

How to Find the Maximum Likelihood Estimation for a Uniform Distribution

Authored by
stats writer

March 1, 2026

RECOMMENDED CITATION

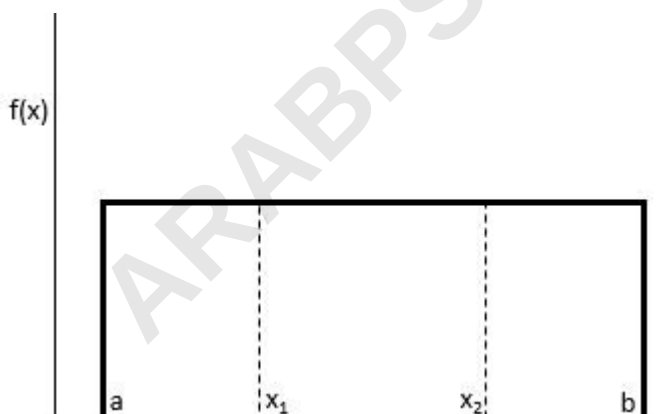
stats writer (2026). *How to Find the Maximum Likelihood Estimation for a Uniform Distribution*. PSYCHOLOGICAL SCALES. Retrieved from <https://scales.arabpsychology.com/?p=133414>

Understanding the Fundamentals of the Continuous Uniform Distribution

In the vast field of **statistics**, the **uniform distribution** represents one of the most fundamental concepts for modeling scenarios where every possible outcome within a specific range is equally likely to occur. Unlike other distributions that may cluster around a central mean, a **uniform distribution** maintains a constant **probability density function** across its entire support interval, typically defined by the **parameters** a and b . These parameters represent the minimum and maximum boundaries of the distribution, respectively, ensuring that any **random variable** falling within this interval has an identical chance of being observed.

The mathematical representation of this probability is elegantly simple. When we seek to determine the probability that a value will fall between two specific points, x_1 and x_2 , within the broader interval from a to b , we utilize a linear ratio based on the lengths of these intervals. Specifically, the formula is expressed as $P(x_1 < X < x_2) = (x_2 - x_1) / (b - a)$. This calculation highlights the inherent symmetry and simplicity of the distribution, making it a critical tool in both theoretical **statistics** and practical applications such as **random number generation** and simulation modeling.

Visually, the **uniform distribution** is often depicted as a rectangular shape, where the height of the rectangle is determined by the inverse of the interval's width, $1/(b - a)$. This ensures that the total area under the curve equals one, satisfying the primary axiom of **probability theory**. By understanding these foundational properties, researchers can begin to explore more complex inferential techniques, such as determining how to best estimate these boundary **parameters** when they are unknown and must be inferred from observed data.



This tutorial provides a comprehensive deep dive into the **maximum likelihood estimation** (MLE) process. We will examine how mathematicians and data scientists derive the most statistically probable **estimate** for the parameters a and b when provided with a sample of data points. This process is essential for building accurate models in fields ranging from economics to engineering, where precise boundary estimation is often required to define the limits of a system or process.

Theoretical Foundations of Maximum Likelihood Estimation

Maximum likelihood estimation is a sophisticated statistical methodology used to estimate the **parameters** of a probability distribution based on observed data. The core philosophy behind MLE is to identify the parameter values that maximize the **likelihood function**, which represents the joint probability of the observed sample as a function of the unknown parameters. In simpler terms, we are looking for the version of the distribution that is most likely to have produced the dataset we currently possess.

The **likelihood function** is central to this endeavor. For a set of independent and identically distributed (i.i.d.) observations, the likelihood is the product of the individual probability density functions for each observation. While many statistical models rely on **gradient descent** or other iterative **optimization** techniques to find the peak of this function, the **uniform distribution** presents a unique challenge. Because the boundaries of the distribution are the very parameters we are trying to estimate, the **likelihood function** is not always differentiable in the traditional sense at its maxima.

In many cases, taking the natural logarithm of the likelihood--resulting in the **log-likelihood**--simplifies the math significantly, converting products into sums. However, for the **uniform distribution**, we must pay close attention to the constraints of the **sample space**. If any observed data point falls outside the range, the likelihood of that specific parameter set becomes zero. Therefore, the MLE process for this specific distribution involves a careful analysis of the order statistics of the sample data rather than just solving derivative equations.

Step 1: Constructing the Likelihood Function for Uniform Data

The initial step in the **maximum likelihood estimation** process is the formal construction of the **likelihood function**. For a **uniform distribution**, the PDF for a single observation x is $1/(b - a)$, provided that $a \leq x \leq b$. If we have a sample of size n , the joint probability (the likelihood) is the product of these individual densities. This assumes that each data point in our set is independent of the others, a standard assumption in frequentist **statistics**.

Mathematically, the **likelihood function** $L(a, b)$ is expressed as the product from $i=1$ to n of $1/(b - a)$. This simplifies to $(b - a)$ raised to the power of negative n . However, there is a critical condition: this formula only holds if all observed values fall within the interval. If even a single observation is smaller than a or larger than b , the entire likelihood drops to zero because such a sample would be impossible to generate from a **uniform distribution** with those specific boundaries.

$$\prod_{i=1}^n f(x_i; a, b) = \prod_{i=1}^n \frac{1}{(b-a)} = \frac{1}{(b-a)^n}$$

Understanding this constraint is vital. The **likelihood function** is effectively a piecewise function. It has a non-zero value only when a is less than or equal to the minimum value in the sample and b is greater than or equal to the maximum value in the sample. This dependency on the data boundaries is what distinguishes the **maximum likelihood estimation** of the uniform distribution from other common distributions like the Gaussian or Poisson distributions, where the **parameters** typically affect the shape rather than the support of the function.

Step 2: Deriving the Log-Likelihood Function

Once the primary **likelihood function** is established, the next logical progression in the **maximum likelihood estimation** workflow is to derive the **log-likelihood** function. By applying the natural logarithm to the likelihood equation, we transform the multiplicative relationship into an additive one, which is generally much easier to manipulate mathematically. For the uniform case, taking the log of $(b - a)^{-n}$ results in $-n \ln(b - a)$.

This transformation is valid because the natural logarithm is a monotonically increasing function. Therefore, the values of a and b that maximize the **log-likelihood** will also maximize the original **likelihood function**. In the context of the **uniform distribution**, the **log-likelihood** clarifies the relationship between the sample size and the width of the interval, showing how the "penalty" for a larger interval grows as more data points are collected.

$$\log \prod_{i=1}^n f(x_i; a, b) = \log \prod_{i=1}^n \frac{1}{(b-a)} = \log ((b-a)^{-n}) = -n \log (b-a)$$

The **log-likelihood** function serves as the basis for the subsequent differentiation steps. It is important to remember that while we are using algebraic manipulation, the underlying constraints regarding the minimum and maximum observations still apply. The **log-likelihood** is only defined and relevant within the region where the likelihood is non-zero. Outside of these boundaries, the function does not yield a valid **estimate**, reinforcing the idea that our **estimator** must be strictly informed by the extreme values present in our dataset.

Step 3: Analyzing Derivatives and Optimization Challenges

In standard **maximum likelihood estimation**, the typical procedure involves taking the partial derivative of the **log-likelihood** with respect to the **parameters**, setting those derivatives to zero, and solving for the variables. This approach works perfectly for smooth, bell-shaped likelihood surfaces. However, for the **uniform distribution**, the derivative of the **log-likelihood** function does not provide a simple "zero-point" that indicates a maximum.

The derivative of the **log-likelihood** with respect to a reveals a monotonically increasing function

within the allowed range. This suggests that to maximize the likelihood, we should make a as large as possible. Conversely, the derivative with respect to b is monotonically decreasing, suggesting that we should make b as small as possible to maximize the function. This behavior indicates that the maximum occurs at the boundaries of the feasible region, rather than at a stationary point where the slope is zero.

$$\frac{n}{(b-a)}$$

Because the derivatives do not vanish, we cannot use **gradient descent** in the traditional way to find an interior solution. Instead, we must look at the constraints imposed by the data. The derivative of the **log-likelihood** with respect to b is illustrated below, further confirming that the likelihood is maximized when the denominator $(b - a)$ is minimized, provided the interval still contains all the data points.

$$-\frac{n}{(b-a)}$$

Step 4: Identifying the Maximum Likelihood Estimators

To finalize the **maximum likelihood estimation**, we identify the values of a and b that satisfy our optimization criteria under the given constraints. As we established, we want to maximize the **likelihood function** by making the interval as small as possible, while still ensuring that every sample point X_i is contained within it. This means that a must be no larger than the smallest value in our data, and b must be no smaller than the largest value in our data.

Therefore, the most likely value for the **parameter** a is the minimum observed value in the sample, denoted as $\min(X_1, X_2, \dots, X_n)$. If we were to choose an a value any larger than this minimum, the likelihood would immediately drop to zero because the minimum value would then fall outside the distribution's range. This **estimator** is considered the most efficient way to define the lower bound of a **uniform distribution** based solely on observed evidence.

Similarly, the **maximum likelihood estimate** for b is the maximum observed value in the sample, denoted as $\max(X_1, X_2, \dots, X_n)$. Since the **likelihood function** is monotonically decreasing with respect to b , the smallest possible value for b that remains valid (i.e., is greater than or equal to all data points) will produce the highest likelihood. Thus, the sample maximum is the definitive **estimator** for the upper boundary of the distribution.

Practical Implications of MLE in Uniform Modeling

Applying **maximum likelihood estimation** to a **uniform distribution** has significant implications in real-world data science. For instance, in quality control, if a machine produces components whose lengths are uniformly distributed, the MLE provides the most statistically sound method for estimating the range of lengths the machine is capable of producing based on a inspected batch. It allows engineers to set realistic tolerances and identify when a process has drifted beyond its intended boundaries.

Furthermore, these **estimates** are used in simulation and **Monte Carlo methods**. When we need to model an uncertain process where only the extremes are known or can be estimated, using the sample minimum and maximum as the distribution's **parameters** provides a model that is perfectly tailored to the observed evidence. While MLE for uniform distributions is known to be slightly biased (as the true range is often slightly wider than the sample range), it remains a "consistent" **estimator**, meaning it converges to the true value as the sample size increases.

In summary, the process of finding the **maximum likelihood estimation** for a **uniform distribution** involves a unique blend of calculus and logical constraint analysis. By focusing on the likelihood of the boundaries, researchers can derive **parameters** that offer deep insights into the underlying nature of their data. Whether you are working in finance, physics, or computer science, mastering these estimation techniques is a vital step in becoming a proficient **statistician** or data analyst.