

How to Easily Interpret the P-value in Regression Output

Authored by
stats writer

November 21, 2025

RECOMMENDED CITATION

stats writer (2025). *How to Easily Interpret the P-value in Regression Output*.

PSYCHOLOGICAL SCALES. Retrieved from <https://scales.arabpsychology.com/?p=98736>

The P-value is arguably one of the most critical metrics derived from any statistical analysis, particularly within regression output. Fundamentally, the P-value serves as a quantitative measure of evidence against the null hypothesis (H_0). It quantifies the probability of observing the current data, or data even more extreme, assuming that the null hypothesis is true--that is, assuming there is no actual relationship between the predictor variables and the response variable in the population being studied. When analyzing regression results, this probability directly informs the researcher about the likelihood that the observed relationship between the independent and dependent variables is merely a product of random sampling variation, rather than reflecting a genuine, underlying connection.

Interpreting the P-value requires a comparison against a predetermined threshold, known as the alpha level (α), which is most commonly set at 0.05. A low P-value (typically $P < 0.05$) signifies that the results are highly unlikely to have occurred purely by chance if the null hypothesis were true. Consequently, a small P-value provides strong evidence to reject H_0 , leading to the conclusion that the relationship is statistically significant. Conversely, a high P-value suggests that the observed data is reasonably likely under the assumption of the null hypothesis, meaning we lack sufficient evidence to claim statistical significance for that particular relationship, and thus the variable or model may not be a useful predictor.

Understanding the Role of the P-value in Statistical Modeling

When we employ Multiple Linear Regression, the primary goal is to quantify and model the intricate relationship between two or more predictor variables (also known as independent variables) and a single response variable (the dependent variable). This powerful statistical technique allows researchers to build predictive models, estimate the magnitude of influence each predictor exerts, and understand the overall predictive power of the combined set of variables. Unlike simple correlation, regression establishes a functional form--a mathematical equation--that links inputs to the output, making it indispensable in fields ranging from economics and finance to social sciences and engineering.

The process of analyzing these relationships is inherently tied to the principles of Hypothesis Testing. Every coefficient generated by the regression model--whether for an individual predictor or for the model as a whole--is tested against the assumption that it is zero in the population. If the coefficient is truly zero, it means the variable has no linear effect on the response variable. The P-value is the tool we use to assess the plausibility of this assumption. Therefore, whenever we perform multiple linear regression, we must scrutinize the P-values in the output to rigorously determine if the observed relationships are scientifically meaningful and possess statistical significance.

A thorough understanding of how to interpret these P-values is paramount for drawing correct

conclusions from the analysis. Misinterpreting a P-value can lead to incorrectly claiming a relationship where none exists (Type I error) or failing to recognize a genuine relationship (Type II error). This tutorial provides a detailed, formal guide explaining how to correctly interpret the various P-values generated in the output of a multiple linear regression model, using Excel as the demonstration environment.

The Hypothesis Testing Framework in Regression Analysis

In the context of regression, hypothesis testing is performed separately for the entire model and for each individual predictor variable. For any regression coefficient (β), the null hypothesis (H_0) always asserts that the population coefficient is zero ($\beta = 0$). This zero coefficient signifies that the corresponding predictor variable has no influence on the dependent variable. Conversely, the alternative hypothesis (H_a or H_1) posits that the coefficient is not zero ($\beta \neq 0$), implying that a real relationship exists. The P-value is the direct result of a statistical test (typically an F-test for the overall model or a t-test for individual coefficients) designed to quantify the evidence against H_0 .

To make a decision about rejecting or failing to reject the null hypothesis, we must define the level of risk we are willing to accept, which is the alpha level (α). When α is set at 0.05 (the standard), we are accepting a 5% risk of incorrectly rejecting the null hypothesis (a false positive). If the calculated P-value is less than or equal to α ($P \leq 0.05$), the observed data is considered sufficiently improbable under H_0 , leading us to reject the null hypothesis and conclude that the variable or model is statistically significant.

It is crucial to remember that the P-value itself is not the probability that the null hypothesis is true, nor is it a measure of the effect size or practical importance. It is merely a measure of how incompatible the data are with the specific assumption outlined in the null hypothesis. A variable can be statistically significant (low P-value) but have a negligible impact (small coefficient), or it can be practically important (large coefficient) but lack statistical significance (high P-value) due to small sample size or high variability. Therefore, the P-value must always be interpreted alongside the coefficient estimates, confidence intervals, and the context of the research question.

Practical Example: Interpreting Regression P-values in Excel

To demonstrate the interpretation process, let us consider a practical scenario involving educational data. Suppose our research objective is to determine whether two factors--the total number of hours a student spends studying and the number of preparatory exams taken--have a measurable impact on the final score a student achieves on a major college entrance exam. This scenario is perfectly suited for Multiple Linear Regression, as we have two or more independent variables influencing one dependent variable.

In this model, we designate **hours studied** and **prep exams taken** as the primary predictor variables, intending to estimate their individual and combined effects on the **exam score**, which serves as our dependent variable. After collecting data from a representative sample of students and running the analysis using the statistical tools available in Excel, we obtain a comprehensive regression summary output. This output contains all the necessary statistics, including coefficients, R-squared values, and, critically, the P-values needed for hypothesis testing.

The following visual representation captures the regression output generated by Excel for this specific educational model. This summary organizes the results into several key tables, but our immediate focus for hypothesis testing must be on the Analysis of Variance (ANOVA) section and the Coefficient table, where the P-values for the overall model and the individual variables are clearly displayed:

D	E	F	G	H	I	J	K
SUMMARY OUTPUT							
<i>Regression Statistics</i>							
Multiple R	0.857						
R Square	0.734						
Adjusted R Square	0.703						
Standard Error	5.366						
Observations	20						
ANOVA							
	<i>df</i>	<i>SS</i>	<i>MS</i>	<i>F</i>	<i>Significance F</i>		
Regression	2	1350.76	675.38	23.46	0.00		
Residual	17	489.44	28.79				
Total	19	1840.20					
	<i>Coefficients</i>	<i>Standard Error</i>	<i>t Stat</i>	<i>P-value</i>	<i>Lower 95%</i>	<i>Upper 95%</i>	
Intercept	67.67	2.82	24.03	0.00	61.73	73.61	
hours	5.56	0.90	6.18	0.00	3.66	7.45	
prep_exams	-0.60	0.91	-0.66	0.52	-2.53	1.33	

Upon reviewing the output structure, we identify three distinct P-values that require careful scrutiny, each addressing a unique statistical question:

The P-value pertaining to the overall model fit (found in the ANOVA table).

The P-value associated with the first predictor variable: hours studied.

The P-value associated with the second predictor variable: prep exams taken.

Understanding the decision rule for each of these values is essential: if $P \leq 0.05$, the result is statistically significant; if $P > 0.05$, the result is not statistically significant at the 5% alpha level. We will now proceed to interpret each of these critical P-values individually.

Interpreting the Overall Model's P-value (Significance F)

The first P-value to assess in the regression output relates to the overall significance of the model. This value is derived from the F-test, which is presented under the column labeled **Significance F**, typically located within the ANOVA summary table of the Excel output. The overall model test challenges the null hypothesis that all regression coefficients (excluding the intercept) are simultaneously equal to zero. Rejecting this null hypothesis means that at least one of the predictor variables in the model significantly contributes to explaining the variation in the response variable.

In the provided example screenshot, we observe that the P-value for the overall model, labeled **Significance F**, is reported as **0.00** (which conventionally represents a value extremely close to zero, often < 0.0001). Since this calculated P-value (0.00) is far less than our established alpha level of 0.05, we possess overwhelming evidence to reject the null hypothesis of the overall model. This leads to the robust conclusion that the regression model as a whole is statistically significant.

Stated in practical terms specific to our scenario, this significant F-test result indicates that the combination of "hours studied" and "prep exams taken" collectively has a statistically significant relationship with the final exam score. The model, as formulated, explains significantly more variance in the exam scores than would be expected if all predictors had zero effect. This confirms that the model is a worthwhile starting point for further analysis and interpretation of the individual predictor effects.

Assessing the Individual Predictor Variable P-values

While the overall model's significance is important, researchers are often more interested in determining which specific variables are driving the observed effects. This is where the individual P-values for each coefficient, found in the Coefficients table of the output, become critical. These P-values are the result of separate t-tests, each testing the null hypothesis that the individual coefficient for that predictor variable is zero, holding all other variables constant.

Let us first examine the P-value associated with the first predictor variable, **hours studied**. Although the exact numerical value is not explicitly shown next to the variable name in the condensed output, if we assume, based on the original analysis context, that this value is less than 0.05, the interpretation is straightforward. Since this P-value falls below the 0.05 threshold, we reject the null hypothesis for this coefficient. We can confidently conclude that the number of hours studied is a statistically significant predictor of the final exam score. In essence, students who study more hours are predicted to achieve higher scores, independent of the number of prep

exams they take.

This determination of statistical significance confirms that the relationship between study hours and exam performance is unlikely to be due merely to random sampling noise. Therefore, policy recommendations or pedagogical adjustments based on increasing study hours would be supported by the statistical evidence derived from this regression coefficient.

Interpreting the P-value for Prep Exams Taken

Next, we turn our attention to the P-value corresponding to the second predictor variable, **prep exams taken**. The output provides the P-value for this variable as **0.52**. This value is considerably higher than the standard alpha level of 0.05. When the P-value (0.52) is greater than α , we must fail to reject the null hypothesis for this specific variable.

The conclusion is that the number of prep exams taken, when considered alongside the hours studied, is not statistically significant. This lack of significance implies that, within the context of this specific multiple regression model, there is insufficient statistical evidence to claim a relationship between the number of prep exams taken and the final exam score that is distinguishable from random chance. In practical terms, the marginal benefit of taking additional prep exams, after accounting for the study hours, does not yield a statistically reliable improvement in the final score.

It is important to understand the nuance here: this does not necessarily mean prep exams are useless entirely, but rather that their unique contribution to predicting the exam score is weak or non-existent in this sample, especially once the effect of study hours is already accounted for. This scenario often arises when two predictor variables are highly correlated (a condition known as multicollinearity), or when one variable simply possesses much stronger predictive power than the other.

The Importance of Model Refinement and Variable Selection

The finding that the variable "prep exams taken" is not statistically significant ($P = 0.52$) has immediate implications for model refinement. Since this variable does not add any statistically reliable improvement to the overall predictive power of the model, analysts may decide to remove it entirely. Retaining non-significant variables can needlessly complicate the model, reduce its parsimony, and potentially inflate the standard errors of the remaining coefficients, making interpretation more challenging.

The standard procedure following the identification of non-significant predictors is to perform a model respecification. In this case, removing the non-significant variable would lead us to run a simpler model: a simple linear regression using only **hours studied** as the sole predictor variable. This refined model, based exclusively on the statistically significant predictors, is often more

robust, easier to interpret, and potentially more accurate when applied to new, unseen data, adhering to the principle of using the simplest model that adequately explains the phenomenon.

This iterative process of model testing and refinement, guided by the P-values, is central to rigorous hypothesis testing and effective statistical modeling. By systematically eliminating variables that do not meet the criteria for significance, we ensure that the final model is both statistically sound and relevant to the research questions being addressed.

ARABPSYCHOLOGY.COM