

Understanding the 10% Condition in Statistics: A Comprehensive Guide

Authored by
Mohammed loot

November 7, 2025

RECOMMENDED CITATION

Mohammed loot (2025). *Understanding the 10% Condition in Statistics: A Comprehensive Guide*. PSYCHOLOGICAL STATISTICS. Retrieved from <https://statistics.arabpsychology.com/?p=12226>

Introduction: Setting the Statistical Stage

In the realm of [statistics](#), many foundational concepts rely on simplified models of chance events. One such fundamental concept is the [Bernoulli trial](#). A Bernoulli trial is defined as an experiment that has only two mutually exclusive outcomes: typically labeled as "success" or "failure." Crucially, the probability of success must remain constant across every iteration of the experiment. This simple framework allows statisticians to model complex scenarios, from quality control checks to medical testing outcomes, where the result is purely binary.

A classic and easily visualized example of a Bernoulli trial is the simple flip of a fair coin. The coin can only land on heads or tails, representing our two possible outcomes. If we designate heads as a "success" and tails as a "failure," the probability of success remains a fixed 0.5 for each flip, assuming the coin is balanced and the previous outcomes do not influence the next. When we string together multiple truly independent Bernoulli trials, we often model the resulting distribution using the [Binomial distribution](#).

However, calculating exact probabilities using the Binomial distribution can become computationally intensive when the number of trials, denoted as n , grows very large. To simplify these calculations, especially when n is substantial, statistical practice often dictates using the [normal distribution](#) as a powerful approximation of the Binomial distribution. This approximation is only valid, however, if certain conditions are met, the most critical being the assumption that the trials are truly [independent](#).

The Critical Role of Independence in Sampling

The concept of [independence](#) is the cornerstone of reliable statistical inference. When we state that trials are independent, we mean that the outcome of one trial has absolutely no bearing or influence on the outcome of any subsequent trial. In the coin flip example, independence is inherent; the coin has no memory of the previous result. But what happens when we sample from a finite group, where the act of selecting an item changes the composition of the remaining group?

Consider a scenario where we are sampling without replacement--meaning once an item (or person) is selected for the sample, it is removed from the [population size](#) and cannot be chosen again. In this situation, the trials are technically **not** independent. The probability of the next success or failure changes slightly with each selection. For example, if we are selecting marbles from a jar containing 10 red and 10 blue marbles, the probability of drawing a red marble first is 10/20. If we draw a red marble and do not replace it, the probability of drawing a second red marble immediately drops to 9/19. This sequential change violates strict independence, necessitating the use of conditional probability.

While this lack of independence is mathematically true, statisticians have recognized that if the

[sample size](#) is extremely small relative to the overall population, the change in probabilities becomes negligible, making the assumption of independence a safe and practical simplification. The difference between sampling with replacement (truly independent) and sampling without replacement (dependent) diminishes as the population size increases relative to the sample size. This observation leads directly to the establishment of a formal guideline known universally as **The 10% Condition**.

Defining and Formalizing The 10% Condition

The **10% Condition** provides a crucial safety threshold for statistical modeling, particularly when dealing with finite populations and sampling without replacement. It allows us to proceed with powerful statistical techniques, such as the Normal approximation to the Binomial distribution, even when the foundational assumption of independence is technically violated due to the finite nature of the population pool. This condition acts as a practical rule of thumb, maintaining the integrity of our calculations while simplifying complex real-world data collection processes.

The 10% Condition: As long as the [sample size](#) (n) is less than or equal to 10% of the total [population size](#) (N), we can confidently proceed under the assumption that the Bernoulli trials are independent for calculation purposes. Mathematically, this is expressed as: $n \leq 0.1N$.

The rationale behind setting the limit at 10% is rooted in the mathematical demonstration that the error introduced by assuming independence when it does not strictly exist remains acceptably small. When the sample constitutes only a tiny fraction of the whole, removing a few items does not significantly alter the underlying probability distribution of the remaining population. For example, removing one person from a population of 100,000 barely changes the proportion, whereas removing one person from a population of 10 radically changes the subsequent probabilities. By limiting the sample to no more than 10% of the population, we ensure that the dependency introduced by sampling without replacement is statistically insignificant for most inference procedures.

Intuition Through Calculation: The Impact of Sampling Ratio

To truly grasp why the 10% limit works, we can examine a tangible example involving conditional probability where sampling without replacement occurs. Let's revisit a classroom scenario. Suppose a class has 20 students ($N=20$), and the true proportion who prefer football over basketball is exactly 50% (10 students prefer football, 10 prefer basketball). We are interested in the probability that four randomly selected students ($n=4$) all prefer football. Here, the ratio n/N is $4/20 = 20\%$. Since 20% is greater than 10%, we should anticipate a noticeable difference between the independent and dependent calculations.

Let X be the [random variable](#) representing the number of students who prefer football. We want

to find $P(X=4)$. We compare two methods: sampling with replacement (truly independent) and sampling without replacement (dependent).

First, if we assume the trials are strictly independent (sampling with replacement, or if the population were infinite), the probability calculation is straightforward multiplication based on the constant proportion:

$$P(\text{All 4 students prefer football}) = (10/20) \times (10/20) \times (10/20) \times (10/20) = 0.5 \times 0.5 \times 0.5 \times 0.5 = \mathbf{0.0625}.$$

Second, if the trials are dependent because we are sampling without replacement (which is typical in real-world surveys), the probability shifts with each selection:

$$P(\text{All 4 students prefer football}) = (10/20) \times (9/19) \times (8/18) \times (7/17) \approx 0.5 \times 0.4737 \times 0.4444 \times 0.4118 \approx \mathbf{0.0433}.$$

The difference between the independent (0.0625) and dependent (0.0433) probability is significant--an error of about 30% relative to the true value. This large discrepancy confirms that when the sample size (4) exceeds 10% of the population size (20), the assumption of independence is invalid and leads to inaccurate results. Now, consider how this difference shrinks as the population size grows, keeping the sample size constant at $n=4$.

The following visualization demonstrates how the proximity of the independent and non-independent probabilities converges as the ratio of sample size to population size decreases. Notice that when the population size is much larger (e.g., $N=1000$, where $n/N = 0.4\%$), the difference between the two calculations becomes imperceptibly small, validating the utility of the 10% threshold.

Classroom Size	Sample Size / Classroom Size	$P(\text{All 4 students prefer football})$ <i>Independent trials</i>	$P(\text{All 4 students prefer football})$ <i>Not Independent trials</i>
20	20.0%	$10/20 \times 10/20 \times 10/20 \times 10/20 = .0625$	$10/20 \times 9/19 \times 8/18 \times 7/17 = .0433$
40	10.0%	$20/40 \times 20/40 \times 20/40 \times 20/40 = .0625$	$20/40 \times 19/39 \times 18/38 \times 17/37 = .0530$
100	4.0%	$50/100 \times 50/100 \times 50/100 \times 50/100 = .0625$	$50/100 \times 49/99 \times 48/98 \times 47/97 = .0587$
1,000	0.4%	$500/1,000 \times 500/1,000 \times 500/1,000 \times 500/1,000 = .0625$	$500/1,000 \times 499/999 \times 498/998 \times 497/997 = .0621$

Practical Applications and Necessary Conditions for Statistical Inference

The 10% Condition is not the sole requirement for utilizing the [Normal approximation](#) for Binomial distributions; it works in tandem with other critical checks. Specifically, for the approximation to be sufficiently accurate, two additional conditions--often referred to as the Success/Failure Condition--must also be satisfied. These ensure that the distribution is roughly symmetrical and bell-shaped,

justifying the use of the Normal model.

The complete list of conditions required for using the Normal approximation to model the sampling distribution of a proportion includes:

Randomization Condition: The data must come from a well-designed random sample or randomized experiment, ensuring the sample is representative of the population.

The 10% Condition: The [sample size](#) n must not exceed 10% of the [population size](#) N ($n \leq 0.10N$). This guarantees that the sampling without replacement does not introduce significant dependency errors.

Success/Failure Condition: Both the number of expected successes (np) and the number of expected failures ($n(1-p)$) must be at least 10. This ensures that the shape of the Binomial distribution is close enough to the symmetric Normal curve.

In practice, when designing a study or analyzing survey data, researchers must rigorously verify all these conditions before proceeding with inferential tests like confidence intervals or hypothesis testing based on the Normal model. Failing to check the 10% Condition, especially when sampling from finite, small populations, introduces unnecessary bias and compromises the validity of the derived conclusions. While the 10% rule is robust, statisticians always prefer a sample size that is significantly less than 10%--perhaps 5% or even 1%--to maximize the accuracy of the independence assumption.

Conclusion: Ensuring Robust Statistical Modeling

The **10% Condition** is an indispensable tool in the statistician's toolkit, serving as a boundary condition that reconciles theoretical purity with practical necessity. It allows us to treat trials as independent--a requirement for many powerful statistical models--even when sampling without replacement from a finite pool. By ensuring that the sample size does not unduly deplete the population, we mitigate the risk of introducing significant statistical error due to dependency.

The core message of the condition is one of statistical safety: when your sample is small relative to the population, the composition of the population changes so slightly with each draw that the shift in conditional probability is negligible. Adherence to this condition is essential for producing reliable and valid statistical inferences, ensuring that our estimates about the larger population remain accurate and unbiased. Always verify $n \leq 0.10N$ before proceeding with calculations that rely on the assumption of independence.

Further Exploration and Resources

For readers interested in deepening their understanding of the statistical distributions and theorems that rely on the concept of independence, we recommend exploring the following related

topics:

[An Introduction to the Normal Distribution](#)

[An Introduction to the Binomial Distribution](#)

[An Introduction to the Central Limit Theorem](#)