

# What is Categorical Distribution?

Authored by  
**stats writer**

December 13, 2025

## RECOMMENDED CITATION

stats writer (2025). *What is Categorical Distribution?*. PSYCHOLOGICAL SCALES.  
Retrieved from <https://scales.arabpsychology.com/?p=107304>

The field of probability theory relies on models that accurately describe the likelihood of various outcomes. Among these models, the **categorical distribution** stands out as a fundamental tool for modeling experiments where the outcome is nominal and discrete. Often referred to as a generalized Bernoulli distribution, the categorical distribution is a type of discrete probability distribution that governs a single trial experiment resulting in one of  $K$  mutually exclusive and exhaustive outcomes.

Formally, a **categorical distribution** describes the probability that a random variable will assume a value corresponding to one of  $K$  specified categories. These categories are typically labeled numerically, such as  $\{1, 2, \dots, K\}$ . The core assumption is that each category  $i$  possesses an associated probability, denoted as  $p_i$ , such that the sum of all these probabilities equals unity.

This distribution is essential in various statistical and machine learning applications, particularly when dealing with data that falls into distinct, non-ordered groups, such as predicting a customer's preferred color, the outcome of a quality control test (pass/fail/rework), or classifying objects into predefined classes.

## Key Characteristics and Formal Criteria

For any statistical model to be legitimately classified as a **categorical distribution**, it must rigorously satisfy a set of four core criteria. These conditions ensure that the model accurately reflects a single trial process with distinct, finite outcomes.

First, the outcomes must be **discrete**. This means the result of the experiment must fall into clearly separated categories without any possibility of intermediate values. The data must be countable, not measurable. For instance, if the categories are defined as types of transport (Car, Bus, Train), an observation cannot be "Car-like-Bus."

The categories are **discrete** and mutually exclusive. The random variable can only take on values corresponding exactly to one of the defined  $K$  categories.

There must be  $K \geq 2$  potential categories. If there were only one category, the outcome would be certain (probability 1), and the concept of a distribution would become trivial.

The probability ( $p_i$ ) associated with the random variable assuming a value in each category must adhere to the fundamental axioms of probability, meaning  $0 \leq p_i \leq 1$  for all categories  $i$ .

The collective probability mass must be fully accounted for. The sum of the probabilities for all  $K$  categories must precisely equal 1. That is,  $\sum p_i = 1$  for  $i = 1$  to  $K$ .

These conditions ensure that the model is mathematically consistent and provides a complete description of the sample space for a single observation or trial. Failure to meet any of these

requirements would mean the distribution is either continuous, unbounded, or fundamentally misdefined.

## The Mathematical Formalism of Categorical Distribution

The **categorical distribution** is parameterized by a vector of probabilities,  $\mathbf{p} = (p_1, p_2, \dots, p_K)$ , where  $p_i$  represents the probability of observing the  $i$ -th category. Since the categories are nominal, a convenient way to represent the outcome of a trial is through the use of indicator variables or the "one-hot" encoding scheme.

Let  $X$  be a random variable that follows a categorical distribution with parameter vector  $\mathbf{p}$ . The probability mass function (PMF) is typically defined using the one-hot encoding, where the outcome vector  $\mathbf{x}$  is a vector of length  $K$  with a single element equal to 1 (indicating the observed category) and the rest being 0. If the outcome is category  $i$ , then  $x_i = 1$  and  $x_j = 0$  for  $j \neq i$ .

The probability mass function (PMF) for an outcome  $\mathbf{x}$  is given by the product form:

$$P(X = \mathbf{x}) = p_1^{x_1} * p_2^{x_2} * \dots * p_K^{x_K}$$

Since only one  $x_i$  can be 1, this formula elegantly reduces to simply  $p_i$ , the probability of the observed category. This formulation is particularly useful when analyzing statistical properties like expected value or variance, although these properties are typically calculated based on the individual components of the one-hot encoded vector rather than the single scalar outcome.

## Illustrative Example: The Standard Die Roll

A classic and highly accessible example of a **categorical distribution** is the simple experiment of rolling a fair, six-sided die. This experiment perfectly encapsulates the properties required by the distribution.

In this scenario, the number of potential outcomes,  $K$ , is 6: {1, 2, 3, 4, 5, 6}. Since the die is assumed to be fair, the probability associated with each outcome is uniform, meaning  $p_i = 1/6$  for all  $i$  from 1 to 6. This setup satisfies the criteria for a discrete probability distribution, as outcomes are non-overlapping and exhaust the entire sample space.

Let us verify how the die roll adheres to the formal requirements:

**Example of a Categorical Distribution:**

Outcomes of rolling a dice

| Value | Probability |
|-------|-------------|
| 1     | 1/6         |
| 2     | 1/6         |
| 3     | 1/6         |
| 4     | 1/6         |
| 5     | 1/6         |
| 6     | 1/6         |

**Discreteness:** The outcomes are strictly integer values (1, 2, 3, 4, 5, 6). A roll cannot result in 3.5, confirming the nature of the random variable as discrete.

**Minimum Categories:** We have  $K = 6$ , which clearly meets the requirement of  $K \geq 2$ .

**Probability Bounds:** Each probability is  $1/6$ , which falls squarely between 0 and 1 ( $0 < 1/6 < 1$ ).

**Summation to Unity:** The sum of all probabilities is  $(1/6) + (1/6) + (1/6) + (1/6) + (1/6) + (1/6) = 6/6 = 1$ . The distribution is properly normalized.

This example demonstrates that the **categorical distribution** is ideal for modeling events where all outcomes are fixed and predetermined, and we are only interested in the likelihood of a single observation falling into one of these specific bins.

## Distinguishing Discrete from Continuous Variables

Understanding the context of the **categorical distribution** necessitates a clear distinction between discrete probability distributions and continuous probability distributions. The defining factor lies in whether the underlying random variable is countable or measurable.

A helpful operational definition, often referred to as a "Rule of Thumb" in introductory statistics, centers on the method of observation:

### Rule of Thumb: Count vs. Measure

If you can **count** the number of potential outcomes, you are dealing with a **discrete random variable**. Examples include the number of defective items in a batch, the count of successes in a series of trials, or the outcome of a single categorical experiment (like rolling a die).

Conversely, if you must **measure** the outcome, you are working with a **continuous random variable**. Continuous variables, such as height, weight, temperature, or elapsed time, can take on any value within a given range, potentially having infinite possibilities.

The **categorical distribution** is fundamentally rooted in the discrete domain. Its categories are separated by gaps, meaning there are no intermediate states. This separation is key to its application in classification problems where observations fall definitively into one predetermined bucket. The ability to count the outcomes is what allows us to assign a specific probability mass to each category, as opposed to requiring a continuous probability density function.

## Real-World Applications of Categorical Models

The utility of the **categorical distribution** extends far beyond simple games of chance; it serves as a powerful foundational model across various scientific and engineering disciplines where outcomes are classified into distinct groups. Analyzing these real-world scenarios highlights the distribution's broad applicability.

### Example 1: Flipping a Coin (The Binary Case)

When we flip a coin, we encounter the simplest form of the categorical distribution, known as the Bernoulli distribution, where  $K = 2$ . The outcomes are defined as Heads (Success) and Tails (Failure). Let  $p$  be the probability of Heads. The resulting distribution is parameterized by  $\mathbf{p} = (p, 1-p)$ .

The trial satisfies all criteria: we have two discrete outcomes, both probabilities are between 0 and 1, and they sum to 1. This scenario is foundational not just for games, but also for modeling any single binary decision or event, such as whether a component fails (0) or succeeds (1) in a reliability test.

#### Example of a Categorical Distribution:

Outcomes of flipping a coin

| Value | Probability |
|-------|-------------|
| Heads | 1/2         |
| Tails | 1/2         |

### Example 2: Selecting Marbles from an Urn (Non-Uniform Probabilities)

Consider a classic urn problem, where the probabilities are not uniform. Suppose an urn contains 5

red marbles, 3 green marbles, and 2 purple marbles. If we randomly select one marble from the urn, the outcomes are discrete categories: Red, Green, or Purple. Here,  $K = 3$ .

The associated probabilities are calculated based on the counts (total 10 marbles):  $p_{\text{Red}} = 5/10 = 0.5$ ,  $p_{\text{Green}} = 3/10 = 0.3$ , and  $p_{\text{Purple}} = 2/10 = 0.2$ . The probability vector is  $\mathbf{p} = (0.5, 0.3, 0.2)$ . Crucially,  $0.5 + 0.3 + 0.2 = 1.0$ , confirming the model structure.

### Example of a Categorical Distribution:

Outcomes of selecting a marble

| Value  | Probability |
|--------|-------------|
| Red    | 5/10        |
| Green  | 3/10        |
| Purple | 2/10        |

This type of distribution is commonly used in sampling theory and serves as a precursor to understanding more complex models like the [multinomial distribution](#) when multiple selections are made.

### Example 3: Selecting a Card from a Deck (High K Value)

When selecting a single card from a standard 52-card deck, we can define the categories based on suit ( $K=4$ : Hearts, Diamonds, Clubs, Spades) or based on rank ( $K=13$ : Ace through King). If we consider the suit, we have  $K=4$ , and for a fair deck,  $p_i = 13/52 = 1/4 = 0.25$  for each suit. This illustrates how the **categorical distribution** handles multiple categories while maintaining the constraints of mutual exclusivity and summation to unity. This model is critical in quality control and risk assessment where a single event is classified into one of many potential outcomes.

**Example of a Categorical Distribution:**

Outcomes of selecting a card

| Value | Probability |
|-------|-------------|
| Ace   | 1/13        |
| 2     | 1/13        |
| 3     | 1/13        |
| 4     | 1/13        |
| 5     | 1/13        |
| 6     | 1/13        |
| 7     | 1/13        |
| 8     | 1/13        |
| 9     | 1/13        |
| 10    | 1/13        |
| Jack  | 1/13        |
| Queen | 1/13        |
| King  | 1/13        |

**Comparison with Related Discrete Distributions**

The **categorical distribution** is often confused with or defined in relation to other major discrete probability distributions. Understanding the subtle differences, particularly regarding the number of possible outcomes ( $K$ ) and the number of independent trials ( $n$ ), is crucial for correct statistical modeling. The categorical distribution acts as the most fundamental building block, describing the outcome of a single event where there are multiple categories. It is defined specifically by having  $K \geq 2$  potential outcomes and  $n = 1$  trial.

Using the notation where  $K$  is the number of possible outcomes and  $n$  is the number of trials, we can establish the following relationships:

**Bernoulli Distribution:** This is a special case of the **categorical distribution** where  $K$  is restricted to exactly 2 outcomes (e.g., success/failure). It models a single trial ( $n = 1$ ) and is parameterized by a single probability,  $p$ , the likelihood of success.

Relationship:  $K = 2$  outcomes,  $n = 1$  trial.

**Binomial Distribution:** This distribution extends the Bernoulli concept by modeling the total **number of successes** over a fixed number of independent trials ( $n \geq 1$ ). Critically, it still requires

only two possible outcomes per trial ( $K = 2$ ). The binomial distribution provides a count, whereas the categorical distribution provides the identity of the single outcome.

Relationship:  $K = 2$  outcomes,  $n \geq 1$  trial.

**Multinomial Distribution:** The multinomial distribution is the generalization of the binomial distribution to cases where there are more than two outcomes ( $K \geq 2$ ). If we perform  $n$  independent trials and record the counts for each of the  $K$  categories, the resulting distribution of counts is multinomial.

Relationship:  $K \geq 2$  outcomes,  $n \geq 1$  trial.

In essence, a single trial sampled from a multinomial distribution is mathematically equivalent to an outcome sampled from a **categorical distribution**. The categorical model is the atomic unit from which the multinomial distribution is constructed.

## Parameter Estimation in Categorical Models

A crucial aspect of applying any discrete probability distribution is determining the parameter vector  $\mathbf{p} = (p_1, p_2, \dots, p_K)$ . In real-world data science, these probabilities are rarely known beforehand (unlike a fair die roll) and must be estimated from observed data.

The most common and statistically efficient method for estimating the parameters of a **categorical distribution** is the technique of **Maximum Likelihood Estimation (MLE)**. Given a dataset of  $N$  independent observations, where  $n_i$  is the observed count of the  $i$ -th category (so that  $N = \sum n_i$ ), the MLE for the probability  $p_i$  is simply the relative frequency of that category:

$$\hat{p}_i = n_i / N$$

For example, if an analysis of website traffic shows that out of 500 visitors, 200 arrived via Search (Category 1), 150 via Social Media (Category 2), and 150 via Direct Link (Category 3), the MLE estimates would be  $p_{\text{Search}} = 0.4$ ,  $p_{\text{Social}} = 0.3$ , and  $p_{\text{Direct}} = 0.3$ . This intuitive approach ensures that the estimated parameters maximize the likelihood of observing the specific sample data.

When categories are sparse or sample sizes are small, leading to zero counts for some categories, simple MLE can sometimes lead to poor predictive performance. In such cases, Bayesian methods utilizing a **Dirichlet prior** (which is the conjugate prior for the categorical distribution) or techniques like **Laplace smoothing** are employed to introduce pseudo-counts, ensuring that all estimated probabilities are non-zero, thereby stabilizing the model.

## Conclusion: Importance in Data Science

The **categorical distribution** serves as a cornerstone of statistical modeling, particularly in disciplines involving classification and discrete choice. Its simplicity--modeling a single observation into one of  $K$  bins--belies its critical importance as the foundation for more complex models, including those used in sophisticated machine learning algorithms like Naive Bayes classifiers and neural network output layers (often via the Softmax function, which estimates the probabilities  $p_i$ ).

By defining the strict criteria for discreteness, mutual exclusivity, and probability normalization, this distribution provides a precise mathematical framework for understanding and predicting nominal outcomes. Whether used directly to model a single decision or indirectly as the elemental trial structure within multinomial and logistic regression models, the categorical distribution remains an indispensable tool for statisticians and data scientists seeking to quantify uncertainty in classified data.