

# Select the First Row by Group Using dplyr

Authored by  
**Mohammed loot**

November 7, 2025

## RECOMMENDED CITATION

Mohammed loot (2025). *Select the First Row by Group Using dplyr*. PSYCHOLOGICAL STATISTICS. Retrieved from <https://statistics.arabpsychology.com/?p=12017>

Data analysis workflows frequently demand specialized techniques to isolate and extract specific observations from large datasets based on criteria defined within subgroups. A fundamental and common requirement for analysts utilizing the [R](#) statistical environment is the precise selection of the first, last, or an arbitrary Nth record belonging to each unique group within their data structure. This operation is crucial for tasks like deduplication, finding the earliest entry, or identifying extremes (minimum or maximum values) across defined categories.

The challenges inherent in this task--managing group boundaries, establishing a meaningful order, and applying positional logic--are elegantly solved by the powerful [dplyr](#) package. As a cornerstone of the [Tidyverse](#) ecosystem, **dplyr** provides a highly efficient, readable, and reproducible framework for data manipulation. By mastering a concise combination of grouping, ordering, and filtering functions, users can precisely target and extract the single row that meets the specified positional criteria for every category present in the dataset.

The core philosophy behind selecting a specific row by group relies on establishing a reliable internal order within each group before applying a positional filter. This structured approach is essential not only for achieving the desired results but also for ensuring that complex data cleaning and transformation pipelines are robust and easily auditable. The following sections will break down the essential syntax and demonstrate practical applications, ensuring clarity in how group boundaries and sorting mechanisms interact to produce the required output.

## Understanding the Core Logic: Grouping, Ordering, and Filtering

To effectively select the "first" observation per group, the process must be carefully structured into three distinct and sequential steps. This pipeline design is characteristic of effective [dplyr](#) workflows, ensuring that each transformation builds logically upon the last. The three stages are: grouping the data, arranging the internal order, and filtering by position. Without a clearly defined order, selecting the "first" row is arbitrary and generally unhelpful, as it merely captures the first row encountered in the dataset's current, possibly random, arrangement.

The first step utilizes the [group\\_by\(\)](#) function to partition the entire [data frame](#) into logical subgroups based on one or more categorical variables. Once the groups are established, the subsequent steps operate independently on each of these partitions. The second, and most critical, step involves using the [arrange\(\)](#) function. This function establishes the specific sort order--either ascending or descending--that dictates which observation will be assigned the first position. For instance, if you want the observation with the lowest score, you arrange the scores in ascending order; if you want the highest score, you arrange descendingly.

The final step leverages the [filter\(\)](#) function in conjunction with the powerful **window function**, [row\\_number\(\)](#). The [row\\_number\(\)](#) function assigns an integer rank (1, 2, 3, ...) to each row within its current group, respecting the order established by [arrange\(\)](#). By setting the filter condition to

`row_number() == 1`, we instruct **dplyr** to retain only the single observation that holds the rank of one within every defined group. This synergy of functions provides a clean and highly expressive solution for positional extraction.

The foundational syntax below illustrates this three-step pipeline structure. It is designed to select the row that ranks first according to the variable specified in the `arrange()` function. It is important to remember that the effectiveness of this technique hinges entirely on the preceding sorting step; the positional rank is meaningless without a defined order.

```
df %>%  
group_by(group_var) %>%  
arrange(values_var) %>%  
filter(row_number()==1)
```

## Constructing a Reproducible Data Example in R

To provide a concrete and easily replicable demonstration of this powerful technique, we will construct a sample **data frame** within the [R](#) environment. This dataset simulates a common scenario in data analysis, representing performance metrics across different teams. Our data structure includes two critical variables: `team`, which serves as the primary **grouping variable**, and `points`, which is the **value variable** we will use for establishing the order and subsequent selection.

Working with this tangible example allows us to solidify the conceptual understanding of how the [dplyr](#) functions interact within the pipeline. We will use this small, controlled dataset to rigorously demonstrate how to correctly identify and extract rows corresponding to extremes, such as the minimum number of points, the maximum number of points, and other specific positional ranks for each unique team identifier. This practical setup ensures that the results of the manipulation commands are immediately observable and verifiable.

Before executing the core manipulation commands, two preliminary steps are necessary: ensuring the [dplyr](#) library is loaded into the session and properly constructing the source data frame. The code block below details the creation of the sample data, which contains ten observations distributed unevenly across three distinct teams (A, B, and C). Notice the variety in the `points` variable, which will allow us to test both ascending and descending sorting logic accurately across the groups.

```
#create dataset  
df <- data.frame(team=c('A', 'A', 'A', 'B', 'B', 'B', 'C', 'C', 'C', 'C'),  
points=c(4, 9, 7, 7, 6, 13, 8, 8, 4, 17))
```

```
#view dataset  
df
```

```
team points
```

```
1 A 4  
2 A 9  
3 A 7  
4 B 7  
5 B 6  
6 B 13  
7 C 8  
8 C 8  
9 C 4  
10 C 17
```

## Case Study 1: Isolating the Minimum Value per Group (Ascending Order)

In this first scenario, our objective is highly specific: to select the row that contains the absolute **minimum value** of `points` for every respective team. When applying the grouped operation, this minimum value must correspond to the first row after the data has been rigorously sorted in ascending order. This is a common requirement in quality control or performance tracking where the lowest metric (e.g., error rate, fastest time) is often the most significant observation to isolate.

The workflow initiates by piping the raw data frame, `df`, into the `group_by()` function, unequivocally specifying `team` as the variable that defines the boundaries of our subgroups. Immediately following the grouping, we invoke `arrange()` on the `points` variable. Crucially, `arrange()` performs an ascending sort by default (smallest to largest). This action ensures that the row containing the minimum `points` value is placed at the top (position 1) within its team group.

The final command, `filter(row_number() == 1)`, acts as the selection mechanism. Because the data has been pre-sorted ascendingly, retaining only the observation assigned the first position effectively extracts the record corresponding to the minimum points for each team. The output confirms that for Team A and Team C, the minimum point score is 4, while for Team B, the minimum is 6, demonstrating the accurate isolation of the lowest value across independent groups.

```
library(dplyr)
```

```
df %>%  
  group_by(team) %>%  
  arrange(points) %>%
```

```
filter(row_number()==1)
```

```
# A tibble: 3 x 2
```

```
# Groups: team
```

```
team points
```

```
1 A 4
```

```
2 C 4
```

```
3 B 6
```

## Case Study 2: Identifying the Maximum Value Using Descending Order

When the analytical requirement shifts from finding the minimum value to identifying the **maximum value** per group, the necessary adjustment is remarkably simple, confined primarily to the ordering step. The underlying structure of the data pipeline remains consistent: we must still group the data, arrange the internal order according to the desired metric, and then filter for the first positional rank. This consistency highlights the robustness and predictability of the [dplyr](#) framework.

To successfully achieve a descending sort--which places the largest values at the head of each group--we must utilize the `desc()` helper function within [arrange\(\)](#). By wrapping the `points` variable inside `desc()`, we invert the default sorting behavior. This inversion ensures that when `row_number()` subsequently assigns ranks, the observation possessing the highest score for `points` within that specific team receives rank number one.

This technique is incredibly versatile because it allows analysts to select the row associated with any aggregate statistic--maximum, minimum, or even a specific median value--simply by defining the appropriate sort order before the positional filter is applied. The resulting output confirms the successful isolation of the highest score for each team: Team C (17), Team B (13), and Team A (9). Furthermore, this approach is highly efficient because it avoids the need to calculate the maximum value separately and then join it back to the original [data frame](#); the selection happens in a single, streamlined operation.

```
df %>%
```

```
group_by(team) %>%
```

```
arrange(desc(points)) %>%
```

```
filter(row_number()==1)
```

```
# A tibble: 3 x 2
```

```
# Groups: team
```

```
team points
```

```
1 C 17
2 B 13
3 A 9
```

## Extending the Technique: Selecting Nth and Last Rows

The true utility and flexibility of using the `row_number()` function within a grouped `filter()` operation extend far beyond merely selecting the first row. This approach enables analysts to select **any arbitrary Nth row**--such as the second-highest observation, the third-lowest, or any positional ranking based on the defined sort order. This generalization is essential for advanced data auditing and outlier detection, where interest might lie in the runners-up rather than just the top performer.

To generalize this concept, the user simply modifies the conditional statement inside the `filter()` command. Instead of setting `row_number() == 1`, the condition is changed to `row_number() == n`, where `n` represents the specific desired rank. For instance, if we maintain the descending sort established previously (highest score ranked first) and wish to extract the second-highest scorer per team, we set `n` to 2. This direct manipulation of the rank index provides granular control over the selection process, regardless of the size or complexity of the groups.

Furthermore, selecting the absolute **last row** of a group is a common requirement, particularly when dealing with time-series or sequential data where the latest record is significant. This is achieved using the special `R` function `n()`. When used inside a `filter()`, `n()` returns the total count of rows within the current group. Therefore, the condition `row_number() == n()` reliably captures the final record of every group, irrespective of whether the group sizes are equal or vary significantly. This capability is invaluable when ordering by a different variable, such as a timestamp, and extracting the final measurement recorded.

Consider the following examples demonstrating how to select the second row, followed by the syntax for selecting the last row. In both examples, the data is still sorted in descending order of points, meaning the second row represents the second-highest score, and the last row represents the minimum score.

To select the 2nd row by group (i.e., the second-highest score), the syntax is adjusted as follows:

```
df %>%
  group_by(team) %>%
  arrange(desc(points)) %>%
  filter(row_number()==2)
```

Alternatively, to select the last row by group (the row with the minimum score, given the initial descending sort):

```
df %>%  
group_by(team) %>%  
arrange(desc(points)) %>%  
filter(row_number()==n())
```

## Summary and Best Practices in Grouped Manipulation

Mastering the grouped data manipulation pipeline provided by the [Tidyverse](#) is an absolutely essential skill for conducting efficient and transparent data analysis. The powerful combination of [group\\_by\(\)](#), [arrange\(\)](#), and [filter\(row\\_number\(\)\)](#) establishes a highly versatile framework for extracting specific positional data points from datasets, whether the goal is to find the minimum, maximum, or a custom Nth rank within a category. This three-step process ensures that the selection logic is clearly defined and reproducible across different data environments.

A critical takeaway for all analysts is the absolute dependence of the selection on the preceding ordering step. If the [arrange\(\)](#) step is omitted from the pipeline, the subsequent `filter(row_number() == 1)` command will simply select the first row based on the arbitrary existing order of the source [data frame](#) rows within each group. This can lead to unpredictable and incorrect results, especially when dealing with data read from external files where row order is not guaranteed. Therefore, defining an explicit order is mandatory when the positional rank must correspond to a specific value (like min or max).

It is also worth noting that for users running modern versions of [dplyr](#), alternative and often more concise helper functions exist, such as `slice_head(n=1)` and `slice_tail(n=1)`. These functions offer a slightly abbreviated syntax specifically for selecting the very first or last row, respectively, assuming the data has already been grouped and ordered correctly. However, the foundational technique using [filter\(row\\_number\(\)\)](#) remains paramount, as it provides the maximum flexibility for complex conditional selections, such as selecting the second or third record, and offers the clearest understanding of the underlying logic of grouped window functions.

## Further Reading and Resources

For further exploration of grouped operations, data ordering, and advanced data manipulation techniques within the [R](#) ecosystem and the **Tidyverse**, consult the following related guides and documentation:

[How to Arrange Rows in R](#)

[How to Count Observations by Group in R](#)

[How to Find the Maximum Value by Group in R](#)