

What is the difference between covariance and variance?

Authored by
stats writer

December 7, 2025

RECOMMENDED CITATION

stats writer (2025). *What is the difference between covariance and variance?*.

PSYCHOLOGICAL SCALES. Retrieved from <https://scales.arabpsychology.com/?p=106569>

In the expansive field of statistics, two fundamental measures—**variance** and **covariance**—are essential tools for analysts and researchers. While they sound phonetically similar and are derived from related mathematical principles, their applications and interpretations are distinctly different. Understanding the precise definition and function of each measure is crucial for accurately analyzing data distribution and relationships between variables.

The core distinction lies in their scope: **Variance** focuses on a single variable, quantifying the internal spread or dispersion of values within a given dataset. It is a measure of how far individual data points deviate from the mean. Conversely, **Covariance** is inherently bivariate; it assesses the directional relationship between two separate variables, indicating whether they tend to increase or decrease together.

This comprehensive guide provides an in-depth exploration of both concepts, detailing their formal definitions, deriving their computational formulas, and illustrating their practical application through numerical examples. By the end, readers will possess a robust understanding of when and how to appropriately use variance versus covariance in quantitative analysis.

Understanding Variance: The Measure of Dispersion

Variance serves as the cornerstone for quantifying the variability or dispersion within a single set of observations. It measures how spread out values are in a given dataset relative to their central tendency, specifically the mean. A small variance indicates that data points cluster tightly around the mean, suggesting high consistency, while a large variance implies that data points are widely scattered, indicating greater volatility or spread. This measure is fundamental because it provides the mathematical basis for many subsequent statistical tests and concepts, such as standard deviation and hypothesis testing.

To calculate variance, we first determine the difference between each individual observation and the mean of the dataset. These differences, known as residuals, are then squared. Squaring the differences serves two critical purposes: first, it eliminates negative signs, ensuring that deviations in both directions (above and below the mean) contribute positively to the measure of spread; second, it heavily penalizes larger deviations, making the variance highly sensitive to outliers. The sum of these squared differences is then divided by the number of observations (or $n-1$ for a sample) to yield the average squared deviation.

It is vital to distinguish between population variance (σ^2) and sample variance (s^2). When working with a complete population, the denominator used is N (the population size). However, in practical statistics, we almost always deal with samples drawn from a larger population. In this case, we use $n-1$ in the denominator—a technique known as Bessel's correction—to provide an unbiased estimator of the true population variance. Failure to apply this correction when analyzing a sample would systematically underestimate the population variability.

Variance Formula and Components

The formula used to find the variance of a **sample** (denoted as **s²**) is structured to capture the average squared distance from the central point:

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

The components of this formula are defined precisely to ensure mathematical accuracy in measuring dispersion:

x: Represents the **sample mean**, which is the arithmetic average of all observations in the dataset. It serves as the central anchor against which deviations are measured.

x_i: Denotes the *i*th observation in the sample. This represents each individual data point being analyzed.

N: Specifies the total sample size, meaning the total count of observations in the dataset.

Σ: This is the Greek capital letter Sigma, signifying the operation of summation. It instructs us to sum all the preceding squared deviations.

Understanding the structure of this formula reveals that **variance** is measured in units squared relative to the original data units. For instance, if the original data represents weight in kilograms (kg), the variance will be expressed in kilograms squared (kg²). This dimensional discrepancy is often why the square root of the variance, known as the **Standard Deviation**, is preferred for practical interpretation, as it returns the measure of spread back into the original units.

Illustrative Example of Sample Variance

Consider two distinct datasets, each containing 10 values, to demonstrate how variance reflects the underlying spread of the data.

Dataset 1: 6, 7, 10, 13, 14, 14, 18, 19, 22, 24

Using a calculator, we can find that the sample variance for this dataset is **36.678**. This numerical result quantifies the average squared distance of these data points from their mean.

Now suppose we had another dataset with 10 values that are more dispersed:

Dataset 2: 6, 13, 19, 24, 25, 30, 36, 43, 49, 55

The variance calculated for the second dataset is substantially larger than the first. This difference in variance values provides immediate insight: the values in the second **dataset** are far more spread out and exhibit higher volatility compared to the tighter clustering observed in the first dataset. This comparative analysis demonstrates the power of variance in quantifying data

heterogeneity.

Defining Covariance: Measuring Relationship Direction

In contrast to variance, which is univariate, **Covariance** is a bivariate measure that quantifies the extent to which two random variables change together. It specifically assesses the nature of the linear relationship between Variable X and Variable Y. It indicates the direction of the relationship—whether they move in tandem (positive covariance) or in opposite directions (negative covariance).

The computation of covariance involves comparing the deviation of X from its mean against the simultaneous deviation of Y from its mean. If, for a given observation, both X and Y are simultaneously above their respective means (positive residuals) or simultaneously below their respective means (negative residuals), their product will be positive, contributing to a positive overall covariance. Conversely, if one variable is above its mean and the other is below, the product is negative, contributing to a negative covariance.

It is important to emphasize that covariance only measures the direction of the linear relationship, not its strength. The magnitude of the covariance value is dependent on the units of the variables involved, making it difficult to compare covariance values across different pairs of variables. This limitation often necessitates the use of the derived metric, the correlation coefficient, for standardized strength assessment.

Covariance Formula and Interpretation

The formula used to find the covariance between two variables, X and Y, involves calculating the average of the cross-products of the deviations from their respective means:

$$\text{COV}(X, Y) = \Sigma(x_i - \bar{x})(y_i - \bar{y}) / n$$

The elements used in this calculation are critical for determining the nature of the association:

x: The **sample mean** of variable X.

x_i: The *i*th observation of variable X.

y: The **sample mean** of variable Y.

y_i: The *i*th observation of variable Y.

n: The total number of pairwise observations, ensuring that each data point for X is correctly paired with a corresponding data point for Y.

Σ: The summation symbol, indicating the aggregate sum of all cross-products.

The interpretation of the **covariance** result hinges entirely on its sign. A **Positive Covariance** implies a direct relationship: as X increases, Y tends to increase, and vice versa. A **Negative**

Covariance implies an inverse relationship: as X increases, Y tends to decrease. A covariance near zero suggests little to no linear relationship between the variables, although they may still be related in a non-linear fashion.

Practical Examples of Covariance Direction

To solidify the understanding of covariance, let us examine two scenarios demonstrating positive and negative relationships.

Suppose we have the following dataset with 10 paired values for X and Y :

X	Y
3	6
5	7
6	7
6	13
7	16
8	15
12	17
14	20
15	24
19	27

Using a calculator, we can find that the covariance between X and Y is **31.8**. Since this value is positive, it tells us that as the values for X increase, the values for Y tend to increase as well, indicating a strong positive linear association.

Now suppose we had another dataset with 10 paired values showing an inverse relationship:

X	Y
3	28
5	25
6	24
6	25
7	19
8	15
12	13
14	15
15	4
19	3

Using a calculator, we can find that the covariance between X and Y is **-38.55**. Since this value is negative, it tells us that as the values for X increase, the values for Y tend to decrease, demonstrating an inverse linear association.

The Fundamental Difference: Univariate vs. Bivariate

The most profound conceptual difference between **variance** and **covariance** is their scope of application. Variance is inherently a **univariate** statistic; it requires only one variable and one **dataset** to be computed, focusing solely on internal data spread. It answers the question: "How dispersed is this set of numbers?" It provides a metric for the risk or uncertainty associated with a single quantity.

Covariance, however, is a **bivariate** statistic, meaning it requires two variables (X and Y) to be calculated. Its purpose is relational: it examines the joint variability of two variables. It answers the question: "How does the movement of X correspond to the movement of Y?" This relational focus makes covariance crucial in fields like finance (measuring portfolio risk) and social **statistics** (analyzing demographic correlations).

Furthermore, variance can be considered a special case of covariance. If you calculate the covariance of a variable with itself, $\text{COV}(X, X)$, the result is the variance of X, $\text{VAR}(X)$. This mathematical identity reinforces that variance is the measure of how a single variable relates to its own mean deviation, establishing a strong link between the two concepts despite their distinct usage scenarios.

Applications and Limitations in Statistical Analysis

Understanding when and where to deploy these metrics is crucial for effective data analysis. **Variance** is indispensable for risk assessment and quality control. In manufacturing, variance helps determine the consistency of product measurements; in finance, the variance of stock returns quantifies investment volatility. The higher the variance, the more spread out values the values are, and the measure can range from zero (no spread at all) to any number greater than zero.

Covariance is primarily utilized in multivariate analysis when we want to quantify the relationship between two variables. A positive value for covariance indicates a positive relationship between two variables, while a negative value indicates a negative relationship between two variables. Its practical application is limited by its non-standardized nature, meaning its magnitude is dependent on the units of the variables involved.

To overcome the scaling issue inherent in covariance, analysts typically normalize the covariance by dividing it by the product of the standard deviations of the two variables. This process yields the **Pearson correlation coefficient** (correlation). Correlation standardizes the measure of association, resulting in a value between -1 and +1, which provides a unit-free measure of the strength and direction of the linear relationship, offering a superior tool for comparison across different datasets and for modeling complex relationships in advanced statistics.