

Learning Simple Linear Regression with R: A Step-by-Step Guide

Authored by
Mohammed loot

November 6, 2025

RECOMMENDED CITATION

Mohammed loot (2025). *Learning Simple Linear Regression with R: A Step-by-Step Guide*. PSYCHOLOGICAL STATISTICS. Retrieved from <https://statistics.arabpsychology.com/?p=11921>

[Simple linear regression](#) (SLR) is a foundational [statistical modeling](#) technique used primarily to investigate and quantify the linear relationship between two continuous variables: a single [explanatory variable](#) (or predictor) and a corresponding response variable (or outcome). Mastering this technique is essential for data analysts seeking to understand how variations in one factor influence another.

The core principle of simple linear regression involves identifying and fitting a straight line that minimizes the total squared distance to all observed data points. This resulting line, which best summarizes the data trend, is formalized by the classic linear equation:

$$\hat{y} = b_0 + b_1x$$

Understanding the interpretation of the coefficients in this equation is critical for accurate reporting of the model results:

$\hat{y}$: Represents the estimated or predicted value of the [response variable](#) based on the model.

b_0 : Denotes the **intercept**, which is the baseline predicted value of $\hat{y}$ when the explanatory variable (x) is zero.

b_1 : Represents the **slope** of the regression line, quantifying the magnitude and direction of the relationship--specifically, the average change in $\hat{y}$ associated with a one-unit increase in x .

If the model demonstrates statistical significance, this equation becomes a powerful tool, enabling researchers not only to describe the association between variables but also to generate reliable [predictions](#). This comprehensive tutorial provides a rigorous, step-by-step guide on how to conduct a robust simple linear regression analysis using the powerful statistical programming environment, [R](#).

Step 1: Data Preparation and Importing the Sample Dataset into R

To demonstrate the practical application of the simple linear regression procedure, we will first establish a sample dataset. This simulated data set captures observations for 15 individuals, tracking two crucial variables: the total number of **hours studied** for an examination and the corresponding **exam score** achieved.

In this context, we hypothesize a link where the time spent studying directly influences the resulting score. Consequently, *hours studied* functions as our [explanatory variable](#) (x), and *exam score* serves as the response variable (y). The following [R](#) code illustrates the necessary commands to define and load this data structure as a data frame object, which is the standard format for statistical analysis in R:

Create the sample data frame containing hours studied and exam scores

```
df <- data.frame(hours=c(1, 2, 4, 5, 5, 6, 6, 7, 8, 10, 11, 11, 12, 12, 14),
```

```
score=c(64, 66, 76, 73, 74, 81, 83, 82, 80, 88, 84, 82, 91, 93, 89))
```

```
# Display the first six observations to verify data loading
```

```
head(df)
```

```
hours score
```

```
1 1 64
```

```
2 2 66
```

```
3 4 76
```

```
4 5 73
```

```
5 5 74
```

```
6 6 81
```

```
# Attach the dataset to make variables globally accessible for convenience
```

```
attach(df)
```

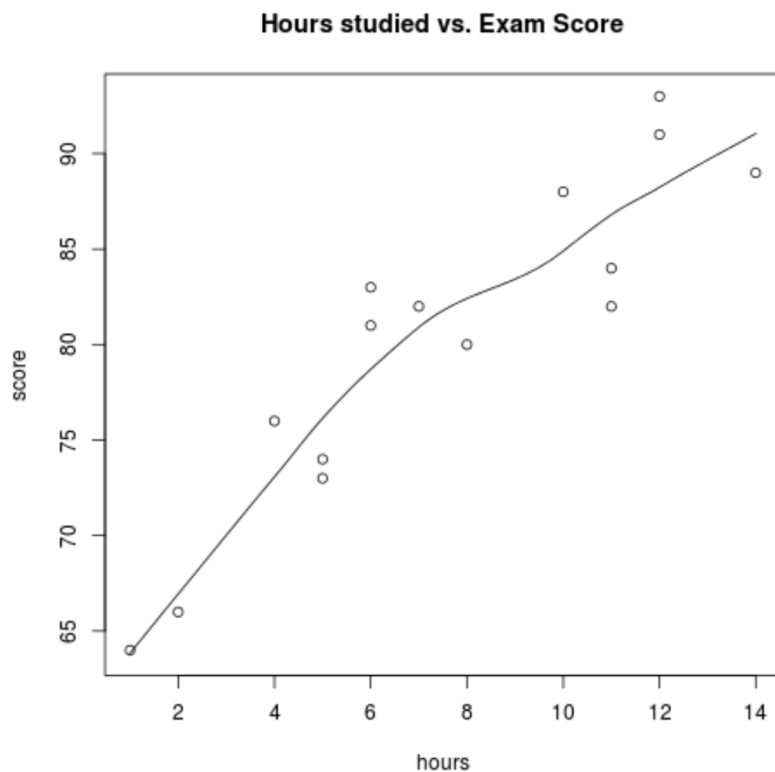
Ensuring the data is correctly structured and loaded is a non-negotiable prerequisite for any statistical analysis. The use of the `attach()` function, while sometimes cautioned against in large projects, simplifies subsequent code in this tutorial by allowing direct reference to the variable names (like `hours` and `score`) without needing to specify the data frame prefix (`df$`) repeatedly.

Step 2: Exploratory Data Analysis and Initial Assumption Checks

Before fitting any statistical model, it is mandatory to perform exploratory data analysis (EDA). This step is crucial for verifying that the data aligns with the fundamental assumptions underlying simple linear regression, most importantly the assumption of linearity. If the relationship between the predictor and response is not linear, the model results will be unreliable and potentially misleading.

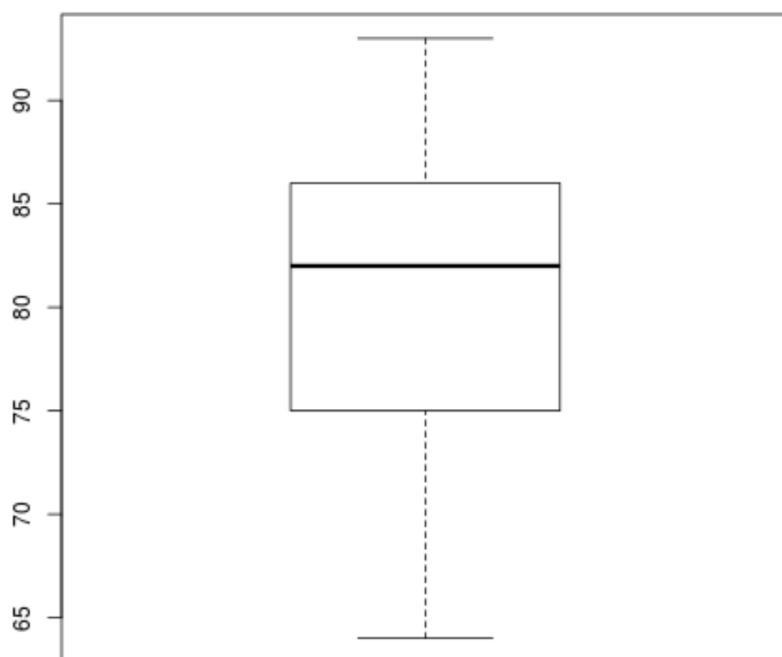
We begin by generating a [scatterplot](#) to visually inspect the association between our two variables. This visualization plots *hours studied* on the x-axis against the *exam score* on the y-axis. A clear, roughly straight-line pattern is necessary to justify the use of a linear model. We utilize the `scatter.smooth()` function in [R](#), which automatically superimposes a smoothed line to better aid in visualizing the underlying trend:

```
scatter.smooth(hours, score, main='Hours Studied vs. Exam Score')
```



The resulting [scatterplot](#) demonstrates a discernible, positive linear trend: as study time increases, exam performance tends to rise consistently along a straight path. This robust visual evidence provides strong preliminary support for proceeding with a linear regression model. Furthermore, we must assess the potential influence of [outliers](#)--extreme observations that can disproportionately skew the resulting regression line and violate assumptions. We utilize a **boxplot** of the response variable (score) to detect any such anomalies.

boxplot(score)



Since the boxplot displays no individual data points (typically represented by small circles) outside the whiskers, we confidently conclude that our dataset is free from significant [outliers](#) in the response variable, allowing us to proceed to the formal model estimation without the need for data transformation or removal.

Step 3: Estimating and Interpreting the Regression Coefficients

With the assumption of linearity visually confirmed and the absence of influential outliers established, we can now fit the formal simple linear regression model. In the [R](#) environment, the primary function for this task is `lm()` (linear model). We specify the model using formula syntax, treating *score* as the dependent variable predicted by *hours*:

Fit the simple linear regression model using score as the response and hours as the predictor

```
model <- lm(score~hours)
```

```
# View the detailed statistical summary of the fitted model  
summary(model)
```

Call:

```
lm(formula = score ~ hours)
```

Residuals:

```
Min 1Q Median 3Q Max
```

-5.140 -3.219 -1.193 2.816 5.772

Coefficients:

Estimate Std. Error t value Pr(>|t|)

(Intercept) 65.334 2.106 31.023 1.41e-13 ***

hours 1.982 0.248 7.995 2.25e-06 ***

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

Residual standard error: 3.641 on 13 degrees of freedom

Multiple R-squared: 0.831, Adjusted R-squared: 0.818

F-statistic: 63.91 on 1 and 13 DF, p-value: 2.253e-06

The pivotal output from the summary relates to the **Coefficients** table, which yields the estimated parameters for the regression line. These coefficients allow us to formulate the specific equation describing the relationship in our sample data:

Predicted Score = 65.334 + 1.982 × (Hours Studied)

This equation offers immediate insight into the model's findings. The **intercept** (65.334) suggests that a student who studies for zero hours is predicted to achieve a score of approximately 65.334 points. More importantly, the **slope coefficient** (1.982) indicates a positive, meaningful association: for every single additional hour a student dedicates to studying, their expected exam score increases by an average of 1.982 points, holding all other factors constant. The utility of this estimated linear relationship lies in its predictive power; for instance, a student studying for 10 hours is predicted to achieve a score of **85.154** points.

Step 4: Assessing Model Fit and Statistical Significance

A robust regression analysis requires evaluating not only the estimated coefficients but also the model's overall statistical validity and goodness-of-fit. The `summary()` output provides several key diagnostic metrics to fulfill this requirement, ensuring the model's conclusions are trustworthy.

We first examine the significance of the predictor variable using the **p-value** ($\text{Pr}(>|t|)$). For the *hours* variable, the p-value is 2.25e-06, which is profoundly lower than the conventional significance level ($\alpha = 0.05$). This outcome leads us to reject the null hypothesis that the slope is zero, confirming that there is a highly significant statistical association between the number of hours studied and the resulting exam score. This verifies that *hours* is a relevant and useful predictor.

Next, the **Multiple R-squared** statistic measures the proportion of variance in the response

variable (score) that is successfully explained by the [explanatory variable](#) (hours). A value of **0.831** (or 83.1%) suggests that 83.1% of the total variation observed in the exam scores can be accounted for by the variation in the number of hours studied. This figure strongly indicates a very good fit for our linear model, demonstrating substantial explanatory power.

Furthermore, the **Residual Standard Error (RSE)**, calculated as **3.641**, provides a tangible measure of the typical deviation, estimating the average distance between the actual observed scores and the scores predicted by the fitted line. Measured in points, a lower RSE suggests higher precision in the model's predictions. Finally, the overall model significance is assessed by the **F-statistic** (63.91) and its corresponding p-value (2.253e-06). Because this p-value is extremely small, we affirm the overall statistical significance of the regression, meaning the model is far superior to a random chance or null model.

Step 5: Validating Core Assumptions Through Residual Analysis

The final, non-optional step in linear regression analysis is thoroughly examining the model's [residuals](#) (the differences between observed and predicted values) to ensure that the method's underlying statistical assumptions have been met. Violations of these assumptions can invalidate the significance tests and coefficient interpretations performed in the previous steps.

Checking for Homoscedasticity: The Residual vs. Fitted Plot: We must first verify the assumption of [homoscedasticity](#), which requires that the variance of the errors remains constant across all levels of the predictor variable. This is assessed using the Residual vs. Fitted Values Plot. The horizontal axis displays the predicted (fitted) values, while the vertical axis displays the [residuals](#). The points should exhibit a purely random scatter around the horizontal zero line, showing no discernible patterns like a widening or narrowing funnel.

Define residuals from the fitted model

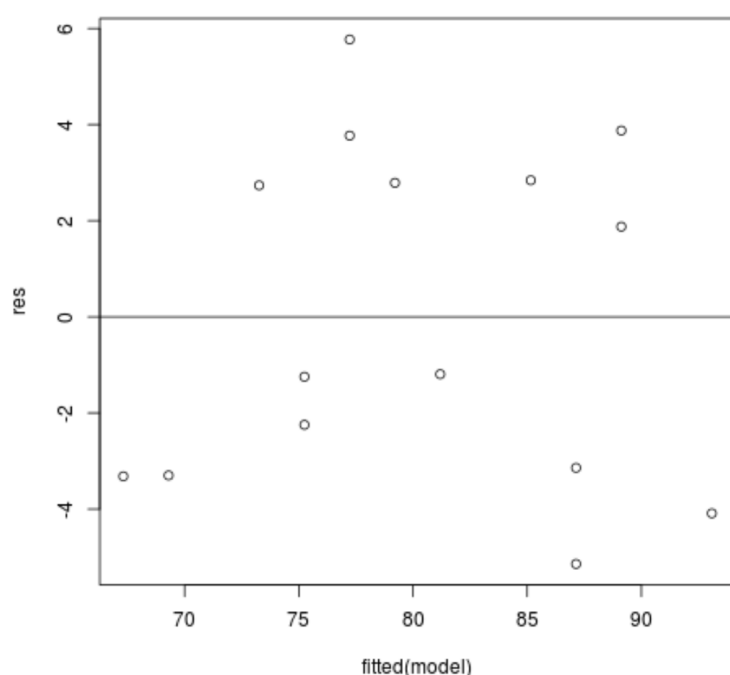
```
res <- resid(model)
```

Produce residual vs. fitted plot

```
plot(fitted(model), res)
```

Add a horizontal reference line at 0

```
abline(0,0)
```



In our plot, the scatter of the residuals appears appropriately random and centered around the zero line, confirming that the variance of the errors is constant across the range of fitted values. Thus, the crucial assumption of [homoscedasticity](#) is satisfied.

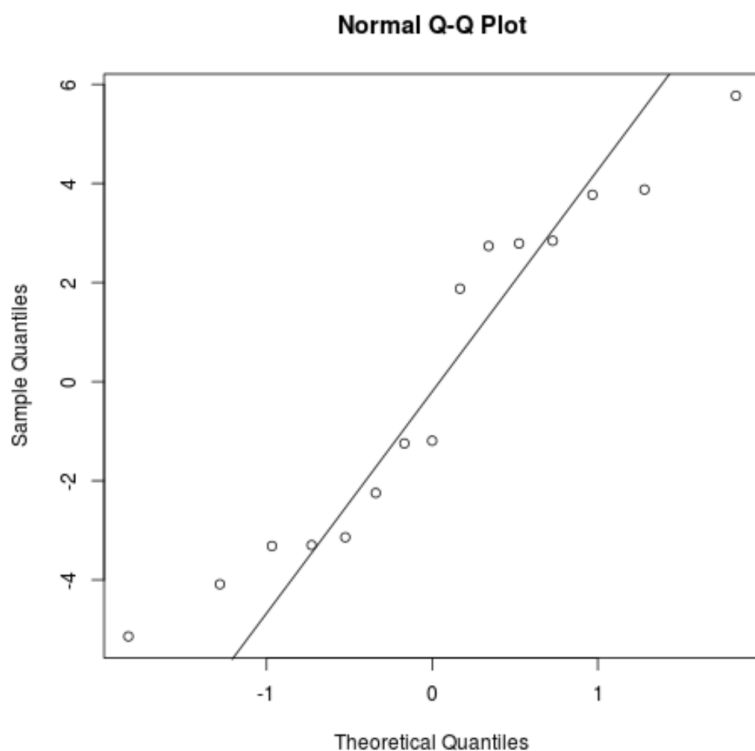
Verifying Normality: The Q-Q Plot: The second critical assumption requires that the [residuals](#) follow an approximate normal distribution. We verify this visually using a **Normal Quantile-Quantile (Q-Q) plot**. If the residuals are normally distributed, the plotted data points should align closely along the theoretical 45-degree diagonal line. Significant deviations, particularly at the tails, would suggest a violation of this assumption.

Create Q-Q plot for residuals

```
qqnorm(res)
```

Add a straight diagonal line to the plot for reference

```
qqline(res)
```

The residuals align closely with the 45-degree reference line, deviating only slightly at the extreme ends. This visual assessment confirms that the residuals are approximately normally distributed. Because we have successfully validated both the linearity and the underlying statistical assumptions of normality and [homoscedasticity](#), the results and interpretations derived from our simple linear regression model are considered statistically reliable and sound.

The complete [R](#) code utilized for all steps and commands demonstrated in this tutorial is publicly available [here](#) for replication and reference.