

Statistical Analysis of Machinery Variance by Python

Joao Gabriel Ostrowski¹, József Menyhárt^{2*}

¹University of Debrecen, Faculty of Engineering, Mechanical Engineering Department, Ótemető u. 2-4, 4028 Debrecen, Hungary
E-mail: ostrowski@mailbox.unideb.hu

^{2*}Corresponding author, University of Debrecen, Faculty of Engineering, Department of Air and Road Vehicles, Ótemető u. 2-4, 4028 Debrecen, Hungary,
E-mail: jozsef.menyhart@eng.unideb.hu

Abstract: Based particularly on data technologies, information is rapidly evolving in engineering. In mechanical engineering, maintenance is benefiting the most from data innovations, the reduction of maintenance costs, and the improvement of system availability. Modern reliability engineering targets uncertainty with advanced statistics and data science. This paper explores maintenance descriptive analytical techniques to determine whether variations between porosity percentage in batches from two similar machines are due to randomness, suggesting technical issues on machinery. The finding is backed by statistical analysis, hypothesis testing, and data mining. Python 3 was used following best practices of data science and data visualization as the main toolset to explore the working data, execute the proposed statistical analysis, and make conclusions. The findings were promising reinforcing the importance of advanced statistics for reliability studies. As a conclusion, we failed to reject the null hypothesis suggesting that there was no apparent difference between sample means. With a confidence of 95%, the difference is explained by natural variations and randomness inherent to the fabrication process.

Keywords: reliability engineering; python programming; statistical analysis; maintenance optimization

1 Introduction

Information is evolving in engineering, based particularly on data technologies that facilitated information flow through departments making processes more efficient. Applied examples include live feedback data from sensors, live feed for stock control, predictive maintenance schemes, etc. Cloud implementations for Big Data streamline and Data Science play a prominent role in the introduction of Industrial Internet of Things in modern factories [1].

In mechanical engineering, maintenance is the most advanced field regarding new data technologies including such giants as IBM and Google worked extensively in applied cases for the industry, aiming at common business goal requirements, minimizing costs and late maintenance while maximizing machinery availability, production output, and customers' satisfaction [2]. Although most of the field development encompasses Maintenance Predictive Analytics through Machine Learning implementations, models fail to predict events when there is no past data detailing a similar occurrence, the adoption of Maintenance Descriptive Analytics in which Modern Reliability Engineering takes part, applying advanced statistics [3]. The focus of this research is to explore some new approaches to the statistical methodologies applied to support optimized maintenance planning in industrial environment.

1.1 The S.M.A.R.T Methodology

The S.M.A.R.T methodology was used as a business best practice to define the goals of this research. First, defined by George T. Doran [4], consultant and former Director of Corporate Planning for Washington Water Power Company, S.M.A.R.T. is an acronym to define goals and objectives in a clear and data-driven manner. According to Paul J. [5], an effective goal must be:

- Specific: A goal should be clear and specific, otherwise it is not possible to focus on efforts or feel truly motivated to achieve them. The five "W" questions (what, why, who, where, which) provide a large scale of support to gain maximum specificity;
- Measurable: It is only possible to plan a timeline of activities if one can perceive when tasks start and come to end. Such need becomes attainable by tracking progress, for tracking, there must be some measurable parameters;
- Achievable: A goal needs to be realistic and attainable to be successful. In other words, it should stretch and push barriers for a challenge while still remaining approachable.
- Relevant: This step ensures that the goal matters to the business are in line with other relevant goals.

Time-based: Every goal needs a target date to have a deadline to focus on and something to work toward.

In line with Peter Drucker's management concepts [6], such methodology started to be widely applied on analytics and data projects due to its attention to measurable results, which clarifies the base for effectivity analysis.

1.2 Goal

According to the previous guideline, the goal of this paper is in connection with reliability engineering objectives according to Connor, P. [7]:

- Goal 1: Evaluation of the variations in porosity percentage of two different machines based on their sensor data set to determine if the difference is due to randomness or some underlying causes. The required confidence level of the analysis is 95%. The analysis should be carried out in a month.

1.3 Foundations of Reliability Engineering

A brief introduction of reliability engineering as a discipline and its roots are rooted in the military due to the necessity of cost-effective systems, and a comparison between reliability and quality frameworks.

1.4 Military Roots of Reliability Engineering

Technology is known to develop exponentially due to military efforts, and the reliability engineering as a discipline was a necessary step by the United States during the 1950s as an attempt to reduce the increasing failure rates and improve the availability of military equipment. Due to the rapid development of electronic devices embedded in systems, the US Department of Defense (DOD) set up the Advisory Group of Reliability of Electronic Equipment (AGREE) report in 1952. The report consisted of strict rules regarding a test of thousands of hours applying high stress, cyclical, high and low temperatures, vibration and switching. [26]

In 1965, the DOD issued the MIL-STD-785 Reliability Programs for Systems and Equipment integrating the standardization for several liable engineering activities, meanwhile, the AGREE routine was accepted by NASA and all major high technology purchasers and suppliers. European countries followed up on the reliability progress by 1990s with a series of Reliability standards that became integrated into the International Standards Organization (ISO), ISO/IEC 60 300 e.g. [26]

1.5 Reliability and Quality

Reliability and quality, despite their distinctive concepts, are both attributes used interchangeably in several cases. Regarding the user level, it is commonly associated with technological products including computers, smartphones, or house appliances. On the other hand, from an industrial point of view, the idea is attributed to machines and parts.

There is no discussion when it comes to the increasing reliability required by modern society, upon a purchase one expected the item to function as described

for the longest time period as possible, avoiding new purchases and substitutions. This is exactly where reliability engineering differs from the old fashioned quality control paradigm, a time dimension was also taken into consideration. In other words, while quality is assessed against a set of specifications, which define conforming or non-conforming parts, a binary distribution of good and bad.

As an uncertain science, the main objective of reliability engineering reveals whether an item will work over a certain period of time, which can be answered by distributions of probabilities. As in its core, reliability is defined as *"the probability that an item will perform a required function without failure understated conditions for a stated period of time"*. [27]

In the words stated by James R. Schlesinger, Former US Secretary of State for Defense, "reliability is engineering in its most practical form". According to Connor, P. the priorities of such science can be listed as follows [30]:

- 1 To apply engineering knowledge and specialist techniques to prevent or to reduce the likelihood or frequency of failures;
- 2 To identify and correct the causes of failures that do occur, despite the efforts to prevent them;
- 3 To determine ways of coping with failures that do occur, if their causes have not been corrected;
- 4 To apply methods for estimating the likely reliability of new designs, and for analyzing reliability data.

2 Mathematical Methods and Tools

The methods used to quantify reliability include the mathematics of probability and statistics. As mentioned above, reliability engineering handles uncertainty. Understanding that variation is part of engineering and inherent to all manufacturing processes is the necessary step for controlling and minimizing unexpected events. Statistical methods provide the means for analyzing, understanding, and controlling variation. [28]

2.1 Box Plot Graph Construction and Analysis

A box and whisker plot often called a box plot displays the five-number summary of a set of data. The five-number summary is the minimum, first quartile, median, third quartile, and maximum. [28]

In a box plot, we draw a box from the first quartile to the third quartile. A vertical line goes through the box at the median. The whiskers go from each quartile to the minimum or maximum.

One common measure of spread for data analysis is the interquartile range (IQR), which is calculated as $Q_3 - Q_1$ and provides a measure of spread that it's independent of the distribution mean.

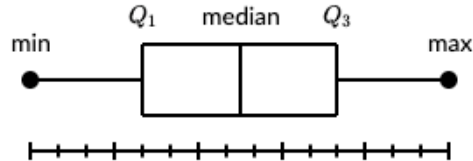


Figure 1
Representation of the five-number summary of a box plot

2.2 Outliers Detection through the Interquartile Range (IQR) Rule

For a rough symmetric dataset, both mean and standard deviation are good descriptive statistics to improve understanding around the dataset, however, for a skewed distribution or distributions that are not normal, the median and the interquartile range are more robust metrics to rely on for extracting insights. Since standard deviation is a function of the distribution mean, it is also heavily influenced by outliers, which are data points that do not fit the overall distribution adding noise to it.

Outliers can be visualized, and boxplots are great in this sense, or one could use mathematical techniques for identifying them. One of the techniques commonly found in the industry and statistical libraries for coding is the interquartile range rule [8].

The IQR rule used shows that a datapoint is an outlier if it is more than $1.5 \times IQR$ above the third quartile or below the first quartile. Said differently, low outliers are below $Q_1 - 1.5 \times IQR$ and high outliers are above $Q_3 + 1.5 \times IQR$.

2.3 Normality Test

An important decision point when working with a sample of data is whether to use parametric or nonparametric statistical methods. [9] [29]

The simpler form would assume that data have known and specific distribution, often a Gaussian, or normal, distribution. If a data sample does not follow a known distribution, then the assumptions of parametric statistical tests are not valid and nonparametric statistical methods must be used. Thus, normality tests are applied to working data sets for checking its correlation to a normal distribution. For the scope of this research, graphical and statistical methods were applied to check the normality of working data [9].

There are two common graphical methods for the normality test, the histogram plot, and the quantile-quantile plot, or QQ plot. The histogram with a Gaussian kernel density estimation was achieved by using the seaborn library under the `.distplot(kde=True)` command. For validation, the kernel density should be visual to a Gaussian distribution.

The QQ plot was executed by the statsmodel library under the command `qqplot(line='s')`. The QQ plot generates an idealized distribution and plots the compared data as a scatterplot above a 45 degrees line representing the desired distribution. There are plenty of excellent Python visualization libraries available, including the built-in matplotlib. However, we can use some alternative libraries. Seaborn is a popular data visualization library for Python programming language. Seaborn is an important Python visualization library built on top of matplotlib. It can give the similar information like ggplot2. It gives us the capability to create amplified data visuals. This helps us understand the data by displaying it in a visual context to unearth any hidden correlations between variables or trends that might not be obvious initially. Seaborn has a high-level interface as compared to the low-level of Matplotlib. [17-20] [24] [25]

Statsmodels is a Python package/library that allows engineers or users to explore data, estimate statistical models, and perform statistical tests during their work. An extensive list of descriptive statistics, tests, plotting functions, and result statistics are available for different types of data and each estimator. It complements SciPy's stats module. Statsmodels is part of the Python programming language. It is oriented towards data analysis, data science and statistics. Statsmodels is built on top of the numerical libraries NumPy and SciPy. The statsmodels integrate with Pandas for data handling and uses Patsy for an R-like formula interface. The graphical functions are based on the Matplotlib library. [17-20] [24] [25]

A perfect approximation using this method implies all datapoints from the sample aligned with the idealized line. When the datapoints tend to go above the idealized normal, distribution line the distribution is right-skewed, contrarily, the distribution is left-skewed when the observations tend to be below the normal line as exemplified in Figure 2 [10].

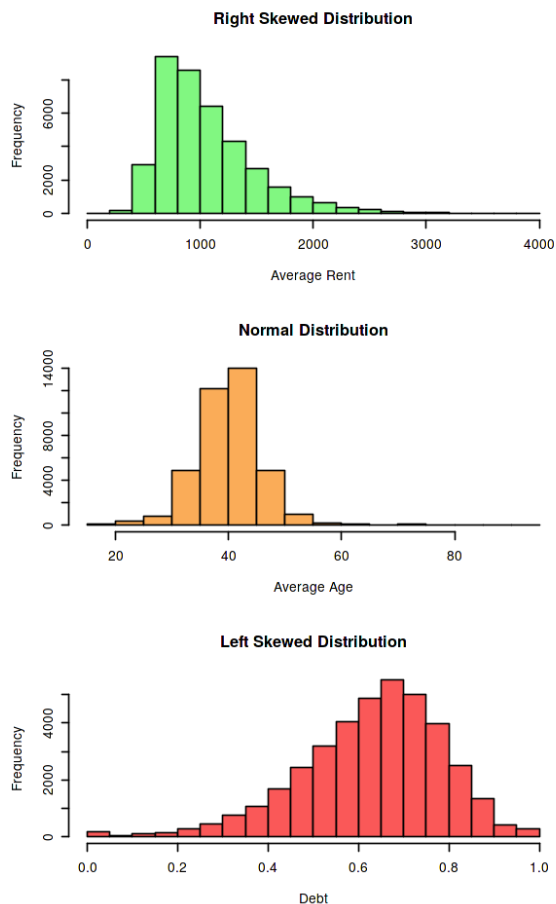


Figure 2
Examples of QQ Plot distributions

The statistical methods are diverse in the coding field, which differs on the preliminary assumptions and considerations upon the tested data. For the research scope, an algorithmic version of the Shapiro-Wilk normality test was applied by using the `scipy.stats` library under the command `shapiro()`.

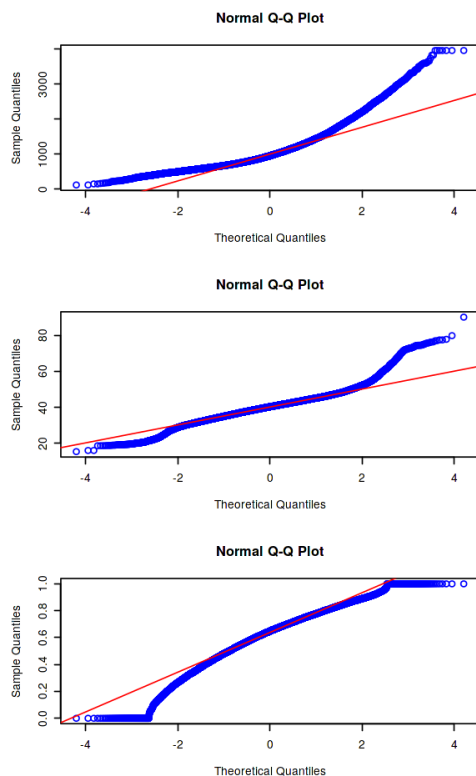


Figure 3
Examples of QQ Plot distributions

2.4 Central Limit Theorem

One of the reasons why normal distribution appears regularly when dealing with applied statistics must complete with the central limit theorem. One of the most fundamental and profound concepts in statistics [11], the central limit theorem is the core for several statistical activities involving using a sample for making inferences over a large and unknown population. It demonstrates that the sampling distribution of the sample means can be approximated as a normal distribution if the number of samples n is large enough, for statistical purposes $n > 30$ [12].

The versatility of the theorem is that it holds its principle for every random variable distribution, being normally distributed or not. Figure 4 illustrates the concept from a non-normal distribution use case.

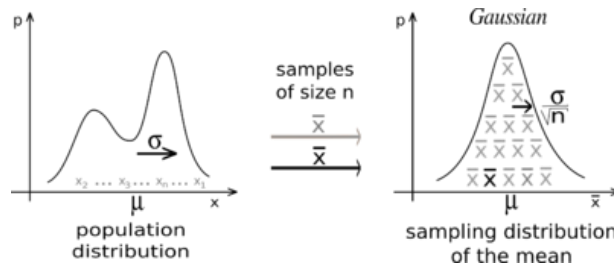


Figure 4

Central limit theorem use case

2.5 T-test for Difference of Two Sample Means

Testing the hypothesis that the mean of a sample is the same as that of an assumed population is a very common statistical analysis use case. When this is the case and when the true population mean (μ) and standard deviation (σ) are known this is a typical z test. On the other hand, when parameters from the original population are not given the test and classified as a t-test using sample statistics as approximation parameters and it has the variation corrected based on the Student-t distribution. [31-34]

Generally, the hypothesis being tested in a t-test suggests the difference between means of two sample means with relatively small sample sizes. The assumptions of the test are the following [31-34]:

- Data points are independent;
- Data points are accurately recorded;
- The sample size is small, usually less than 30.

The statistical test is then evaluated according to Equation 1.

$$Ttest = \frac{(x_1 - x_2) - d}{SE_{\bar{x}}} \quad (1)$$

$$SE = \sqrt{\frac{s_1^2}{n_1} + \frac{s_2^2}{n_2}} \quad (2)$$

The t-statistic is then compared to a t-table or applied to a t Distribution Calculator [13] together with the calculated degrees of freedom (DF) associated with the sample sizes for obtaining the p-value which resembles the probability of observing a result as extreme as the one observed in the experiment. The obtained p-value is then compared to the experimental confidence level to accept or reject the null hypothesis and reach conclusions. [31-34]

3 Methods

The technology used to process analysis varies according to the purpose and can be summarized as follows:

- Python 3.2 was used by specific libraries such as Pandas and Numpy for structured data handling and summary statistics, Matplotlib and Seaborn for visualizations, Statsmodels and Scipy for specific statistical tests;

3.1 Goal 1 - Porosity Analysis

Briefly reviewing the goal definition, for the given data set of machine porosity, a study is necessary to understand whether the porosity occurrence is due to processual randomness. The action plan consists of:

- Exploratory data analysis for better understanding the working data;
- Validation of the underlying distribution of the sample;
- State null hypothesis and alternative hypothesis for the experiment;
- Evaluation of sample confidence intervals (CI) at 95% confidence level;
- Two-sample t-test statistical test.

Table 1
First 4 samples of Machine Porosity data set [29]

Sample Number	Machine A - Percent porosity [%]	Machine B - Percent Porosity [%]
Sample 1	1.72	1.78
Sample 2	1.73	1.99
Sample 3	1.74	1.68
Sample 4	1.53	1.69

As a good practice of data analysis, the initial focus is on EDA execution or Exploratory Data Analysis for getting a better sense of the working data. The working data is a small machine sensor dataset with 20 observations for each machine. Summary statistics are presented below as well as the visualization of population distribution for Machine A and Machine B as a frequency plot (Figure 4) and a box plot (Figure 5).

Table 2
Summary statistics for Machine A and Machine B working data [29]

	Machine A	Machine B
Mean [%]	1.730	1.771
Median [%]	1.745	1.760
Max [%]	1.790	1.990
Min [%]	1.530	1.680
Std Dev [%]	0.063246	0.074302

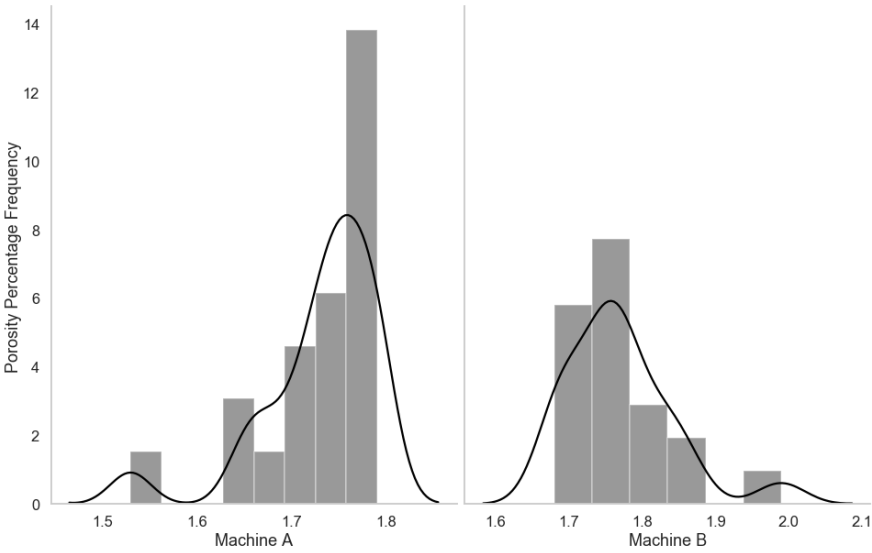


Figure 5
Porosity percentage frequency distribution

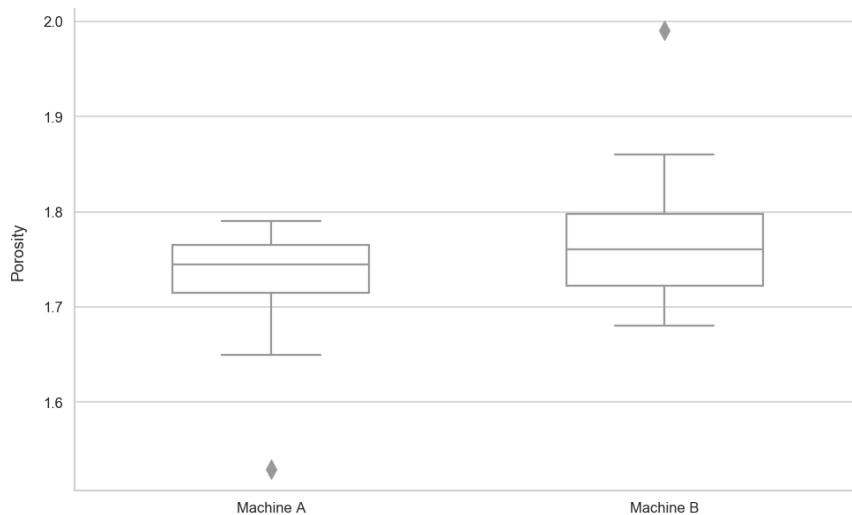


Figure 6

Box plot visualization from Machine A and B working data

It is a common assumption for engineering processes that the population of a production process follows a normal distribution, and provided that n is large enough (>30) it is possible to assume that the means of samples coming from such population are also normally distributed. A normality test was executed for both distributions in order to confirm if the dataset could be approximated and treated as a normal distribution since such consideration allow simplified statistical tests for testing the hypothesis.

From Figure 5 two different outliers are detected by using the $1.5 \times IQR$ rule, which made both samples to fail upon a Shapiro-Wilk test due to the small n of 20. However, after treating the outliers, overwriting both with the mean value of each sample, both samples passed the normality test. Figure 6 shows the resulting QQ plot graphical method for normality estimation for both machines after removing outliers.

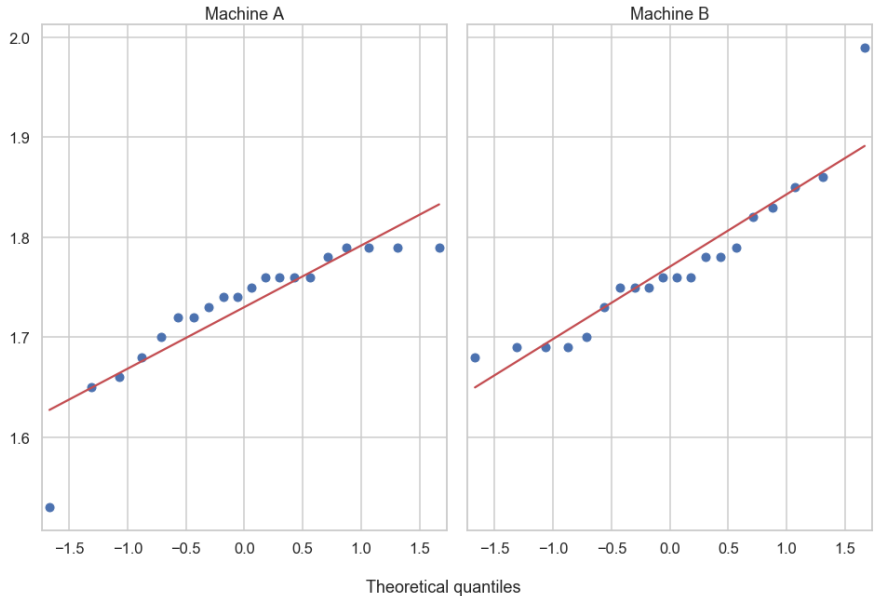


Figure 7
QQ plot for machines A and B after removing outliers

For the Two-sample t-test, the null hypothesis (H_0) under analysis is that both samples come from a common distribution, and therefore are similar having their mean difference is zero. In other words, the different machines have no impact on the porosity percentage of produced parts. Additionally, the alternative hypothesis(H_a)is that the mean difference between samples is different from zero, meaning that samples come from different populations. Hypothesis statements summarized in Table 3.

Table 3
Summarized hypothesis

Set	Null Hypothesis	Alternative Hypothesis
1	$\mu_1 - \mu_2 = 0$	$\mu_1 - \mu_2 \neq 0$

As an experiment best practice, the first assessment is completed by looking at an eventual overlap of CI at 95% confidence level, which would reinforce the null hypothesis. Figure 8 presents a clear overlap between sample means at 95% confidence level.

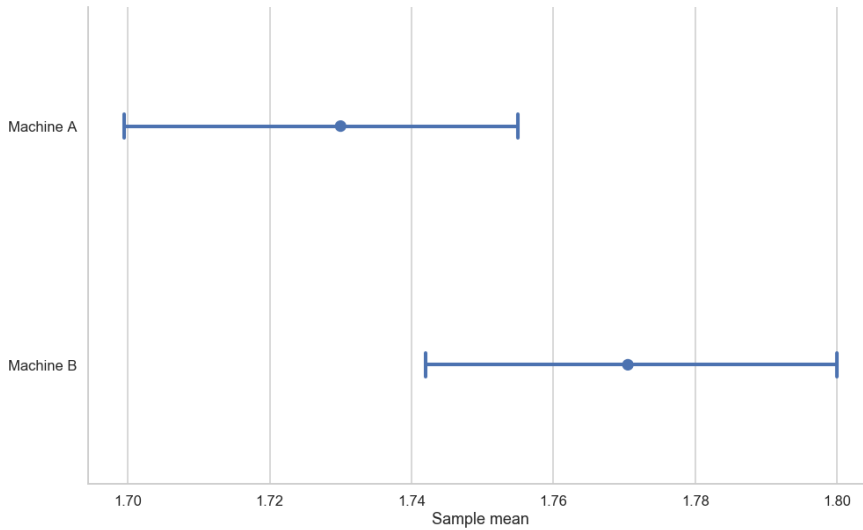


Figure 8

The confidence interval for sample means

The standard error of the mean difference is calculated and evaluated at a Z score of 1.96, covering 95% confidence level. Results are summarized in Figure 9 which are obtained by using Miller's A/B test calculator [14].

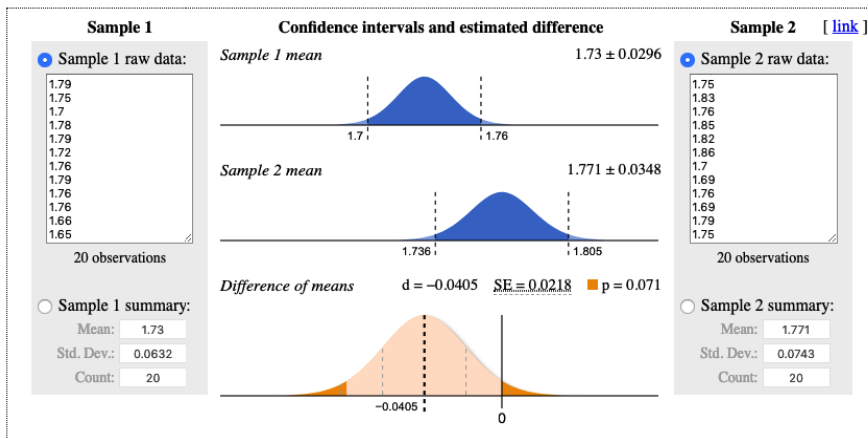


Figure 9

T-test summary from Evan's Awesome A/B Tools

Similar results were obtained using statsmodel T-test function considering equal variance and n number of samples under the command `stats.ttest_ind(equal_var=False)`.

Conclusion

Since the advent of data technologies, engineers need to be prepared to face data integration and analysis tasks which will require a lot more cognitive skills and deep work [15] to generate intelligence and insights as outputs to improve maintenance processes. [21-23]

Most of the meaningful insights by the analysis is derived from adjacent areas of engineering, supporting the idea of a more complete professional who is able to design, implement, carry out an audit, and analyze data.

The purposed goal emphasized statistical and data analysis, covering the entire pipeline from data extraction, transformation, and load (ETL) to p-value determination for reaching conclusions and business suggestions. The intuition built refines inspection and maintenance planning, where better-informed decisions are made aided by data-driven processes. Understanding such a mindset will gradually direct the industry towards modernization as it is already happening with modern industrial revolutions, such as the Industrial Internet of Things (IIoT) [17].

Understanding the mathematics of reliability allow leaders not to see future events, but the distribution of the most likely and desired outcomes in a way companies can prioritize better actions.

In conclusion, in a world where processes are growing complex with many different human and processual interactions understanding how to deal with uncertainty is an important skill for the modern engineer.

The method can be further refined by other soft computing methods. The Support Vector Machine method is well suited for analysing and grouping large amounts of data. The most promising method seems to be fuzzy logic, which provides an opportunity to examine so-called fuzzy sets. This is essential for a modern technical system. There are several operational and maintenance cases where operators have unequivocal data or perceptions. In this case, statistics-based machine learning is needed to make decisions.

Modern maintenance systems run on a real-time basis these days providing the fastest data reporting that helps maintainers take the fastest action. Modern enterprise management systems such as SAP and Oracle provide the ability to store and analyse data. The solution in this article can be implemented in any enterprise management system provided there is sufficient hardware capacity to perform complex calculations. Increased IT demand requires financial investment, which is not worth for every enterprise. This decision is always an economic decision and may be different for every company.

References

- [1] Boyes, H., et al. (2018) The industrial internet of things (IIoT): An analysis framework. *Computers in Industry*, Volume 101, 2018, pp. 1-12
- [2] Ageeva, Y. (2018, June 17) Predictive Maintenance Scheduling with IBM Watson Studio Local and Decision Optimization. In *Towards Data Science*. Retrieved 08:48, March 12, 2019, <https://towardsdatascience.com/predictive-maintenance-scheduling-with-ibm-data-science-experience-and-decision-optimization-25bc5f1b1b99>
- [3] Karim, R. et al. (2016) Maintenance Analytics – The New Know in Maintenance. *IFAC-PapersOnLine*, Volume 49, Issue 28, pp. 214-219
- [4] Doran, G. T. (1981) There's a S.M.A.R.T. Way to Write Management's Goals and Objectives. *Management Review*, 70, 35-36
- [5] Meyer, P. J. (2003) *Resources Attitude Is Everything: If You Want to Succeed Above and Beyond.*, Published January 1st 2003 by Paul J. Meyer Resources, ISBN: 9780898113044
- [6] Bogue, R. (2018) "Use S.M.A.R.T. goals to launch management by objectives plan". *TechRepublic.*, Retrieved February 13, 2020, <https://www.techrepublic.com/article/use-smart-goals-to-launch-management-by-objectives-plan/>
- [7] Connor, P. D. T. (2012) *Practical Reliability Engineering*. Wiley, Fifth Edition
- [8] Upton, G., Cook, I.: *Understanding Statistics*, Publisher: OUP Oxford, 1996, ISBN 9780199143917
- [9] Brownlee, J. (2018) A Gentle Introduction to Normality Tests in Python. In *Machine Learning Mastery*. Retrieved 17:29, February 18, 2019, <https://machinelearningmastery.com/a-gentle-introduction-to-normality-tests-in-python/>
- [10] Bachmann, J. (2018, September 01) Statistical Analysis: A frequentist approach. In *Kaggle Kernels*. Retrieved 20:47, March 11, 2019, <https://www.kaggle.com/janiobachmann/statistical-analysis-a-frequentist-approach>
- [11] Brooks, D. G.: *The Sampling Distribution and Central Limit Theorem*, 2nd edition, Amazon Books, CreateSpace Independent Publishing Platform (December 31, 2012) ISBN-13: 978-1530441617
- [12] Stat Trek. (2019, March 06) Central Limit Theorem. In *Statistics Dictionary Stat Trek*. Retrieved 20:42, March 06, 2019, https://stattrek.com/statistics/dictionary.aspx?definition=central_limit_theorem

- [13] Stat Trek. (2019, March 07) T-distribution calculator. In Statistics Dictionary Stat Trek. Retrieved 19:03, March 06, 2019, <https://stattrek.com/online-calculator/t-distribution.aspx>
- [14] Miller, E. (2019, March 03) Evan's Awesome A/B Tools. In [evanmiller.org](http://www.evanmiller.org). Retrieved 13:36, March 03, 2019, <http://www.evanmiller.org/ab-testing/t-test.html>
- [15] Newport, C. (2016) Deep Work: Rules for Focused Success in a Distracted World. Grand Central Publishing, Fifth Edition
- [16] Boyes, H., et al. (2018) The industrial internet of things (IIoT): An analysis framework. Computers in Industry, Volume 101, 2018, pp. 1-12
- [17] Singh, S.: Become a Data Visualization Whiz with this Comprehensive Guide to Seaborn in Python, Retrieved: February 13, 2020, <https://www.analyticsvidhya.com/blog/2019/09/comprehensive-data-visualization-guide-seaborn-python/>
- [18] statsmodels 0.11.0: Statistical computations and models for Python, Retrieved: February 13, 2020, <https://pypi.org/project/statsmodels/>
- [19] statsmodels 0.11.0: statsmodels, Retrieved: February 13, 2020 <http://www.statsmodels.org/stable/index.html>
- [20] patsy: Describing statistical models in Python, Retrieved: February 13, 2020, <https://patsy.readthedocs.io/en/latest/index.html>
- [21] Mankovits T., Szabó T., Kocsis I., Páczelt I.: Optimization of the Shape of Axi-Symmetric Rubber Bumpers Strojnicki Vestnik -Journal of Mechanical Engineering 60:(1) pp. 61-71 (2014)
- [22] Deák K., Mankovits T., Kocsis I.: Optimal wavelet selection for manufacturing defect size estimation of tapered roller bearings with vibration measurement using Shannon Entropy Criteria, Strojnicki Vestnik -Journal of Mechanical Engineering 63 : 1 pp. 3-14, 12 p. (2017)
- [23] Sziki G. Á., Sarvajcz K., Szántó A., Mankovits T.: Series Wound DC Motor Simulation Applying MATLAB SIMULINK and LabVIEW Control Design and Simulation Module. Periodica Polytechnica - Transportation Engineering & pp. 1-5, 5 p. (2019)
- [24] Kordic, B., Popovic, M., Ghilezan S.: Formal Verification of Python Software Transactional Memory Based on Timed Automata, Acta Polytechnica Hungarica, Vol. 16, No. 7, 2019, ISBN 1785-8860
- [25] Horváth Á.: The Cxnet Complex Network Analyser Software, Acta Polytechnica Hungarica, Vol. 10, No. 6, 2013, ISBN 1785-8860
- [26] O'Connor, P. D. T.; Kleyner, A.: Introduction to Reliability Engineering, Practical Reliability Engineering, Fifth Edition. 2012 John Wiley & Sons, Ltd. Published 2012 by John Wiley & Sons, Ltd.,

https://media.wiley.com/product_data/excerpt/28/04709798/0470979828-54.pdf

- [27] Huff, A. M.: Reliability Of a Lunar, Western Kentucky University, Retrieved February 13, 2020, https://digitalcommons.wku.edu/cgi/viewcontent.cgi?article=1302&context=stu_hon_theses
- [28] HSC: Numeracy for Science - Mathematics is fundamental to science, Retrieved February 13, 2020, <https://investigatingsciencehsc.com/numeracy-in-science/>
- [29] DiazEIE419Q3 - Percent Porosity Machine A Machine B Sample, Retrieved February 13, 2020, <https://www.coursehero.com/file/23869487/DiazEIE419Q3/>
- [30] Chowdavaram, GUNTUR – 522019., R.V.R. & J. C. COLLEGE OF ENGINEERING, Regulations (R-17), Scheme of Instructions, Examinations and Syllabi For Two Year M.Tech. Degree, Programme in MACHINE DESIGN
- [31] NCSS Statistical Software: Chapter 206, Two-Sample T-Test, Retrieved February 13, 2020, https://www.ncss.com/wp-content/themes/ncss/pdf/Procedures/NCSS/Two-Sample_T-Test.pdf
- [32] NCSS Statistical Software: Chapter 207, Two – Sample T-Test from Means and SD's, Retrieved February 13, 2020, https://ncss-wpengine.netdna-ssl.com/wp-content/themes/ncss/pdf/Procedures/NCSS/Two-Sample_T-Test_from_Means_and_SDs.pdf
- [33] Demetriadis, S.: Data Analysis and Hypothesis Testing Using the Python ecosystem, t-Test & ANOVAs, Aristotle University of Thessaloniki, Retrieved February 13, 2020, <http://pytolearn.csd.auth.gr/dnld/sdemetri@UVa2016-StatTests.pdf>
- [34] plotly – Graphing Libraries: T-Test in Python/v3, Retrieved February 13, 2020, <https://plot.ly/python/v3/t-test/#two-sample-t-test>