

Learning PySpark: Calculating the Mean of a DataFrame Column

Authored by
Mohammed looti

November 11, 2025

RECOMMENDED CITATION

Mohammed looti (2025). *Learning PySpark: Calculating the Mean of a DataFrame Column*. PSYCHOLOGICAL STATISTICS. Retrieved from <https://statistics.arabpsychology.com/?p=16535>

Calculating **descriptive statistics** is an essential initial phase in nearly every modern data analysis and machine learning workflow. When handling truly massive datasets, standard Python libraries often become insufficient, necessitating the use of distributed computing frameworks. [PySpark](#), the Python API for Apache Spark, offers highly efficient methods for performing these complex calculations across large, distributed clusters. A particularly frequent requirement is determining the [mean](#) (or average) of a numeric column. Understanding how to correctly implement aggregation functions using the [DataFrame](#) API is crucial for leveraging Spark's power efficiently and effectively. This comprehensive guide details the two most common and optimized approaches for calculating the mean of one or more columns within a PySpark DataFrame, complete with practical code examples for immediate application.

Choosing the Right Method for Mean Calculation

In the PySpark ecosystem, calculating the mean can be approached in two primary ways. The choice between methods largely depends on the desired output: whether you need a single, scalar result aggregated across the entire DataFrame, or if you aim to calculate multiple column means simultaneously and present them neatly within a resultant summary DataFrame. Both techniques utilize built-in functions specifically optimized for Spark's underlying distributed architecture, guaranteeing robust performance even when processing petabytes of data. It is vital to select the method that aligns precisely with your specific analytical objective and the subsequent data processing steps.

The first strategy focuses on obtaining a single metric using the dedicated aggregation pipeline, ideal for immediate programmatic use. The second strategy leverages projection and selection, which is better suited for producing human-readable summary reports. We will examine both approaches in detail, ensuring that data practitioners can confidently choose the appropriate tool for their task.

Method 1: Calculating the Mean for a Single, Specific Column using Aggregation

This approach relies heavily on the powerful [aggregation](#) function, which is typically employed when the requirement is a single, scalar result for a defined column, such as the overall average score. This method is highly efficient, minimizing overhead and providing a streamlined way to extract a specific statistical metric directly. It is the preferred method when the resulting mean value is needed immediately as a Python variable for calculations outside the Spark pipeline.

```
from pyspark.sql import functions as F
```

```
# Calculate the mean of the column named 'game1'  
df.agg(F.mean('game1')).collect()
```

Method 2: Calculating Means for Multiple Columns using Selection Transformation

Conversely, when the objective shifts to calculating averages across several numerical columns and presenting these results within a new, compact DataFrame structure, the use of the `.select()` transformation combined with the `mean()` function is generally the cleaner and more practical choice. This method excels at generating summary tables where multiple averages are needed simultaneously, offering immediate visual insight into the central tendency of various metrics within the dataset.

```
from pyspark.sql.functions import mean
```

```
# Calculate mean for game1, game2, and game3 columns  
df.select(mean(df.game1), mean(df.game2), mean(df.game3)).show()
```

Initializing the Environment and Sample Data

Before proceeding with the actual calculation demonstrations, it is imperative that we establish a functional [PySpark](#) environment and construct a sample [DataFrame](#) to simulate typical structured data analysis scenarios. This foundational preparation involves initializing the `SparkSession`, which serves as the primary entry point for all underlying Spark functionality, and clearly defining both the raw data and the schema that will define our DataFrame's structure. For clarity in these examples, we utilize a simple dataset tracking numerical scores achieved by various teams across three distinct 'games'.

The setup process begins by importing the necessary components from `pyspark.sql`, most critically the [SparkSession](#) object. We then structure our raw input as a standard Python list of lists, where each inner list represents a single row of observations. Following this, we explicitly define the column names (the schema). The final step involves using the `spark.createDataFrame()` method to instantiate the distributed DataFrame object, which is now ready for efficient analytical processing.

The following initialization block executes the necessary setup code. It creates the environment, loads the sample data, and crucially, displays the resulting DataFrame structure. It is essential to confirm that the target columns for mean calculation (`game1`, `game2`, `game3`) are correctly inferred as numerical types, as statistical averages are mathematically valid only on fields containing numeric data.

```
from pyspark.sql import SparkSession  
spark = SparkSession.builder.getOrCreate()
```

```
# Define data representing team scores across three games
```

```

data = ,
,
,
,
,
]

# Define column names for the dataset
columns =

# Create dataframe using the defined data and column names
df = spark.createDataFrame(data, columns)

# View the resulting dataframe structure and content
df.show()

```

```

+-----+-----+-----+-----+
| team|game1|game2|game3|
+-----+-----+-----+
| Mavs| 25| 11| 10|
| Nets| 22| 8| 14|
| Hawks| 14| 22| 10|
| Kings| 30| 22| 35|
| Bulls| 15| 14| 12|
| Blazers| 10| 14| 18|
+-----+-----+-----+

```

Method 1 Deep Dive: Single Column Aggregation

When the analytic requirement dictates calculating the average score exclusively for one column, such as `game1`, the most efficient and direct strategy is the application of the `.agg()` function in conjunction with `functions.mean()` (which is typically aliased as `F.mean()` for brevity and cleaner code). The **aggregation function** applies the specified statistical operation--be it mean, sum, count, or standard deviation--across every row of the designated column, ultimately returning a new, very small DataFrame containing only the calculated result. This approach is superior in terms of performance because it is designed to minimize data shuffling and focus the computation on the specific subset of data required.

The result of executing `df.agg(F.mean('game1'))` is a DataFrame consisting of a single row and a single column. To move beyond the DataFrame abstraction and extract the raw numerical mean

value into a standard Python variable, we must invoke the `.collect()` action. This is a critical step, as `.collect()` triggers the execution of the calculation on the distributed cluster and transfers the resultant data back to the driver program as a list of `Row` objects. Subsequently, the index slicing notation is required to correctly access the first row's data and then the first (and only) element within that row, isolating the scalar mean value.

We execute the following syntax to calculate the precise [mean](#) of values contained solely within the **game1** column of our sample DataFrame:

```
from pyspark.sql import functions as F
```

```
# Calculate mean of column named 'game1'  
df.agg(F.mean('game1')).collect()
```

```
19.333333333333332
```

The calculated mean score for the **game1** column is revealed to be approximately **19.333**. This result can be easily verified using a simple manual calculation: $(25 + 22 + 14 + 30 + 15 + 10) / 6 = 19.333...$ This validation confirms the accuracy of the PySpark function implementation. This methodology remains the standard choice when the final output required is a single numerical metric intended for subsequent programmatic consumption or integration into non-Spark processing pipelines.

Method 2 Deep Dive: Multiple Column Selection

When the analytical objective is to compute the average across several distinct columns--for instance, `game1`, `game2`, and `game3`--and display those aggregated results clearly within a new, consolidated summary table, the `.select()` transformation, paired with the imported `mean` function, provides a far cleaner and more expressive syntax. Crucially, `.select()` differs from `.agg()`: while `.agg()` is designed primarily for dataset reduction, `.select()` allows the user to project a new [DataFrame](#) where the resultant columns are defined by the statistical operations applied to the original columns. This creates a highly readable summary.

This approach yields a result format that is immediately digestible when paired with the `.show()` action. To optimize clarity, we import the specific `mean` function directly from `pyspark.sql.functions`. Within the `.select()` statement, we then explicitly list each column whose mean we wish to calculate. A key point to remember is that PySpark automatically prefixes the resulting column names in the summary [DataFrame](#) with `avg()`, clearly indicating the transformation applied.

We apply the following syntax to calculate and display the mean values simultaneously for the

game1, **game2**, and **game3** columns across the entire DataFrame:

```
from pyspark.sql.functions import mean
```

```
# Calculate mean for game1, game2 and game3 columns
```

```
df.select(mean(df.game1), mean(df.game2), mean(df.game3)).show()
```

```
+-----+-----+-----+
| avg(game1)| avg(game2)| avg(game3)|
+-----+-----+-----+
|19.333333333333332|15.166666666666666| 16.5|
+-----+-----+-----+
```

The resultant summary DataFrame offers an immediate and concise overview of the central tendency for all three analyzed score columns. Based on this output, data analysts can quickly deduce the following averaged scores:

The mean score in the **game1** column is approximately **19.333**.

The mean score in the **game2** column is approximately **15.167**.

The mean score in the **game3** column is exactly **16.5**.

Handling Missing Data and Performance Optimization

When executing statistical calculations within any distributed computing environment, the robust management of missing data--specifically **null values** or NaN entries--is a paramount concern. In [PySpark](#), both the `mean()` function utilized within the `.select()` transformation and the `F.mean()` function employed within the [aggregation](#) method adhere to a crucial default behavior: they will automatically skip and ignore any `null` or NaN values encountered within the column during the mean calculation process. This behavior aligns with best practices in common statistical software, ensuring that missing data points do not erroneously distort the average derived from the valid entries.

It is critical for practitioners to understand this default setting. If a mean calculation is required that necessitates treating **null values** as zero (a form of data imputation), or if the project requires the explicit removal of entire rows containing nulls (known as listwise deletion) before the calculation occurs, then an explicit data cleaning step must be executed beforehand. This typically involves using DataFrame methods such as `.fillna(0)` or `.na.drop()` prior to applying the aggregation function. Neglecting to account for this default null handling can lead to significant discrepancies if automatic data imputation was erroneously expected.

From a performance perspective, both the `.agg()` and `.select(mean(...))` patterns are highly

optimized operations within the PySpark framework. However, a general rule of thumb for large-scale data processing is that when calculating a multitude of statistics (e.g., mean, minimum, maximum, standard deviation) simultaneously across numerous columns, structuring the computation within a single, unified `.agg()` call that includes all desired functions is generally the most efficient pattern. This strategic approach minimizes the number of required passes over the distributed data partitions, thereby drastically reducing I/O operations and network overhead, leading to faster overall execution times. For simpler tasks focusing only on the mean, either method provides excellent, high-speed performance and strong code readability.

Conclusion and Next Steps in PySpark

Mastering statistical aggregation within [PySpark](#) is a fundamental skill essential for effective large-scale data processing and analysis. Whether the immediate need is a single scalar result for a specific column using the high-performance `.agg()` method, or a neatly summarized DataFrame displaying multiple column averages via the expressive `.select()` transformation, PySpark offers clear, highly functional, and optimized methods to achieve the desired result. Data professionals must always be cognizant of the framework's default handling of **null values** to ensure the statistical output is interpreted correctly and accurately reflects the underlying data distribution.

For those dedicated to expanding their [PySpark](#) expertise beyond simple column averages, the robust framework provides deep support for advanced features such as window functions, complex grouping operations (e.g., calculating the mean score per team using `.groupBy()`), and integrated, powerful machine learning capabilities. We strongly encourage further exploration of the official documentation and community resources to unlock the full potential of distributed, big data analysis.

The following resources detail how to perform other common statistical and transformation tasks within the PySpark environment: