

Calculate Mean for Multiple Columns Using dplyr

Authored by
Mohammed looti

November 16, 2025

RECOMMENDED CITATION

Mohammed looti (2025). *Calculate Mean for Multiple Columns Using dplyr*.
PSYCHOLOGICAL STATISTICS. Retrieved from
<https://statistics.arabpsychology.com/?p=2652>

Streamlining Data Aggregation with [dplyr](#)

Effective data manipulation is the foundational requirement for rigorous statistical analysis and empirical research. When working within the powerful statistical environment of [R](#), the [dplyr](#) package stands out as an essential component of the [Tidyverse](#), providing a highly consistent and expressive grammar for data wrangling. This package utilizes a core set of functions, often called "verbs," that significantly enhance code readability and operational efficiency. A frequent analytical need across disciplines involves calculating descriptive statistics, such as the [mean](#), not just across groups, but across a specific selection of columns for each individual observation--a process known as row-wise operation. This capability is paramount when summarizing performance metrics or aggregated characteristics on a case-by-case basis.

This comprehensive guide provides an expert walkthrough of the precise methodology required to calculate the arithmetic [mean](#) value across multiple specified columns within an [R data frame](#). We will harness the specialized functionalities of the [dplyr](#) package, focusing specifically on techniques designed for efficient row-level computation. We will systematically break down the essential syntax, demonstrate a practical, real-world example using sample data, and conclude with advanced strategies for selecting columns dynamically, ensuring your data aggregation methods are both robust and scalable for integration into professional data analysis workflows.

The Foundational Syntax for Row-Wise Calculations

To successfully compute the average of several designated columns on a row-by-row basis using [dplyr](#), analysts must employ a strategic sequence of functions. The typical operational flow begins by passing your [data frame](#) through the [pipe operator](#) (`%>%`), which chains commands together seamlessly. The crucial initial step involves invoking [rowwise\(\)](#), which fundamentally changes the context of subsequent calculations, instructing [dplyr](#) to treat each row as a distinct group. Following this, you use [mutate\(\)](#) to generate and define the new column containing the average, and finally, you utilize the specialized function [c_across\(\)](#) to accurately specify the columns that will contribute to the [mean](#) calculation for that specific row.

The foundational code structure that orchestrates this powerful row-wise aggregation is concise and highly readable. It is essential, however, to ensure that the [dplyr](#) library is loaded into your R session prior to execution, a step typically accomplished using the command `library(dplyr)`. This ensures all necessary verbs are available for use in the pipeline.

`library(dplyr)`

```
df %>%  
rowwise() %>%  
mutate(game_mean = mean(c_across(c('game1', 'game2', 'game3')), na.rm=TRUE))
```

Let us systematically analyze the components of this pipeline to understand their precise roles. The [pipe operator](#) (`%>%`) facilitates the seamless flow of the data frame named `df` through the subsequent functions. The inclusion of [rowwise\(\)](#) is critical because it fundamentally shifts the operational focus from the default column-centric operations to a row-centric perspective, ensuring that operations are applied individually to each row. The [mutate\(\)](#) function is then responsible for creating and populating our new derived variable, which we have aptly named `game_mean`. Within [mutate\(\)](#), the standard [mean\(\)](#) function calculates the average of the subset of values defined by [c_across\(\)](#).

The function [c_across\(\)](#) is specifically designed for use within row-wise operations, allowing the user to select multiple columns (in this case, **game1**, **game2**, and **game3**) to be aggregated into a vector for computation. Furthermore, the argument `na.rm = TRUE`, which is nested within the [mean\(\)](#) function call, represents a vital consideration for ensuring data quality. This argument guarantees that any [missing values \(NA\)](#) encountered in the specified columns are automatically excluded from the calculation. By excluding these gaps, we prevent the resulting row [mean](#) from incorrectly evaluating to `NA` itself, thereby significantly ensuring the robustness and integrity of the final derived metric even when dealing with imperfect or incomplete source data.

Preparing the Data for a Practical Example

To fully grasp the mechanics and practical utility of the syntax outlined above, we transition now to a concrete, easily relatable example. Imagine a scenario where we are analyzing an [R data frame](#) that systematically logs the points scored by various players across three distinct competitive basketball games. This dataset structure is perfectly suited for demonstrating row-wise calculations, as our primary goal is to determine the average score achieved by each individual player across these three performance metrics, regardless of their team affiliation.

We will begin by constructing a sample [data frame](#), which we will name `df`. This structure will include a categorical variable identifying the **team** and three quantitative columns: **game1**, **game2**, and **game3**. Critically, we have intentionally introduced instances of [missing data](#), represented by the value `NA`, within the scores. This common feature mirrors the imperfect nature of real-world datasets and highlights the necessity of implementing effective mechanisms for handling these gaps--a core strength of the [dplyr](#) approach when paired with the `na.rm` argument.

#create data frame

```
df <- data.frame(team=c('A', 'A', 'A', 'B', 'B', 'B', 'C', 'C'),  
game1=c(10, 12, 17, 18, 24, 29, 29, 34),  
game2=c(8, 10, 14, 15, NA, 19, 18, 29),  
game3=c(4, 5, 5, 9, 12, 12, 18, 20))
```

#view data frame

```
df
```

```
team game1 game2 game3
```

```
1 A 10 8 4
```

```
2 A 12 10 5
```

```
3 A 17 14 5
```

```
4 B 18 15 9
```

```
5 B 24 NA 12
```

```
6 B 29 19 12
```

```
7 C 29 18 18
```

```
8 C 34 29 20
```

This generated data structure provides a clear, row-level record of individual player scores across the three games. Our analytical objective is straightforward: to compute a single, representative average score for each player, combining their performance from **game1**, **game2**, and **game3**. The resulting average, which will be housed in a newly created column, will serve as a standardized, easily interpretable metric, enabling direct comparison of player performance across all measured events.

Executing the Row-Wise Mean Calculation

With our sample [data frame](#) properly initialized and ready for transformation, we can now proceed to apply the robust [dplyr](#) syntax to calculate the crucial row-wise [mean](#) for our specified game columns. This execution will seamlessly integrate into our data structure, augmenting the existing columns by appending a new column that contains the calculated average performance for every player in the dataset.

The following code snippet demonstrates the precise application of the [dplyr](#) pipeline discussed in the previous section. It explicitly targets the columns **game1**, **game2**, and **game3** for the averaging process. Notice the continued, vital use of [na.rm = TRUE](#), which guarantees that [missing values](#) are handled gracefully, preventing erroneous results and ensuring the integrity of all calculated averages across the board.

library(dplyr)

```
#calculate mean value in each row for game1, game2 and game3 columns
```

```
df %>%
```

```
rowwise() %>%
```

```
mutate(game_mean = mean(c_across(c('game1', 'game2', 'game3')), na.rm=TRUE))
```

```
# A tibble: 8 x 5
```

```
# Rowwise:
team game1 game2 game3 game_mean

1 A 10 8 4 7.33
2 A 12 10 5 9
3 A 17 14 5 12
4 B 18 15 9 14
5 B 24 NA 12 18
6 B 29 19 12 20
7 C 29 18 18 21.7
8 C 34 29 20 27.7
```

The resulting output clearly demonstrates the successful creation of the new column, `game_mean`, which holds the precise calculated [mean](#) score for each player across the three games. To ensure complete clarity and verify the function's reliability, let us review the calculations for a few specific rows, paying close attention to how the function handles the critical issue of [missing data](#):

Row 1: For the first player on Team A, the scores are 10, 8, and 4. The arithmetic mean is calculated as $(10 + 8 + 4) / 3$, correctly yielding **7.33**.

Row 2: The second player on Team A recorded scores of 12, 10, and 5. The calculation confirms the mean as $(12 + 10 + 5) / 3$, resulting in exactly **9**.

Row 5: This row is highly informative, featuring the fifth player on Team B with scores 24, `NA`, and 12. Because we correctly specified `na.rm = TRUE`, the `NA` value is automatically excluded from the numerator and the count. The average is calculated based only on the two valid entries: $(24 + 12) / 2$, which results in precisely **18**. This elegantly demonstrates the robustness of the [dplyr](#) pipeline in efficiently managing real-world data imperfections.

Scaling Up: Dynamic Column Selection with [starts_with\(\)](#)

While manually listing column names, such as `c('game1', 'game2', 'game3')`, is sufficient for small and unchanging datasets, this approach quickly becomes inefficient, tedious, and highly error-prone when dealing with large datasets that contain numerous columns sharing a common naming structure. To solve this critical scalability challenge, [dplyr](#) offers powerful mechanisms known as [select helpers](#), which facilitate dynamic column selection based on patterns. Utilizing these helpers dramatically improves the flexibility, maintainability, and overall conciseness of your R code.

One of the most valuable and frequently employed [select helpers](#) is [starts_with\(\)](#). This function is specifically designed to select all columns whose names commence with a defined string prefix.

In the context of our basketball example, every column intended for averaging (**game1**, **game2**, and **game3**) shares the consistent prefix "game". Therefore, substituting the manual list with [starts_with\('game'\)](#) provides an efficient, clean, and highly scalable alternative for the row-wise [mean](#) calculation.

library(dplyr)

```
#calculate mean value in each row for columns that start with 'game'
df %>%
  rowwise() %>%
  mutate(game_mean = mean(c_across(c(starts_with('game'))), na.rm=TRUE))

# A tibble: 8 x 5
# Rowwise:
  team game1 game2 game3 game_mean
  <fct> <dbl> <dbl> <dbl> <dbl>
1 A    10    8    4    7.33
2 A    12   10    5    9
3 A    17   14    5   12
4 B    18   15    9   14
5 B    24   NA   12   18
6 B    29   19   12   20
7 C    29   18   18  21.7
8 C    34   29   20  27.7
```

As demonstrated by the output, employing this modified syntax yields results functionally identical to the previous, manually-defined example. However, the true advantage lies in its scalability potential. If your underlying [data frame](#) were to expand to include numerous game columns--say, **game1** through **game20**--using [starts_with\('game'\)](#) would instantly select all relevant variables using a single, concise, and unchanging expression. This robust approach makes your analytical code highly resilient to structural changes in the data, provided that the foundational column naming convention remains consistent.

Advanced Column Selection Strategies

The effectiveness of [starts_with\(\)](#) is merely the gateway to the comprehensive suite of tools provided by [dplyr](#) for selecting variables. The package offers a rich variety of [select helpers](#) that grant analysts exceptional flexibility in targeting columns based on complex patterns, criteria, or numerical ranges, moving beyond the limitation of explicit enumeration. Gaining proficiency with these helpers is fundamental to streamlining data preparation and analysis tasks, especially when

managing large or highly dynamic datasets where manual column indexing is simply impractical.

These specialized functions can be seamlessly integrated within [c_across\(\)](#) or other core [mutate\(\)](#) operations. Some of the most valuable and commonly used selection aids include:

[ends_with\(\)](#): Selects columns whose names terminate with a specified string (e.g., `ends_with('_score')`).

[contains\(\)](#): Selects columns whose names incorporate a specified literal string anywhere within the name (e.g., `contains('total')`).

[matches\(\)](#): Selects columns whose names adhere to a provided regular expression pattern, offering the highest level of detailed control (e.g., `matches('$')` for columns ending in a number).

[num_range\(\)](#): Specifically designed for selecting columns that follow a numbered naming convention, such as `var01`, `var02`, etc. (e.g., `num_range('game', 1:5)`).

[everything\(\)](#): A utility function that selects all remaining columns not yet specified. It is often employed strategically after manually arranging important columns to ensure all variables are included in the final output.

By strategically integrating these [select helpers](#) into your data manipulation workflows, you gain significant operational flexibility, ensuring you can target specific subsets of columns with precision without relying on error-prone manual enumeration. This capability is paramount for writing adaptable and robust analytical code that can easily withstand evolving requirements and changes in source data structure.

Conclusion: Mastering Row-Wise Aggregation

The essential ability to calculate row-wise [means](#) across multiple columns is a pervasive and fundamentally important requirement in nearly all quantitative data analysis efforts, especially when the objective is to summarize individual observations across diverse metrics. As extensively demonstrated throughout this guide, the [dplyr](#) package within the [R](#) environment offers an exceptionally elegant, highly efficient, and syntactically readable solution for achieving this complex aggregation task.

The power of this technique lies in the harmonious combination of three core functions: [rowwise\(\)](#) to set the calculation context, [mutate\(\)](#) to generate the result column, and [c_across\(\)](#) to dynamically specify the target variables for aggregation. When this core structure is augmented with robust [select helpers](#), such as [starts_with\(\)](#), analysts can execute sophisticated, scalable data aggregations with minimal code overhead. This methodological advantage not only simplifies the resulting analytical code but also significantly bolsters its flexibility, clarity, and ease of

maintenance, ensuring the code can be easily adapted to accommodate evolving analytical needs and much larger datasets.

We strongly encourage data practitioners to actively integrate these specialized [dplyr](#) functions into their regular workflow. By mastering these techniques--especially the integration of row-wise operations and dynamic column selection--you will undoubtedly elevate your data manipulation skills in [R](#), enabling you to extract deeper, more precise insights from your datasets with unparalleled efficiency and reliability.