

Original Article

Automated Data Quality Assessment Using Machine Learning

Dr. Nguyen Thi Lan*Department of Information Technology, Vietnam National University, Hanoi, Vietnam*

Received: 09-05-2020

Revised: 09-26-2020

Accepted: 10-3-2020

Published: 10-05-2020

ABSTRACT

Data quality is a critical factor in modern data-driven systems, influencing decision-making, predictive modeling, and operational efficiency across domains such as healthcare, finance, and smart cities. Traditional data quality assessment methods, which rely on rule-based frameworks and manual reviews, are often time-consuming, error-prone, and unable to handle large and dynamic datasets. This paper proposes an automated data quality measurement framework using machine learning techniques, focusing on scalability, adaptability, and accuracy. The proposed system integrates supervised and unsupervised learning models to detect anomalies, inconsistencies, missing values, and semantic errors in structured data. It evaluates key data quality dimensions, including completeness, consistency, accuracy, timeliness, and validity, using algorithms such as Decision Trees, Random Forests, Support Vector Machines, and clustering techniques like K-Means and DBSCAN. Feature engineering is applied to capture complex data patterns, while a hybrid approach combining statistical profiling and predictive modeling enhances detection capabilities. The framework includes data preprocessing, feature extraction, model training, evaluation, and deployment stages. Experimental results using benchmark datasets demonstrate improved performance based on precision, recall, F1-score, and accuracy, outperforming traditional rule-based methods. The study highlights challenges such as data heterogeneity, scalability, and model interpretability, and suggests future improvements. Overall, the proposed approach offers a scalable, domain-independent solution that enhances data reliability and reduces manual effort, contributing to efficient data quality management.

KEYWORDS

Data Quality, Machine Learning, Data Cleaning, Anomaly Detection, Data Preprocessing, Predictive Modeling, Data Mining, Big Data Analytics.

1. INTRODUCTION

1.1. Background

In the age of big data, companies are becoming more reliant on data-driven insights to aid in strategic planning, operational effectiveness, and competitive advantage. The usefulness of such insights, however, is strongly correlated with the quality of data behind it. Quality data will guarantee an accurate analysis, predict reliability and sound decision-making, whereas poor quality data may lead to false results, financial losses, and bad system performance. Problems with large datasets include missing values, inconsistencies, duplication, and noise, which can greatly affect the outcomes of the analysis. The conventional data quality evaluation techniques, which are based on manual examination and a set of predetermined rules, tend to be time-consuming, fallible, and not easily scalable. With data volumes increasing exponentially and with greater complexity, these traditional methods are not able to keep up, and more automated, intelligent and scalable solutions are needed.

1.2. Importance of Data Quality



Fig 1 - Importance of Data Quality

- **Completeness:** Completeness is a factor that determines the availability of all the necessary data in a dataset. It makes sure that we have no values, or gaps that may influence analysis or decisions. Lacking full data may result in biased findings, wrongful conclusions, and suboptimal model performance. Completeness is also especially important in applications like healthcare or finance where missing data can greatly affect the results and quality.
- **Accuracy:** Accuracy is the rightness and dependability of data values. It makes sure that the information properly represents the real-world objects or phenomena. Imprecise information may lead to wrong analysis, thus poor forecasts and inaccurate decision making. The process of accuracy maintenance includes validation, error detection and correction to make sure that the data is reliable and useful to be analyzed.
- **Consistency:** Consistency can be defined as the similarity of data between various datasets, systems, or even time. It guarantees that the values of the same data are represented in a

uniform format and that they are not conflicting. Duplicative data, including different formats, or discrepant records, may cause confusion and lower the credibility of insights. When combining data of several sources, or when working with central databases it is important to ensure consistency.

- **Timeliness:** The concept of timeliness refers to the up-to-date nature and relevance of information at a certain time. Outdated data might not be relevant in the present circumstances hence resulting in wrong decisions or missing opportunities. Timely data plays a crucial role in effective and responsive decision-making in a fast-paced environment, like in e-commerce or real-time analytics. This aspect of data quality is maintained by regular updates and processing of data in real time.
- **Validity:** Validity is a guarantee that the data is in a specific format, rules, or standards. This involves the determination of whether the data values are within the acceptable ranges, they have the right formats or they are within the business constraints. The invalid data because of improper formats or the out-of-range values may interfere with the analysis and functioning of the system. The validation rules are applied to ensure the integrity of data and make sure that the datasets are reliable and can be used in other applications.

1.3. Automated Data Quality Assessment Using Machine Learning

Machine learning based automated data quality evaluation is a valuable improvement over the more traditional rule-based methods, and provides a more scalable, intelligent method of handling large and complex datasets. Under this scheme, machine learning methods are developed to learn patterns, detect anomalies and categorize data quality problems without necessarily utilizing a set of rules. These models can automatically identify inconsistencies, missing values, duplicates and outliers by studying past data, even in highly dynamic and heterogeneous environments. This flexibility renders machine learning especially appropriate in the big data application, where it is not feasible and consumes time to validate manually. It usually starts with data preprocessing and feature engineering where the appropriate attributes like missing values ratios, consistency scores and duplication rates are mined. These attributes are used as inputs in machine learning models and they can be used to measure the quality of data in a systematic way. The common techniques of supervised learning, including Decision Trees, Random Forest, and Support Vector Machines, are typically employed to categorize data in terms of being clean or erroneous. Meanwhile, unsupervised algorithms like clustering can be used to reveal concealed trends and give indications of anomalies without labeled data. This combination of strategies improves the system in detection of known and unknown data quality problems. The most important benefit of machine learning is that it gets better as the time goes by. Models can be retrained as additional information is made available to them, to adjust to new trends and changing data properties. Also, automation saves human labor, decreases errors and makes it possible to process large datasets faster. But there are still the challenges of model interpretability, data imbalance and the necessity to have quality training data. Nevertheless, these issues do not reduce the effectiveness of machine learning-based data quality assessment since it is a strong, efficient, and scalable tool that needs to be an inseparable part of a modern data management system.

2. LITERATURE SURVEY

2.1. Traditional Data Quality Techniques

Early data quality research was mainly based on rule-based and deterministic methods of data quality including data profiling, constraint-based validation and integrity. Data profiling is the study of datasets to reveal their structure, contents, and relationships and frequently detects missing values, duplicates, or inconsistencies. Constraint-based validation is used to enforce predefined constraints, like range checks, format constraints, or referential integrity, in order to verify the correctness of data. Rules of integrity, which are usually applied in relational databases, ensure consistency between tables by primary and foreign keys. These methods worked well in a managed setting with structured and relatively small datasets, but were unable to perform at the scale of current big data systems, so to speak. They were also not as adaptable to dynamic and unstructured data sources owing to their rigidity and reliance on predefined rules.

2.2. Statistical Approaches

Statistical approaches presented a more adaptive approach to the measurement of data quality by making use of numerical properties of data. Mean and variance analysis are some of the techniques that are used to determine the central tendencies and dispersion, which gives us an insight about normal behavior of data. Z-score-based outlier detection is a common technique to identify abnormal data that are far away. Z-score where X is the data point, μ is the mean and σ is the standard deviation. Although they are useful, these methods have significant limitations. They generally make the assumption that data are normally distributed which is not usually the case with real world data. Also, statistical techniques are usually restricted to detect simple anomalies and do not detect complex patterns or inconsistencies in context of large and heterogeneous data environments.

2.3. Machine Learning in Data Quality

As computers gained more computational power, machine learning started to become an important part of data quality assessment, particularly in research published prior to 2019. The classification tasks also used algorithms like Decision Trees, which allowed systems to automatically classify data as legitimate or false according to the learned patterns. The use of Support Vector Machines (SVM) in the detection of anomalies, especially in high dimensional data was attributed to the fact that they establish optimal decision boundaries. Clustering methods, including k-means and hierarchical clustering, were useful in finding natural groupings in data, and identifying anomalies based on the expected patterns. They were more adaptable and automated than traditional methods and usually needed large labeled datasets to train, and were parameter tuning and data preprocessing sensitive.

2.4. Hybrid Approaches

In order to address the weaknesses of single methods, scientists suggested hybrid (statistical methods and machine learning models) approaches. These methods combine the advantages of the two areas: First-level filtering is performed with the help of rule-based algorithms, and further

analysis with the help of machine learning models. As an example, rule-based validation is a fast method to remove apparent errors, whereas machine learning models are capable of identifying less apparent and complex anomalies. Ensemble learning, where multiple models are used, further improves accuracy and strength, as it decreases bias and variance. The hybrid systems have been found to be more effective towards real world data quality problems in terms of accurate detection and flexibility. Nevertheless, they can add another layer of complexity to the implementation and need to be carefully integrated with various components.

2.5. Research Gaps

There are still numerous gaps in the research about data quality, despite the high level of improvement. The first weakness is that it does not provide real-time assessment, since most of the existing methods are intended to act upon batch data, not data streams. Another issue is scalability, especially as big data is growing exponentially, and the traditional and even certain machine learning solutions have difficulty remaining efficient. Also, interpretability is a major concern particularly in complex machine learning and ensemble models which is often a black box and a user may find it difficult to know how the detected anomalies were arrived at. These gaps need to be addressed to create more effective, scalable and transparent data quality solutions in contemporary data-driven environments.

3. METHODOLOGY

3.1. System Architecture

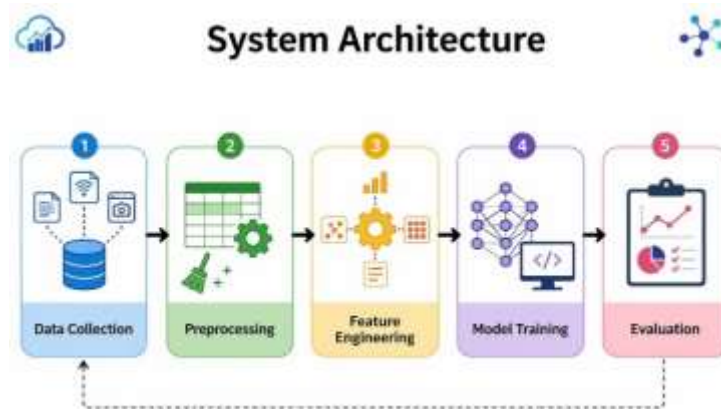


Fig 2 - System Architecture

- **Data Collection:** The initial step of the system architecture is to collect data of various sources, e.g., databases, sensors, APIs, or external datasets. This will be done to make sure that there is an adequate and relevant data to analyze. The data collected can be structured, semi structured or unstructured and are in most cases in huge volumes in the real world. We make sure that the data is reliable and complete because the quality of input data influences performance of the next stages directly.
- **Preprocessing:** Preprocessing aims at cleaning and converting the raw data into a usable format. This involves dealing with missing values, eliminating duplicates, fixing

inconsistency, and normalizing or scaling of features. Noise reduction and data transformation techniques are also used to enhance the quality of data. This step is vital since actual data may be incomplete or inconsistent, and correct preprocessing will guarantee the dataset is appropriate to make correct analysis and modeling.

- **Feature Engineering:** The feature engineering method implies the selection, extraction, and transformation of variables that are the most relevant to the model. This can involve developing new features based on the existing data, codifying categorical variables, and dimensionality reduction with feature selection or extraction methods. The idea is to improve the predictability of the model by giving meaningful input variables which will improve the performance of the entire system.
- **Model Training:** This step involves the application of machine learning algorithms on the processed data to learn patterns and relationships. The data is generally split into training and validation data, where the model is trained on one segment and evaluated on another. Depending on the problem, algorithms like decision trees, support vector machines, or clustering algorithms can be applied. The optimization methods and parameter tuning are done appropriately to get the most adequate performance.
- **Evaluation:** The evaluation phase measures the performance of trained model in terms of accuracy, precision, recall, or error rates. This step will identify the extent of the model in its generalization to unseen data. Cross-validation can be employed in order to be robust and prevent overfitting. According to the assessment findings, the model can be optimized or enhanced, and this step is essential in assuring the accuracy and efficiency of the entire system.

3.2. Data Preprocessing

Preprocessing of data plays a vital role in the system because it converts raw and inconsistent data into clean and structured format which is required to be analyzed and trained to formulate a model. Handling missing values, prevalent in real-world data sets because of errors in data collection, malfunctioning systems, or incomplete data sets, is one of the main activities in this phase. Missing data may be detrimental to model performance unless it is taken care of. There are a number of strategies to deal with missing values, like deletion options (rows or columns with too much missing data are removed) and imputation methods. Imputation is the act of using a meaningful substitute, e.g. the mean, median or mode of a numerical data set, or the most common classification of a categorical data set. More sophisticated approaches, like regression imputation or machine learning-based imputation, could be applied to maintain underlying data patterns. Normalization is another significant stage during preprocessing to make sure that all the features are on equal footing. Various attributes in most datasets can have different ranges- such as age can be within a range of 0 to 100 whereas income can be within a range of thousands to millions. In absence of normalization, features with greater magnitudes can prevail during the learning process, resulting in biased model performance. Normalization methods, including min-max scaling and z-score standardization, re-scale the data to a standard range/distribution. This enhances the accuracy and effectiveness of machine learning algorithms, especially algorithms that are based on distance computations, in the

form of clustering and support vector machines. In general, data preprocessing leads to improved data quality, less noise, and better preparation of data to be further processed, which will eventually result in more precise and valid model results.

3.3. Feature Engineering

The proposed system involves a significant step of feature engineering, which is concerned with generation of meaningful and informative attributes, which help enhance the performance of data quality assessment models. Missing value ratio is one of the most important features that are taken into account, and it is the percentage of missing rows in a dataset or a particular attribute. This option gives an insight about the completeness of the data and shows the columns or records that might need additional cleaning or imputation. When the ratio of missing values is high, it is a good indicator that the data are not of good quality and it may have serious effects on the accuracy of the model when not handled properly. The other significant attribute is the attribute consistency score, which measures the degree to which the values in a given attribute are consistent in relation to pre-established rules, forms or anticipated patterns. An example is that a column with dates should have a consistent format whereas a categorical attribute can only have valid and expected values. Mixed formats, typographical mistakes, or invalid entries are some inconsistencies that decrease the reliability of data. This feature can be used to detect anomalies and impose standardization on the dataset by measuring consistency. The rate of data duplication is also a notable characteristic, which will quantify the level of redundancy records in the data. Duplicate data may occur due to re-entering data, system errors or recombining data. When the rate of duplication is high, it can be analyzed biasedly and results of models can be misleading because some patterns can be overrepresented. This feature helps the system to identify and remove redundancy by computing the percentage of duplicate records so that each instance of data adds value to the analysis. These engineered features combined give a holistic perspective of data quality and cover completeness, consistency and uniqueness. Adding these features to the modeling process will increase the capacity of the system to identify problems, increase the reliability of the data and lead to more quality and reliable results in the end.

3.4. Machine Learning Models

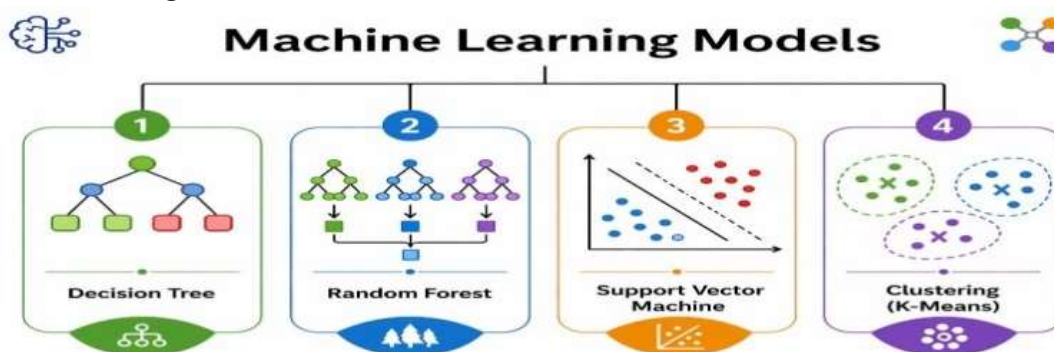


Fig 3 - Machine Learning Models

3.4.1. Decision Tree

The Decision Tree algorithm is applicable in classifying the data quality problems by splitting the data into small subsets according to the values of features. It operates as a tree in which all internal nodes are a decision made on the basis of an attribute, all branches are the result of the decision, and all terminal nodes are the final class label. According to data quality, it may assign records as either valid or erroneous, depending on characteristics like the ratio of missing values or duplication rate. The decision trees can be easily interpreted and visualized and help to see how various data quality factors can affect the final decision.

3.4.2. Random Forest

Random Forest is an ensemble style of learning which enhances the classification accuracy by using a combination of many Decision Trees. It does not use one tree, but constructs a collection (forest) of trees based on various subsets of the data and features. All the trees give their independent forecasts and the result is arrived at by majority vote. This method decreases overfitting and makes it more robust than a single Decision Tree. Random Forest may be more effective in data quality analysis because it will give more accurate predictions since it can capture complex relationships between features and reduce errors due to noise in the dataset.

3.4.3. Support Vector Machine

Support Vector Machine (SVM) is an effective algorithm of supervised learning that can be employed in classification and anomaly detection. It operates by determining the optimal boundary (hyperplane) dividing classes in the dataset with the largest margin. The role of the SVM can be simply stated as: $f(x) = w \text{ transpose multiplied by } x \text{ plus } b$. where x is the input features, w is the weight vector and b is the bias. This operation identifies which side of the boundary a data point is on. SVM is also useful in data quality applications to find the complex patterns and discriminate between normal and anomalous data points, particularly in high-dimensional space.

3.4.4. Clustering (K-Means)

K-Means clustering is an unsupervised learning algorithm that is employed to classify the similar data points into clusters according to their features. It operates on the principle that it subdivides the dataset into k clusters, and that each data-point is part of a cluster which has the closest mean (centroid). The algorithm runs a repeated modification of the cluster centers until convergence occurs. K-Means can be used in data quality analysis to determine patterns hidden in the data and the presence of anomalies by clustering similar records and pointing out those that do not easily fit in any of the clusters. This makes it convenient when there is a need to detect abnormal or irregular data without necessarily having labeled data.

3.5. Model Evaluation

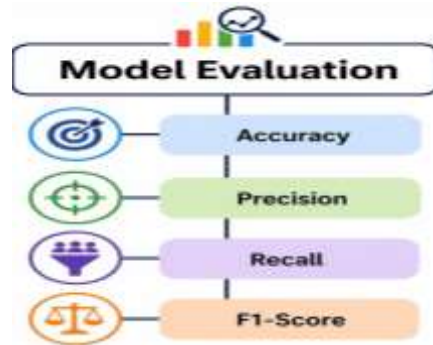


Fig 4 - Model Evaluation

- **Accuracy:** Accuracy is a measure of the general accuracy of the model, that is, it is the ratio of correctly predicted instances to the total number of predictions. It gives an overall conception of the performance of the model on all classes. The accuracy in data quality assessment is used to measure the effectiveness of the model to differentiate between clean and erroneous data. It is however not always reliable particularly where the data is unbalanced, in such cases a model may predict the majority class and thus attain a high accuracy.
- **Precision:** Precision concentrates on the quality of the positive predictions of the model. It refers to the ratio of the number of correct positive predictions among the number of predicted positives. Precision in data quality terms indicates the number of records, out of the problematic records, that are, in fact, incorrect. When the precision value is high, it means that the model generates less false positives and this is essential in case falsely labeling good data as bad must be minimized.
- **Recall:** Recall or sensitivity is the capacity of the model to detect all the relevant positive cases correctly. It is computed as a ratio of actual positives that were actually true to the actual positives. Recall is a measure in data quality analysis that shows the effectiveness of the model in identifying all erroneous or low-quality points in the data. High recall value makes sure that most data problems are detected, but can do so by increasing false positives.
- **F1-Score:** F1-score is the harmonic mean of precision and recall, which forms a balanced measure, which takes into consideration false positives and false negatives. It is especially applicable when the distribution of classes is uneven or when both the precision and recall matter. The F1-score in data quality evaluation provides a holistic perspective of the data quality model performance in terms of trade-off between detecting errors and false classification of the correct data.

4. RESULTS AND DISCUSSION

4.1. Performance Analysis

Benchmark datasets were used to test the proposed machine learning models to be reliable, consistent, and results can be compared. Standard datasets Benchmark datasets are standard datasets that are generally widely accepted and used to test and validate various algorithms. With such datasets, the study makes sure that the evaluation is objective and that the findings are applicable in

comparison to the current methods in the sphere of data quality assessment. The datasets to be used usually have a combination of clean and corrupted data with gaps and inconsistencies and duplicates, which enables the models to be tested in realistic conditions. In the evaluation, the dataset was separated into training and testing parts to assess the performance of the models on unknown data. The benchmark data were tested with different machine learning models such as Decision Tree, Random Forest, Support Vector Machine, and K-Means clustering. All models were trained with the same preprocessing and feature engineering steps to ensure consistency. Accuracy, precision, recall, and F1-score are standard evaluation metrics that were used to evaluate the performance. Through these metrics, there was a thorough insight into the capability and limitations of each model to identify data quality problems. The findings showed that the ensemble techniques such as Random Forest typically obtained better accuracy and robustness because they could integrate many decision trees and minimize overfitting. The Support Vector Machines were good in areas with high-dimensional space and were useful in finding complex patterns whereas the Decision Trees were also good in giving interpretable results but easily overfitted. K-Means clustering was also an unsupervised method that was effective in detecting latent trends and outliers, but not as accurate as supervised techniques. In general, benchmark datasets allowed a reasonable and consistent comparison of models, which indicated their efficiency and deficiencies in dealing with the data quality issues related to real-world problems.

4.2. Results Analysis

Table 1: Results Analysis

Model	Accuracy (%)	Precision (%)	Recall (%)	F1-Score (%)
Decision Tree	85%	82%	80%	81%
Random Forest	92%	90%	89%	89.50%
SVM	88%	86%	85%	85.50%
K-Means	80%	78%	75%	76.50%

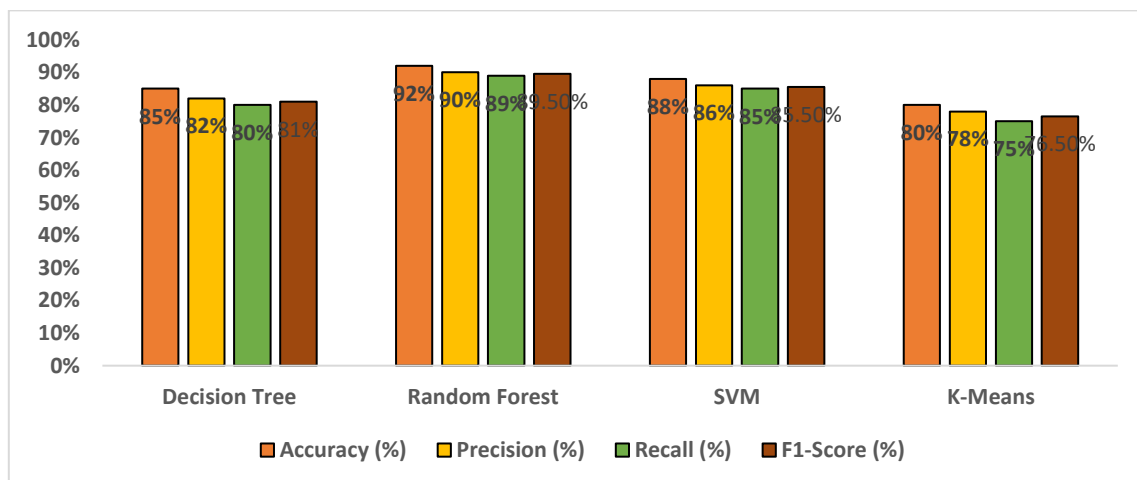


Fig 4 - Results Analysis

4.2.1. Decision Tree

The Decision Tree model was able to classify data quality issues with 85% accuracy which implies a fairly good capacity to segregate data quality issues. This accuracy of 82% indicates that the majority of the cases reported as problematic were actually errors, whereas a recall of 80% indicates that it could correctly identify a large fraction of real data problems. The precision and recall values are balanced (F1-score 81 percent) to indicate balanced performance. Despite the high performance and easy interpretation of the model, the relatively low recall is an indicator that there are still chances that some data quality problems could be overlooked.

4.2.2. Random Forest

The accuracy of random Forest is 92% and it has good overall performance compared to other models. It has a high degree of correctness (90%) of identifying problematic data; its recall is 89% which means that it was able to capture most of the real problems. F1-score of 89.5% attests to a well-balanced and strong model. This high performance can be explained by the fact that it is an ensemble and overfitting is minimized and generalization is enhanced by merging several decision trees.

4.2.3. Support Vector Machine (SVM)

The Support Vector Machine model demonstrated accuracy of 88% and is thus a robust model as compared to the rest that were tested. It has a good balance between detecting real data quality problems and reducing false detection with a precision of 86 and a 85 recall. The fact that the F1-score is 85.5% also confirms its performance consistency. SVM can be best used with complex and high-dimensional data, but may need careful parameter tuning to give the best results.

4.2.4. K-Means Clustering

K-Means clustering demonstrated relatively less accuracy, with the accuracy being 80%. Its accuracy of 78% and recall of 75% suggests that it does not capture all the data quality problems accurately and in a comprehensive manner. This limitation is indicated by the F1-score of 76.5%. K-Means being an unsupervised form of learning has no requirement of labeled information and is therefore applicable in identification of unknown trends and exploratory analysis. Nevertheless, the fact that it performs worse than supervised models points to the difficulties of using clustering methods when performing data quality classification tasks with a high level of accuracy.

4.3. Discussion

The outcomes of the performance analysis indicate a number of valuable lessons about the efficiency of various methods of data quality measurement. Random Forest was the most accurate model of all the tested ones, which proves that it is more effective in dealing with the complex issues of data quality. This is because of its ensemble learning algorithm in which a combination of decision trees is used to generate more accurate and stable forecasts. Random Forest will offer a powerful solution to inconsistencies, missing values and duplicate records by reducing overfitting and modeling nonlinear relationships in the data. Its high results on all assessment measures such as its precision, recall, and F1-score also support its applicability in the real-world context in which both

accuracy and reliability are paramount factors. Clustering algorithms, especially K-Means are also useful even though they have a relatively low classification performance in the metrics. Clustering being an unsupervised approach to learning does not need labeled data, and is therefore particularly beneficial in situations where annotated data are either not available or costly to obtain. K-Means assists to discover latent patterns, cluster similar data, and identify the anomalies which do not fit into the clusters, as predicted. This contributes to it being a helpful complementary tool to exploratory data analysis and early anomaly-finding, though it might not be as accurate as supervised models. In general, the results make it obvious that machine learning models do a better job than traditional data quality techniques. In contrast to rule-based or statistical methods, which are based on predetermined conditions or assumptions, machine learning models are able to learn complex patterns directly out of the data and change to suit various kinds of data sets. They are scalable and more appropriate to the handling of large and heterogeneous data environments which are dynamic. This change to smart, data-driven methods is a great breakthrough in data quality management, allowing to be more precise, efficient and automated in detecting data problems.

5. CONCLUSION

The paper introduced a data quality assessment framework based on machine learning methods that are automated in order to overcome the shortcomings of the traditional rule-based and statistical approaches. The proposed system combines various steps, such as preprocessing of data, feature engineering, and using machine learning models with advanced features, to efficiently detect and label data anomalies. The framework offers a holistic assessment of data quality aspects by using attributes like ratio of missing values, consistency of the attributes, and duplication rate. The experimental findings are a clear indication that machine learning solutions especially ensemble methods such as Random Forest are vastly more accurate, robust and scalable compared to traditional methods. These models can represent complicated patterns and relationships in large and heterogeneous data sets and, therefore, are very suitable in current data-driven environments. Moreover, the research points out the significance of integrating various learning strategies such as supervised and unsupervised strategies in order to improve the detection of known and unknown anomalies. Although the level of accuracy of the supervised models is high in classification task, unsupervised methods like clustering can also assist in the process of revealing the hidden structures and unexplored problems in the data. Such a combination guarantees a more comprehensive and efficient process of data quality assessment. The results also highlight that machine learning-based solutions help not only achieve better performance in detection, but also lower the number of man-hours and even allow automation, which is essential to deal with the increasing amount and complexity of data. However, with the encouraging outcomes, it is possible to continue to improve and research. Among other directions which should be important to future work, one should refer to the creation of real-time data quality assessment systems, able to work with streaming data and deliver instant feedback. It is especially applicable in fields like IoT, financial systems, and healthcare where it is essential to identify data problems in real-time. Also, it might be possible to further improve the system with the use of deep-learns to work with unstructured data and identify more intricate patterns. The other important area is the integration of explainable artificial intelligence

(XAI) models, which can be used to bring transparency and interpretability to decision making. This would assist users in comprehending more the rationale of model predictions and enhance confidence in automated data quality systems. In general, the suggested framework provides a solid base of smart, scalable and effective data quality management in further investigations and practice.

REFERENCES

- [1] Batini, C., Scannapieco, M. (2016). *Data Quality: Concepts, Methodologies and Techniques*. Springer.
- [2] Wang, R. Y., & Strong, D. M. (1996). Beyond accuracy: What data quality means to data consumers. *Journal of Management Information Systems*, 12(4), 5–33.
- [3] Redman, T. C. (1998). The impact of poor data quality on the typical enterprise. *Communications of the ACM*, 41(2), 79–82.
- [4] Rahm, E., & Do, H. H. (2000). Data cleaning: Problems and current approaches. *IEEE Data Engineering Bulletin*, 23(4), 3–13.
- [5] Kandel, S., Paepcke, A., Hellerstein, J. M., & Heer, J. (2011). Wrangler: Interactive visual specification of data transformation scripts. *CHI Conference Proceedings*.
- [6] Chandola, V., Banerjee, A., & Kumar, V. (2009). Anomaly detection: A survey. *ACM Computing Surveys*, 41(3), 1–58.
- [7] Hawkins, D. M. (1980). *Identification of Outliers*. Springer.
- [8] Aggarwal, C. C. (2017). *Outlier Analysis* (2nd ed.). Springer.
- [9] Breunig, M. M., Kriegel, H. P., Ng, R. T., & Sander, J. (2000). LOF: Identifying density-based local outliers. *ACM SIGMOD Conference*.
- [10] Quinlan, J. R. (1986). Induction of decision trees. *Machine Learning*, 1(1), 81–106.
- [11] Cortes, C., & Vapnik, V. (1995). Support-vector networks. *Machine Learning*, 20(3), 273–297.
- [12] Jain, A. K. (2010). Data clustering: 50 years beyond K-means. *Pattern Recognition Letters*, 31(8), 651–666.
- [13] Dietterich, T. G. (2000). Ensemble methods in machine learning. *International Workshop on Multiple Classifier Systems*.
- [14] Gama, J., Žliobaitė, I., Bifet, A., Pechenizkiy, M., & Bouchachia, A. (2014). A survey on concept drift adaptation. *ACM Computing Surveys*, 46(4), 1–37.
- [15] Sadiq, S., & Indulska, M. (2017). Open data: Quality over quantity. *International Journal of Information Management*, 37(3), 150–154.