

How to Perform a Repeated Measures ANOVA in Excel: A Step-by-Step Guide

Authored by
stats writer

March 10, 2026

RECOMMENDED CITATION

stats writer (2026). *How to Perform a Repeated Measures ANOVA in Excel: A Step-by-Step Guide*. PSYCHOLOGICAL SCALES. Retrieved from <https://scales.arabpsychology.com/?p=135037>

Understanding the Fundamentals of Repeated Measures ANOVA

The **Repeated Measures Analysis of Variance (ANOVA)** is a sophisticated statistical procedure employed by researchers to evaluate whether the means of three or more related groups exhibit a **statistically significant** difference. Unlike a standard **One-Way ANOVA**, which compares independent groups, the repeated measures version is specifically designed for scenarios where the same subjects are measured multiple times under different conditions or across various time intervals. This longitudinal approach allows researchers to account for individual variability among participants, thereby increasing the **statistical power** of the test by reducing the error variance associated with subject-specific differences.

In a typical research setting, **ANOVA** serves as a critical tool for hypothesis testing. By examining the variance within and between groups, the test determines if observed differences in means are likely due to the experimental manipulation or simply the result of random chance. When utilizing a **Repeated Measures ANOVA**, the focus is on the "within-subjects" factor, meaning we are looking at how the same person changes from condition A to condition B to condition C. This methodology is indispensable in fields such as psychology, medicine, and marketing, where tracking the progress of a single cohort is often more insightful than comparing disparate groups of individuals.

The primary advantage of this statistical model lies in its **efficiency** and **sensitivity**. Because every participant acts as their own control, the "noise" or background variation caused by individual differences--such as genetic makeup, personality traits, or baseline skill levels--is effectively partitioned out of the analysis. This makes it significantly easier to detect subtle effects of a treatment or intervention that might otherwise be obscured in an independent samples design. Understanding the underlying logic of this test is the first step toward successfully implementing it within a data processing environment like **Microsoft Excel**.

Determining the Appropriate Use Cases for this Statistical Procedure

Before proceeding with the technical execution in **Excel**, it is vital to identify when a **Repeated Measures ANOVA** is truly necessary. The most common application is in **longitudinal studies**, where data points are collected at multiple intervals, such as pre-test, post-test, and a six-month follow-up. For instance, a clinician might measure the **blood pressure** of patients before starting a new medication, one month into the treatment, and three months after completion. In this case, the time points represent the "related groups," and the goal is to see if the medication consistently lowers pressure over time.

Another prevalent use case is the **crossover design** or **within-subjects experiment**. In these scenarios, participants are exposed to every level of a categorical independent variable. For example, a researcher might want to test the impact of different background noises (silence, white

noise, and classical music) on a student's ability to focus. Each student would perform a task under all three conditions. Because the same subjects provide data for every category, the observations are not independent, necessitating the use of a **Repeated Measures ANOVA** to maintain statistical validity and avoid violating the assumption of independence required by other tests.

Finally, this test is frequently used in **marketing research** and **product development**. Companies often conduct sensory testing where a panel of experts rates several different formulations of a product. Since each judge tastes every sample, the ratings are inherently linked to the specific preferences and sensitivities of the individual judge. By applying a **Repeated Measures ANOVA**, the company can determine which formulation is objectively superior while accounting for the fact that some judges may be generally more critical or lenient than others, thus ensuring a more accurate **mean** comparison across the product line.

Essential Statistical Assumptions and Data Prerequisites

To ensure the results of a **Repeated Measures ANOVA** are reliable, several **statistical assumptions** must be met. The first and perhaps most critical is the **assumption of sphericity**. Sphericity requires that the variances of the differences between all possible pairs of related groups be equal. If this assumption is violated, the **F-statistic** may be positively biased, leading to an increased **Type I error** rate. While Excel does not automatically test for sphericity using Mauchly's test, it is a factor that researchers must consider, potentially applying corrections like the Greenhouse-Geisser adjustment if the data appears problematic.

The second major assumption is **normality**. This implies that the residuals (the differences between the observed values and the group means) should follow a **normal distribution** for each level of the within-subjects factor. While **ANOVA** is generally considered robust to minor deviations from normality, especially with larger sample sizes, significant **skewness** or **kurtosis** can undermine the accuracy of the **p-value**. Users should ideally screen their data using histograms or descriptive statistics before running the final analysis to ensure the distribution is approximately bell-shaped.

Lastly, the **independence of observations** must be maintained across different subjects. While the measurements within a single subject are related, the subjects themselves must be independent of one another. This means that the response of one participant should not influence the response of another. Additionally, the dependent variable should be measured on a **continuous scale** (interval or ratio). If your data is categorical or ordinal, alternative non-parametric tests, such as the **Friedman test**, should be explored instead of the standard **Repeated Measures ANOVA**.

Formatting and Structuring Data within the Excel Spreadsheet

Proper data organization is the foundation of any successful analysis in **Microsoft Excel**. To perform a **Repeated Measures ANOVA**, you must structure your spreadsheet so that each row represents a unique subject (or participant) and each column represents a different measurement point or condition. This format is often referred to as "wide format" data. For example, if you are tracking the reaction times of five patients across four different drugs, your first column should list the patient identifiers, while the subsequent four columns should contain the specific reaction time values for Drug A, Drug B, Drug C, and Drug D.

When entering the data, ensure that there are no empty cells within the range you intend to analyze. Missing data points can cause the **Data Analysis Toolpak** to produce errors or lead to an imbalanced design that **Excel** cannot handle using its standard functions. It is also good practice to include clear headers for your columns. These labels will not only help you keep track of your variables but can also be included in the analysis range to make the final output table easier to read and interpret. Avoid including any non-numeric characters or special symbols within the data cells, as **Excel** requires pure numerical input for **statistical analysis**.

Step 1: Enter the data.

Enter the following data, which shows the response time (in seconds) of five patients on the four drugs:

	A	B	C	D	E	F	G
1	Patient	Drug 1	Drug 2	Drug 3	Drug 4		
2	Patient 1	30	28	16	34		
3	Patient 2	14	18	10	22		
4	Patient 3	24	20	18	30		
5	Patient 4	38	34	20	44		
6	Patient 5	26	28	14	30		
7							
8							
9							
10							
11							
12							
13							
14							
15							

Once the data is entered as shown above, double-check for any typos. Even a single incorrectly entered decimal point can significantly alter the **standard deviation** and the resulting **F-statistic**.

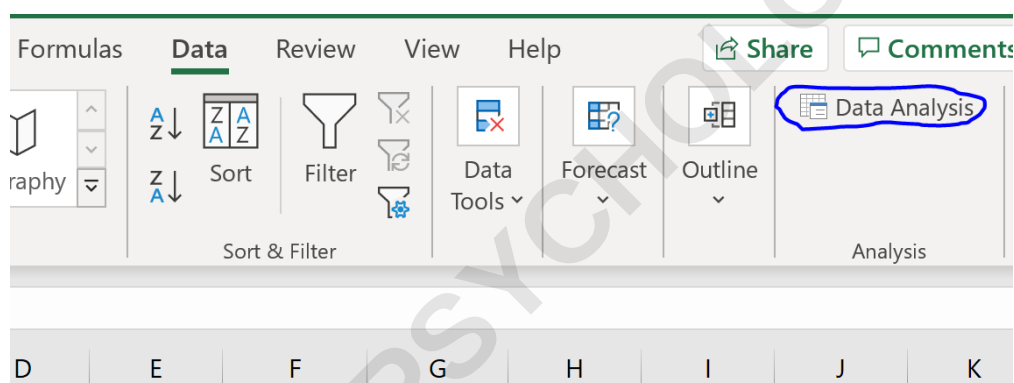
Organizing your data meticulously at this stage ensures that the subsequent calculations are performed on a clean and accurate dataset, which is the cornerstone of **quantitative research**.

Activating and Utilizing the Excel Data Analysis Toolpak

The **Data Analysis Toolpak** is a powerful **add-in** for **Excel** that provides a variety of advanced statistical functions, including **ANOVA**, **regression**, and **correlation**. By default, this toolpak is not active in a fresh installation of **Excel**. To enable it, you must navigate to the "File" tab, select "Options," click on "Add-ins," and then choose "Excel Add-ins" from the "Manage" dropdown menu. After clicking "Go," check the box for "Analysis Toolpak" and click "OK." Once activated, the "Data Analysis" button will appear in the "Analysis" group on the "Data" tab.

Step 2: Perform the repeated measures ANOVA.

To perform the repeated measures ANOVA, go to the **Data** tab and click on **Data Analysis**. If you don't see this option, you need to first enable the add-in as described above.



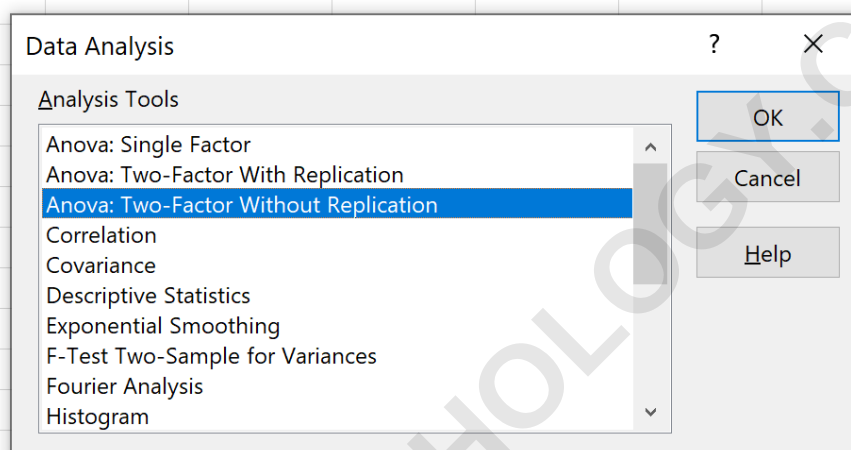
Activating this tool is a one-time process; once it is enabled, it will remain available in your **Excel** environment for all future projects. This toolpak is essential because **Excel's** standard worksheet functions, like **AVERAGE** or **STDEV**, cannot perform the complex variance partitioning required for a **Repeated Measures ANOVA**. By utilizing the **Analysis Toolpak**, you gain access to a professional-grade interface that automates the generation of comprehensive summary tables and statistical metrics.

Step-By-Step Execution of the ANOVA Procedure

Once you have clicked on **Data Analysis**, a list of available tools will appear. To conduct a **One-Way Repeated Measures ANOVA**, you should select **Anova: Two-Factor Without Replication**. Although the name might seem counterintuitive, this specific tool is the correct choice because a **Repeated Measures ANOVA** is mathematically equivalent to a **Two-Way ANOVA** where one

factor is the "Treatment" (the conditions or time points) and the second factor is the "Subjects" (the individual participants).

	A	B	C	D	E	F	G
1	Patient	Drug 1	Drug 2	Drug 3	Drug 4		
2	Patient 1	30	28	16	34		
3	Patient 2	14	18	10	22		
4	Patient 3	24	20	18	30		
5	Patient 4	38	34	20	44		
6	Patient 5	26	28	14	30		
7							
8							
9							
10							
11							
12							
13							
14							
15							
16							
17							
18							
19							
20							



After selecting **Anova: Two-Factor Without Replication** and clicking "OK," a new dialog box will appear. You must define the **Input Range** by highlighting the entire block of data, including the labels in the first row and column. Ensure that the "Labels" checkbox is ticked if you have included headers. The "Alpha" level is set to 0.05 by default, which is the standard threshold for **statistical significance** in most scientific research. Finally, specify an **Output Range**--a blank area on your worksheet where **Excel** will place the results--and click "OK."

	A	B	C	D	E	F	G	H	I	J	K
1	Patient	Drug 1	Drug 2	Drug 3	Drug 4						
2	Patient 1	30	28	16	34						
3	Patient 2	14	18	10	22						
4	Patient 3	24	20	18	30						
5	Patient 4	38	34	20	44						
6	Patient 5	26	28	14	30						
7											
8											
9											
10											
11											
12											
13											
14											
15											
16											
17											
18											
19											
20											

Anova: Two-Factor Without Replication

Input Range:

☐ Labels

Alpha:

Output options

☒ Output Range:

☐ New Worksheet Ply:

☐ New Workbook

OK Cancel Help

Selecting the "Two-Factor Without Replication" option effectively tells **Excel** to partition the total variance into three sources: the variance between the different drugs, the variance between the individual patients, and the residual error. In the context of a **Repeated Measures ANOVA**, we are primarily concerned with the variance between the drugs, as this tells us if the treatment had an effect. The variance between patients is accounted for but is not the primary focus of our **hypothesis** testing.

Analyzing the Components of the Generated ANOVA Output

After clicking "OK," **Excel** will generate a detailed output consisting of two main sections: a **Summary** table and an **ANOVA** table. The Summary table provides descriptive statistics for each subject (rows) and each condition (columns), including the count, **sum**, average, and **variance**. These values are useful for a preliminary check of your data; for instance, you can quickly see which drug resulted in the fastest or slowest average response time across all five patients.

The core of the analysis is found in the **ANOVA** table, which breaks down the **Sum of Squares (SS)**, **Degrees of Freedom (df)**, **Mean Square (MS)**, and the **F-statistic**. The table is divided into "Rows," "Columns," and "Error." In our specific case of testing four drugs across five patients, the "Columns" row represents the effect of the drugs, while the "Rows" row represents the differences between the patients. The "Error" row represents the random variation that cannot be explained by either the drugs or the patients.

	A	B	C	D	E	F	G	H
1	Patient	Drug 1	Drug 2	Drug 3	Drug 4			
2	Patient 1	30	28	16	34			
3	Patient 2	14	18	10	22			
4	Patient 3	24	20	18	30			
5	Patient 4	38	34	20	44			
6	Patient 5	26	28	14	30			
7								
8		Anova: Two-Factor Without Replication						
9								
10		<i>SUMMARY</i>	<i>Count</i>	<i>Sum</i>	<i>Average</i>	<i>Variance</i>		
11		Row 1	4	108	27	60		
12		Row 2	4	64	16	26.66667		
13		Row 3	4	92	23	28		
14		Row 4	4	136	34	104		
15		Row 5	4	98	24.5	51.66667		
16								
17		Column 1	5	132	26.4	76.8		
18		Column 2	5	128	25.6	42.8		
19		Column 3	5	78	15.6	14.8		
20		Column 4	5	160	32	64		
21								
22								
23		ANOVA						
24		<i>Source of Variation</i>	<i>SS</i>	<i>df</i>	<i>MS</i>	<i>F</i>	<i>P-value</i>	<i>F crit</i>
25		Rows	680.8	4	170.2	18.10638	5.07E-05	3.259167
26		Columns	698.2	3	232.7333	24.75887	1.99E-05	3.490295
27		Error	112.8	12	9.4			
28								
29		Total	1491.8	19				
30								

Focusing on the "Columns" row is essential for answering our research question. The **Mean Square** for Columns is calculated by dividing the Sum of Squares by the Degrees of Freedom. The **F-statistic** is then derived by dividing the MS for Columns by the MS for Error. This ratio effectively compares the systematic variance (the effect of the drugs) to the unsystematic variance (the error). A high F-value suggests that the differences between the drugs are much larger than what would be expected by chance alone, pointing toward a significant result.

Interpreting Statistical Significance and the F-Ratio

The final step in interpreting the **Excel** output is to examine the **p-value** and the **F critical value**. The p-value indicates the probability of obtaining an F-statistic as extreme as, or more extreme than, the one calculated, assuming that the **null hypothesis** (that all means are equal) is true. If the p-value is less than your chosen Alpha level (typically 0.05), you have sufficient evidence to **reject the null hypothesis**. In the example provided, the p-value is **0.0000199**, which is

substantially lower than 0.05, indicating a highly significant result.

Alternatively, you can compare the calculated **F-statistic** to the **F critical value** (F crit) provided by **Excel**. The F crit is the threshold that the calculated F-statistic must exceed for the results to be considered significant at the specified Alpha level. In our results, the F-statistic of **24.75887** is much larger than the F critical value. This leads to the same conclusion: there is a statistically significant difference in the mean reaction times across the four different drugs tested on the patients.

It is important to remember that a significant **Repeated Measures ANOVA** only tells you that at least one group mean is different from the others; it does not specify *which* pairs of groups are different. To determine the specific differences, researchers typically perform **post-hoc analysis**, such as Tukey's HSD or Bonferroni corrections. Unfortunately, **Excel** does not include built-in post-hoc tests in the Toolpak, so these must be calculated manually or using external **statistical software** to pinpoint the exact nature of the differences.

Formally Reporting the Results for Academic or Professional Use

Once the analysis is complete, the findings must be reported in a clear and standardized format, often following **APA style** for academic publications. A formal report should include the type of test performed, the number of participants, the purpose of the study, the **F-statistic**, the degrees of freedom for both the effect and the error, and the p-value. This level of detail allows other researchers to verify your work and understand the strength of your findings.

For the example of the drug reaction study, a formal report might look like this: "A **One-Way Repeated Measures ANOVA** was conducted on five individuals to investigate the influence of four different drugs on response time. The analysis revealed a statistically significant effect of drug type on response time, **$F(3, 12) = 24.76, p < 0.001$** ." This concise statement captures all the essential statistical information needed for an audience to evaluate the study's outcome.

In addition to the text report, creating a visual representation of the means can be highly beneficial. A **line chart** or a bar graph with **error bars** (representing the standard error or confidence intervals) can help illustrate the trend across the conditions. While **Excel** provides the raw data needed for these charts in the Summary table, the user must manually create the visualization to accompany the statistical report, providing a comprehensive overview of the research results.

Considering the Constraints of Excel for Comprehensive Data Analysis

While **Microsoft Excel** is an excellent tool for basic statistical tasks and is widely accessible, it has several limitations when it comes to complex analyses like the **Repeated Measures ANOVA**. One major drawback is the lack of automated checks for assumptions. Professional software like **SPSS**,

R, or **SAS** will automatically run Mauchly's test for sphericity and provide adjusted p-values (like Greenhouse-Geisser) if the assumption is violated. In **Excel**, the user must be aware of these issues and handle them manually, which increases the risk of error.

Furthermore, as previously mentioned, the absence of integrated **post-hoc testing** is a significant hurdle. After finding a significant F-test, most researchers need to know exactly which conditions differ. Performing these comparisons in **Excel** requires manual calculation of **t-tests** with manual adjustments for multiple comparisons (like the **Bonferroni correction**), which is time-consuming and prone to clerical mistakes. For studies involving a large number of conditions or complex interactions, **Excel** can quickly become overwhelmed.

Despite these limitations, **Excel** remains a valuable starting point for many researchers and students due to its intuitive interface and the fact that it requires no coding knowledge. It is perfect for quick exploratory data analysis or for smaller datasets where the assumptions are likely to be met. However, for high-stakes research or complex experimental designs, it is highly recommended to consult a **statistician** or utilize more robust statistical software to ensure that the analysis is both comprehensive and technically sound.