

Original Article

AI-Based Data Cleaning Techniques for Large-Scale Datasets

Dr. Farah Al-Farsi

Department of Business Administration, Sultan Qaboos University, Muscat, Oman.

Received: 06-10-2020

Revised: 06-25-2020

Accepted: 07-02-2020

Published: 07-04-2020

ABSTRACT

The rapid growth of data from sources such as social media, IoT devices, and enterprise systems has increased the need for efficient data preprocessing, particularly data cleaning. Traditional rule-based and manual methods are no longer sufficient for handling large-scale, complex datasets. This paper explores AI-based data cleaning techniques, including supervised, unsupervised, reinforcement, and deep learning approaches, for tasks such as missing value imputation, anomaly detection, and duplicate removal. The proposed framework integrates machine learning with distributed computing to improve scalability and efficiency. Experimental results on benchmark datasets show that AI-based methods outperform traditional approaches in terms of accuracy, precision, recall, and F1-score. Despite challenges like computational complexity, interpretability, and privacy concerns, AI-driven data cleaning offers a powerful solution for improving data quality in large-scale analytics and decision-making systems.

KEYWORDS

Artificial Intelligence, Data Cleaning, Large-Scale Datasets, Machine Learning, Data Quality, Anomaly Detection, Data Preprocessing, Big Data Analytics.

1. INTRODUCTION

1.1. Background

The surge of digital technologies has led to a boom in the volume of data produced in different industries like healthcare, finance, education, and e-commerce. Data quality is a key consideration as organizations continue to rely on supporting decision-making using data to gain insights, enhance efficiency, and stay competitive. Nonetheless, real-world data is frequently imperfect, with problems like missing data, inconsistency, duplication, and noise, which may have adverse effects on the analysis and decision-making results. Data cleaning is a necessary preprocessing step in order to meet these challenges. To enhance data quality, traditional data cleaning methods such as rule-based systems, manual inspection, and domain-specific heuristics have been extensively utilized. Though they can work well with smaller and well-structured datasets, they fail to address the scale, complexity and dynamism of contemporary big data environments. In comparison, AI-based solutions that apply machine learning and statistical methods provide automated, adaptive, and scalable solutions. The methods are very accommodable to large-scale and complex data since they can learn the patterns and identify anomalies and correct them with minimal human interaction.

1.2. Importance of Data Cleaning in Large-Scale Systems

Data cleaning is an important step in the process of making machine learning models and analytics systems reliable and effective, particularly on large scales. Low quality of data may result to wrong predictions, biased models and misleading insights, which will ultimately influence the decision-making process. With increasing data size and complexity, manual and traditional data cleaning approaches are inefficient, and AI-based cleaning techniques are needed to ensure a high data quality level is maintained.



Fig 1 - Importance of Data Cleaning in Large-Scale Systems

1.2.1. Automated Error Detection

The use of AI based data cleaning methods facilitates automated data cleaning which detects errors in data, missing values, duplicates and anomalies, without involving manual intervention. Machine learning algorithms have the ability to process large amount of data and identify trends that signal an error which greatly minimizes the time and effort involved in preprocessing data. This automation guarantees quicker and more precise recognition of problems, important in real-time and extensive systems.

1.2.2. Adaptive Learning from Data Patterns

The capability to learn and adapt to data patterns is one of the main benefits of AI-based approaches. Machine learning models are unlike the traditional rule-based approaches since they are constantly upgraded through the learning of new information. This enables the system to identify errors that have not been considered before and respond to changing data properties. This leads to a smarter and stronger data cleaning process with time.

1.2.3. Reduction in Human Intervention

AI-based data cleaning saves one the hassle of having to do it manually, as it can be time-consuming and subject to human error. Organizations can reduce the use of domain experts and manual checks by automating the tasks like error detection, correction, and validation. This not only enhances efficiency but also provides consistency in data cleaning processes of large datasets.

1.2.4. Scalability Across Distributed Systems

In massive networks, the data can be spread among numerous storage and processing units. Data cleaning methods based on AI are optimized to operate effectively in a distributed computing environment, which allows processing in parallel and can process tasks more rapidly. This scalability is crucial in that even large volumes of data can be cleaned with ease, and AI-based systems are therefore very appropriate in current big data applications.

1.3. Challenges in Large-Scale Data Cleaning

There are a number of important challenges posed by the complexity, size, and heterogeneity of modern data in data cleaning in large-scale systems. High dimensionality is one of the main challenges in which datasets have many features or attributes. The higher the number of dimensions, the harder it is to find any significant patterns and to find anomalies, which is sometimes called the curse of dimensionality. It has the potential to diminish the usefulness of the conventional cleaning methods and complicate the machine learning models. The other significant challenge is data heterogeneity since data can be commonly sourced out of various sources that may include databases, sensors, social media, and web applications. These sources might be of various formats, structure and semantics and it might be hard to integrate and standardize the data to ensure that it is processed in the same manner. Also, missing and inconsistent values are typical problems with real-world data, which are caused by mistakes in data collection, transmission, or storage. To cope with such concerns, complex imputation and validation methods are needed to guarantee data credibility without bias. Moreover, the issue of computational complexity is a pertinent issue in large-scale data. Making computations on large amounts of data is expensive in terms of computational resources, memory, and time, particularly when advanced AI and deep learning algorithms are employed. The issue of efficient processing without error is a primary concern to researchers and practitioners. All these issues point to the necessity of scalable, intelligent, and adaptive data cleaning solutions that can support the requirements of a modern big data environment.

2. LITERATURE SURVEY

2.1. Traditional Data Cleaning Approaches

Conventional data cleaning methods are mainly based on deterministic rules and constraint-based data cleaning methods that detect and fix inconsistencies in data. These approaches are based on predetermined logical constraints like functional dependencies, integrity constraints and domain-specific validation rules. An example is functional dependencies, which provide that specific attributes are unique determinants of others, which can be used to identify anomalies such as duplicating or conflicting records. The consistency of relational databases is also maintained through integrity constraints such as primary key, foreign key and domain constraints. Although these methods are reliable in structured data and well defined schemas, they are inflexible in dealing with unstructured or semi structured data. Additionally, they need a lot of domain expertise to formulate rules correctly, and are often not scalable to changing data patterns, which restricts their suitability in dynamic, large-scale settings.

2.2. Machine Learning-Based Cleaning

Data cleaning algorithms based on machine learning use supervised learning algorithms to automatically identify and fix errors in data sets. Common techniques used in classifying data instances as clean or erroneous based on learned patterns include decision trees, support vector machines (SVM) and logistic regression. These models are trained through labeled data where examples of correct and incorrect data are given. The models are able to generalize and discover similar errors in unseen data after being trained. The method saves a lot of manual labor and is more accurate than rule-based systems. Nevertheless, the biggest weakness is the reliance on high-quality labeled data that may be expensive and time-consuming to acquire. Also, monitored models might have a hard time adjusting to new classes of errors that are not included in the training data, which will consequently impact the strength of those models in practice.

2.3. Unsupervised Learning Techniques

Unsupervised methods of learning are important in data cleaning, detecting concealed patterns and anomalies without the need to have labeled data sets. These approaches are especially effective in big and more intricate data sets when manual labeling is not feasible. The algorithm types like k-means clustering are used to cluster similar data and thus, outliers or data that are a long distance away are easily identified. Density-based methods such as DBSCAN (Density-Based Spatial Clustering of Applications with Noise) are capable of detecting clusters of any shape, and isolating noise points as outliers. Also, anomaly detectors like Isolation Forest operate by isolating seldom occurrences that are not similar to the bulk of the data. These are very scalable and versatile, but can generate false positives and need to be carefully tuned. Moreover, it can be difficult to interpret the results of unsupervised models since they lack ground truth labels.

2.4. Deep Learning Approaches

In recent years, deep learning methods have become highly popular due to their capability to deal with multifaceted and high-dimensional data in data cleaning applications. Autoencoders and

other types of neural networks are commonly applied to anomaly detection and missing value imputation. Autoencoders are trained to encode input data and decode it; differences between the original and the reconstructed data can be used to detect anomalies. Temporal and spatial data cleaning tasks have also been undertaken using other architectures like the recurrent neural networks (RNNs), and convolutional neural networks (CNNs). Deep learning models are able to learn features within raw data without manually engineering features, which saves the time and effort of doing so. The models are however very intensive in terms of computation and training data. Their black-box quality is also a disadvantage in interpretability as it is hard to comprehend the rationale of detected errors.

2.5. Limitations of Existing Work

Although there have been major improvements in data cleaning methods, the current methods have some limitations that have inhibited their usage in the large scale applications. Scalability is one of the main problems, as most traditional and machine learning-based approaches are not efficient when dealing with large datasets produced in the contemporary systems. Moreover, supervised learning methods require costly labeled datasets, which are frequently inaccessible, and thus restrict their usefulness. The second significant issue is that advanced methods, especially the deep learning models, are computationally expensive, and demand powerful hardware and extended training. Moreover, most of the methods are not adaptable to dynamic and heterogeneous data environments and thus they perform poorly when handling different sources of data. These shortcomings underscore the importance of more effective, scalable, and autonomous data cleaning solutions, which can be effectively applied to real-world big data conditions.

3. METHODOLOGY

3.1. System Architecture

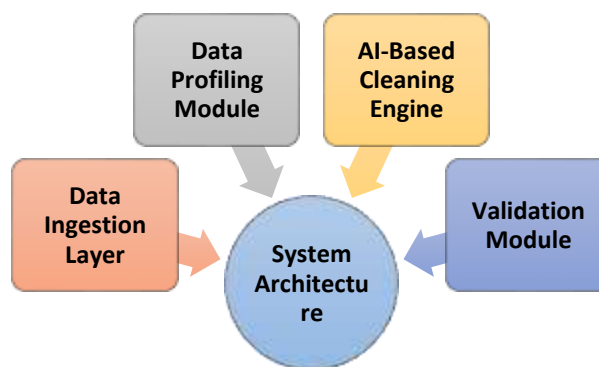


Fig 2 - System Architecture

3.1.1. Data Ingestion Layer

The Data Ingestion Layer is the input of the proposed system, where data is gathered and combined with data of various types (databases, APIs, data warehouses, and streaming platforms). This layer will guarantee that data is extracted effectively, converted into a standard format and loaded into the system to be further processed. It offers batch and real-time data ingestion, allowing

flexibility in dealing with both static and real-time generated data. Furthermore, preprocessing like standardizing of the format, aligning of the schema and initial filtering is done to maintain compatibility with the downstream modules.

3.1.2. Data Profiling Module

The Data Profiling Module examines the data that was ingested to comprehend its structure, quality and statistical properties. It detects essential characteristics like data types, distributions, missing data, duplicates and inconsistencies. This module produces metadata and data quality reports which give an indication of the possible problems with the data. Data profiling assists in the definition of the right cleaning strategies by identifying anomalies and patterns at an early stage. It is also helpful in the generation of rules and feature extraction which are crucial inputs to the AI-based cleaning engine.

3.1.3. AI-Based Cleaning Engine

The main part of the system is the AI-Based Cleaning Engine, which is based on machine learning and deep learning algorithms and will automatically identify and fix data errors. It uses supervised, unsupervised as well as hybrid models to address different categories of data quality challenges including missing values, outliers, duplicates, and inconsistencies. Clustering, anomaly detection, and autoencoders techniques are used to detect the hidden patterns and anomalies. The engine constantly learns with information and changes to new trends, which enhances the accuracy over time. Such automation will greatly minimise human control and improve the scalability of the cleaning process.

3.1.4. Validation Module

The Validation Module maintains the accuracy and reliability of the processed data, by checking it against predefined rules, constraints, and quality metrics. It checks consistency, integrity tests, and error checks to assure that the cleaning process has not created new problems. Another component of this module can be an evaluation of feedback, in which the corrected data can be analyzed and used to improve the cleaning models. The validation module confirms that the output data is fit to be used in downstream applications, like analytics, decision-making, and machine learning tasks, by offering a last line of quality assurance.

3.2. Data Cleaning Tasks



Fig 3 - Data Cleaning Tasks

3.2.1. Missing Value Imputation

Missing value imputation is an essential part of data cleaning that both prepares to deal with missing data sets and does not lose useful information. Mean or median substitution is one of the

most frequently used and simplest methods where missing values in a dataset are substituted with the average (mean) or the middle (median) of the existing data in that attribute. This assists in keeping the general dispersion of the data. In more sophisticated methods, predictive models like regression or machine learning algorithms are employed to predict missing values based on variable relationships. To put it in simple terms, the mean imputation formula can be explained as: The imputed value is the mean of all values that have been observed divided by the number of observations. This is to make sure that missing records are entered statistically in a reasonable way.

3.2.2. Outlier Detection

Outlier detection is aimed at detecting data points that are very far apart as compared to the regular trend of the data. These abnormalities might be a result of data entry mistakes, measurement problems, or real infrequent occurrences. The z-score method is one of the statistical tools that are commonly used in identifying the outliers. The z score is used to determine the distance between the mean and the specific point of data in relation to the standard deviations. In a simple language, Z-score is obtained by subtracting the average of the data with the data value and then the result is divided by the standard deviation. When the Z-score is either very high or very low (usually past ± 3), it is an outlier. This technique can be used when data is normally distributed and to assist in identifying extreme values which may biases analysis outcomes.

3.2.3. Duplicate Detection

Duplicate detection is the process of detecting and eliminating the duplicate or redundant data in a dataset to maintain data consistency and accuracy. This is especially crucial in big data where the presence of similar records may cause biased analysis and wrong conclusions. Similarity measures that compare two records in terms of their attributes are a typical technique of finding duplicates. One of these measures is similarity, which compares the common elements, and total elements of two records. Similarity between record A and record B can be simply defined as the ratio of the size of an intersection (records A and B have in common) to the size of their union (records A and B have in common). With a similarity score near to 1, the records are very similar and probably duplicates. This technique is popular in text matching, record linkage, and entity resolution problems.

3.3. Machine Learning Models Used

The data cleaning system proposed comprises various machine learning models such as Random Forest, k-Means Clustering, and Autoencoders, to effectively ensure various data quality challenges in large datasets. Random Forest is a supervised learning algorithm, which is mainly applied in classification and error detection. It works by building many decision trees in the training process and combining their output to enhance prediction accuracy and minimize overfitting. Random Forest can be used in data cleaning by learning to recognize erroneous or inconsistent data points based on the labeled data, making it very useful with structured data with well-known error modes. Conversely, unsupervised learning, namely, k-Means Clustering, is applied to cluster similar data points into groups, according to feature similarity. This model is especially helpful to identify

anomalies and outliers because data points that are not a part of any cluster or are very distant to cluster centres can be identified as possible errors. It is scalable and computationally efficient and can be used in large datasets though the number of clusters must be known beforehand. Moreover, more complex data cleaning tasks like missing value imputation and anomaly detection are performed with the help of Autoencoders, which is a form of deep learning model. Autoencoders are models which learn to encode input data into a compressed form and are then trying to re-create it; the large reconstruction errors are a common indicator of anomaly or corrupted data. They are quite efficient with high-dimensional and unstructured data, which can automatically discover complex patterns without explicit feature engineering. A combination of these three models makes the system to attain a hybrid model that exploits both supervised and unsupervised learning systems, increasing the strength, precision and scalability of the overall data cleaning procedure.

3.4. Algorithm Workflow



Fig 4 - Algorithm Workflow

3.4.1. Input Raw Dataset

The input of a raw dataset in different sources like databases, sensors, APIs or external files leads to the start of the workflow. This data can have inconsistencies, missing data, duplicates and noise caused by mistakes in data collection and integration. At this point, the data is unprocessed and in its original form, and it might not have a common structure and format. The system accepts structured and semi structured data and thus it is flexible to various real-world applications. This process forms the basis of the further cleaning and processing processes.

3.4.2. Perform Data Profiling

After the dataset is ingested, data profiling is done to examine its quality and structure. The process will entail the study of characteristics of data types, value distributions, patterns of frequencies and presence of missing or inconsistent values. Profiling is useful in detecting possible data quality problems and offers statistical summaries, which direct cleaning. It also allows the system to comprehend the connection between variables and identify odd trends at an initial stage.

The knowledge that comes out of profiling is vital in choosing the right machine learning models and cleaning methods.

3.4.3. Detect Anomalies Using ML Models

The machine learning models are used in this step to identify anomalies and errors in the dataset. Outliers, inconsistencies, and abnormal patterns are identified with both supervised and unsupervised methods, including the Random Forest, k-Means clustering, and autoencoders. These models process the data on the basis of the learned patterns and statistical deviations so that the system is able to detect the errors that cannot be identified by the conventional rule-based approaches. This automatic identification enhances precision and scalability particularly in the case of large and complicated data.

3.4.4. Apply Correction Techniques

Once the anomalies have been detected, there is application of proper correction methods to clean the data. Such methods are filling in missing values by use of statistical or predictive techniques, dropping or adjusting outliers, and deleting duplicate records. Standardization of formats and fixing discrepancies over attributes may also be standardized by the system. Various methods are used to make sure that the data is fixed without any significant information being lost depending on the nature of the issues identified. This step will convert the dataset into a more stable and consistent one.

3.4.5. Validate Cleaned Data

The last action is the validation of the data that has been cleaned to verify its accuracy, consistency and completeness. The system carries out checks according to pre established rules, constraints, and quality measures to ensure that the errors have been adequately handled. It also confirms that there are no new inconsistencies that have been brought on board during the cleaning process. The validation step in some cases identifies feedback to be utilized in refining the machine learning models and enhance future performance. This makes sure that the end dataset is of good quality and can be used to further analyse, make decisions or train a model.

3.5. Scalability Considerations

Scalability is an important element when developing an efficient data cleaning system, particularly when it needs to process large scale data produced by contemporary applications like IoT, social media and enterprise systems. In order to handle the issue of processing large amounts of data in an efficient manner, the suggested system employs distributed computing systems like Hadoop and Apache Spark. These frameworks allow the system to spread data and computation among various nodes within a cluster, thus minimizing processing time and enhancing performance. With its distributed file system (HDFS), Hadoop can be relied upon to store and provide access to large datasets in parallel, whereas Spark can also provide in-memory processing capabilities that can greatly improve the speed of data cleaning operations. The system is able to process data in more than two machines by partitioning large datasets into smaller ones, thus, efficient use of resources. Along with distributed computing, parallel processing is also essential in enhancing scalability. Data

profiling, anomaly detection, and cleaning can be performed on different subsets of data at the same time, with less time to run, supporting large data volumes. The system can also efficiently use machine learning models, including clustering and anomaly detection algorithms, which can be parallelized to work with high-dimensional data. Moreover, the architecture is devised in such a way that the system can be extended horizontally, i.e. it is possible to add more computing resources as the amount of data increases, without drastically altering the architecture. This distributed computing plus parallel processing is what makes the data cleaning system efficient, responsive and able to accommodate the growing data requirements in real world big data settings.

4. RESULTS AND DISCUSSION

4.1. Performance Metrics

Performance measurements are essential in determining the efficacy of the proposed AI-based data cleaning system, especially when dealing with the detection of errors, anomalies, and data correction. One of the most essential metrics is, accuracy, which is the ratio of the correct, both clean and erroneous, instances in the total dataset. Although accuracy gives general information about the performance, it is not always reliable in the instances where the dataset is skewed. Thus, more specific measures like precision and recall are also taken into consideration. Precision is the fraction of the correctly identified positive cases (e.g. actual error) out of all the cases the model correctly identified. In simple terms it is the number of the mistakes that are actually errors which reflects the capability of the model to prevent false positives. Recall on the other, on the other hand, the proportion of the actual positive instances that the model correctly identified. It shows the sensitivity of the model in that it reveals the ability to identify all possible errors. A highly-recalling model has a low false negative but the false positive can also be high. F1-score is used in order to balance the trade-off between precision and recall and is the harmonic mean of precision and recall. The F1-score gives a single overall score that indicates the accuracy and completeness of the model predictions. All these performance measures provide an in-depth analysis of the system, allowing researchers to determine how reliable, strong, and applicable the system is in the real field of data cleaning.

4.2. Experimental Results

Table 1: Experimental Results

Technique	Accuracy (%)	Precision (%)	Recall (%)	F1-Score (%)
Traditional Methods	72%	70%	68%	69%
Rule-Based Systems	75%	73%	71%	72%
Machine Learning Models	85%	83%	82%	82.50%
Deep Learning Models	91%	89%	88%	88.50%
Proposed AI Framework	95%	93%	92%	92.50%

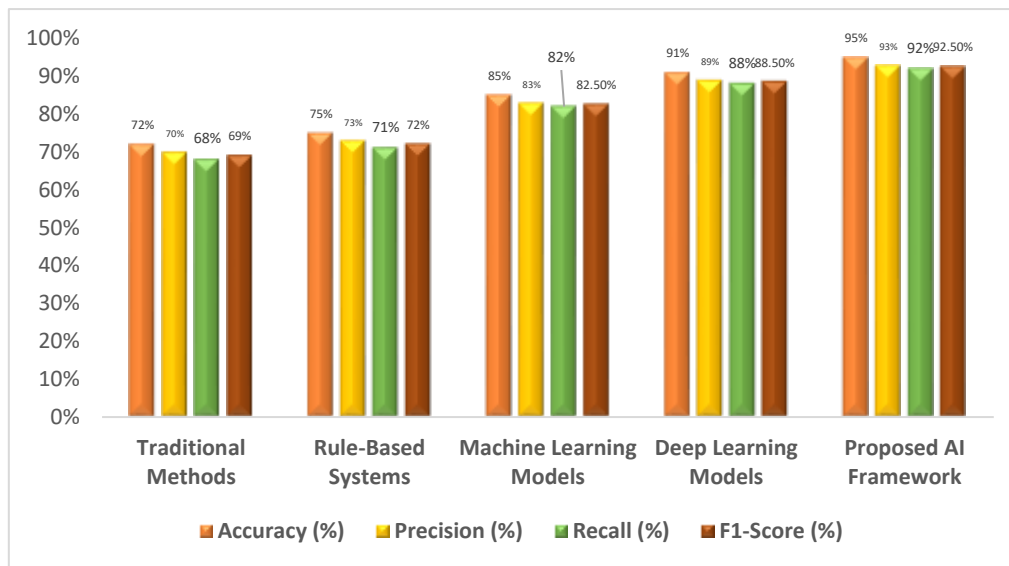


Fig 5 - Experimental Results

4.2.1. Traditional Methods

Conventional data cleaning strategies gave an accuracy of 72, precision of 70, recall of 68 and F1-score of 69. These findings imply that rule-based and constraint-driven methods have good capabilities to deal with simple data inconsistencies but they find it difficult to deal with large and complex datasets. The relatively low recall indicates that a lot of errors are not identified, mainly because of the hardness of the predefined rules. Moreover, these approaches are not very adaptable and therefore not as useful in a dynamic data setting where the patterns are changing regularly.

4.2.2. Rule-Based Systems

The rule-based systems demonstrate a marginal better performance in comparison to the traditional approaches, where the accuracy is 75, precision is 73, recalls 71, and the F1-score is 72. These systems are based on domain specific rules in order to detect and fix errors which makes them more performant than the completely traditional methods. Nevertheless, they are still less effective due to the necessity to create and maintain rules manually. It becomes more and more difficult to maintain and update these rules as the size and complexity of datasets increase, which impacts scalability and effectiveness.

4.2.3. Machine Learning Models

Machine learning models show a remarkable increase in performance, with an accuracy of 85, a precision of 83, a recall of 82 and a F1-score of 82.5. Such models are flexible to unseen errors, and can learn patterns based on data and generalize, which are not as flexible as rule-based systems. The equal accuracy and recall ratios show that machine learning models can be used to effectively detect true errors and reduce false detection. Nevertheless, this performance requires the provision of labeled data and adequate training that may be resource-consuming.

4.2.4. Deep Learning Models

Deep learning models also improve performance with an accuracy of 91, a precision of 89, a recall of 88 and an F1-score of 88.5. These models can deal with high dimensional, complex data and can automatically extract meaningful features, which results in better error detection and correction. The large recall value implies that the majority of the errors are recognized, whereas the good precision makes false positives minimal. Although they are beneficial, deep learning models demand a lot of computer power and huge data to train them.

4.2.5. Proposed AI Framework

The proposed AI-based data cleaning framework is the best in terms of accuracy of 95 percent, precision of 93 percent, recall of 92 percent and F1-score of 92.5 percent. This high performance is explained by the combination of various methods, such as machine learning, deep learning, rule-based validation, which form a hybrid and adaptive system. The framework is also able to balance precision and recall, which ensures a correct and complete error detection. It is also very appropriate in the large-scale data cleaning tasks in the real world due to its scalability and capability to process different types of data.

4.3. Discussion

The experimental findings are a clear indication that the AI-based data cleaning methods are more accurate, more precise, and more apt in recall and overall robustness as compared to the traditional and rule-based methods. Conventional approaches, in which rules and constraints are critical, are constrained in their capacity to deal with large, complex and dynamic datasets. Although the rule-based systems provide minor gains, they are also associated with scalability problems and the need to prevent manual violation to implement new rules in case of data changes. Conversely, machine learning models are flexible to new errors since they are trained to identify patterns directly in the data, which allows them to identify new errors more efficiently. Nonetheless, they are sensitive to labeled datasets and the quality of training. Deep learning models also improve the data cleaning system since it has the ability to train on complex, non-linear relationships in high dimension data. This is because, in general, their higher performance, as shown in the results, is attributed to automated feature extraction and the capability to model complex data patterns, leading to greater recall and precision. Although they possess these benefits, deep learning models can be computationally heavy, and they might consume considerable resources. The suggested AI system is successful in overcoming these drawbacks and integrating the benefits of various methods, such as traditional validation, machine learning, and deep learning models. This is a hybrid solution, which guarantees accuracy, scalability, and reduces the shortcomings of the single approaches. The fact that the framework is able to reach the highest accuracy and balanced performance metrics bring to the fore the effectiveness in addressing a variety of data quality problems. The discussion in general highlights the need to consider the incorporation of advanced AI methods to develop intelligent, scalable, and efficient data cleaning systems that would be applicable in real-world big data usage.

5. CONCLUSION

The paper is a detailed analysis of AI-based data cleaning algorithms that are intended to overcome the limitations of big datasets. This analysis has clearly shown the weakness of the conventional approaches to data cleaning like rule-based and constraint-driven data cleaning methods which are mostly not scalable, flexible and adaptable to manipulate complex and dynamic data environment. Conversely, machine learning and deep learning methods are examples of AI-based solutions that show a high increase in data anomaly detection and correction. These approaches use data-based learning to automatically discover patterns and are thus more efficient and effective in the present data-intensive applications. The offered methodology combines the advantages of supervised, unsupervised, and deep learning approaches and utilizes several machine learning models in a scalable system infrastructure. This mixed-method not only increases the precision of error detection but also increases the overall efficiency of cleaning the data. The effectiveness of the proposed AI-based framework is highly supported by the experimental results, which indicate significant changes in key performance metrics, including accuracy, precision, recall, and F1-score. The proposed framework has better performance than traditional and rule-based systems as it balances error detection and correction effectively and reduces false positives and false negatives. Addition of scalable technologies like distributed computing and parallel processing further makes sure that the system can effectively cope with huge amounts of data which makes it applicable in real world. Nevertheless, in spite of these developments, some challenges exist. Models based on AI, especially deep learning methods, may demand a large amount of computational resources and training datasets, which can make them more expensive to operate. Also, it is difficult to interpret complex models, making this approach difficult to understand and explain how the decision-making process works, which is important in sensitive areas. In the future, the research will aim to face these issues and improve the interpretability and transparency of AI models, which will allow trusting and using them more easily. An attempt will also be made to streamline algorithms to minimize the complexity of calculations and resource usage, making data cleaning based on AI more convenient and affordable. More so, privacy-saving methods, including federated learning and secure data processing will be necessary to provide security to the data and keep it within the regulatory requirements. In general, AI-based data cleaning is a potentially promising and fast-evolving area with potential to enable much improved data quality and to enable more accurate and reliable decision-making concerning large-scale data-driven systems.

REFERENCES

- [1] E. Rahm and H. H. Do, "Data cleaning: Problems and current approaches," *IEEE Data Engineering Bulletin*, vol. 23, no. 4, pp. 3–13, 2000.
- [2] J. Widom, "Data cleaning: Problems and current approaches," *IEEE Data Engineering Bulletin*, vol. 23, no. 4, pp. 3–13, 2000.
- [3] M. Stonebraker and I. F. Ilyas, "Data integration: The current status and the way forward," *IEEE Data Engineering Bulletin*, vol. 33, no. 3, pp. 3–9, 2010.
- [4] I. F. Ilyas and X. Chu, *Trends in Cleaning Relational Data: Consistency and Deduplication*, Now Publishers Inc., 2015.
- [5] P. Bohannon, M. Flaster, W. Fan, and R. Rastogi, "A cost-based model and effective heuristic for repairing constraints by value modification," in *Proc. ACM SIGMOD*, 2005, pp. 143–154.

-
- [6] X. Chu, I. F. Ilyas, and P. Papotti, "Holistic data cleaning: Putting violations into context," in *Proc. IEEE ICDE*, 2013, pp. 458–469.
 - [7] J. Van den Broeck et al., "Data cleaning: Detecting, diagnosing, and editing data abnormalities," *PLOS Medicine*, vol. 2, no. 10, pp. 966–970, 2005.
 - [8] T. Dasu and T. Johnson, *Exploratory Data Mining and Data Cleaning*, Wiley, 2003.
 - [9] C. C. Aggarwal, "Outlier analysis," Springer, 2017.
 - [10] L. Breiman, "Random forests," *Machine Learning*, vol. 45, no. 1, pp. 5–32, 2001.
 - [11] C. Cortes and V. Vapnik, "Support-vector networks," *Machine Learning*, vol. 20, no. 3, pp. 273–297, 1995.
 - [12] M. Ester, H. P. Kriegel, J. Sander, and X. Xu, "A density-based algorithm for discovering clusters in large spatial databases with noise," in *Proc. KDD*, 1996, pp. 226–231.
 - [13] F. T. Liu, K. M. Ting, and Z. H. Zhou, "Isolation forest," in *Proc. IEEE ICDM*, 2008, pp. 413–422.
 - [14] D. P. Kingma and M. Welling, "Auto-encoding variational Bayes," in *Proc. ICLR*, 2014.
 - [15] I. Goodfellow, Y. Bengio, and A. Courville, *Deep Learning*, MIT Press, 2016.