

Learning Label Encoding for Multiple Columns in Scikit-Learn

Authored by
Mohammed loot

October 27, 2025

RECOMMENDED CITATION

Mohammed loot (2025). *Learning Label Encoding for Multiple Columns in Scikit-Learn*. PSYCHOLOGICAL STATISTICS. Retrieved from <https://statistics.arabpsychology.com/?p=4459>

In the expansive and complex world of [machine learning](#), the initial and often most time-consuming phase is data preparation. This stage, known as [preprocessing](#), is crucial because raw data rarely conforms to the requirements of analytical models. A common challenge arises when dealing with [categorical data](#)--variables that represent distinct groups or labels (such as colors, geographical regions, or boolean flags). Since the vast majority of [machine learning](#) algorithms are fundamentally mathematical and require [numerical data](#) inputs, these descriptive categories must be systematically transformed into a quantitative format.

One powerful and widely adopted technique for this transformation is [Label Encoding](#). This method simplifies the data by assigning a unique integer identifier to every distinct category within a feature column. While the concept is simple, applying it correctly and efficiently, especially across multiple features simultaneously, is a core skill in effective [feature engineering](#). This comprehensive guide will explore the mechanisms of [Label Encoding](#) and provide clear, practical instructions on how to leverage the [Scikit-learn](#) library in Python to streamline this process across selected columns within a [pandas DataFrame](#).

Understanding the Mechanics of Label Encoding

[Label Encoding](#) is an integral [preprocessing](#) technique specifically designed to handle [categorical variables](#) by converting them into an integer representation. Essentially, every unique string or category found in a column is mapped to a distinct integer value. For instance, if a dataset contains a 'Size' column with entries "Small," "Medium," and "Large," the encoder might map these to 0, 1, and 2, respectively.

The primary advantage of this conversion is that it renders the [categorical data](#) compatible with algorithms that exclusively process [numerical data](#). However, data practitioners must be acutely aware of the major conceptual drawback: [Label Encoding](#) inadvertently introduces an artificial sense of order or hierarchy. When an algorithm processes the encoded values (e.g., 0, 1, 2), it may incorrectly assume that 2 is numerically superior or quantitatively "more" than 1, and 1 is "more" than 0.

Due to this implicit ordering, [Label Encoding](#) is optimally suited for [ordinal data](#)--categories that possess a natural, meaningful rank (e.g., educational levels, satisfaction ratings). Conversely, for [nominal data](#) (categories without inherent rank, such as country names or colors), introducing a numerical hierarchy can severely mislead models. In such cases, alternative methods like [One-Hot Encoding](#) are typically preferred to maintain feature independence and prevent the model from learning non-existent ordinal relationships.

Visualizing the Label Encoding Transformation

To solidify the understanding of this transformation, let us examine how a simple [categorical](#)

[variable](#), such as "Team" names ('A', 'B', 'C', 'D'), is processed. [Label Encoding](#) takes these unique string values and assigns a unique, non-negative integer to each one.

In practice, the assignment of these integers is generally determined by the alphabetical or lexicographical order of the categories within the data column. Therefore, 'A' would be mapped to 0, 'B' to 1, 'C' to 2, and so on. This mechanism successfully converts qualitative information into a quantitative format suitable for computation, while preserving the distinction between the original categories.

The following visual aid demonstrates this direct mapping, illustrating how each unique categorical element, like a 'Team' name, is systematically converted into its corresponding integer label by the encoder:

Original Data		Label Encoded Data	
Team	Points	Team	Points
A	25	0	25
A	12	0	12
B	15	1	15
B	14	1	14
B	19	1	19
B	23	1	23
C	25	2	25
C	29	2	29

Implementing Label Encoding Across Multiple Features

Real-world [machine learning](#) datasets frequently contain numerous [categorical columns](#) requiring encoding. Manually initializing and applying a separate [LabelEncoder](#) instance for every single column is inefficient and prone to errors. Fortunately, the robust integration between [Scikit-learn](#) and the [pandas DataFrame](#) structure allows for a highly efficient and elegant solution to encode multiple features simultaneously.

The core of this efficiency lies in the powerful [pandas DataFrame](#) method, `apply`. By selecting a specific subset of columns (which itself functions as a DataFrame), we can use the `apply` method to execute a function--in this case, the `fit_transform` method of the [LabelEncoder](#)--across all selected features in one operation. The `fit_transform` method is ideal here, as it first learns the mapping (fitting) and then applies the transformation to the data.

The generalized Python syntax provided below demonstrates how to perform [Label Encoding](#) concisely across any designated list of columns within a [pandas DataFrame](#):

```
from sklearn.preprocessing import LabelEncoder
```

```
#perform label encoding on col1, col2 columns  
df = df.apply(LabelEncoder().fit_transform)
```

This powerful snippet begins by importing the necessary `LabelEncoder` from [Scikit-learn's preprocessing](#) module. It then isolates the target columns (`col1`, `col2`) and uses the `apply` function to execute the encoding instance across them. The results are immediately written back into the original DataFrame columns, replacing the string categories with their integer representations. The following section brings this syntax to life with a complete, executable demonstration.

Practical Demonstration: Encoding a Sports Data DataFrame

Setting up the DataFrame

To effectively demonstrate multi-column [Label Encoding](#), we will utilize a simulated [pandas DataFrame](#) containing player statistics. This dataset is designed to include a mix of descriptive (categorical) features and measurable (numerical) features, making it an excellent realistic scenario for data [preprocessing](#).

We begin by constructing the sample DataFrame using the Python [pandas](#) library:

```
import pandas as pd
```

```
#create DataFrame  
df = pd.DataFrame({'team': ,  
  'position': ,  
  'all_star': ,  
  'points': })
```

```
#view DataFrame  
print(df)
```

```
team position all_star points  
0 A G Y 11  
1 A F N 8  
2 B G Y 10
```

```

3 B F Y 6
4 B F Y 6
5 C G N 5
6 C G Y 9
7 D F N 12

```

The 'team', 'position', and 'all_star' columns clearly contain string-based [categorical values](#). Our goal is to convert these three features into [numerical data](#) that can be used directly by a [machine learning](#) model.

Applying Label Encoding to Multiple Columns

We now apply the optimized encoding method to transform the 'team', 'position', and 'all_star' columns simultaneously into their respective integer labels. This is achieved by selecting the subset of columns and applying the `LabelEncoder().fit_transform` method in a single line of code, demonstrating the efficiency of this [Scikit-learn](#) and pandas integration.

The following Python code executes the multi-column encoding:

```
from sklearn.preprocessing import LabelEncoder
```

```
#perform label encoding across team, position, and all_star columns
df = df.apply(LabelEncoder().fit_transform)
```

```
#view updated DataFrame
print(df)
```

```

team position all_star points
0 0 1 1 11
1 0 0 0 8
2 1 1 1 10
3 1 0 1 6
4 1 0 1 6
5 2 1 0 5
6 2 1 1 9
7 3 0 0 12

```

The resulting output confirms that the concise code block successfully applied the transformation. All string values within the specified categorical columns have been replaced by their integer-encoded counterparts, preparing the data for model consumption.

Interpreting the Encoded Output and Best Practices

By examining the final DataFrame, we can clearly see the mapping performed by the [LabelEncoder](#). The integer assignment is based on the alphabetical order of the unique values found within each feature. Understanding this mapping is essential for interpreting model results:

For the **team** column, the alphabetical conversion yielded:

"A" → 0

"B" → 1

"C" → 2

"D" → 3

For the **position** column:

"F" (Forward) → 0

"G" (Guard) → 1

And for the binary **all_star** column:

"N" (No) → 0

"Y" (Yes) → 1

This transformation is robust and applicable to any size or number of [categorical columns](#) within a [pandas DataFrame](#) using the demonstrated syntax. The efficiency of this method makes it a staple in any data [preprocessing](#) pipeline.

Considerations for Choosing the Right Encoding Strategy

While [Label Encoding](#) offers simplicity, its application must be governed by careful consideration of the data type to prevent undesirable outcomes in [machine learning](#) models. The central issue remains the automatic creation of an ordinal hierarchy.

As previously established, [Label Encoding](#) should be reserved primarily for [ordinal data](#). When applied to [nominal data](#), the numerical differences between the encoded integers (e.g., 0, 1, 2) can trick models sensitive to magnitude, such as linear and distance-based algorithms, into believing that one category is quantitatively superior to another, even when no such relationship exists. This false hierarchy can introduce significant bias.

Therefore, for non-ordered [nominal data](#), [One-Hot Encoding](#) is the generally preferred technique. Furthermore, it is critical for data scientists to handle any [missing values](#) present in the categorical columns before proceeding with encoding. Whether through imputation or removal, complete data

ensures the encoder learns the correct set of unique labels and performs a flawless transformation. Always align your chosen encoding strategy with the specific requirements of your target [machine learning](#) algorithm and the intrinsic characteristics of your dataset.

Conclusion

Mastering the conversion of [categorical data](#) is a foundational requirement for any successful [machine learning](#) project. We have demonstrated that [Label Encoding](#), when integrated with [Scikit-learn](#)'s preprocessing tools and the efficient ``apply`` method of a [pandas DataFrame](#), offers a streamlined approach to transforming multiple features into a [numerical format](#).

This technique is particularly valuable for encoding [ordinal data](#), where maintaining an inherent order aids model performance. By systematically and efficiently transforming [categorical features](#), data scientists can ensure their datasets are properly structured for robust model training and ultimately achieve improved analytical results. Always prioritize understanding the data type--nominal versus ordinal--when selecting your encoding method.

Additional Resources

For those looking to deepen their expertise in data [preprocessing](#) techniques and other essential Python tasks for [machine learning](#), the following tutorials offer further valuable guidance and examples: