

Advantages & Disadvantages of Using Standard Deviation

Authored by
stats writer

November 18, 2025

RECOMMENDED CITATION

stats writer (2025). *Advantages & Disadvantages of Using Standard Deviation*.
PSYCHOLOGICAL SCALES. Retrieved from <https://scales.arabpsychology.com/?p=95299>

Understanding Standard Deviation

The standard deviation (SD) is one of the most fundamental measures in statistics. It quantifies the amount of variation or dispersion within a set of values. Specifically, it measures the typical deviation of individual data points from the dataset's central tendency, often represented by the sample mean. A low standard deviation indicates that the data points tend to be tightly clustered around the mean, while a high standard deviation suggests that the data points are spread out over a wider range of values.

To fully appreciate its utility and limitations, it is essential to understand the underlying mathematical framework. For a sample dataset, the formula for the standard deviation (denoted as s) is calculated as the square root of the variance. This calculation incorporates differences between each observation and the mean, squaring those differences to eliminate negative values before averaging them:

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

This formula relies on several key statistical components:

Σ : The summation symbol, indicating the sum of all calculated values.

x_i : Represents the i th individual observation in the dataset.

$\bar{x}$: Denotes the sample mean (Link 2/5 for mean), which is the arithmetic average of the dataset.

n : Represents the total size of the sample, where $n - 1$ is used for calculating the sample standard deviation to ensure an unbiased estimate of the population standard deviation.

Overview of Key Advantages and Disadvantages

The standard deviation is a cornerstone of statistical analysis, yet its application involves specific strengths and weaknesses. Selecting the appropriate measure of dispersion for any statistical challenge requires a clear understanding of these properties.

The primary advantages of using the standard deviation (Link 2/5 for SD) center around its computational integrity and its intuitive output:

Advantage #1: Comprehensive Data Utilization. The methodology for calculating SD ensures that every observation contributes to the final measure of spread.

Advantage #2: Interpretive Clarity. The result is expressed in the original units of measurement, making it easy to understand the typical distance from the mean.

Conversely, the main limitation of the standard deviation is a direct consequence of the squaring operation in its formula:

Disadvantage #1: Sensitivity to Outliers. Extreme values, or outliers ([Link 1/5 for outlier](#)), exert a highly disproportionate influence on the SD, potentially leading to a misleading representation of the data's central spread.

Advantage #1: The Standard Deviation Utilizes All Observations

A significant strength of the standard deviation is its highly comprehensive approach to data analysis. Unlike simpler measures of spread that might rely only on two data points (like the range) or the central 50% (like the Interquartile Range), the SD calculation involves processing every single observation in the dataset. This is statistically favorable because using all available data ensures that the calculation is based on the full information content of the sample, providing a richer and generally more reliable measure of variability.

Consider an initial dataset representing exam scores from a small college class:

Scores: 68, 70, 71, 75, 78, 82, 83, 83, 85, 90, 91, 91, 92

Calculating the descriptive statistics for this set reveals that the **sample standard deviation is approximately 8.46**. This single value is the result of measuring the squared distance of every one of the thirteen scores from the sample mean ([Link 3/5 for mean](#)), thereby guaranteeing that even the scores at the absolute extremes contribute to the final summary of dispersion.

To highlight the importance of this comprehensive data utilization, we can compare the standard deviation's response to that of the Interquartile Range (IQR) ([Link 1/5 for IQR](#)). The IQR measures the spread of the middle 50% of the data by calculating the difference between the third quartile (Q3) and the first quartile (Q1), effectively disregarding the lower and upper 25% of the values. While the IQR serves as an excellent measure of spread when data is skewed, it inherently overlooks information contained in the dataset's tails.

Let us modify the initial dataset by changing only the lowest score to an extreme low value, introducing a severe outlier:

Scores: 22, 70, 71, 75, 78, 82, 83, 83, 85, 90, 91, 91, 92

Upon recalculating the standard deviation for this new set, we observe that the **sample standard deviation jumps significantly to 18.37**. This dramatic increase accurately reflects the increased variability caused by the new extreme value (22). However, **the interquartile range remains 17.5**, unchanged, because the middle 50% of values were completely unaffected by the score change at the lower end. This comparison decisively illustrates that the standard deviation is sensitive to, and incorporates, all observations in its calculation, whereas other measures may not.

Advantage #2: The Standard Deviation Offers Clear Interpretation

Beyond its computational completeness, the [standard deviation](#) (Link 3/5 for SD) is highly favored in practical applications for its intuitive and unambiguous interpretation. SD is always expressed in the same units as the original data, meaning the resulting value directly represents the typical distance or magnitude of deviation an observation has from the central [mean](#) (Link 4/5 for mean). This direct relationship simplifies communication and makes the statistic readily accessible, aiding quick decision-making based on statistical output.

If we revisit our first set of exam scores, where the calculated sample standard deviation was **8.46**, the interpretation is straightforward: a typical student's exam score deviates by approximately 8.46 points from the class average. If the data had measured product defects, the standard deviation would be in the number of defects. This adherence to the original scale provides immediate context and prevents misinterpretation of the magnitude of spread.

By contrast, other measures of dispersion designed for relative comparisons can be less intuitive when assessed in isolation. A key example is the [Coefficient of Variation](#) (CV) (Link 1/5 for CV). The CV provides a standardized, unit-less measure of dispersion, calculated as the ratio of the standard deviation (s) to the [sample mean](#) ($\bar{x}$) (Link 5/5 for mean):

Coefficient of Variation: $s / \bar{x}$

In our example, the mean exam score is 81.46. The coefficient of variation is thus $8.46 / 81.46$, yielding a value of **0.104** (or 10.4%). While the CV is invaluable for comparing the relative volatility between two entirely different datasets--such as comparing the spread of salaries in dollars versus the spread of ages in years--the raw value of 0.104 does not inherently convey the absolute magnitude of score dispersion within the class as clearly as 8.46 points does. For descriptive statistics aimed at summarizing a single dataset, the direct interpretability of the standard deviation is generally superior.

Disadvantage #1: High Sensitivity to Outliers and Extreme Values

The most significant drawback associated with the use of the [standard deviation](#) (Link 4/5 for SD) is its tendency to be disproportionately affected by extreme values. As previously noted, the calculation requires squaring the difference between each data point and the mean. When an [outlier](#) (Link 2/5 for outlier) is far from the mean, this large distance is squared, resulting in an exponentially greater contribution to the overall variance and, consequently, a substantially inflated standard deviation. This inflation can render the standard deviation meaningless as a descriptor of the "typical" spread among the core data points.

To demonstrate this sensitivity, consider a dataset tracking the annual salaries (in thousands of

dollars) of 10 employees at a company, all within a reasonable range:

Salaries (in thousands): 44, 48, 57, 68, 70, 71, 73, 79, 84, 94

Given this distribution, the salaries are relatively consistent, and the calculated **sample standard deviation is approximately 15.57** thousand dollars. This SD value accurately conveys the measure of variability present in this homogeneous group.

Now, suppose we introduce one single, highly extreme outlier (Link 3/5 for outlier) by replacing the largest salary with a massive executive compensation, while keeping the other nine salaries the same:

Salaries (in thousands): 44, 48, 57, 68, 70, 71, 73, 79, 84, 895

The effect of this one extreme value is profound. The calculated **sample standard deviation for this new dataset explodes to approximately 262.47** thousand dollars. This dramatically inflated SD value is now completely misleading, failing to represent the actual variability experienced by the nine employees who earn between 44k and 84k. The standard deviation is heavily pulled toward the extreme value, distorting the perception of the data's true dispersion.

Consequently, when datasets are known to contain or are suspected of containing outliers, statisticians often turn to robust measures of dispersion, such as the Interquartile Range (IQR) (Link 2/5 for IQR). Since the IQR focuses only on the central quartiles, it is unaffected by these extreme values, offering a much more stable and reliable measure of spread in heavily skewed or outlier-ridden distributions.

Conclusion and Further Study

The standard deviation (Link 5/5 for SD) remains an indispensable tool in statistics, valued for its rigorous use of all data points and its highly interpretable output in the context of the original units. It provides a powerful absolute measure of data dispersion.

However, analysts must remain judicious regarding its primary weakness: its profound sensitivity to outliers (Link 4/5 for outlier). When data quality is uncertain or distributions are highly skewed, alternative robust measures like the Interquartile Range (IQR) (Link 3/5 for IQR) should be utilized to ensure that statistical conclusions accurately reflect the typical spread of the data.

For those interested in deepening their understanding of variability and statistical measures, the following tutorials provide additional information about using the standard deviation and related concepts in statistics:

Exploring Variance and Mean Absolute Deviation as complementary measures of spread.

A comparative analysis of SD versus the Coefficient of Variation (Link 2/5 for CV) for comparing data sets.

Best practices for identifying and handling outliers (Link 5/5 for outlier) in large statistical datasets.

ARABPSYCHOLOGY.COM