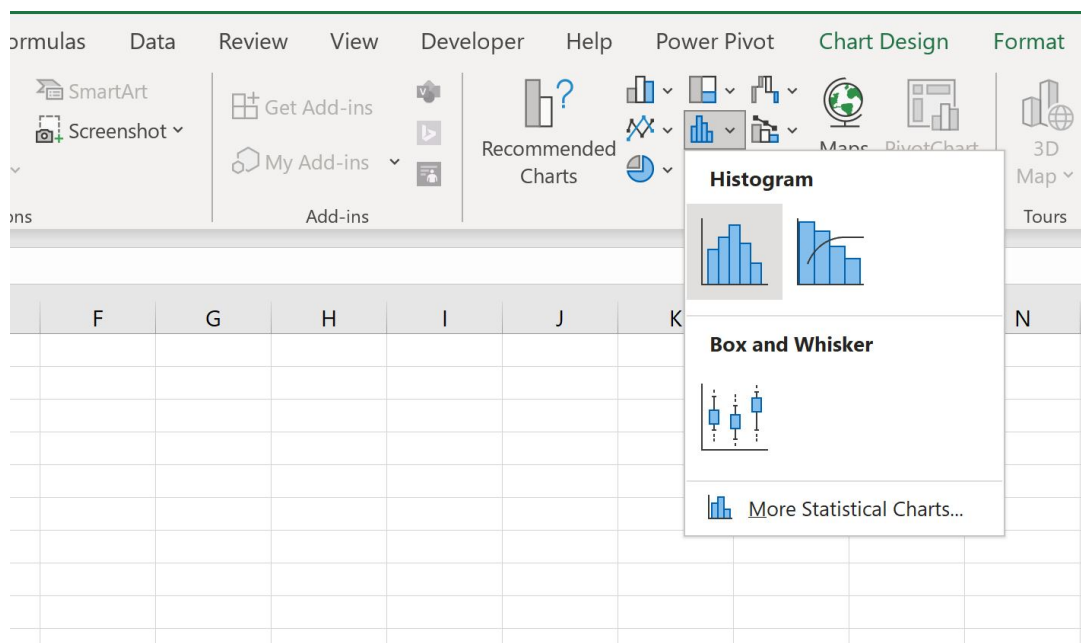
To demonstrate this, we will visualize the frequency distribution of points scored by players using a [histogram](#). The histogram is the optimal choice for understanding how a continuous numerical variable is distributed. To generate this visualization for our 'Points' variable, follow these precise steps:

Begin by selecting the relevant data range for the 'Points' variable, which spans cells **B2:B11** in our current dataset.

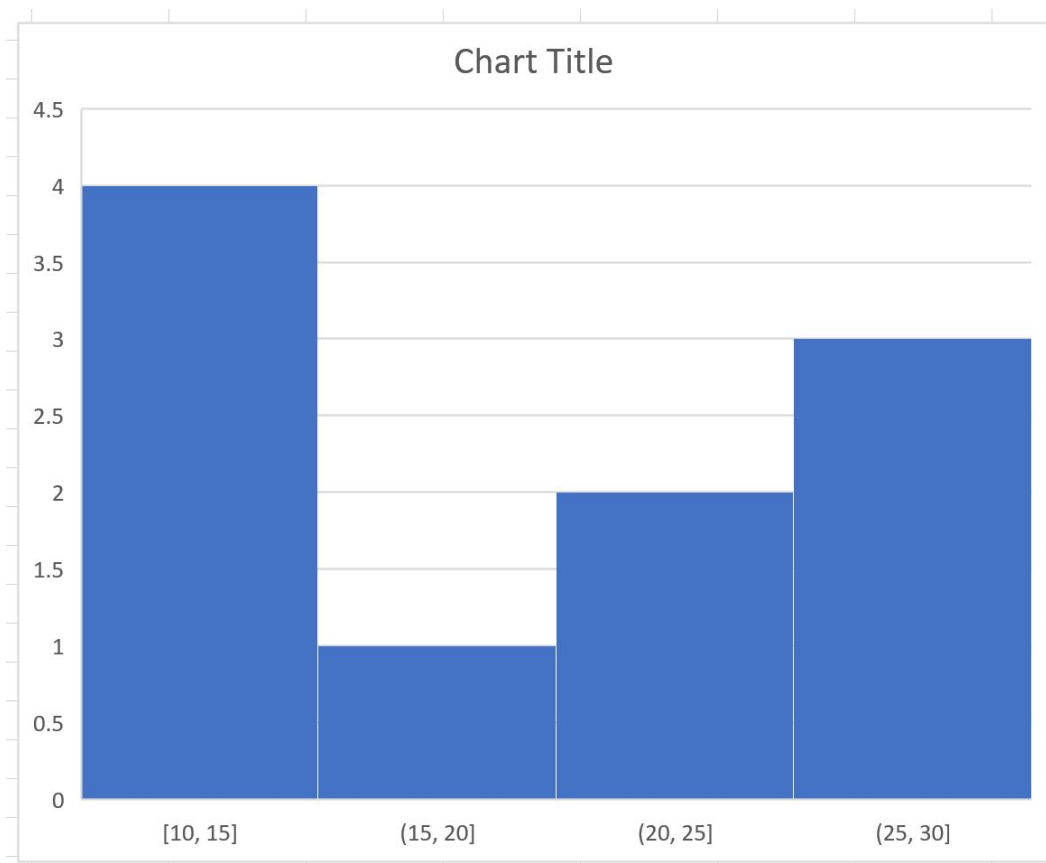
Locate and click the **Insert** tab situated on the main Excel ribbon interface.

Within the **Charts** section, click on the dedicated **Histogram** chart icon.

The graphical representation of these steps, guiding you to the correct menu selection, is clearly illustrated in the screenshot provided below:



Once these instructions are executed, Excel will instantly generate the [histogram](#). This chart offers an immediate, intuitive graphical analysis of how the observed points are clustered and spread across the group of basketball players:



Analyzing this resulting histogram provides a clear visual breakdown of performance frequency across defined bins (point ranges). By observing the height of each bar, we can quickly derive frequency counts:

The largest single group, consisting of **4 players**, scored within the 10 to 15 point bracket.

Only **1 player** recorded scores in the subsequent, lower-frequency range of 15 to 20 points.

2 players achieved scores falling between 20 and 25 points.

A second significant cluster, comprising **3 players**, is found in the highest scoring bracket, ranging from 25 to 30 points.

Interpreting this [histogram](#) allows us to characterize the overall shape of the data [distribution](#)--we can observe whether it is symmetric, positively or negatively skewed, or, as suggested here, potentially bimodal (indicating two main clusters of performance). This type of visualization is crucial for identifying common performance levels and flagging any unusually high or low outliers. This visualization process can, and should, be efficiently repeated for all other quantitative [variables](#), such as 'Rebounds' and 'Assists', to ensure a comprehensive understanding of their respective distributions.

Step 4: Precisely Identifying and Quantifying Missing Values

The culmination of our [Exploratory Data Analysis](#) requires a meticulous assessment of data quality, achieved by identifying and precisely quantifying the count of [missing values](#) present within our [dataset](#). Since missing data presents a universal challenge in real-world statistical applications, understanding its location and magnitude is paramount for preserving the integrity and accuracy of subsequent modeling or analysis. The quantification step guides decisions regarding how best to handle these gaps, whether through sophisticated imputation techniques or simple row exclusion.

While Excel doesn't have a single dedicated function for counting non-numeric blanks in a range, we can combine several powerful functions to achieve this accurate count. To determine the number of [missing values](#) specifically within column B (the 'Points' variable), we employ the following complex, yet highly effective, array formula construction:

=SUMPRODUCT(--NOT(ISNUMBER(B2:B11)))

We dissect this critical formula to appreciate the logic behind its effectiveness:

ISNUMBER(B2:B11): This function evaluates every cell in the defined range, generating an array of TRUE/FALSE results. TRUE means the cell contains a legitimate number; FALSE means it is empty, text, or an error.

NOT(...): This logical function immediately reverses the array results. Cells that were FALSE

(missing data) now become TRUE, effectively isolating the non-numeric entries we wish to count.

--: The double unary operator is used to convert the logical TRUE/FALSE values into their numeric equivalents (1 for TRUE, 0 for FALSE). Consequently, all identified missing data points are converted to the value 1.

SUMPRODUCT(...): This function concludes the process by efficiently summing the array of 1s and 0s, yielding the total count of non-numeric, or [missing values](#), within the chosen range.

To apply this, enter the formula into a summary cell, such as **B19**. The formula can then be dragged horizontally to cells **C19** and **D19**. Excel's automatic cell reference adjustment ensures that you swiftly calculate the missing data count for every subsequent [variable](#) in the dataset, completing the data quality audit.

B19								
	A	B	C	D	E	F	G	H
1		Points	Rebounds	Assists				
2		21	8	0				
3		15	7	4				
4		15	9 NA					
5		30		5				
6		29	5	5				
7		26	7	6				
8		22	4	7				
9		10		12				
10		13	2	13				
11		16	1	11				
12								
13	Mean	19.7	5.375	7				
14	Median	18.5	6	6				
15	Q1	15	3.5	5				
16	Q3	25	7.25	11				
17	Min	10	1	0				
18	Max	30	9	13				
19	Missing	0	2	1				
20								
21								
22								
23								

The resulting summary clearly establishes the distribution of data gaps across our performance metrics:

The **Points** column is complete, registering **0** missing values.

The **Rebounds** column contains **2** missing values, indicating a significant gap in collection for this

metric.

The **Assists** column exhibits **1** missing value.

Equipped with this precise understanding of data gaps, analysts can confidently determine the optimal next steps for data cleaning and preparation, whether that involves removal of incomplete records, sophisticated imputation, or utilizing methods robust to [missing values](#). Successfully concluding this step marks the end of a thorough [Exploratory Data Analysis](#) on the sample [dataset](#).

Conclusion and Additional Resources for Excel Analysis

The completion of these four structured steps provides a robust and comprehensive groundwork for interpreting any dataset. By systematically applying [descriptive statistics](#), generating insightful [visualizations](#), and meticulously quantifying data quality issues like missingness, you achieve a profound preliminary understanding of your data's intrinsic structure. This knowledge is not optional; it is foundational, guiding all subsequent analytical decisions and guaranteeing the scientific reliability of any models or conclusions derived from the data. By harnessing the readily available, powerful features of [Microsoft Excel](#), even entry-level data analysis tasks can be approached with rigor and professionalism, paving the way for more advanced data science undertakings.

To continue developing your analytical proficiency in Excel, we recommend exploring these related tutorials that cover other common and powerful spreadsheet tasks:

[Mastering What-If Analysis Techniques in Excel](#)

[Advanced Data Sorting and Organization in Excel](#)

[Efficiently Filtering and Subsetting Data within Excel](#)