

Learn How to Calculate the Median of a Column in PySpark DataFrames

Authored by
Mohammed looti

November 11, 2025

RECOMMENDED CITATION

Mohammed looti (2025). *Learn How to Calculate the Median of a Column in PySpark DataFrames*. PSYCHOLOGICAL STATISTICS. Retrieved from <https://statistics.arabpsychology.com/?p=16530>

The Importance of the Median in Large-Scale Data Processing

The [Median](#) is a fundamental statistical measure integral to effective data analysis, primarily used to ascertain the central tendency of a dataset. Unlike the arithmetic mean, which is highly susceptible to skewing by extreme outliers, the median robustly identifies the exact middle value once a dataset has been sorted. This characteristic makes it a superior measure of typical performance, especially when dealing with data distributions that are not perfectly normal or symmetrical. When handling big data challenges, utilizing specialized, distributed processing tools becomes mandatory. This is where high-performance frameworks like Apache Spark, accessed via [PySpark DataFrames](#), prove indispensable for efficient computation across massive, distributed environments.

PySpark is equipped with a rich suite of highly optimized statistical functions housed within the [pyspark.sql.functions](#) module. These functions are designed to simplify complex statistical calculations, such as calculating the median, eliminating the need for complex, manual implementations. The choice of methodology--whether employing aggregation or selection techniques--depends entirely on the analytical goal: extracting a single scalar median value or generating a tabular summary of medians across multiple columns. Understanding these distinct approaches is vital for writing Spark code that is not only functional but also highly performant and scalable when processing gigabytes or terabytes of data.

In this comprehensive guide, we will meticulously detail two primary, high-performance methods for calculating the median of columns within a PySpark environment. Both techniques rely on native Spark SQL functionalities, ensuring maximum throughput and computational accuracy. We will initiate this exploration by setting up a practical, structured dataset, followed by detailed demonstrations of each method. Furthermore, we will include a rigorous, manual verification step to confirm the precision of the calculated [median](#) results, providing complete confidence in the analytical output.

PySpark Setup: Initializing the Environment and Sample Data

Before we can begin the statistical analysis, it is essential to establish a fully operational PySpark session and construct the sample [DataFrame](#) that will serve as our source data. This prerequisite setup involves initializing the **SparkSession**, which acts as the entry point to Spark functionality, and defining the structure of the data. Our example dataset is designed to simulate performance metrics--specifically, game scores--for several fictional teams, providing numerical columns suitable for immediate statistical computation.

The following code block outlines the standard procedure for initializing the Spark driver and creating a sample DataFrame, named `df`. This DataFrame is structured with one categorical column (`team`) and three numerical columns (`game1`, `game2`, `game3`), representing the scores that

require central tendency analysis. The use of `SparkSession.builder.getOrCreate()` is the recommended practice, ensuring that an existing session is reused or a new one is created only if necessary, thereby optimizing resource usage.

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

```
#define data
```

```
data = ,
```

```
,
```

```
,
```

```
,
```

```
,
```

```
]
```

```
#define column names
```

```
columns =
```

```
#create dataframe using data and column names
```

```
df = spark.createDataFrame(data, columns)
```

```
#view dataframe
```

```
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|
```

```
+-----+-----+-----+-----+
```

Once executed, the resulting DataFrame is immediately available for statistical computation. The structure clearly presents six distinct records and their corresponding scores across the three game metrics. With the data successfully loaded and structured, we can now transition to applying the specialized median calculation techniques, commencing with the optimized method designed for extracting the median from a single, specific column.

Approach 1: Calculating a Single Median Value using df.agg()

For scenarios requiring the calculation of a single aggregate statistic across an entire column, the `df.agg()` function is the most robust and commonly utilized method in PySpark. The `agg()` transformation is highly flexible, allowing analysts to apply one or multiple aggregation functions--such as `median`, `sum`, or `count`--in a single, unified operation across the DataFrame. This approach is highly efficient when the desired output is a singular, scalar representation of the central tendency for a particular metric.

Implementation requires importing the `functions` module, conventionally aliased as `F`, and subsequently passing the `F.median()` function along with the name of the target column to the `agg()` method. It is crucial to remember that `agg()` inherently produces a new, single-row DataFrame containing the result. Therefore, to retrieve the final numerical value for immediate use within subsequent Python logic or reporting, the command chain `.collect()` must be applied. This sequence isolates the scalar value from the DataFrame's structured output.

The following code snippet demonstrates this precise technique, focusing exclusively on calculating the median score for the `game1` column. This technique is specifically optimized for scenarios where memory efficiency is paramount, as it avoids generating large intermediate DataFrames and instead focuses on returning a precise, single statistical measure, making it ideal for integration into streamlined data pipelines.

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

```
#calculate median of column named 'game1'  
df.agg(F.median('game1')).collect()
```

Detailed Walkthrough: Single-Column Median Verification

We now apply the `agg()` method to calculate the median for the **game1** column, leveraging our initialized DataFrame. The syntax requires a clear specification of the PySpark function, `F.median`, and passing the column name as its string argument. The subsequent command, `.collect()`, is the vital final step that extracts the raw numerical result from the structured DataFrame produced by the aggregation operation.

Upon execution, the system yields the median score for the **game1** metric. This calculated value provides an essential benchmark for understanding the middle performance level within the dataset, ensuring the baseline measure is not disproportionately affected by unusually high or low individual scores. The resulting precise median value for the six scores recorded in `game1` is **18.5**.

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

```
#calculate median of column named 'game1'  
df.agg(F.median('game1')).collect()
```

18.5

To solidify our understanding and confirm the computational accuracy, we manually verify this result. The scores contained within the **game1** column are (25, 22, 14, 30, 15, 10). To determine the [median](#), these values must first be arranged in ascending order: 10, 14, **15, 22**, 25, 30. Since the dataset contains an even number of observations (six records), the median is correctly derived by averaging the two innermost values, 15 and 22. The exact calculation is $(15 + 22) / 2 = \mathbf{18.5}$. This manual computation confirms the flawless accuracy of the PySpark aggregate function.

Approach 2: Calculating Medians Across Multiple Columns using df.select()

While the [df.agg\(\)](#) method is perfect for obtaining a single, scalar output, real-world data analysis frequently demands the simultaneous calculation of medians for multiple columns, with the results presented in a clear, comparative table. For this specific requirement, the optimal method involves leveraging the [df.select\(\)](#) function, paired with the direct import of the `median` function from `pyspark.sql.functions`.

The `select()` method serves to transform the DataFrame structure by applying specific functions to designated columns and returning a new DataFrame that contains only the results of those transformations. By iteratively passing multiple instances of `median(df.column_name)` to the `select()` function, we instruct Spark to calculate the median for every specified column across the entire dataset. This operation efficiently produces a single-row DataFrame where each column header corresponds to the original metric, and the cell contains its calculated median.

This technique offers exceptional utility for comparative analysis, as it neatly organizes all required central tendency measures into a single, easily interpretable output table. Analysts can thus rapidly compare the middle values of different metrics--such as **game1**, **game2**, and **game3**--without the need for executing and managing multiple, separate aggregation queries.

```
from pyspark.sql.functions import median
```

```
#calculate median for game1, game2 and game3 columns  
df.select(median(df.game1), median(df.game2), median(df.game3)).show()
```

Example 2: Multiple Column Median Calculation (Detailed Walkthrough)

Executing the `select()` approach to determine the medians for **game1**, **game2**, and **game3**

results in PySpark generating a new DataFrame. This output clearly maps the calculated [median](#) back to its original column name, providing a much more structured and immediately reportable result compared to the single scalar extracted via the `agg()` method.

The resulting single-row table provides an immediate statistical overview of the central performance level across all three games. The output confirms that `game1` has a significantly higher median score (18.5) compared to `game2` and `game3`, which exhibit lower central tendencies. This swift comparative analysis aids data scientists in understanding the general distribution and potential difficulty differences across the various metrics represented in the source dataset.

from pyspark.sql.functions import median

```
#calculate median for game1, game2 and game3 columns
df.select(median(df.game1), median(df.game2), median(df.game3)).show()
```

```
+-----+-----+-----+
|median(game1)|median(game2)|median(game3)|
+-----+-----+-----+
| 18.5| 14.0| 13.0|
+-----+-----+-----+
```

The calculated results explicitly establish the following central measures for each game score column:

The [median](#) of values in the **game1** column is **18.5**.

The median of values in the **game2** column is **14.0**.

The median of values in the **game3** column is **13.0**.

An essential feature of using the built-in [median](#) function in PySpark is its intrinsic ability to manage missing data. Regardless of whether the `agg()` or `select()` method is employed, the function is specifically engineered to automatically disregard any records containing **null values** during the calculation process. This robust behavior guarantees that the final median calculation is derived solely from valid, non-null data points, thereby preventing inaccurate measurements of central tendency that could result from data sparsity or corruption.

Advanced Considerations: Null Handling and Performance Optimization

In the context of production-level data engineering, encountering missing values (nulls) is inevitable, and PySpark's handling of these values is paramount for accurate statistical outcomes. As discussed, the native `median` function within `pyspark.sql.functions` inherently excludes nulls from the computation. This default behavior is typically advantageous, as it avoids the

introduction of bias that might arise from arbitrary imputation, such as replacing nulls with zero. However, if the analytical requirements dictate a specific imputation strategy--for example, replacing nulls with the column's mean or a previous observation--that data preparation step must be rigorously executed **prior** to initiating the median aggregation.

Regarding performance, both the [df.agg\(\)](#) and [df.select\(\)](#) methods, when utilizing Spark's native functions, are highly optimized for distributed execution. PySpark leverages the Catalyst Optimizer to translate these high-level Python commands into highly efficient physical execution plans on the underlying Java Virtual Machine (JVM). For datasets of substantial size, relying on these built-in functions is drastically more performant than implementing median calculations through custom User Defined Functions (UDFs) or attempting to retrieve and process the entire dataset on the driver node, which invariably leads to severe performance degradation and potential out-of-memory exceptions.

For extremely large datasets, particularly those where latency is a concern, it is worth noting Spark's capability for approximate quantile calculations. While the standard `F.median()` function typically strives for an exact result, analysts should be aware of related tools like `F.approx_quantile()`. This function allows for a calculated trade-off, sacrificing a minor degree of mathematical accuracy in exchange for substantial performance improvements. Choosing the exact `F.median()` function, as demonstrated throughout this guide, guarantees absolute precision, a necessity unless performance challenges on truly massive scales necessitate the use of approximation techniques.

Summary and Next Steps in PySpark Statistical Analysis

Mastering the calculation of the median is fundamental to data exploration and statistical reporting within a big data environment. PySpark offers two powerful and highly efficient methods to fulfill this requirement: the `df.agg()` method, which is best suited for extracting a single, scalar median value, and the `df.select()` function, which is ideal for performing simultaneous calculations across multiple metrics and presenting the results in a structured, tabular format for immediate comparison.

By consistently utilizing the native functions available within `pyspark.sql.functions`, data engineers and analysts ensure that their processing remains fully scalable and optimally efficient within the distributed computing architecture of Apache Spark. This practice maximizes performance and leverages the full power of Spark's underlying optimization engines.

To further expand your proficiency in data aggregation and statistical computation using the PySpark framework, consider exploring the following related functionalities and tutorials:

Calculating the Mean (Average) and other Moments of Distribution in PySpark DataFrames.

Understanding and Implementing Advanced PySpark Window Functions for row-level analysis.
Handling Null Values and Data Imputation Strategies in PySpark for robust data cleaning.