

Learning to Impute Missing Data: Replacing NA Values with the Median in R

Authored by
Mohammed looti

October 28, 2025

RECOMMENDED CITATION

Mohammed looti (2025). *Learning to Impute Missing Data: Replacing NA Values with the Median in R*. PSYCHOLOGICAL STATISTICS. Retrieved from <https://statistics.arabpsychology.com/?p=4844>

Introduction: Handling Missing Data and Median Imputation in R

Missing data, often represented as [NA values](#) in [R](#), is a common challenge in data analysis. These gaps can arise from various reasons, such as data entry errors, equipment malfunctions, or survey non-responses. If not handled appropriately, missing data can lead to biased results, reduced statistical power, and inaccurate models. Therefore, addressing missing values is a crucial step in any robust data preprocessing workflow.

One effective strategy for dealing with missing numerical data is [imputation](#), where missing values are replaced with substituted estimates. Among various imputation techniques, replacing missing values with the [median](#) is a widely adopted method, particularly when the data contains [outliers](#) or is skewed. The median, unlike the mean, is less sensitive to extreme values, making it a robust choice for maintaining the central tendency of a dataset without undue influence from unusual observations.

In this comprehensive guide, we will explore practical methods for replacing NA values with the median using powerful functions from the [dplyr](#) and [tidyr](#) packages in R. These packages, part of the [tidyverse](#) ecosystem, streamline data manipulation tasks, offering an intuitive and efficient approach to data cleaning. We will walk through three distinct scenarios: imputing a single column, multiple specified columns, and all numeric columns within a [data frame](#).

Preparing Your Data: A Practical Example

Before diving into the imputation methods, let's establish a working example using a sample data frame. This will allow us to demonstrate each technique clearly and observe its impact on the data. Our example data frame, named `df`, simulates a small dataset containing player statistics, where some observations are intentionally marked as missing (NA) to illustrate the imputation process.

The following R code snippet initializes our sample data frame. It includes a character column for player names and three numeric columns: `points`, `rebounds`, and `blocks`, which will be the focus of our imputation efforts. Notice the presence of NA values in these numerical columns, representing the data points we aim to fill.

```
#create data frame
```

```
df <- data.frame(player=c('A', 'B', 'C', 'D', 'E'),  
points=c(17, 13, NA, 9, 25),  
rebounds=c(3, 4, NA, NA, 8),  
blocks=c(1, 1, 2, 4, NA))
```

```
#view data frame
```

```
df
```

```
player points rebounds blocks
```

```
1 A 17 3 1
```

```
2 B 13 4 1
```

```
3 C NA NA 2
```

```
4 D 9 NA 4
```

```
5 E 25 8 NA
```

Upon inspecting the data frame, you can clearly see the missing values in the `points`, `rebounds`, and `blocks` columns. Our goal is to systematically replace these NAs with the median value calculated from the non-missing observations within their respective columns, ensuring our dataset is complete and ready for further analysis.

Method 1: Replacing NA Values with Median in a Single Column

When you need to address missing data in a specific numeric column, the combination of `dplyr`'s [mutate](#) and [across](#) functions, along with `tidyr`'s [replace_na](#), provides a clean and efficient solution. This method is ideal for targeted imputation, allowing you to precisely control which variables are modified.

The core of this approach lies in the `mutate()` function, which is used to add new variables or transform existing ones. We pair it with `across()` to apply a function to selected columns. Within `across()`, we specify the target column (e.g., `points`) and then define the transformation using an anonymous function (`~`). This anonymous function first calculates the median of the column, explicitly handling existing NA values by setting `na.rm=TRUE`, and then uses `replace_na()` to substitute any missing values in that column with the calculated median.

Let's demonstrate this by replacing the NA values in the `points` column of our `df` data frame with the median of the non-missing `points` values.

```
library(dplyr)
```

```
library(tidyr)
```

```
#replace NA values in points column with median of points column
```

```
df <- df %>% mutate(across(points, ~replace_na(., median(., na.rm=TRUE))))
```

```
#view updated data frame
```

```
df
```

```
player points rebounds blocks
```

```
1 A 17 3 1
```

```
2 B 13 4 1
```

```
3 C 15 NA 2
4 D 9 NA 4
5 E 25 8 NA
```

In this example, the median of the `points` column (excluding the NA) is 15 (calculated from 17, 13, 9, 25). Consequently, the single NA value in the `points` column for player 'C' has been successfully replaced with 15. It is important to note that only the specified `points` column was affected by this operation; the `rebounds` and `blocks` columns, which also contain missing values, remain unchanged. This confirms the precise, column-specific nature of this imputation method.

Method 2: Imputing NA Values Across Multiple Specific Columns

Often, you may encounter scenarios where multiple columns within your data frame require imputation. Instead of applying the single-column method repeatedly, `dplyr`'s `across()` function allows you to specify several columns simultaneously, streamlining your code and enhancing efficiency. This approach is particularly useful when a group of related variables all need the same imputation strategy.

To implement this, you simply provide a [vector](#) of column names to the `across()` function. The rest of the logic remains consistent: `replace_na()` is used in conjunction with `median(..., na.rm=TRUE)` to ensure that each specified column has its missing values filled with its own respective median. This ensures that the imputation is performed independently for each column, preserving their individual statistical properties.

Let's expand our imputation to cover both the `points` and `blocks` columns. We will replace any NA values in these columns with their individual medians.

```
library(dplyr)
```

```
library(tidyr)
```

```
#replace NA values in points and blocks columns with their respective medians
df <- df %>% mutate(across(c(points, blocks), ~replace_na(., median(., na.rm=TRUE))))
```

```
#view updated data frame
```

```
df
```

```
player points rebounds blocks
```

```
1 A 17 3 1.0
2 B 13 4 1.0
3 C 15 NA 2.0
4 D 9 NA 4.0
```

```
5 E 25 8 1.5
```

After executing this code, you'll observe that the NA in the `points` column is replaced by 15 (as before), and the NA in the `blocks` column (for player 'E') is now replaced by 1.5. The median of the `blocks` column (1, 1, 2, 4) is 1.5. This demonstrates how easily you can apply the same imputation logic across a user-defined subset of your data frame's columns, significantly reducing repetitive coding and potential errors.

The `rebounds` column, however, still retains its missing values, as it was not included in the `c(points, blocks)` vector. This highlights the precise control offered by specifying columns explicitly.

Method 3: Broad Imputation for All Numeric Columns

For datasets with numerous numeric columns that all require missing value imputation, manually listing each column can be cumbersome and prone to error. Fortunately, `dplyr` provides an elegant solution using ``where(is.numeric)`` within the `across()` function. This powerful construct allows you to apply the imputation logic to all columns that meet a specific criterion, such as being of a numeric data type.

This method is particularly advantageous for exploratory data analysis or initial cleaning steps on large datasets. By targeting all numeric columns, you ensure a consistent imputation strategy across all quantitative variables, while safely leaving non-numeric columns (like character or factor variables) untouched. The `is.numeric` predicate acts as a filter, dynamically selecting columns for the operation.

Let's apply this broad imputation method to our data frame, filling NA values in all numeric columns with their respective medians.

```
library(dplyr)
```

```
library(tidyr)
```

```
#replace NA values in all numeric columns with their respective medians
```

```
df <- df %>% mutate(across(where(is.numeric), ~replace_na(., median(., na.rm=TRUE))))
```

```
#view updated data frame
```

```
df
```

```
player points rebounds blocks
```

```
1 A 17 3 1.0
```

```
2 B 13 4 1.0
```

```
3 C 15 4 2.0
4 D 9 4 4.0
5 E 25 8 1.5
```

Upon reviewing the updated data frame, you will notice that all NA values within the `points`, `rebounds`, and `blocks` columns have been replaced. Specifically, the NA in `rebounds` (for players 'C' and 'D') is now replaced with 4, which is the median of the `rebounds` column (3, 4, 8). The `player` column, being non-numeric, remains entirely unchanged, preserving its original structure and content.

This method provides a robust and scalable solution for ensuring data completeness across all quantitative variables, making it an indispensable tool for data preprocessing in R. It significantly reduces the effort required for imputation in wide datasets, ensuring consistency and accuracy across your analytical pipeline.

Considerations and Best Practices for Imputation

While median imputation is a valuable and robust technique for handling missing data, especially in the presence of outliers, it is essential to understand its implications and consider best practices. Simple imputation methods like median replacement can reduce variance and potentially distort relationships between variables if not applied judiciously. Therefore, it's crucial to acknowledge the assumptions and limitations.

One key consideration is the "Missing At Random" (MAR) assumption. Median imputation often performs well when data are MAR, meaning the probability of a value being missing depends only on observed data, not on the missing value itself. If data are "Missing Not At Random" (MNAR), where missingness is related to the unobserved value, median imputation might introduce bias. Always strive to understand the mechanism behind your missing data.

For more complex missing data patterns or when higher accuracy is paramount, consider advanced imputation techniques. These might include K-Nearest Neighbors (KNN) imputation, multiple imputation (e.g., using packages like `mice` or `VIM`), or model-based imputation. These methods often provide more sophisticated estimates for missing values by leveraging relationships within the dataset, offering a more nuanced approach than single-value imputation.

Regardless of the chosen method, always document your imputation strategy. Transparency in your data preprocessing steps is vital for reproducibility and for others to understand your analytical choices. It is also good practice to perform sensitivity analyses by comparing results obtained from different imputation methods or by analyzing the data both with and without imputation to assess the robustness of your findings.

Additional Resources

To further enhance your data wrangling and imputation skills in R, consider exploring the following resources and tutorials. These can provide deeper insights into the capabilities of the tidyverse packages and introduce you to more advanced techniques for data cleaning and preparation.

The official [tidyverse](#) website for comprehensive documentation and tutorials.

Detailed guides on [dplyr](#) for data manipulation and transformation.

Documentation for [tidyr](#) focusing on data tidying and handling missing values.

Tutorials on [imputation](#) methods in R for a broader understanding of various techniques.