

Lack of Fit Test in R: A Step-by-Step Guide to Model Evaluation

Authored by
Mohammed looti

November 5, 2025

RECOMMENDED CITATION

Mohammed looti (2025). *Lack of Fit Test in R: A Step-by-Step Guide to Model Evaluation*. PSYCHOLOGICAL STATISTICS. Retrieved from <https://statistics.arabpsychology.com/?p=10871>

The [lack of fit test](#) is an essential statistical tool within [regression analysis](#), specifically designed to assess the adequacy of a proposed statistical model. Its core function is to rigorously evaluate whether the structural form of the model--such as assuming linearity versus curvilinearity--is appropriate for describing the observed data. A successful analysis hinges on choosing a model that is **complex enough** to capture the underlying trend but **simple enough** to remain interpretable and avoid the pitfalls of overfitting.

This powerful test addresses a crucial question in model selection: Does introducing additional complexity into a statistical model significantly improve its explanatory power over a simpler, baseline version? By comparing two nested models--a **Full Model**, which incorporates the additional terms, and a **Reduced Model**, which is a structural subset of the full model--we determine if the variance explained by the extra terms justifies their inclusion. If the difference in the residual fit between the two models is statistically significant, we conclude that the simpler structure exhibits a "lack of fit," meaning it systematically fails to capture the relationship adequately.

To illustrate this concept, consider a common research problem: analyzing the relationship between the *number of hours a student spends studying* and their resulting *exam score*. While a basic assumption might be that this relationship is straight (linear), real-world data often suggests a more nuanced, non-linear pattern. Perhaps the benefit of studying diminishes after a certain point, or scores increase exponentially with sustained effort. Identifying this curvilinear pattern requires a more sophisticated model structure.

We propose the following two competing [regression models](#) to describe this relationship, ensuring the reduced model is structurally nested within the full model:

Full Model (Quadratic Structure): $\text{Score} = \beta_0 + B_1(\text{hours}) + B_2(\text{hours})^2$

Reduced Model (Linear Structure): $\text{Score} = \beta_0 + B_1(\text{hours})$

The following comprehensive tutorial provides a step-by-step guide on how to execute a formal lack of fit test using the powerful and flexible [R statistical environment](#). Our objective is to objectively determine if the quadratic term ($B_2(\text{hours})^2$) in the full model contributes a statistically significant improvement in fit compared to the baseline linear model, thus validating the need for a non-linear approach.

Step 1: Establishing the Dataset and Initial Visualization

The foundation of any robust statistical investigation is the quality and structure of the underlying data. For this demonstration, we will generate a synthetic dataset in R that closely simulates the relationship between study hours and exam performance for 50 hypothetical college students.

Utilizing a synthetic approach offers two key advantages: it ensures that the example is perfectly reproducible by any user running the exact same code, and it allows us to intentionally design a subtle non-linear relationship to clearly highlight the necessity and utility of the lack of fit test.

We begin by setting a seed, which guarantees the consistent generation of random numbers across different sessions. We then construct the data frame, ensuring the underlying mathematical relationship incorporates a non-linear component (a cubic term) along with random noise. This intentional design makes the quadratic model (our full model) a statistically better candidate than the simpler linear model, setting up an ideal scenario for validating the test procedure. The following sequence of R commands generates and previews the required dataset:

Ensure reproducibility of random number generation

```
set.seed(1)
```

```
# Create the synthetic dataset with 50 observations
```

```
df <- data.frame(hours = runif(50, 5, 15), score=50)
```

```
df$score = df$score + df$hours^3/150 + df$hours*runif(50, 1, 2)
```

```
# Display the initial six rows of the resulting data frame
```

```
head(df)
```

```
hours score
```

```
1 7.655087 64.30191
```

```
2 8.721239 70.65430
```

```
3 10.728534 73.66114
```

```
4 14.082078 86.14630
```

```
5 7.016819 59.81595
```

```
6 13.983897 83.60510
```

Once the data frame is generated, the critical next step is the visualization of the relationship between the independent variable (hours studied) and the dependent variable (exam score). A simple [scatterplot](#) provides invaluable intuitive evidence regarding the appropriate model form. If the data points clearly deviate from a straight line, it strongly supports the necessity of testing a non-linear model. We employ the robust `ggplot2` package in R for generating a clear graphical representation, which helps us anticipate the statistical outcome.

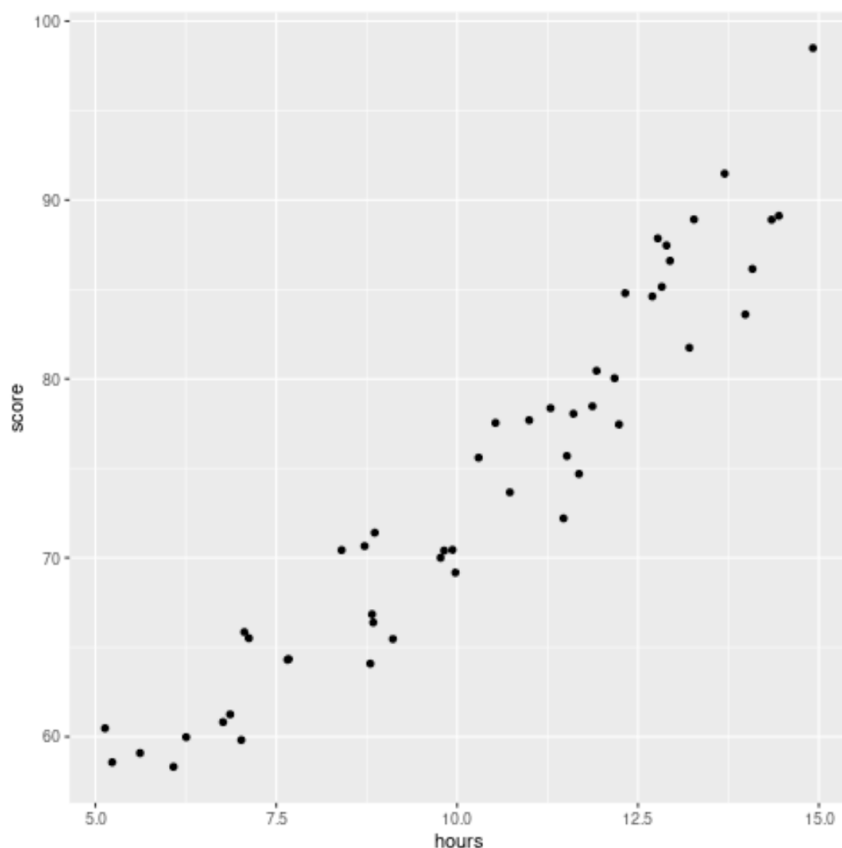
Load the specialized ggplot2 visualization package

```
library(ggplot2)
```

```
# Generate the scatterplot of hours vs. score
```

```
ggplot(df, aes(x=hours, y=score)) +
```

```
geom_point()
```



A careful inspection of the generated scatterplot reveals that the arrangement of the data points is indeed not perfectly linear. There is a discernible upward curve, suggesting that the incremental gain in score associated with each additional hour of study is increasing, or at least not constant, across the range of study times. If we were to apply only a simple linear model, it would systematically misestimate the scores, particularly underestimating performance at the higher end of the study hours range. This visual conformation serves as the initial, compelling justification for proceeding with the formal lack of fit analysis.

Step 2: Specification and Fitting of Competing Regression Models

Following the exploratory data analysis, the next technical phase involves fitting the two necessary [statistical models](#) to the dataset. The fundamental requirement for conducting a valid lack of fit test is that the models must be strictly **nested**. This nesting means that the reduced model must be mathematically equivalent to the full model when one or more coefficients in the full model are constrained (typically set to zero). In our example, the reduced linear model is perfectly nested within the full quadratic model by setting the quadratic coefficient (B_2) to zero.

We utilize R's core function for linear modeling, `lm()`, to estimate the parameters for both structures. For the full model, we employ the `poly(hours, 2)` function, which fits a second-degree [polynomial regression](#). This function is particularly useful as it generates orthogonal polynomial terms, which can improve numerical stability and simplify interpretation compared to manually calculating the squared terms.

Fit the Full Model: Includes both linear and quadratic terms

```
full <- lm(score ~ poly(hours,2), data=df)
```

Fit the Reduced Model: Includes only the linear term

```
reduced <- lm(score ~ hours, data=df)
```

The resulting `full` model object now contains the parameter estimates capturing both the linear and quadratic effects, providing the flexibility needed to model the observed curvature. Conversely, the `reduced` model assumes a strictly linear relationship, serving as the baseline structure against which the additional complexity of the full model is rigorously judged. The subsequent step will use the differences in the models' residual sums of squares (RSS) to calculate the test statistic.

Step 3: Executing the Formal Lack of Fit Test via ANOVA

The definitive step in this procedure is the formal statistical comparison of the two nested models using the standard [Analysis of Variance \(ANOVA\)](#) framework. The R function `anova()` is specifically designed to efficiently compare the fit of nested linear models by partitioning the total variation in the response variable.

The lack of fit test is fundamentally framed around comparing the model error under two competing assumptions, defined by the following statistical hypotheses:

Null Hypothesis (H₀): The reduced model is statistically sufficient, implying that the additional terms (the quadratic term) in the full model contribute no significant explanatory power. We hypothesize that the added coefficient β_2 is zero.

Alternative Hypothesis (H_A): The full model provides a significantly better fit to the data than the reduced model, meaning the added complexity accounts for substantial unexplained variation. We hypothesize that the added coefficient β_2 is not zero.

We execute the test by providing both fitted model objects to the `anova()` function. For this comparison, the order in which the models are specified is important, as the function calculates the reduction in residual variance achieved by moving from the simpler model to the more complex one.

```
# Perform the lack of fit test comparing the full and reduced models
anova(full, reduced)
```

Analysis of Variance Table

Model 1: score ~ poly(hours, 2)

Model 2: score ~ hours

Res.Df RSS Df Sum of Sq F Pr(>F)

1 47 368.48

2 48 451.22 -1 -82.744 10.554 0.002144 **

Signif. codes: 0 '***' 0.001 '**' 0.01 '*' 0.05 '.' 0.1 ' ' 1

The interpretation of this ANOVA table is critical for concluding the analysis:

Residual Degrees of Freedom (Res.Df): Model 1 (Full) has 47 degrees of freedom, and Model 2 (Reduced) has 48. The difference (Df = -1, or 1 degree of freedom added) precisely corresponds to the single parameter added in Model 1--the quadratic term--which is the focus of the test.

Residual Sum of Squares (RSS): This represents the amount of variation in the exam scores that remains unexplained by the model. Model 1 (Full) has a substantially smaller RSS (368.48) than Model 2 (Reduced, 451.22). The difference (Sum of Sq = 82.744) quantifies the specific amount of variance explained solely by the quadratic term.

F-Statistic and P-value: The resulting F-statistic, 10.554, measures the ratio of the variance explained by the added term to the residual variance. This F-value yields a corresponding p-value of 0.002144. Comparing this p-value to the conventional significance threshold (typically $\alpha = 0.05$), we find that 0.002144 is highly significant.

Since the p-value is significantly smaller than 0.05, we decisively reject the [null hypothesis](#) (H_0). The statistical conclusion is clear: the full model, which includes the quadratic effect of study hours, provides a significantly better fit to the observed data than the simpler linear model. This mandates the selection of the more complex, appropriate quadratic structure for all subsequent predictive and inferential tasks.

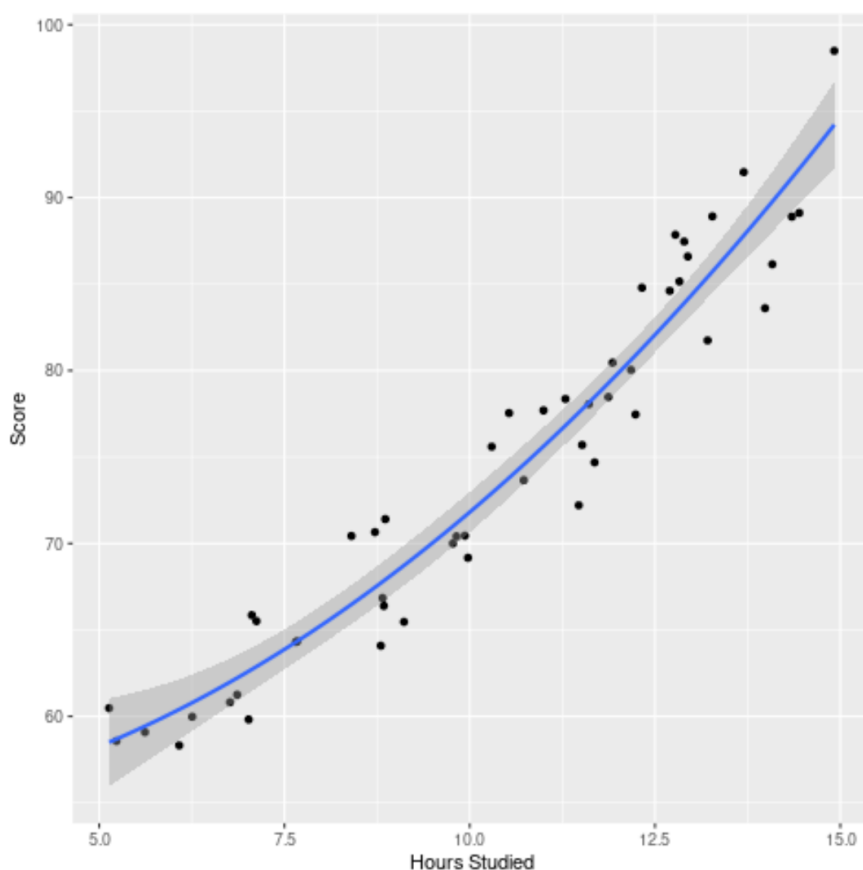
Step 4: Confirming the Model Fit Through Visualization

While the ANOVA test provides rigorous statistical proof, the final step involves visualizing the chosen model fit against the raw data points. This visual confirmation is essential for validating the statistical finding and ensuring that the selected model behaves logically and predictably across the entire range of the predictor variable. It serves as a necessary diagnostic check against any

potential misinterpretation of the numerical results.

We return to the `ggplot2` framework, enhancing our initial scatterplot by overlaying the fitted regression curve corresponding to the statistically superior quadratic model. We utilize the `stat_smooth()` function, explicitly defining the method as linear modeling (`method='lm'`) and ensuring the formula matches our selected second-degree polynomial structure (`formula = y ~ poly(x, 2)`).

```
ggplot(df, aes(x=hours, y=score)) +  
geom_point() +  
stat_smooth(method='lm', formula = y ~ poly(x,2), size = 1) +  
xlab('Hours Studied') +  
ylab('Score')
```



The visualization provides powerful confirmation of the statistical finding. The smooth, curved regression line generated by the quadratic model adheres closely to the central tendency of the data points across the entire range of hours studied. Crucially, it captures the upward sweep that the linear model would have failed to accommodate. If we had mistakenly chosen the reduced

model, the resulting straight line would have led to systematically biased residuals and poor predictive accuracy. The lack of fit test provided the objective, statistical evidence needed to select this better-fitting, more appropriate model structure.

Additional Resources for Advanced Model Diagnostics

Mastering the lack of fit test is a fundamental step toward robust [regression analysis](#). Researchers frequently need to explore various functional forms and diagnostic checks to ensure the validity and accuracy of their statistical models. To further enhance your proficiency in fitting and evaluating complex models in R, we strongly recommend exploring the following related resources and techniques:

[How to Perform Polynomial Regression in R](#)

Utilizing rigorous methods like the lack of fit test ensures that your chosen statistical model accurately reflects the underlying scientific relationship rather than imposing an overly simplistic or convenient structure. Understanding when and how to reject a reduced model is key to producing reliable statistical findings and sound scientific conclusions.