

Learning Guide: Conducting a One Proportion Z-Test in Python

Authored by
Mohammed looti

November 7, 2025

RECOMMENDED CITATION

Mohammed looti (2025). *Learning Guide: Conducting a One Proportion Z-Test in Python*. PSYCHOLOGICAL STATISTICS. Retrieved from <https://statistics.arabpsychology.com/?p=11973>

The [one proportion z-test](#) stands as a cornerstone in inferential statistics, providing a robust mechanism for comparing the observed success rate derived from a sample against a specific, predetermined [population proportion](#). This test is indispensable across numerous quantitative fields, including epidemiology, market analysis, and stringent quality control processes, because it allows researchers to rigorously assess whether a performance metric or rate observed in a small subset of data is statistically aligned with a larger theoretical standard or expectation. When decisions hinge on binary outcomes--such as pass/fail, yes/no, or success/failure--the one proportion z-test offers the mathematical framework required to move beyond mere observation and make definitive statistical inferences about the entire population.

Before launching into computation, it is vital to understand the underlying prerequisites that validate the use of this specific statistical tool. The test assumes that the sample data was collected randomly and that the observations are independent of one another. Furthermore, to ensure that the sampling distribution of the sample proportion can be reasonably approximated by the normal distribution, the sample size must be sufficiently large. Specifically, statisticians typically require that both the number of expected successes ($n \times p_0$) and the number of expected failures ($n \times (1 - p_0)$) are greater than or equal to 10, ensuring the accuracy of the subsequent Z-statistic calculation and P-value interpretation. Meeting these criteria transforms the process from a simple comparison of percentages into a powerful, verifiable hypothesis test.

The entire structure of the [one proportion z-test](#) is centered on resolving a fundamental dichotomy: Is the observed difference between the sample proportion and the hypothesized population proportion a genuine, statistically significant finding, or is it merely the product of natural, random sampling variability? By standardizing this difference, the test allows researchers to quantify the magnitude of discrepancy and assign a precise probability to the null hypothesis being true. This systematic approach ensures that conclusions are not based on subjective judgment but on quantifiable evidence, leading to data-driven decision-making in high-stakes environments.

The Foundation: Defining the Null and Alternative Hypotheses

The critical first step in any formal hypothesis testing procedure involves the rigorous definition of the competing hypotheses. This framework provides the logical structure for the entire statistical argument. The [null hypothesis](#), denoted as H_0 , always represents the status quo, the assumption of no effect, or the idea that the true population parameter is exactly equal to the hypothesized value. In the specific context of the one proportion z-test, the null hypothesis asserts that the true, but unknown, population proportion (p) is mathematically identical to the specific hypothesized population proportion (p_0) that is being tested against the sample data.

H_0 : $p = p_0$ (The true population proportion is precisely equal to the established hypothesized proportion p_0). This statement is the foundation upon which all subsequent calculations are

based; the Z-statistic is calculated under the explicit assumption that this equality holds true.

In contrast to the conservative stance of the null hypothesis, the alternative hypothesis, H_1 , encapsulates the researcher's claim or suspicion--the outcome they are actively seeking evidence for. The formulation of H_1 is dictated entirely by the underlying research question and determines the nature of the test, leading to either a two-tailed test, where deviation in either direction is significant, or a one-tailed test, which focuses scrutiny on a specific direction of change. Proper formulation here is not a trivial matter; it defines the scope of the statistical investigation and directly influences how the final P-value is interpreted.

H1 (Two-tailed): $p \neq p_0$ (The [population proportion](#) is not equal to the hypothesized value p_0). This is the most common approach, utilized when the researcher is interested in any significant difference, whether the true proportion is higher or lower than the benchmark.

H1 (Left-tailed): $p < p_0$ (The population proportion is less than the hypothesized value p_0). This specific direction is employed when the primary concern is focused solely on a reduction, such as testing if a new manufacturing process has resulted in a lower defect rate than the standard.

H1 (Right-tailed): $p > p_0$ (The population proportion is greater than the hypothesized value p_0). This is appropriate when the research hypothesis only concerns an increase, for example, verifying if a new marketing campaign has successfully boosted the conversion rate above a historical benchmark.

The rigorous process of defining these hypotheses ensures that the statistical conclusion directly and unambiguously addresses the initial research question. A failure to clearly articulate whether the interest lies in a difference, an increase, or a decrease can lead to misinterpretation of the test results, particularly when comparing the calculated Z-statistic to the critical values or when interpreting the resulting [P-value](#). Thus, investing time in accurate hypothesis formulation is the cornerstone of responsible statistical practice, establishing the critical threshold for evidence required to overturn the assumed status quo represented by H_0 .

Quantifying Evidence: Calculating the Z-Test Statistic

The [Z-test statistic](#) serves as the quantifiable link between the observed sample data and the theoretical assumption laid out in the null hypothesis. It standardizes the difference between the observed sample proportion ($\hat{p}$) and the hypothesized population proportion (p_0) by expressing this distance in units of standard error. By transforming the raw difference into a standard score, we can utilize the well-established properties of the standard normal distribution (the Z-distribution) to determine how unusual our sample result would be if the null hypothesis were truly correct.

The mathematical expression for this standardization is precise and foundational to the test:

$$z = (p - p_0) / \sqrt{p_0(1-p_0)/n}$$

Analyzing this formula reveals its conceptual elegance. The numerator, $(p - p_0)$, measures the raw magnitude of the difference observed in the sample. This is the evidence of deviation. The denominator represents the standard error of the sample proportion, calculated using p_0 (not the sample proportion p). This is a crucial distinction: since we assume the null hypothesis ($H_0: p = p_0$) is true for the purpose of the test, we use p_0 to estimate the variability (standard deviation) of the sampling distribution. This denominator effectively measures the expected random variation, providing the necessary scale to judge whether the observed difference in the numerator is large enough to be considered statistically significant.

The components of this calculation are derived directly from the study design and the hypothesis:

p: The **observed sample proportion**, derived from the count of successes (x) divided by the total sample size (n). This is the primary piece of empirical evidence.

p_0 : The **hypothesized population proportion**, which is the fixed value specified in the null hypothesis (H_0).

n: The **sample size**, representing the total number of independent trials or observations utilized in the study. A larger n generally leads to a smaller standard error, making it easier to detect a true difference.

Interpreting the calculated Z-statistic requires placing it on the standard normal curve. A Z-statistic close to zero suggests that the sample proportion (p) is very near the hypothesized proportion (p_0), indicating that the observed data strongly supports the null hypothesis. Conversely, a large positive or large negative Z-statistic signifies that the sample proportion is many standard deviations away from p_0 . This large deviation suggests that the observed result is unlikely to have occurred purely by chance if H_0 were true, thus providing strong evidence against the null hypothesis and necessitating a closer look at the corresponding P-value.

Decision Making: Interpreting Results Using the P-Value and Alpha

The [P-value](#) is arguably the single most important output of any hypothesis test, translating the calculated Z-statistic into a probability that is easy to interpret. The P-value quantifies the evidence against the null hypothesis by defining the probability of observing a sample result as extreme as, or more extreme than, the one actually obtained, assuming that the null hypothesis (H_0) is entirely true. A small P-value suggests that the observed data is highly improbable under the assumption of H_0 , thereby casting serious doubt on the null hypothesis itself.

The statistical decision--whether to reject H_0 or fail to reject H_0 --is reached by comparing this P-value to the predetermined [significance level](#), conventionally denoted by the Greek letter α (alpha). The significance level is the threshold probability chosen by the researcher before

the data collection begins, representing the maximum acceptable risk of making a Type I error--that is, the risk of mistakenly rejecting a true null hypothesis. Common choices for α are 0.05 (5%), 0.01 (1%), or 0.10 (10%), with 0.05 being the standard threshold in many scientific disciplines, signifying that there is only a 5% chance of observing the data if H_0 is true.

The decision rule is straightforward and applies universally across frequentist hypothesis testing frameworks:

If $P\text{-value} < \alpha$: We **reject H_0** . This conclusion means that the observed data is statistically significant at the α level, providing sufficient evidence to conclude that the true population proportion differs from the hypothesized value p_0 .

If $P\text{-value} \geq \alpha$: We **fail to reject H_0** . This indicates that the sample results are plausible under the assumption that H_0 is true. There is insufficient statistical evidence from the sample to conclude that the true population proportion is different from p_0 .

It is essential for researchers and analysts to use precise language when communicating these results. Failing to reject H_0 should never be equated with "accepting" or "proving" the null hypothesis. It simply means that the data collected was not strong enough to definitively disprove the status quo. Furthermore, the selection of the [significance level](#) must be justified by the context of the study; in fields where false positives (Type I errors) carry high risks, such as drug testing, a lower α (like 0.01) is often preferred, demanding much stronger evidence before concluding significance.

Implementing the Z-Test Efficiently in Python

In modern data science and statistical practice, manual calculation of the [Z-test statistic](#) and subsequent P-value lookup is both inefficient and prone to error. To streamline this process and ensure computational accuracy, leveraging specialized libraries within programming languages like [Python](#) is the industry standard. Python's rich ecosystem of data analysis tools makes it an ideal environment for complex statistical modeling and hypothesis testing.

For performing the one proportion z-test, the robust and widely utilized [statsmodels](#) library is the tool of choice. Designed specifically for the estimation of statistical models and the execution of comprehensive statistical tests, `statsmodels` provides the function `proportions_ztest()`. This function automates the entire analytical pipeline, taking the raw sample counts and the hypothesized proportion as input, and immediately returning the calculated Z-statistic and the corresponding P-value, adjusted correctly for the chosen alternative hypothesis.

A clear understanding of the function signature is necessary for correct implementation. The core function and its mandatory parameters are:

proportions_ztest(count, nobs, value=None, alternative='two-sided')

These parameters meticulously map the statistical requirements onto the Python code, ensuring that the test is executed accurately according to the research design:

count: This parameter requires an integer or array representing the number of "successes" (x) observed within the sample. This is the empirical evidence of the desired outcome.

nobs: This specifies the total number of observations or trials (n). This figure is crucial for calculating the standard error and determining the test's statistical power.

value: This parameter sets the hypothesized population proportion (p_0). If omitted, the function defaults to testing the count against zero, but for the one proportion z-test, this should explicitly match the p_0 defined in H_0 .

alternative: This string parameter dictates the structure of the alternative hypothesis (H_1). The default is 'two-sided' ($p \neq p_0$). Researchers must change this to 'smaller' ($p < p_0$, left-tailed) or 'larger' ($p > p_0$, right-tailed) if the research question demands a directional test.

By relying on `proportions_ztest()`, analysts minimize computational errors and ensure that the powerful mathematical engine of the [statsmodels](#) library provides consistent and reliable results, facilitating quick transition from data preparation to formal statistical inference within the [Python](#) environment.

Practical Example: Analyzing Survey Data and Executing the Test

To illustrate the practical application of this test, consider a scenario where a political consulting firm needs to verify if the support level for a candidate in a key district deviates from a previously established benchmark of 64%. The firm sets the hypothesized population proportion (p_0) at 0.64. The research question is framed non-directionally: Does the current support level differ significantly from 64%? This mandates a two-tailed test.

To gather evidence, the firm conducts a random sample survey of 100 eligible voters ($n=100$). The survey results indicate that 60 residents expressed support for the candidate ($x=60$). The hypotheses are formally stated as $H_0: p = 0.64$ and $H_1: p \neq 0.64$. The sample proportion is calculated as $\hat{p} = 60/100 = 0.60$. The goal is to determine if the 4 percentage point difference between the sample (60%) and the hypothesis (64%) is due to genuine population change or simply random chance.

The necessary parameters are input directly into the `proportions_ztest` function, ensuring the count, sample size, and hypothesized value align perfectly with the defined problem:

Hypothesized Population Proportion (p_0): 0.64

Number of Successes (count): 60

Sample Size (nobs): 100

Alternative: 'two-sided' (default)

Executing the test in [Python](#) using the `statsmodels` library yields the following structured output, which is the immediate result of applying the Z-test formula to the provided data:

```
# Import proportions_ztest function from statsmodels
```

```
from statsmodels.stats.proportion import proportions_ztest
```

```
# Perform one proportion z-test comparing 60 successes in 100 trials against p0 = 0.64
```

```
proportions_ztest(count=60, nobs=100, value=0.64)
```

```
(-0.8164965809277268, 0.41421617824252466)
```

Drawing the Final Conclusion and Interpretation

The execution of the `proportions_ztest` function provides the two critical values required for the statistical decision, which must now be interpreted against the established [significance level](#) (α):

Z Test-Statistic: -0.8165. The negative sign indicates that the sample proportion (0.60) is below the hypothesized proportion (0.64). The magnitude (-0.8165) tells us that the observed sample proportion is less than one standard deviation below the mean of the hypothesized sampling distribution.

P-value: 0.4142. This value means that if the true support level were indeed 64% (H_0 is true), there would be a 41.42% chance of observing a sample proportion as far away as 60% (or further) purely due to random variation.

To reach a definitive conclusion, we compare this P-value to the standard [significance level](#), typically $\alpha = 0.05$. Since the calculated [P-value](#) (0.4142) is significantly greater than 0.05, the decision rule dictates that we must fail to reject the [null hypothesis](#) (H_0). The high P-value suggests that the observed data is entirely consistent with the assumption that the true [population proportion](#) is 64%.

The final conclusion, stated clearly and formally, is that there is insufficient statistical evidence, based on this survey data, to conclude that the true proportion of voters supporting the candidate is statistically different from the hypothesized 64%. The observed 4% deficit (60% vs. 64%) is well within the realm of what would be expected from typical sampling variability. The firm should proceed with caution, recognizing that while the sample did not hit the 64% mark, the evidence is not strong enough to scientifically declare the 64% benchmark inaccurate.

Additional Resources for Statistical Testing

For those seeking to deepen their understanding of proportion testing, including related methodologies and variations, the following resources offer valuable context and tools:

[An Introduction to the One Proportion Z-Test](#)

[One Proportion Z-Test Calculator](#)

[How to Perform a One Proportion Z-Test in Excel](#)

[How to Perform a One Proportion Z-Test in R](#)