

Understanding Normality Tests in R: A Practical Guide to Four Methods

Authored by
Mohammed looti

November 2, 2025

RECOMMENDED CITATION

Mohammed looti (2025). *Understanding Normality Tests in R: A Practical Guide to Four Methods*. PSYCHOLOGICAL STATISTICS. Retrieved from <https://statistics.arabpsychology.com/?p=8501>

In the expansive realm of statistical analysis, the proper verification of underlying assumptions is paramount to generating trustworthy results. Many powerful [parametric tests](#), including the ubiquitous t-test and Analysis of Variance (ANOVA), operate under the fundamental premise that the data sample is drawn from a population that follows a [normal distribution](#). If this critical assumption is violated, the resulting p-values and confidence intervals can be misleading, potentially leading to incorrect conclusions regarding the rejection or acceptance of a hypothesis.

Therefore, any rigorous statistical workflow demands a systematic assessment of whether a dataset is truly [normally distributed](#) before proceeding with inferential testing. Fortunately, the **R** programming environment, a cornerstone of modern data science and statistical computing, offers a suite of highly effective and easily implementable tools for conducting this crucial diagnostic assessment.

This comprehensive tutorial delves into four distinct methods for evaluating the normality assumption within **R**. We will explore two intuitive visual techniques and two robust, formal statistical tests, providing the necessary code and interpretation guidelines for each approach.

Summary of Normality Testing Methods in R

Assessing the normality of a dataset typically involves combining visual inspection with objective statistical confirmation. The following four approaches provide a comprehensive strategy for verifying the distributional properties of your data using **R**:

(Visual Method) Create a histogram.

This method involves plotting the frequency distribution of the data. If the resulting visualization displays a symmetrical, mound-shaped curve, often referred to as a "bell-shaped" curve, the data distribution is visually assumed to be normal.

(Visual Method) Create a [Q-Q plot](#) (Quantile-Quantile plot).

This graphical technique compares the observed quantiles of the data against the theoretical quantiles of a normal distribution. If the plotted points closely follow the straight diagonal reference line, it strongly suggests that the data's quantiles align with the theoretical expectations, supporting the assumption of normality.

(Formal Statistical Test) Perform a [Shapiro-Wilk Test](#).

This is one of the most powerful statistical tests for normality, especially for smaller sample sizes. If the resulting [p-value](#) is greater than the chosen significance level (commonly $\alpha = 0.05$), we fail to reject the null hypothesis, indicating that the data is assumed to be normally distributed.

(Formal Statistical Test) Perform a [Kolmogorov-Smirnov Test](#).

This test is less sensitive than the Shapiro-Wilk test for assessing normality but remains a valid option, especially for very large datasets. Similar to the Shapiro-Wilk test, if the [p-value](#) exceeds the significance level ($\alpha = 0.05$), the data is considered consistent with a normal distribution.

The subsequent sections provide a deep dive into the implementation and interpretation of these four powerful methods, complete with practical code examples demonstrating their application within R.

Method 1: Assessing Normality Using a Histogram

The [histogram](#) stands as the most accessible and rapid visual tool for gaining an initial understanding of a dataset's distribution shape. By grouping observations into contiguous bins and displaying the frequency count for each bin, the histogram allows analysts to quickly identify key characteristics such as central tendency, symmetry, and the presence of skewness or multiple modes. A true normal distribution should yield a visualization that is perfectly symmetrical, centered around the mean, and exhibits the classic tapering tails of the bell curve.

While visual inspection via a histogram offers invaluable preliminary insights, it is inherently subjective. The appearance of the histogram can be influenced significantly by the choice of bin size and the overall sample size. For very small samples, the histogram may appear erratic even if the underlying population is normal. Conversely, for large samples, even minor deviations from perfect normality might be visually smoothed over. Therefore, while a symmetrical histogram suggests normality, this finding must always be corroborated by more objective methods, particularly formal statistical tests.

To demonstrate the utility of the histogram in R, we compare two simulated datasets: one generated to be perfectly [normally distributed](#) using the `rnorm()` function, and another designed to be non-normal (exponential distribution) using `rexp()`. The code below sets a seed for reproducibility and plots both distributions side-by-side using the base R plotting functions.

```
#make this example reproducible  
set.seed(0)
```

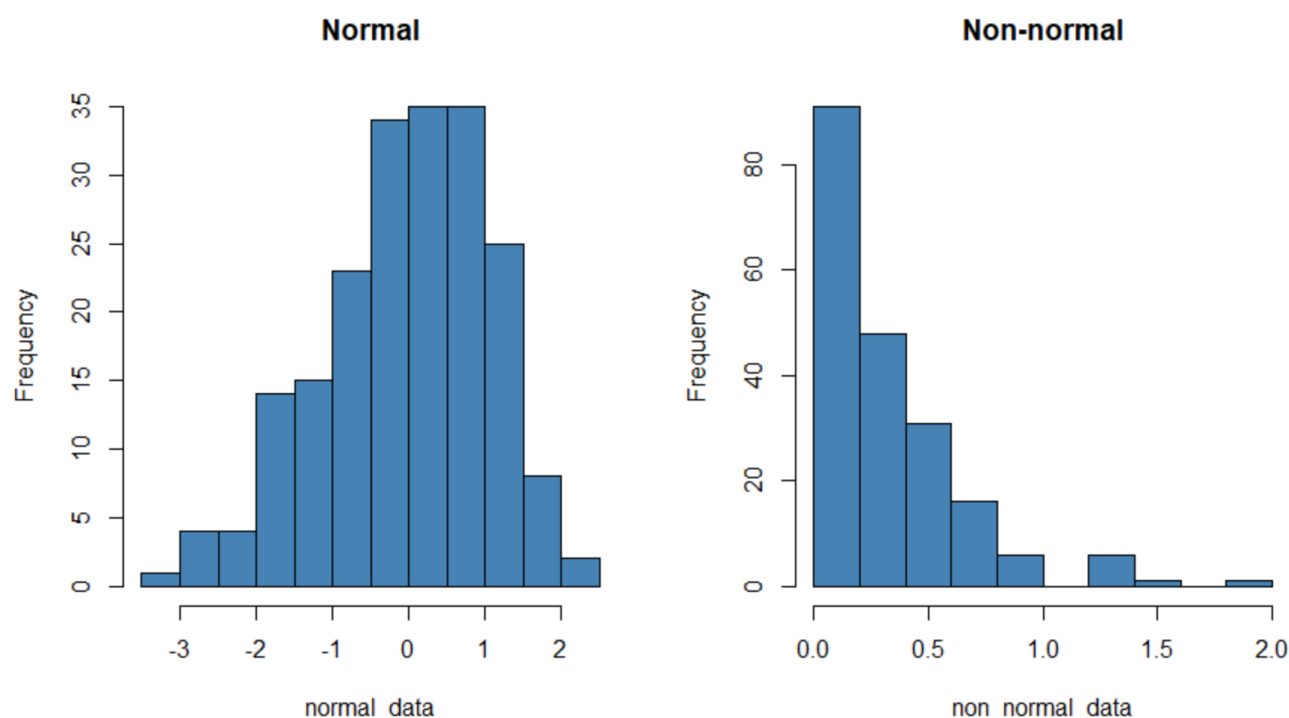
```
#create data that follows a normal distribution  
normal_data <- rnorm(200)
```

```
#create data that follows an exponential distribution  
non_normal_data <- rexp(200, rate=3)
```

```
#define plotting region
```

```
par(mfrow=c(1,2))

#create histogram for both datasets
hist(normal_data, col='steelblue', main='Normal')
hist(non_normal_data, col='steelblue', main='Non-normal')
```



The visual results are striking: the histogram on the left clearly exhibits the desired symmetry around the central peak and the familiar bell-shape, characteristic of a **normally distributed** dataset. In stark contrast, the histogram on the right displays severe right-skewness, where the majority of the data points are clustered near zero and a long tail extends towards positive values. This asymmetry immediately confirms the non-normal nature of the exponentially generated data.

Method 2: Utilizing the Quantile-Quantile (Q-Q) Plot

Moving beyond the general assessment provided by histograms, the [Q-Q plot](#) offers a statistically rigorous graphical method for checking normality. A Q-Q plot works by comparing the quantiles derived from your observed data against the theoretical quantiles expected if the data truly came from a specified distribution (in this case, the normal distribution). If the empirical distribution matches the theoretical distribution, the points generated by the plot should fall precisely along a straight, diagonal reference line.

The interpretation of a Q-Q plot focuses on deviations from this reference line, especially at the

extremes, or "tails," of the distribution. Any significant curvature, gaps, or points straying far from the line suggest a departure from normality. For instance, data that forms an S-shape often indicates a distribution with heavier or lighter tails (kurtosis issues) compared to the ideal normal curve. Similarly, a plot where points curve up or down consistently suggests skewness in the underlying distribution.

In **R**, generating a Q-Q plot is straightforward, utilizing the specialized functions `qqnorm()` and `qqline()`. The `qqnorm()` function plots the data against the theoretical normal quantiles, while `qqline()` adds the critical reference line, facilitating immediate visual comparison and assessment of alignment. This tool is often preferred over the histogram because it provides a standardized benchmark (the straight line) against which all sample sizes can be judged.

#make this example reproducible
set.seed(0)

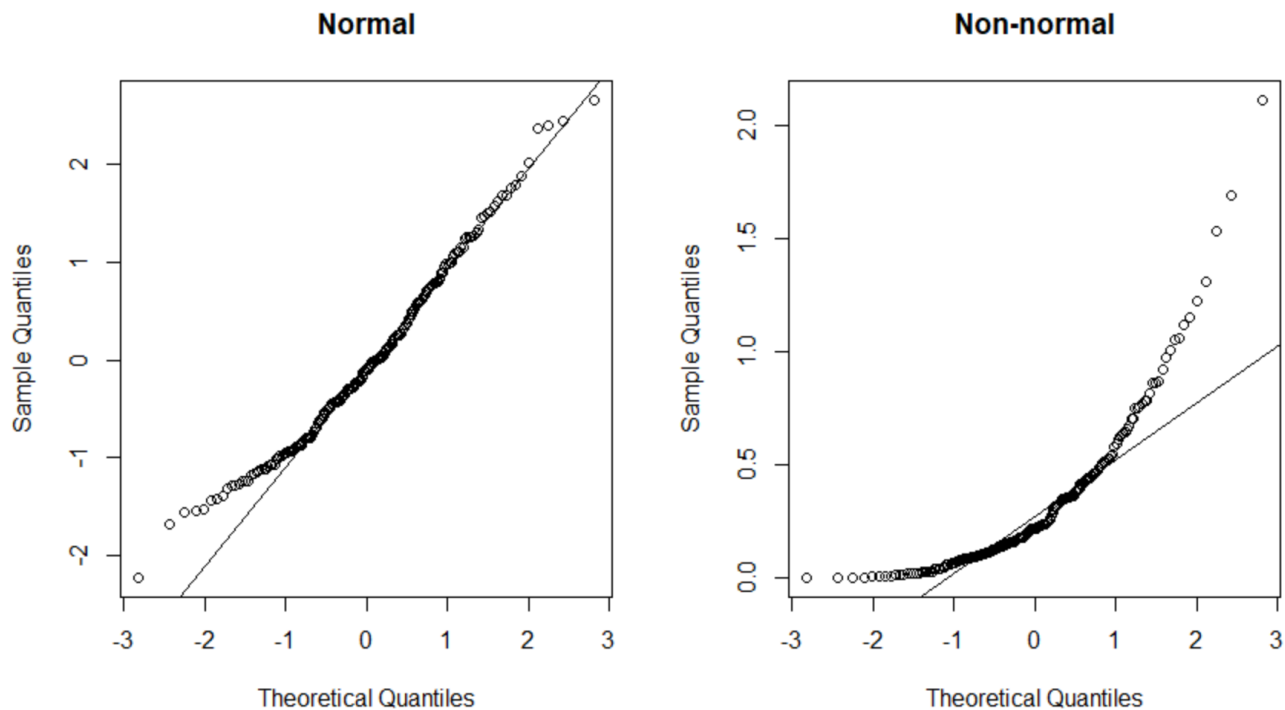
```
#create data that follows a normal distribution
normal_data <- rnorm(200)
```

```
#create data that follows an exponential distribution
non_normal_data <- rexp(200, rate=3)
```

```
#define plotting region
par(mfrow=c(1,2))
```

```
#create Q-Q plot for both datasets
qqnorm(normal_data, main='Normal')
qqline(normal_data)
```

```
qqnorm(non_normal_data, main='Non-normal')
qqline(non_normal_data)
```



The results of the [Q-Q plot](#) confirm the visual findings from the histogram. The plot on the left shows data points lying almost perfectly on the diagonal reference line, providing strong graphical evidence that the dataset is highly consistent with a normal distribution. Conversely, the plot on the right displays a severe curvature, particularly noticeable in the lower quantiles. This pattern is highly characteristic of skewed data, confirming the significant departure from normality found in the exponentially distributed sample.

Method 3: Performing the Shapiro-Wilk Test

While visual methods are excellent for preliminary screening, formal hypothesis tests are indispensable for objective confirmation, especially in high-stakes research. The [Shapiro-Wilk Test](#) is widely recognized as the gold standard among statisticians for assessing normality, offering superior power compared to many alternatives, particularly when dealing with small to moderate sample sizes (typically below 2000 observations). This test evaluates the null hypothesis that the data is normally distributed by comparing the empirical cumulative distribution function (CDF) of the data against the theoretical CDF of a normal distribution.

The interpretation of the Shapiro-Wilk test revolves around its hypothesis structure: the **Null Hypothesis (H₀)** posits that the data is sampled from a normally distributed population, while the **Alternative Hypothesis (H_a)** asserts that the data is not normally distributed. The decision to reject H₀ rests entirely on the resulting [p-value](#). If the p-value falls below the predefined significance level (α , customarily set at 0.05), we reject H₀, concluding that the data is non-normal. Conversely, a p-value greater than 0.05 indicates insufficient evidence to reject H₀, thereby

lending support to the assumption of normality for subsequent [parametric tests](#).

We apply the `shapiro.test()` function in **R** to both our synthetic datasets. The function outputs the test statistic (W), which measures the correlation between the data and the corresponding normal scores, and the critical p-value, which informs our hypothesis decision. Analyzing this output is essential for deriving a conclusive statement regarding the data's distribution.

```
#make this example reproducible  
set.seed(0)
```

```
#create data that follows a normal distribution  
normal_data <- rnorm(200)
```

```
#perform shapiro-wilk test  
shapiro.test(normal_data)
```

Shapiro-Wilk normality test

```
data: normal_data  
W = 0.99248, p-value = 0.3952
```

```
#create data that follows an exponential distribution  
non_normal_data <- rexp(200, rate=3)
```

```
#perform shapiro-wilk test  
shapiro.test(non_normal_data)
```

Shapiro-Wilk normality test

```
data: non_normal_data  
W = 0.84153, p-value = 1.698e-13
```

For the first dataset (`normal_data`), the resulting p-value (0.3952) is considerably larger than the conventional significance threshold of $\alpha = 0.05$. Consequently, we must fail to reject the null hypothesis, confirming that the data is assumed to be [normally distributed](#). Conversely, the p-value for the second test (1.698e-13) is infinitesimally small, falling far below the 0.05 threshold. This compels us to reject the null hypothesis, providing conclusive statistical proof that the exponentially distributed data is **not normally distributed**.

Method 4: Utilizing the Kolmogorov-Smirnov Test

The [Kolmogorov-Smirnov Test](#) (K-S test) is another classical non-parametric goodness-of-fit test.

While often used to compare two empirical distributions, it can also be employed as a one-sample test to determine if a sample distribution significantly differs from a specified theoretical distribution, such as the normal distribution. Although it is generally considered less powerful than the Shapiro-Wilk test for detecting specific departures from normality, particularly in smaller samples, it remains a robust and versatile diagnostic tool.

The core mechanism of the K-S test involves calculating the maximum absolute difference, known as the D statistic, between the empirical cumulative distribution function (ECDF) derived from the sample and the theoretical cumulative distribution function (CDF) of the normal distribution. A larger D statistic indicates a greater divergence between the two distributions. As with the Shapiro-Wilk test, the ultimate decision relies on the [p-value](#): if this value is low (e.g., less than 0.05), the difference is statistically significant, requiring the rejection of the null hypothesis of normality.

To implement this test in **R**, we use the `ks.test()` function. Crucially, when testing against a normal distribution, we must specify the target theoretical distribution using the argument `'pnorm'`. The function requires the standard normal distribution parameters (mean=0, standard deviation=1) unless the data has been standardized prior to the test, or if the mean and standard deviation of the sample are supplied as arguments.

#make this example reproducible
set.seed(0)

```
#create data that follows a normal distribution
normal_data <- rnorm(200)
```

```
#perform kolmogorov-smirnov test
ks.test(normal_data, 'pnorm')
```

One-sample Kolmogorov-Smirnov test

```
data: normal_data
D = 0.073535, p-value = 0.2296
alternative hypothesis: two-sided
```

```
#create data that follows an exponential distribution
non_normal_data <- rexp(200, rate=3)
```

```
#perform kolmogorov-smirnov test
ks.test(non_normal_data, 'pnorm')
One-sample Kolmogorov-Smirnov test
```

```
data: non_normal_data
```

D = 0.50115, p-value < 2.2e-16
alternative hypothesis: two-sided

Examining the output for the normally distributed data, the p-value (0.2296) is significantly above 0.05. This result confirms that the data does not deviate significantly from the theoretical normal distribution, compelling us to retain the null hypothesis of normality. Conversely, the exponentially distributed data yields an extremely small p-value (less than 2.2e-16). This strong statistical evidence dictates the rejection of the null hypothesis, confirming that the distribution is fundamentally different from a normal distribution.

Strategies for Handling Non-Normal Data

If the comprehensive assessment--combining visual evidence from the [histogram](#) and Q-Q plot with the objective results of the Shapiro-Wilk and Kolmogorov-Smirnov tests--consistently indicates that a dataset is **not** normally distributed, the researcher is faced with a choice. While switching immediately to non-parametric tests (which do not assume a specific distribution) is an option, it often comes at the cost of statistical power and interpretability. A preferred initial approach, especially when the violation is due to skewness, is to apply a mathematical transformation to the data.

Data transformation seeks to rescale the original variable so that the resulting distribution is more symmetrical and approximates the normal curve closely enough to satisfy the assumptions of [parametric tests](#). Choosing the correct transformation depends on the nature and severity of the non-normality, but transformations are highly effective for correcting positive (right) skewness.

The three most common and effective mathematical transformations used to normalize positively skewed data include:

- 1. Log Transformation:** This technique involves transforming the original values from x to $\log(x)$. This is highly aggressive and proves most effective when the data spans several orders of magnitude or exhibits extreme positive skewness.
- 2. Square Root Transformation:** A less severe adjustment, this involves transforming the original values from x to $\sqrt{x}$. It provides a moderate correctional effect and is frequently applied to count data, where values are typically non-negative integers.
- 3. Cube Root Transformation:** This method transforms the original values from x to $x^{1/3}$. It offers a transformation intensity that sits between the log and square root methods, providing a reliable corrective measure that is less disruptive than the log function.

By judiciously selecting and applying one of these mathematical transformations, researchers can often rescue their dataset, allowing them to proceed with the preferred parametric statistical

analysis while maintaining robustness and statistical power.

Read [this tutorial](#) to see how to perform these transformations in R.

Additional Resources for Statistical Assumptions

A thorough understanding and verification of all statistical assumptions form the bedrock of accurate and reliable quantitative research. Beyond normality, researchers should also familiarize themselves with related critical prerequisites for linear modeling, such as the assumptions of homoscedasticity (equality of variances) and linearity (a straight-line relationship between variables). Mastering these foundational concepts ensures the validity of subsequent statistical modeling and hypothesis testing.