

Learning Linear Regression in R: Verifying Key Assumptions for Accurate Modeling

Authored by
Mohammed looti

November 12, 2025

RECOMMENDED CITATION

Mohammed looti (2025). *Learning Linear Regression in R: Verifying Key Assumptions for Accurate Modeling*. PSYCHOLOGICAL STATISTICS. Retrieved from <https://statistics.arabpsychology.com/?p=23856>

The process of [Linear Regression](#) is a foundational [statistical method](#) used widely across fields like economics, social sciences, and engineering. Its primary goal is to model the relationship between a response variable (Y) and one or more explanatory variables (X). Specifically, it seeks to fit a straight line that minimizes the sum of squared differences between the observed data points and the values predicted by the model. When utilized correctly, it provides powerful insights into how changes in predictor variables influence the outcome variable. However, the reliability of the inferences drawn from a linear regression model--such as confidence intervals and p-values--is entirely dependent upon meeting a specific set of critical statistical [assumptions](#).

Before attempting to interpret the coefficients or draw conclusions about variable significance, data analysts must meticulously verify that the underlying assumptions of the Ordinary Least Squares (OLS) method have not been violated. Failure to validate these conditions can lead to biased estimates, incorrect standard errors, and ultimately, flawed conclusions that misrepresent the true relationship between the variables. This comprehensive guide details the four primary assumptions inherent in linear regression and provides practical steps for checking each of them within the [R statistical computing environment](#).

The four fundamental assumptions that must be satisfied for a robust and reliable linear regression model are:

Linear Relationship: The relationship between the independent variable(s) and the dependent variable must be approximately linear.

Independence of Residuals: The errors (or [residuals](#)) of the model must be independent of each other. This is particularly critical in time series data where autocorrelation is common.

Homoscedasticity: The variance of the [residuals](#) must remain constant across all levels of the predictor variable(s).

Normality of Residuals: The [residuals](#) of the model must be approximately normally distributed.

If any of these fundamental assumptions are violated, the standard errors calculated by the model will be inaccurate. This means that hypothesis tests (like t-tests) and confidence intervals will be unreliable, potentially leading to incorrect decisions regarding the significance of the predictor variables. Throughout this tutorial, we will use R's built-in **mtcars** dataset as an illustrative example, focusing on modeling the relationship between fuel efficiency (**mpg**, the response variable) and engine displacement (**disp**, the predictor variable).

Evaluating Assumption 1: Verifying the Linear Relationship

The first and most intuitive assumption of linear regression dictates that a straight-line relationship must exist between the independent variable (X) and the dependent variable (Y). If the true relationship is non-linear (e.g., quadratic or exponential), fitting a linear model will inevitably result in a poor fit and highly unreliable predictions. Verifying this assumption visually is often the most

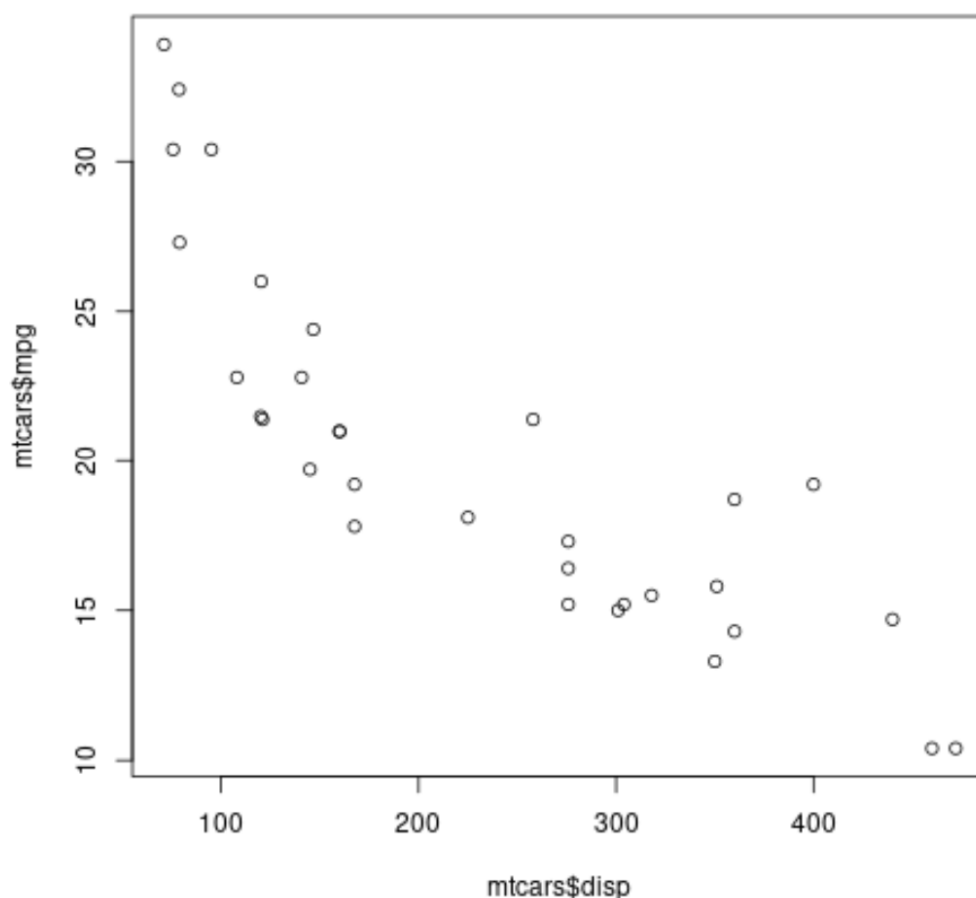
straightforward approach, providing immediate diagnostic feedback on the suitability of the linear model.

The simplest and most effective way to detect whether this assumption is met is by generating a [scatter plot](#) of the independent variable against the dependent variable. In R, this involves plotting the predictor variable on the x-axis and the response variable on the y-axis. If the resulting cloud of data points generally follows a straight line, the assumption of linearity is likely satisfied. Conversely, if the points exhibit a clear curvature (U-shape, S-shape, or exponential curve), transformation of the variables or the use of a non-linear model may be required.

To create the necessary scatter plot for our example using the **mtcars** dataset, we use the base R plotting function. We plot displacement (`disp`) against miles per gallon (`mpg`) to visually inspect their relationship:

```
#create scatter plot of disp vs. mpg  
plot(mtcars$disp, mtcars$mpg)
```

Executing this code produces the raw visualization of the data points, allowing for direct assessment of linearity, as shown in the output graph below.



Upon visual inspection of the scatter plot, we observe that the data points tracing the relationship between engine displacement and fuel efficiency fall approximately along a straight, descending line. While there is some natural variation, there is no pronounced curvature that would suggest a significant violation. Therefore, for this specific relationship, we can reasonably conclude that the **linear relationship** assumption is met, and we can proceed to fit the linear model confidently.

Evaluating Assumption 2: Ensuring Independent Residuals

The second critical assumption of linear regression requires that the error terms, or [residuals](#), are statistically independent of one another. Independence means that the error associated with one observation does not influence the error associated with any other observation. Violation of this assumption, known as [autocorrelation](#) or serial correlation, is most often encountered when dealing with time series data (where an error at time T is correlated with an error at time $T-1$) or spatial data. When residuals are correlated, the standard errors of the coefficients are typically underestimated, leading to inflated t-statistics and an increased likelihood of Type I errors (false positives).

While simple cross-sectional data (like the **mtcars** dataset used here) often meets this assumption

automatically, it is crucial to check in sequential data. The most common formal test used to check for serial correlation is the [Durbin-Watson statistic](#). This test measures the correlation between adjacent errors. A value close to 2 suggests no serial correlation; values significantly lower than 2 indicate positive autocorrelation, and values significantly higher than 2 suggest negative autocorrelation.

Alternatively, particularly when analyzing time series data, an Autocorrelation Function (ACF) plot can provide a visual diagnostic. Ideally, most of the residual autocorrelations should fall within the 95% confidence bands around zero. These bands are typically calculated at approximately ± 2 divided by the square root of n , where n is the sample size. If many spikes in the ACF plot extend outside these bands, it suggests that the independence assumption has been violated, necessitating the use of time series models (like ARIMA) or specialized estimation techniques.

Evaluating Assumption 3: Assessing Homoscedasticity

The third core assumption, known as [Homoscedasticity](#) (meaning "same variance"), requires that the variance of the [residuals](#) remains constant across all predicted values of the dependent variable. In simpler terms, the spread of the errors should be uniform across the entire range of the fitted values. When this condition is not met--meaning the variance of the errors increases or decreases as the predictor variable changes--the residuals are said to suffer from [heteroscedasticity](#).

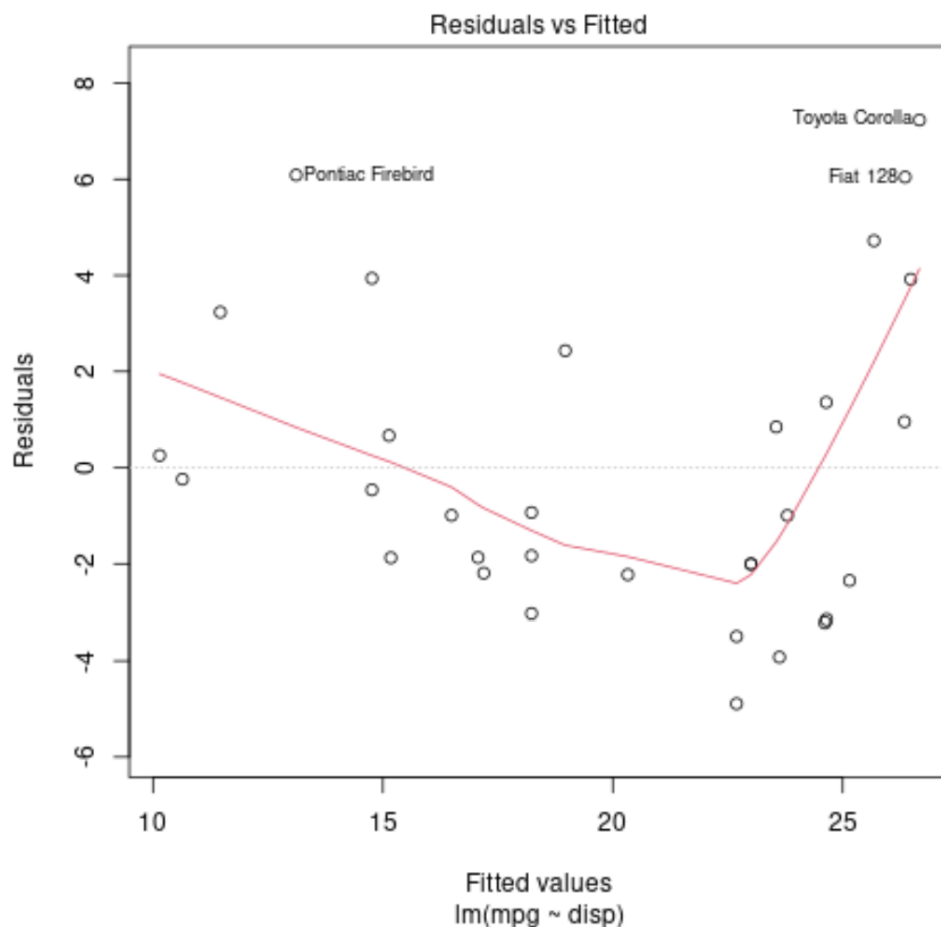
The presence of [heteroscedasticity](#) does not bias the coefficient estimates (they remain consistent), but it severely biases the standard errors. Similar to autocorrelation, this leads to incorrect hypothesis testing. The standard errors tend to be underestimated in regions where the variance is high, making variables appear more significant than they truly are. The most reliable way to detect [heteroscedasticity](#) is by creating a plot of the fitted values against the standardized [residuals](#).

In R, once the linear model is fitted, the built-in plotting functions can generate this diagnostic plot quickly. We first fit the model and then use the ``plot()`` function, specifying index 1, which corresponds to the "Residuals vs Fitted" plot:

```
#fit regression model
model <- lm(mpg ~ disp, data=mtcars)

#create fitted values vs residuals plot
plot(model, 1)
```

This command executes the fitting and plotting sequence, resulting in the diagnostic plot shown below, which maps the predicted values (x-axis) against the corresponding residuals (y-axis).



For the assumption of [Homoscedasticity](#) to hold, the spread of the [residuals](#) on the y-axis should appear randomly scattered around the horizontal line at zero, with no noticeable pattern or fan shape. In the plot above, while the residuals appear to fan out slightly as the fitted values increase, the deviation is not severe enough to indicate a strong presence of concerning heteroscedasticity. If a clear cone or fan shape were present, suggesting non-constant variance, techniques such as Weighted Least Squares (WLS) or robust standard errors would be necessary to correct for the violation.

Evaluating Assumption 4: Checking for Residual Normality

The final primary assumption is that the [residuals](#) of the linear regression model must be approximately normally distributed. It is important to clarify that it is the distribution of the errors, not the outcome variable itself, that needs to be normal. While the OLS coefficient estimates themselves (betas) remain unbiased even if the normality assumption is violated, the normality condition is necessary for the validity of the p-values and confidence intervals, particularly in smaller samples. If the residuals are not normally distributed, the statistical tests may lack power or produce misleading significance levels.

The most common and effective visual method for checking this assumption is to create a Normal Quantile-Quantile Plot, commonly referred to as a [Q-Q plot](#). This plot compares the observed standardized residuals against the theoretical quantiles of a standard normal distribution. If the residuals are normally distributed, the points on the plot should fall roughly along a straight diagonal line. Deviations from this line, particularly at the tails, suggest departures from normality, such as skewness or kurtosis.

To generate the [Q-Q plot](#) for the fitted model in R, we use the `plot()` function again, specifying index 2 this time, which calls the Normal Q-Q diagnostic plot:

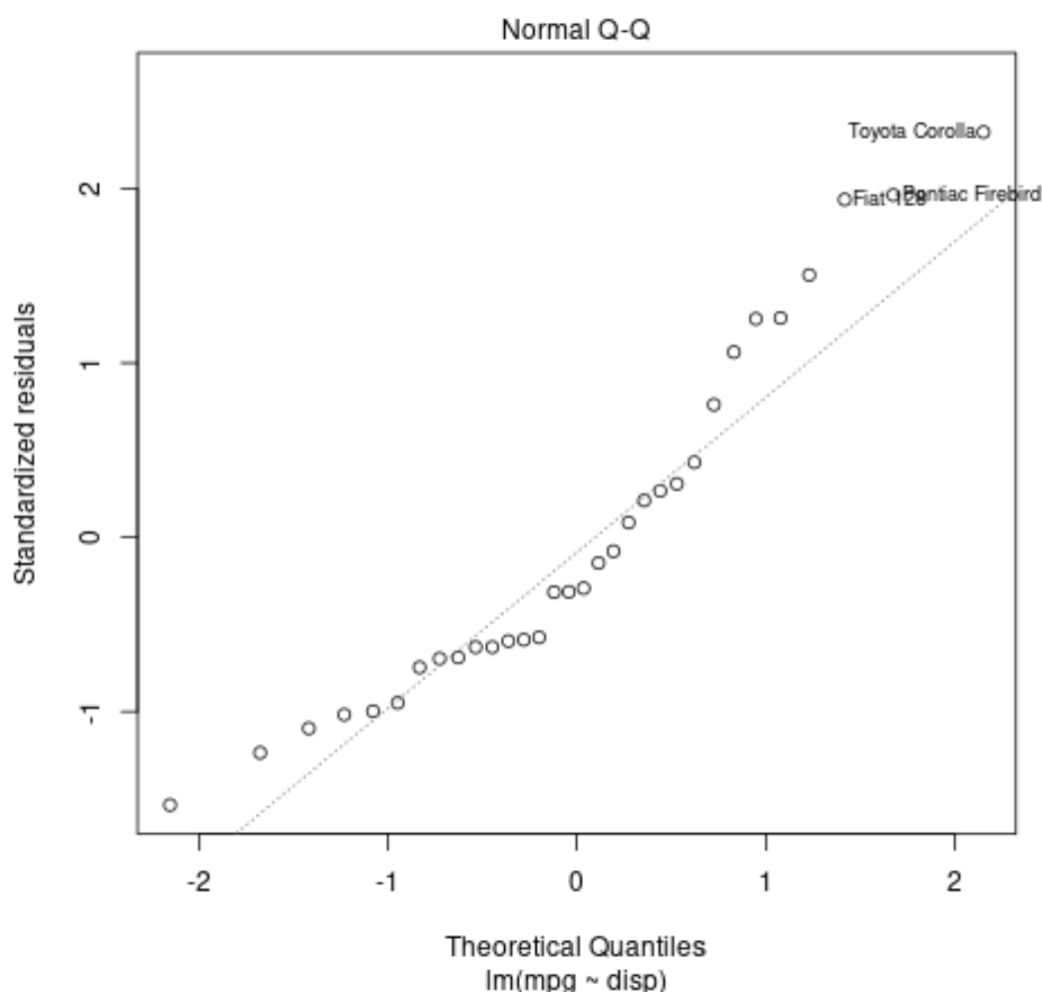
#fit regression model

model <- lm(mpg ~ disp, data=mtcars)

#create Q-Q plot

plot(model, 2)

Running the necessary code produces the following visualization, allowing us to assess the alignment of the standardized residuals against the theoretical normal distribution line.



Examining the [Q-Q plot](#) for our `mpg` vs. `disp` model, we observe that the vast majority of the points cluster tightly around the straight diagonal line. There are minor deviations, particularly noticeable at the extreme upper and lower tails, but these are generally considered minor and often acceptable, especially in smaller sample sizes. Since the overall pattern closely follows the diagonal line, we would assume that the normality assumption is not severely violated, validating the use of standard inference procedures for this model. For more severe violations, transformation of the outcome variable or the use of non-parametric methods might be considered.

Conclusion and Additional Resources

Successfully fitting a linear regression model requires more than just running the `lm()` function in R; it demands a thorough diagnostic process to ensure the model adheres to its underlying statistical [assumptions](#). By systematically checking for linearity, independence of residuals, [homoscedasticity](#), and normality of errors, practitioners can ensure that their coefficient estimates and statistical inferences are both unbiased and reliable. The diagnostic plots provided by R, such

as the scatter plot, the residuals vs. fitted plot, and the [Q-Q plot](#), are essential tools in this validation process, providing clear visual evidence of potential model shortcomings.

The process demonstrated here, using the well-known **mtcars** dataset, confirms that a basic linear relationship between engine displacement and fuel efficiency is appropriate and that the resulting model meets the key requirements for statistical inference. Understanding how to diagnose and address assumption violations is crucial for any analyst seeking to build robust and trustworthy models based on the principles of Ordinary Least Squares regression.

The following tutorials explain how to perform other common tasks in R: