

# How to Easily Conduct Repeated Measures ANOVA in SAS

Authored by  
**stats writer**

December 1, 2025

## RECOMMENDED CITATION

stats writer (2025). *How to Easily Conduct Repeated Measures ANOVA in SAS*.  
PSYCHOLOGICAL SCALES. Retrieved from <https://scales.arabpsychology.com/?p=103427>

The repeated measures ANOVA is a powerful statistical technique essential for analyzing experimental designs where the same subjects are measured under multiple conditions or across various time points. This design is particularly prevalent in fields like psychology, medicine, and human performance studies, offering greater statistical power compared to independent measures designs because it inherently controls for between-subject variability. When implementing this analysis within the SAS statistical software environment, the primary tool utilized is the PROC GLM (General Linear Model) procedure.

Using PROC GLM to conduct a repeated measures analysis requires careful specification of the linear model, ensuring that both the dependent variable and the crucial within-subject factors are correctly identified. The structure of the analysis focuses on partitioning the total variance into components attributable to the treatment effects (the repeated measures) and the error. The resulting output is critical for determining the main effect of the within-subject factor, allowing researchers to assess whether the differences observed across the levels of measurement are statistically significant. This comprehensive tutorial provides a high-level, step-by-step guide to performing and interpreting this essential analysis using the SAS system.

Fundamentally, a repeated measures ANOVA is employed to test for differences between the means of three or more related groups, where the key characteristic is that every participant contributes data to every level of the independent variable. This design minimizes confounding variables and is optimized for situations demanding precise measurement of change or differential response. By following the upcoming steps, you will learn the exact syntax and interpretation required to execute this analysis effectively within your statistical practice.

## Understanding the Theoretical Framework of Repeated Measures ANOVA

Before diving into the coding aspects within SAS, it is crucial to establish a solid understanding of why the repeated measures ANOVA is the appropriate choice for certain experimental designs. This model is essentially an extension of the paired-samples t-test, accommodating three or more levels of a factor. The primary benefit lies in the fact that subject variability, which typically inflates the error term in between-subjects designs, can be separated out and removed from the denominator in the F-ratio calculation. This typically results in a more sensitive test, increasing the likelihood of detecting a true effect if one exists.

The statistical theory underpinning this procedure relies on the assumption of sphericity. Sphericity is an assumption specific to repeated measures designs, requiring that the variances of the differences between all pairs of within-subject conditions are equal. If this assumption is violated, the F-statistic will be positively biased, leading to an increased **Type I error rate** (falsely rejecting the null hypothesis). While the basic PROC GLM setup illustrated here is a fundamental approach, more advanced SAS procedures or the use of the **REPEATED** statement within PROC GLM would

provide automated tests for sphericity, such as the Mauchly's test, and offer adjustments (like Greenhouse-Geisser or Huynh-Feldt corrections) if the assumption is not met. Understanding these theoretical prerequisites is vital for accurate data interpretation.

The null hypothesis (**H<sub>0</sub>**) for the repeated measures ANOVA states that the population means for all levels of the within-subjects factor are equal. Conversely, the alternative hypothesis (**H<sub>A</sub>**) posits that at least two of the population means are significantly different. The analysis aims to generate an F-statistic by comparing the variance due to the treatment (the numerator) against the remaining error variance (the denominator, after removing the subject variance). A large F-statistic, coupled with a small P-value, provides evidence to reject the null hypothesis and conclude that the experimental manipulation had a statistically significant effect.

## Defining the Research Scenario and Data Structure

To illustrate the application of the repeated measures ANOVA in a practical setting, let us consider a typical experimental scenario from psychopharmacology. Suppose a dedicated researcher is interested in determining whether four distinct pharmacological agents (Drug 1, Drug 2, Drug 3, and Drug 4) elicit differential effects on human reaction time. Crucially, to control for individual differences in baseline reaction speed, the researcher employs a **within-subjects design**, meaning the same group of five patients is tested once under the influence of each of the four drugs. This approach ensures that the comparisons drawn are truly reflective of the drug's effect, minimizing the noise introduced by varying patient characteristics.

In this specific design, the **dependent variable** is the measured reaction time (**Value**), which is a continuous metric. The **within-subject factor** is the Drug type (**Drug**), which has four distinct levels (1, 2, 3, 4). The **Subject** identification serves as a crucial blocking factor, accounting for individual differences in the overall linear model. It is imperative that the data are structured in a **long format** for input into PROC GLM, meaning each row represents a single observation (a subject tested under one specific drug condition), rather than a wide format where each subject is represented by a single row containing measurements for all four drugs.

The hypothetical reaction times recorded for the five patients across the four drug conditions are detailed in the table below. This structured data visualization aids in understanding how the measurements are related across the different conditions for each unique participant. Notice how Patient 1 has four separate measurements, corresponding to the four drug levels, a characteristic defining the repeated measures structure:

| Patient   | Drug 1 | Drug 2 | Drug 3 | Drug 4 |
|-----------|--------|--------|--------|--------|
| Patient 1 | 30     | 28     | 16     | 34     |
| Patient 2 | 14     | 18     | 10     | 22     |
| Patient 3 | 24     | 20     | 18     | 30     |
| Patient 4 | 38     | 34     | 20     | 44     |
| Patient 5 | 26     | 28     | 14     | 30     |

## Step 1: Creating the Dataset in SAS

The first practical step in SAS involves transforming the raw measurement data into a usable dataset format. We utilize the **DATA step** to define the variables and input the observations directly using the **DATALINES** statement. This process is fundamental to all statistical analysis in SAS. For our repeated measures scenario, we must ensure three variables are clearly defined: the categorical **Subject** identifier, the categorical **Drug** level, and the continuous dependent variable, **Value** (Reaction Time).

The code segment below demonstrates the implementation of the long data format. Each block of four lines corresponds to a single subject, listing their reaction time (Value) for each sequential Drug (1 through 4). This explicit structure is non-negotiable for correctly specifying the within-subject design later in the PROC GLM procedure. Using descriptive comments, as shown in green, is a best practice for maintaining code clarity and documentation, especially when working with complex statistical programming.

Here is the precise SAS code used to input and create the required dataset named **my\_data**:

```
/*create dataset*/
data my_data;
input Subject Drug Value;
datalines;
1 1 30
1 2 28
1 3 16
1 4 34
2 1 14
2 2 18
2 3 10
2 4 22
3 1 24
```

```
3 2 20
3 3 18
3 4 30
4 1 38
4 2 34
4 3 20
4 4 44
5 1 26
5 2 28
5 3 14
5 4 30
;
run;
```

After executing this DATA step, SAS holds the data in memory, structured appropriately for the upcoming analysis. It is important to verify the data integrity at this stage, perhaps using **PROC PRINT** to ensure that the data input matches the original source table. Any discrepancies here will compromise the validity of the final statistical results. The use of numeric coding (1, 2, 3, 4) for the categorical factor 'Drug' is common practice, though using formats could enhance readability in the output.

## Step 2: Executing the Repeated Measures ANOVA using PROC GLM

The heart of the repeated measures analysis in SAS relies on the PROC GLM procedure. This procedure operates by fitting a linear model to the data, effectively partitioning the variance based on the factors specified. The key to successful implementation lies in correctly defining the categorical variables and the statistical model itself. The structure we use here is a simplified but effective method for one-way repeated measures when the data is in the long format.

Within the PROC GLM block, two primary statements are essential for this analysis. First, the **CLASS** statement is used to explicitly declare which variables are categorical or grouping factors, rather than continuous predictors. In our example, both **Subject** and **Drug** must be listed here. Second, the **MODEL** statement defines the structure of the analysis. For a one-way repeated measures design, the model must include the dependent variable (**Value**) on the left side, and both the Subject and the within-subject factor (**Drug**) on the right side.

The inclusion of **Subject** in the model statement is what mathematically defines the repeated measures structure for this specific type of analysis in SAS. By including **Subject**, we are instructing SAS to remove the variance associated with individual patient differences before calculating the error term used to test the within-subject factor (**Drug**). This distinction is critical for

correctly calculating the F-ratio, ensuring the denominator (error variance) is reduced, thereby increasing the power of the test. Note that the error term for the **Drug** factor is implicitly the residual variance remaining after accounting for Subject and Drug effects, which is equivalent to the Subject\*Drug interaction in a traditional ANOVA partitioning of sums of squares.

The code snippet below executes the procedure on the previously created **my\_data** dataset:

```
/*perform repeated measures ANOVA*/  
proc glm data=my_data;  
class Subject Drug;  
model Value = Subject Drug;  
run;
```

### Step 3: Interpreting the ANOVA Output Table

The core objective of performing the ANOVA is to evaluate the null hypothesis. After running PROC GLM, SAS provides a detailed ANOVA table that summarizes the results of the variance decomposition. When interpreting this output for a repeated measures design, the researcher must focus specifically on the row corresponding to the within-subject factor--in our case, the variable **Drug**. While the results for the **Subject** effect indicate the extent of variability between individuals, this variance is essentially controlled for in this type of design and is not the primary focus for answering the research question.

The critical output generated by SAS, which must be carefully analyzed, is presented below:

## The GLM Procedure

Dependent Variable: Value

| Source          | DF | Sum of Squares | Mean Square | F Value | Pr > F |
|-----------------|----|----------------|-------------|---------|--------|
| Model           | 7  | 1379.000000    | 197.000000  | 20.96   | <.0001 |
| Error           | 12 | 112.800000     | 9.400000    |         |        |
| Corrected Total | 19 | 1491.800000    |             |         |        |

| R-Square | Coeff Var | Root MSE | Value Mean |
|----------|-----------|----------|------------|
| 0.924387 | 12.31302  | 3.065942 | 24.90000   |

| Source  | DF | Type I SS   | Mean Square | F Value | Pr > F |
|---------|----|-------------|-------------|---------|--------|
| Subject | 4  | 680.8000000 | 170.2000000 | 18.11   | <.0001 |
| Drug    | 3  | 698.2000000 | 232.7333333 | 24.76   | <.0001 |

| Source  | DF | Type III SS | Mean Square | F Value | Pr > F |
|---------|----|-------------|-------------|---------|--------|
| Subject | 4  | 680.8000000 | 170.2000000 | 18.11   | <.0001 |
| Drug    | 3  | 698.2000000 | 232.7333333 | 24.76   | <.0001 |

To draw conclusions regarding the differential effects of the drugs, we specifically examine the F-statistic and the associated P-value (Pr > F) for the **Drug** factor. These values quantify the likelihood that the observed differences in mean reaction times between the four drug levels occurred simply by chance, assuming the null hypothesis is true. A P-value less than the predetermined alpha level (typically set at  $\alpha = 0.05$ ) provides statistically robust evidence to reject the null hypothesis. Based on the provided table, we extract the necessary statistics for the effect of the pharmacological treatment:

The F Value for Drug: **24.76**

The p-value (Pr > F) for Drug: **<.0001**

## Drawing Statistical Conclusions

The interpretation of the statistical output directly relates back to the formulation of the hypotheses for this specific experiment. Recall the fundamental hypotheses driving the repeated measures ANOVA:

**H0 (Null Hypothesis):** The mean reaction times are equal across all four drug conditions ( $\mu_1 = \mu_2 = \mu_3 = \mu_4$ ).

**HA (Alternative Hypothesis):** At least one pair of drug conditions exhibits a statistically significant



difference in mean reaction time.

With an observed F-statistic of **24.76** and an associated P-value of **<.0001**, we compare this P-value against the conventional significance threshold of  $\alpha = 0.05$ . Since 0.0001 is substantially smaller than 0.05, the statistical decision is unequivocally to reject the null hypothesis (**H0**). This rejection signifies that the differences in reaction times observed among the four drug types are highly unlikely to be due merely to random sampling variability.

Therefore, we conclude that there is compelling statistical evidence to support the claim that the mean response time is not uniform across the four different drugs. In other words, the type of drug administered has a significant main effect on the reaction time outcome. This highly significant finding indicates that the treatment, represented by the four levels of the Drug factor, has successfully induced differential responses in the subjects. The success of this finding is partially attributable to the efficiency of the repeated measures design in controlling for subject-specific baseline characteristics.

### Considerations for Further Analysis (Post Hoc Testing)

Although the primary ANOVA confirms a significant effect, the researcher's interest usually extends beyond this general finding to determine the specifics: Which drug is the most effective, and which specific pairs of drugs yield differential results? Because the overall F-test is omnibus, failing to perform follow-up comparisons when the null hypothesis is rejected is an incomplete analysis. These subsequent tests, known as **post hoc tests** or multiple comparisons, are necessary to identify the source of the significant variance while maintaining control over the familywise error rate.

In SAS, post hoc testing is seamlessly integrated within the PROC GLM environment using the **MEANS** statement. This statement allows for various comparison methods to be specified, often including options like **LSD** (Least Significant Difference), **Tukey's HSD**, or **Bonferroni** adjustments, depending on the desired level of stringency and the specific hypotheses being tested. For instance, modifying the PROC GLM statement to include **means Drug / tukey;** would automatically perform all pairwise comparisons between the four drug means, adjusting the P-values to account for the multiple testing problem.

Furthermore, should the researcher have implemented a more complex repeated measures design (e.g., one involving multiple between-subjects factors or testing the sphericity assumption), the use of the **REPEATED** statement would be essential for calculating multivariate tests and sphericity corrections. The methods outlined in this tutorial provide the fundamental basis using the long dataset format; however, researchers are encouraged to explore advanced options like PROC MIXED for handling missing data or unbalanced designs common in longitudinal studies. The proper application of post hoc tests ensures that the final conclusions about which drugs differ from



one another are statistically sound and reliable.

ARABPSYCHOLOGY.COM