

Learning to Calculate Sample and Population Variance with Python

Authored by
Mohammed looti

November 2, 2025

RECOMMENDED CITATION

Mohammed looti (2025). *Learning to Calculate Sample and Population Variance with Python*. PSYCHOLOGICAL STATISTICS. Retrieved from <https://statistics.arabpsychology.com/?p=8699>

Understanding the spread or dispersion of data points is arguably the most fundamental concept in modern [statistics](#) and advanced data analysis. The primary quantitative measure used to capture this dispersion is the [variance](#). It offers indispensable insight into how individual data points deviate from the central tendency, specifically the arithmetic mean.

While frequently associated with standard deviation (its positive square root), **variance** serves a critical, foundational role in statistical inference, robust hypothesis testing, and the construction of reliable machine learning models. Crucially, calculating this measure accurately demands a careful distinction based on whether the dataset under examination encompasses the entire collective group (the **population**) or is merely a representative subset (a **sample**) drawn from that larger population.

This comprehensive guide delves into the precise mathematical requirements for calculating both population and sample variance. Furthermore, we will illustrate how to implement these calculations seamlessly and efficiently using Python's powerful, built-in [statistics library](#), providing practical examples for data scientists and analysts.

Defining and Calculating Population Variance

The population variance, mathematically symbolized by the Greek letter sigma squared (σ^2), quantifies the average squared deviation of every data point from the population mean (μ). This specific calculation is only appropriate when the dataset being scrutinized includes every single member of the group of interest--the complete **population**. Since we possess all possible observations, this calculation yields an exact, deterministic value for the true dispersion.

Because the dataset is exhaustive, there is no need for estimation or correction; the resulting variance is a definitive measure of how widely the population data is scattered around its mean. Recognizing when a dataset constitutes a complete population versus a sample is the prerequisite step before applying this statistic.

The standard formula used for calculating the **population variance** (σ^2) is defined as follows:

$$\sigma^2 = \sum (x_i - \mu)^2 / N$$

To ensure clarity, here is a detailed breakdown of the mathematical symbols inherent in the population variance calculation:

Σ : This Greek capital letter denotes the mathematical operation of summation, requiring the addition of all preceding terms.

μ : This symbol represents the [Population mean](#), which is the arithmetic average calculated across

all elements within the population.

x_i : This term denotes the i th element or individual observation recorded within the comprehensive population dataset.

N : This factor represents the total count of elements, known precisely as the **Population size**.

The Necessity of Sample Variance and Bessel's Correction

In contrast to the population calculation, **sample variance** (s^2) is employed when an analyst is only able to examine a subset (a **sample**) drawn from a much larger, often unknowable, population. Since a sample inherently fails to capture the full spectrum of variability present in the entire population, the sample variance serves a crucial role as an **estimator** of the true population variance. This estimation process introduces statistical complexities that must be addressed.

If we were to incorrectly use the population variance formula (dividing by n) on a limited sample, the resulting variance would consistently underestimate the true spread of the population. This systematic bias occurs because the sample mean tends to be closer to the data points within the sample than the true population mean would be. To correct for this inherent underestimation and provide an accurate, unbiased estimate, statisticians apply a technique known as [Bessel's correction](#).

This critical correction involves replacing the simple sample size (n) in the denominator with the degrees of freedom, defined as $(n-1)$. This adjustment ensures that the sample variance is deemed an unbiased estimator of the population variance, thereby providing a statistically reliable measure of dispersion even when working with incomplete data. The formula is structured specifically to account for the uncertainty introduced by sampling.

The accepted standard formula for calculating **sample variance** (s^2) is defined below:

$$s^2 = \frac{\sum (x_i - \bar{x})^2}{(n-1)}$$

The variables utilized specifically within the sample variance formula are defined as follows:

$\bar{x}$: This symbol represents the **Sample mean**, which is the calculated average of all observations contained within the specific sample.

x_i : This denotes the i th element, or individual observation, retrieved from the limited sample dataset.

n : This represents the total count of elements in the sample, commonly referred to as the **Sample size**.

$(n-1)$: This is the degrees of freedom, which critically implements [Bessel's correction](#) to yield an unbiased estimator.

Leveraging Python's Standard Statistics Module

While a theoretical understanding of manual variance calculation is essential for statistical literacy, practical data science workflows necessitate the use of optimized, established libraries. Python's robust standard library provides the excellent `statistics` module, which includes specialized, dedicated functions designed to calculate both types of variance with efficiency and precision.

For calculating the sample variance (which uses the $n-1$ denominator), we utilize the standard `variance()` function. Conversely, when calculating the population variance (which uses the N denominator), we employ the `pvariance()` function--the 'p' explicitly denoting the population context. Both functions accept any iterable dataset, such as a list or tuple, and return the calculated variance value based on the statistical assumption made.

It is paramount for accurate scripting that analysts explicitly import and utilize the correct function based on whether their data is assumed to be a population or a sample. The following basic usage example illustrates how to import and structure the functions for computation:

```
from statistics import variance, pvariance
```

```
# Define a hypothetical dataset 'x'
```

```
# x =
```

```
# calculate sample variance (divides by n-1)
```

```
variance(x)
```

```
# calculate population variance (divides by N)
```

```
pvariance(x)
```

Practical Implementation: Sample vs. Population Variance in Python

To truly grasp the distinction between these two statistical measures, we will apply both functions to an identical dataset. These detailed examples demonstrate how the resulting outputs differ solely based on the statistical assumption made--that is, whether the data is treated as a limited sample or as a complete population.

In the first scenario, we assume our dataset, `data`, represents a **sample** drawn from a much larger collection of possible observations. Consequently, we must utilize the `variance()` function. This function automatically implements Bessel's correction by dividing the sum of squared deviations by $(n-1)$, providing an unbiased estimate of the true population [variance](#).

The code below initializes the sample data and applies the appropriate function. Notice the

calculated result, 22.067, which is slightly inflated compared to the population calculation, reflecting the statistical uncertainty inherent in using a subset of data.

from statistics import variance

```
# define data (assumed to be a sample)
data =

# calculate sample variance (divides by n-1)
variance(data)

22.067
```

For the second scenario, we treat the exact same array, `data`, as if it constitutes the **entire population** of interest. Because we have certainty that all possible data points are included, we use the `pvariance()` function. This function divides by the total number of observations (N), discarding the degrees of freedom correction.

The calculation reflects this change in assumption. The resulting population variance is determined to be **20.596**. This value is consistently smaller than the sample variance calculated previously, illustrating the mathematical effect of removing the uncertainty adjustment required for incomplete data sets.

from statistics import pvariance

```
# define data (assumed to be the entire population)
data =

# calculate population variance (divides by N)
pvariance(data)

20.596
```

Contextual Decision Making: Choosing the Right Variance

The choice between using sample variance or population variance is fundamentally contextual, not mathematical. It hinges entirely upon the scope of your data collection, the integrity of your dataset, and the statistical inferences you intend to draw. Selecting the appropriate calculation is paramount for maintaining the integrity and validity of subsequent statistical tests and model training.

The most common error in variance calculation is misidentifying the data source. If you mistakenly

treat a sample as a population, you introduce bias by systematically underestimating the true spread. Conversely, if you treat a known population as a sample, you unnecessarily inflate the variance, potentially leading to overly conservative statistical conclusions.

Keep these critical guidelines in mind when analyzing your dataset to ensure you calculate the appropriate measure of spread:

You must calculate the **population variance** when the dataset comprehensively includes every value of interest--the complete **population**. Typical examples include analyzing the annual sales figures of all divisions within a corporation or the complete census data for a small town.

You must calculate the **sample variance** when the dataset represents a subset (a **sample**) taken from a larger, often inaccessible, population. This approach is standard practice in survey research, quality control testing, and large-scale social science experiments.

It is a statistical truth that, for any given array of numbers, the sample variance will nearly always be mathematically larger than the population variance. This occurs because dividing by $(n-1)$ instead of n results in a larger quotient, reflecting the statistical necessity to provide an unbiased estimate of the true population [variance](#) when working with limited observations.

Expanding Beyond Variance: Other Measures of Spread

Variance provides a powerful, foundational measure of data dispersion, but it is rarely used in isolation. For a complete and robust statistical picture of data spread, analysts frequently need to examine other related metrics. Python continues to offer powerful tools for calculating these measures quickly and reliably through its standard library and specialized packages like NumPy and Pandas.

To further enrich your understanding of data distribution and variability, consider exploring these complementary measures of spread, which are crucial for comprehensive data analysis:

[Calculating Standard Deviation in Python](#): As the square root of the variance, standard deviation returns the measure of spread back into the original units of the data, making it highly interpretable.

[Determining Interquartile Range \(IQR\)](#): This measure focuses on the middle 50% of the data, offering a robust measure of spread that is significantly less sensitive to extreme outliers than variance or standard deviation.

[Exploring the Coefficient of Variation](#): This is a normalized measure of dispersion, particularly useful for comparing the variability between two different datasets that may have vastly different units or scales.