

How to Normalize NumPy Arrays for Machine Learning: A Step-by-Step Guide

Authored by
stats writer

December 2, 2025

RECOMMENDED CITATION

stats writer (2025). *How to Normalize NumPy Arrays for Machine Learning: A Step-by-Step Guide*. PSYCHOLOGICAL SCALES. Retrieved from <https://scales.arabpsychology.com/?p=103555>

In the realm of data science and machine learning, data normalization is a fundamental technique used to standardize the range of independent variables or features. When working with numerical data structures in Python, specifically those handled by the NumPy library, normalizing a matrix involves scaling its elements so that the vector associated with each row or column achieves a unit norm. This transformation is critical for ensuring that feature magnitudes do not skew model training, particularly in algorithms that rely on distance calculations, such as K-Nearest Neighbors or Support Vector Machines.

The core objective of normalization is to guarantee that all input features contribute equally to the distance metrics and similarity measures used by the model. If raw data features possess vastly different scales--for instance, housing prices ranging from thousands to millions--the features with larger scales would disproportionately influence the final results. By normalizing the data, typically so that the values of each vector sum to 1 (using the L1 norm) or so that the length of the vector is 1 (using the Euclidean norm), we mitigate this potential bias and improve convergence speed and model stability. We will explore how to achieve this crucial data preprocessing step using Python and NumPy, focusing on practical examples provided by the powerful scikit-learn library.

To **normalize** a matrix, or any data vector, means to rescale the values such that they conform to a specified distribution or standard scale. While normalization often implies bringing values between 0 and 1 (Min-Max scaling), in the context of vector normalization using norms, it means scaling the vector so its length (magnitude) equals one. This process ensures dimensional consistency and is a mandatory step before feeding data into many complex machine learning architectures, providing mathematical coherence necessary for reliable predictions and interpretations.

Understanding Vector Normalization and Norm Types

When we discuss normalizing a NumPy matrix, we are essentially normalizing the individual vectors contained within that matrix, which can be defined either by the rows or the columns. Each row or column represents a distinct data observation or feature set, respectively. The mathematical definition of normalization dictates that after the transformation, the magnitude or length of the resulting vector should be equal to one, ensuring a standardized representation irrespective of the original scale of the data.

The choice of normalization technique depends heavily on the downstream task. Two of the most common norms used for this purpose are the L1 norm (Manhattan distance) and the L2 norm (Euclidean distance). The L1 norm calculates the sum of the absolute values of the vector elements, while the L2 norm calculates the square root of the sum of the squared vector elements. When normalizing using the L1 norm, the resulting vector elements will sum up to 1, which is useful when the relative contribution of each feature is important. Conversely, L2 normalization is

beneficial when you want to minimize the impact of outliers and calculate distances in a Euclidean space.

For most classification and clustering problems where data sparsity is not a concern, L2 normalization is often the default choice. However, in applications like text classification, where feature vectors are often sparse, L1 normalization is frequently preferred because it results in a normalized vector whose elements represent the proportion of the total sum, making it highly interpretable in terms of probability distributions or frequency weights. Understanding the implications of choosing between the L1 and L2 norms is vital for effective data preprocessing.

The Standard Method: Leveraging scikit-learn for Efficient Normalization

While normalization can be performed manually using NumPy operations (dividing elements by the calculated norm), the industry standard, and most computationally efficient method, involves using the specialized tools provided by scikit-learn (or `sklearn`). Specifically, the `normalize` function within the `sklearn.preprocessing` module is designed to quickly apply vector normalization to input matrices, handling large datasets with optimized performance.

The simplicity of the `normalize` function lies in its parameters, which allow precise control over the normalization dimension and the type of norm applied. The key parameters are `X` (the input NumPy array or matrix), `axis` (defining whether normalization occurs along rows or columns), and `norm` (specifying the mathematical norm to use, such as 'l1', 'l2', or 'max'). For general purposes, using this library drastically reduces the likelihood of introducing mathematical errors compared to manual implementations.

The fundamental syntax for utilizing this powerful preprocessing tool is demonstrated below. Note the use of the `axis` parameter: `axis=1` instructs the function to normalize each row independently (row vectors), while `axis=0` dictates normalization across columns (column vectors). For the following examples, we will prioritize the l1 norm, which ensures that the sum of the absolute values of the elements in the normalized vector equals one.

```
from sklearn.preprocessing import normalize
```

```
# normalize rows of matrix (L1 norm applied row-wise)
normalize(x, axis=1, norm='l1')
```

```
# normalize columns of matrix (L1 norm applied column-wise)
normalize(x, axis=0, norm='l1')
```

The subsequent sections provide detailed walkthroughs demonstrating how to initialize a sample NumPy matrix and apply the L1 normalization technique along both axes, illustrating the practical

effects of these crucial parameters.

Setting Up the Sample Data: Creating a NumPy Matrix

To practically demonstrate matrix normalization, we must first establish a representative dataset using the NumPy library. This matrix, which we will call `x`, will serve as our unnormalized raw data input. We will use NumPy's `arange` function to create a sequence of numbers and then apply the `reshape` method to transform this sequence into a two-dimensional matrix structure suitable for row and column normalization.

For consistency across our examples, we generate a 3x3 matrix containing linearly spaced integer values. This matrix, defined by the array of numbers from 0 up to (but not including) 36, stepped by 4, provides clear, manageable values that allow us to easily verify the resulting normalized figures. This initial step is fundamental to any data science pipeline, ensuring that the input data is correctly formatted before proceeding to the preprocessing phase.

The code block below outlines the creation and initial display of the matrix `x`. Observe the raw values and how they are distributed across the rows and columns, as this structure will be dramatically altered by the normalization process.

```
import numpy as np
```

```
# Create a 3x3 matrix using arange and reshape
x = np.arange(0, 36, 4).reshape(3,3)
```

```
# Display the original matrix
print(x)
```

```
]
```

Example 1: Row-Wise L1 Normalization (`axis=1`)

The most common requirement for normalization is scaling features associated with each individual observation, which corresponds to scaling across the rows of the matrix. When we set the parameter `axis=1`, the `normalize` function treats each row as an independent vector. The L1 norm (the sum of absolute values) for each row is calculated, and then every element within that specific row is divided by its row's total L1 norm. This ensures that the sum of the resulting elements in any given row equals exactly 1.

We apply this technique to our previously defined matrix `x`. This type of row normalization is particularly useful in machine learning tasks where you need the features of a single data point to

be represented as proportions or probabilities relative to one another. For instance, in topic modeling, normalizing document vectors ensures that longer documents do not exert undue influence merely because they contain a higher frequency of words overall.

Review the output carefully. The calculation for the first row, , involves dividing each element by the row sum ($0 + 4 + 8 = 12$). Thus, the normalized elements become $0/12$, $4/12$, and $8/12$, which simplifies to 0, 0.333..., and 0.666.... This proportional scaling is repeated independently for all three rows, yielding the normalized matrix `x_normed` shown below.

```
from sklearn.preprocessing import normalize
```

```
# Normalize matrix by rows (axis=1) using the L1 norm
```

```
x_normed = normalize(x, axis=1, norm='l1')
```

```
# View the resulting normalized matrix
```

```
print(x_normed)
```

```
]
```

Upon inspection, we can confirm that the values in each row of the normalized matrix now accurately sum to one, confirming the successful application of L1 row normalization:

Sum of first row: $0.0 + 0.333... + 0.666... \approx 1$

Sum of second row: $0.25 + 0.333... + 0.416... \approx 1$

Sum of third row: $0.2857 + 0.3333 + 0.3809 \approx 1$

Example 2: Column-Wise L1 Normalization (`axis=0`)

Alternatively, the normalization requirement might demand that we scale the features across all observations. This is achieved by normalizing column-wise, treating each column as an independent feature vector. By setting the parameter `axis=0`, the `normalize` function computes the L1 norm for each column and uses that specific column sum to scale all elements within it. This operation is particularly relevant when ensuring that no single feature dominates the model due to its overall magnitude across the dataset.

Using the same initial matrix `x`, we now instruct the function to normalize along the vertical axis. Unlike row normalization, where scaling is performed per observation, column normalization scales the features themselves. If the columns represent different features (e.g., height, weight, age), this ensures that regardless of the initial variability or scale of 'weight' compared to 'age', both features contribute equally to the downstream model training process based on their proportional distribution within that feature column.

Consider the first column of the original matrix: . The column sum (L1 norm) is $0 + 12 + 24 = 36$. The normalized values are calculated as $0/36$, $12/36$, and $24/36$, resulting in 0, 0.333..., and 0.666.... This guarantees that the sum of the elements in the new first column is 1. This process is repeated independently for the second and third columns, yielding the following results:

```
from sklearn.preprocessing import normalize
```

```
# Normalize matrix by columns (axis=0) using the L1 norm
```

```
x_normed = normalize(x, axis=0, norm='l1')
```

```
# View the resulting normalized matrix
```

```
print(x_normed)
```

```
]
```

As anticipated, the normalization applied along `axis=0` ensures that the vertical sums equal one. The consistency of these sums confirms the mathematical transformation, providing standardized feature vectors ready for machine learning consumption:

Sum of first column: $0.0 + 0.333... + 0.666... \approx 1$

Sum of second column: $0.083 + 0.333... + 0.583... \approx 1$

Sum of third column: $0.133 + 0.333... + 0.533... \approx 1$

Expanding Normalization: Using the L2 (Euclidean) Norm

While the L1 norm is highly effective for sparsity and probability interpretation, the L2 norm (also known as the Euclidean norm) is arguably the most frequently used form of vector normalization in deep learning and general machine learning. L2 normalization scales the vector such that the sum of the squares of its components equals one, meaning the Euclidean length of the resulting vector is one. This standardizes the vector length while preserving the direction of the vector in space, a property critical for many geometric algorithms.

To implement L2 normalization using scikit-learn, we simply change the `norm` parameter from `'l1'` to `'l2'`. For our original row vector , the L2 norm (magnitude) is calculated as $\sqrt{0^2 + 4^2 + 8^2} = \sqrt{0 + 16 + 64} = \sqrt{80} \approx 8.944$. Each element is then divided by 8.944. The resulting vector will have a length of exactly one. The L2 normalization process effectively projects the data onto a unit hypersphere, which is beneficial because it avoids amplifying the impact of small differences in feature values when dealing with high-dimensional data.

If we were to normalize our matrix `x` using L2 normalization row-wise (`axis=1`), the syntax would be: `normalize(x, axis=1, norm='l2')`. Selecting the appropriate norm (L1, L2, or even the

max norm) is a crucial decision during the data preprocessing phase, directly influencing the performance and behavior of algorithms like K-Means clustering or neural network training, where input stability is paramount. Data scientists typically experiment with both L1 and L2 normalization to determine which best suits their specific model and dataset characteristics.

Summary of Normalization Techniques and Best Practices

The ability to accurately and efficiently normalize matrices is a foundational skill in advanced data analysis using Python. As demonstrated, the `sklearn.preprocessing.normalize` function provides an accessible, robust method for applying standard vector normalization. Crucially, the differentiation between `axis=1` (row-wise, normalizing observations) and `axis=0` (column-wise, normalizing features) dictates the interpretation and effect of the scaling process on the overall dataset.

When incorporating normalization into a production pipeline, there are several best practices to follow. First, always ensure that normalization parameters (the calculated sums or norms) are derived only from the training data, not the test data, to prevent data leakage. Second, remember that normalization is distinct from standardization (Z-score scaling), though both are forms of scaling. Normalization scales data to a unit norm or specific range (like 0 to 1), while standardization rescales data to have a mean of 0 and a standard deviation of 1. The choice between these two depends on whether your data benefits more from bounded ranges or zero-mean distribution.

Finally, while we focused on L1 and L2 norms, scikit-learn also supports the 'max' norm (also known as L_{∞}), which scales each element by the maximum value in that vector, ensuring all resulting values fall between 0 and 1. Mastering these techniques ensures that your NumPy data structures are properly conditioned, leading to more stable, reliable, and faster convergence for complex analytical models. Consistent and appropriate preprocessing is the bedrock of high-performing machine learning systems.

The following tutorials explain how to perform other common operations in Python: