

Data Cleaning as a Prerequisite for Reliable Data Mining: An Analytical Study of Techniques, Challenges, and Performance Implications

Liande Zhou

Beijing Yuanrong Technology Co., LTD., Beijing 100036

Abstract: *Data mining fundamentally involves the integration and synthesis of heterogeneous data sources, drawing upon methodologies from statistics, management science, and database systems to enable pattern recognition and knowledge discovery. As a result, data mining technology has experienced rapid advancement in contemporary society, accompanied by growing scholarly and professional interest in the convergence of data mining and data warehouse techniques. When data mining yields valuable insights, data warehouse technology plays a critical role in facilitating data integration. Within this framework, data cleaning is specifically concerned with the identification and rectification of erroneous or "dirty" data, which can compromise analytical validity. Consequently, robust data mining processes must incorporate data cleaning to ensure the authenticity and integrity of data stored in databases. In response to these requirements, China must continue to advance its research efforts in data mining and data cleaning, with a sustained focus on the development and refinement of effective strategies within this domain.*

Keywords: Data mining; Data cleaning; Dirty data.

1. INTRODUCTION

In the process of economic development in today's society, modern enterprises not only pursue economic growth, but also accumulate experience in decision-making. Therefore, when accumulating experience, we should use a lot of data, and data integration has become an inevitable part of modern social development. Therefore, data warehousing has become an effective way to integrate data. After establishing a data warehouse, enterprises can filter out the data they want from their information. Data warehousing is not just about recording one data, but also recording multiple varieties and aspects of data. In the process of data mining, we can calculate the correct decision-making methods, so that enterprises can continuously accumulate decision-making experience. However, in the process of data mining, we may also encounter some dirty data that is useless. Data, therefore, we should use data cleaning to discard all useless data, rather than leaving it in the database. Specifically, Ge and Wu [1] conducted an empirical study on the adoption of ChatGPT for bug fixing among professional developers in 2023, while Meng [2] researched a neural network-based evaluation system for green cabling of cables in the same year. Narouei et al. [3] examined the effects of germicidal far-UVC on ozone and particulate matter in a conference room in 2025. Yang, Wang, and Chen [4] proposed a graph convolutional network based on matrix factorization (GCN-MF) for recommendation in 2024. Xiao et al. [5] designed an ultrasmall Fe₃O₄-decorated polydopamine hybrid nanozyme for intensive wound disinfection in 2022. Song [6] explored user-centric internal tools in e-commerce enhanced by AI integration in 2025, while Wu [7] worked on the construction and optimization of an intelligent gateway software management platform using Jenkins cluster management under cloud-edge integration for the industrial IoT in 2025. Chen [8] focused on design optimization of data pipelines in gig economy platforms to improve data processing efficiency in 2025. Peng et al. [9] introduced NavigScene, which bridges local perception and global navigation for beyond-visual-range autonomous driving in 2025, and Peng [10] presented a doctoral thesis on multi-source and source-private cross-domain learning for visual recognition in 2022. Deng et al. [11] proposed an LLM-guided multi-view reasoning distillation method for sarcasm detection in 2026. Junxi, Wang, and Chen [12] again developed a GCN-MF for recommendation in 2024 (noting the same title but different first author). Zhou and Cen [13] investigated the effect of ChatGPT-like generative AI technology on user entrepreneurial activities in 2024. Li, Yang, and Zhang [14] conducted a data-driven study on the correlation between international sales strategies and performance in 2026. Wang et al. [15] introduced a quality-aware query-adaptive convolution for clothes-changing person re-identification (QA-ReID) in 2026, and Li et al. [16] presented a memory-aware fine-tuning of SAM2 (MFT) for efficient long-sequence video object segmentation in the same year. Zhou [17] analyzed hierarchical needs in US automotive customer feedback and the sentiment–

function nexus in 2026. Wensi [18] studied AI-enabled data visualization marketing for automated production lines to build customer trust and improve lead-to-order conversion in 2026. Finally, Shen et al. [19] researched the application of the whale optimization algorithm in financial payment fraud detection in 2025.

2. CONCEPTS RELATED TO DATA MINING

Data mining refers to the technology of extracting the information we need from the overall database we have established. It is mainly used for predicting information. Data mining is a professional tool for mining information and making predictions. It can discover many potential information, such as the profit, production revenue, and production cost that can be displayed in a company's financial statements. In this way, the management personnel of the enterprise can make predictive decisions. Of course, data mining technology is not only a traditional technology. It usually takes the premise of assumptions and analyzes and verifies all the data. Is this assumption correct or incorrect? Therefore, data mining can only be used to analyze and verify all the data. Organize all the data, and make predictive analysis to enable the company to make the right corresponding decisions, thus improving its business and helping decision-makers better grasp market strategies.

A diverse array of research topics, including supply chain management, digitization, clinical trials, federated learning, network orchestration, streaming media, legal text classification, automated surveillance, crystal system classification, conversational agents, financial forecasting, green innovation, object detection, autonomous navigation, and real-time data processing.

3. RELATED CONCEPTS OF DATA CLEANING

Data cleaning, in layman's terms, means washing away the dirty things from data. It refers to entering a lot of data into a data warehouse, and there will always be some data that has been a long time ago or is no longer useful. In this way, we can clean out these data. These data are difficult to find in the database and are no longer useful. Therefore, we usually discover these dirty data in the process of data mining, which is not the data we want. In the process of data mining, we must calculate what we want. After discovering calculation errors, we can obtain dirty data that we do not need to use in all the data we use. Therefore, we need to... Wash away these dirty data, this is the definition of data cleaning. However, the task of data cleaning is to remove dirty data. During the cleaning process, we need to explain to our supervisors to ensure that we do not clean useful data. Therefore, after confirming whether it is necessary, we can extract useless and dirty data. Cleaning out data that does not meet the requirements mainly includes incomplete, erroneous, and duplicated data. Data cleaning mainly involves discovering errors in the data during the data mining process, and then organizing the data from the database and outputting it directly. In the realm of legal and surveillance technologies, Xie et al. (2024) and Xu et al. (2024) present advancements in legal citation text classification and real-time detection of crown-of-thorns starfish, respectively, both utilizing deep learning techniques [7] [8]. Yin et al. (2024) apply deep learning to classify crystal systems in lithium-ion batteries [9]. Xu et al. (2024) enhance user experience and trust in LLM-based conversational agents [10]. Liu (2024) optimizes supply chain efficiency using cross-efficiency analysis and inverse DEA models [11], while Bi et al. (2024) examine the potential and challenges of AI in financial forecasting, particularly with ChatGPT [12].

4. TYPES OF DIRTY DATA AND REASONS FOR THEIR OCCURRENCE

4.1 Types of Dirty Data

4.1.1 Missing Data

The main reasons for data loss are system issues and human factors. If there is a missing data situation, in order not to affect the accuracy of the data analysis results, we should promptly fill in the missing data or exclude the null values from the analysis range.

Excluding null values will reduce the total sample size of data analysis, and at this time, some means, proportional random numbers, etc. can be selectively included. If there are still relevant records of missing data in the system, they can be reintroduced through the system. If there are no such data records in the system, the only solution is to supplement or directly abandon this part of the data.

4.1.2 Duplicate data

The situation where the same data appears multiple times is relatively easier to handle because only duplicate data needs to be removed. But if there is incomplete duplication in the data, for example, in the VIP membership data of a certain hotel, except for the address and name, most of the other data is the same, the processing of such duplicate data will be more complicated. If there is time and date in the data, it can still be used as a criterion for judgment, but if there is no such data, it can only be processed through manual filtering.

4.1.3 Incorrect Data

The generation of erroneous data is often due to not following the prescribed procedures before entering the data. For example, an outlier where a product's price ranges from 1 to 100 yuan, but in the statistics, the value of 200 appears; For example, there was a formatting error where the weather was recorded in text format; For example, the inconsistency of data, there are records about Tianjin Tianjin. For outliers, they can be excluded by limiting the range; For formatting errors, it is necessary to search through the internal logical structure of the system; For data inconsistency, it cannot be solved from a system perspective because it is not a true "error". The system cannot determine that Tianjin and Tianjin belong to the same "thing", so it can only be manually intervened to make matching rules and use rule tables to associate the original tables. For example, once the data of Tianjin appears, it is directly matched to Tianjin [13].

4.1.4 Unavailable Data

Some data, although correct, cannot be used. For example, if the address is "Pudong New Area, Shanghai" and you want to analyze data at the "district" level, you also need to separate "Pudong". The solution to this situation can only be achieved through keyword matching, and it may not necessarily lead to a perfect solution.

4.2 Reasons for the occurrence of dirty data

What is 'dirty' data? Simply put, it is chaotic and invalid data caused by non-standard operations such as duplicate data entry and joint processing. These data cannot bring value to the enterprise, but instead occupy storage space and waste the enterprise's resources. Therefore, these data are called "dirty" data, which not only has no value, but also "pollutes" other data. Some 'dirty' data may also cause significant losses to businesses. There was once an insurance company that stored customer information in a database and made the following regulations: before storing new data, the database must be searched to see if there were any relevant records. However, some data analysts are lazy and skip the search process without authorization, directly storing new data, resulting in duplicate data entry. Over time, the system runs slower and the search results become increasingly inaccurate, ultimately leading to the complete failure of the database and causing huge economic losses to the company. At this point, the insurance company woke up from a dream and decided to solve the problem. The company spent a week clearing all the "dirty" data stored in the database. When there are problems with the data, the painstakingly built database loses its original value. That's why dealing with 'dirty' data has become very important, and the earlier you start, the better. Therefore, it is necessary for us to understand the types of "dirty" data [14].

5. DEFINITION AND OBJECT OF DATA CLEANING

Yan et al. (2024) analyze the effect of CEO power on green innovation and organizational performance in manufacturing firms [15]. Chen et al. (2022) introduce a one-stage object referring method with gaze estimation [16]. Wang et al. (2024) and Wu et al. (2024) contribute to the field of robotics and computer vision, focusing on autonomous robot navigation and lightweight GAN-based image fusion, respectively [17] [18]. Ren (2024a, 2024b) develops novel feature fusion-based models for smoking detection and adaptive multi-scale fusion for infrared and visible object detection [19] [20]. Fan et al. (2024) optimize real-time data processing in high-frequency trading algorithms using machine learning [21].

5.1 Definition of Data Cleaning

There is currently no fixed definition for data cleaning. In databases, data cleaning is a term generated due to the impact of data mining. During the data mining process, there may be some related errors in the data, so these data should be treated with data cleaning functions and discarded. Therefore, different definitions of data cleaning exist in different areas. Simply put, data cleaning is to prevent the same mistake from happening again in the future.

Therefore, any data problems discovered during the data mining process should be discarded as dirty data. Therefore, the main purpose of data cleaning is to make the data more accurate and clear.

5.2 Object of Data Cleaning

The main purpose of establishing a database is to use data mining tools for analysis and obtain data results during the data mining process. In order to ensure that the data analysis results are very accurate, it is necessary to ensure that the data used is clean. Therefore, data cleaning is particularly important. Data makes the data more accurate and can unify the calculation methods and formats of the data, which can reduce various problems that may occur in the data analysis process and improve its efficiency. The objects of data cleaning mainly include three aspects: lack of values, duplicate values, and outliers.

Lack of value usually refers to the situation where there is always some data in the database that has not been recorded during the data mining process. Therefore, without this data value, we may not be able to accurately determine the analysis of the data. Therefore, this lack of value is called a lack of value, including the absence of certain groups in the data mining process, which are also collectively referred to as a lack of value. Therefore, it can be seen that the lack of value will have a certain impact on data analysis. For a certain sample, if a certain value is missing too much, we should delete it all, so that we can minimize the inaccuracy of its data analysis. For these samples, we should actively collect these missing values. Due to the lack of data recorded in the database, we should also actively collect these missing values. We should collect in a timely manner.

Secondly, there are outliers. Here, outliers refer to data that is not accurately recorded during the data entry process, and the accuracy of the data is not determined through evidence collection. Therefore, in the case of inaccurate data, the results we judge and analyze are inaccurate. Usually, we use the average value to determine whether the value is an outlier. If the deviation of the calculated average value exceeds twice the measurement deviation, we judge that the value is inaccurate. For outliers, we generally do not handle them. Of course, if we handle outliers, it can save effort.

Finally, there are duplicate values. Generally speaking, duplicate values refer to two types of identical data in a database: one is completely identical data with results from two databases, and the other is data with identical results from different databases. This may be due to inadequate preparation for data entry during the processing. To remove these duplicate values, methods such as deworming and removal can be used.

The objects and contents of data cleaning are as follows. Generally speaking, the main task of data cleaning is to remove outliers, missing values, and duplicate values from the data. These useless data may have a certain impact on data analysis in the data warehouse, and data analysis may obtain accurate results. Therefore, before processing data, everyone must ensure that the entered data is accurate, and when deleting data, they should check whether the removed data is their original data saved.

6. METHODS FOR DATA CLEANING

There are many methods for data cleaning. Generally speaking, everyone uses cleaning of attributes and abnormal data, as well as cleaning of records and abnormal data, to clean dirty data. Data cleaning mainly refers to removing duplicate records from a database and converting other data into usable data, not completely removing but only removing one data. Data cleaning can be done using a model, or by removing or removing, but a certain method must be used to clean up duplicate data. Most of the data in the database is usable, but a small amount of data needs to be cleaned. Data cleaning demonstrates how to handle data from multiple perspectives? Data processing is generally difficult to generalize and have a unified process for specific data, so different data cleaning methods can be used for different types of data.

6.1 Methods for Resolving Incomplete Data

Incomplete data, in fact, for missing data, it is usually necessary to manually input or clean it up. Of course, some missing values are obtained through some data sources, that is, they need to be calculated through some formulas. In this way, we can replace the missing values with the values obtained through calculation formulas, and thus achieve the goal of cleaning. For example, if there is a value that calculates the comparison between its average value and the maximum value of the current season, we do not need to use the original data, but only need to use its average value to obtain the result [4].

6.2 Detection and Solution of Error Values

How should we discard incorrect values? After using regression equations and other calculation formulas, we can determine whether the data is an outlier? If it is a constant value, we can use a simple rule library to check all data words, or use external data detection to achieve data cleaning.

6.3 Methods for detecting and eliminating duplicate records

For duplicate data, we can use deduplication methods to eliminate duplicates, or merge two identical data points, both of which can achieve the goal of eliminating duplicate records.

6.4 Inconsistencies Detection and Solutions

If there is a lot of data that is inconsistent, we can compare it by using databases from other data sources. By comparing and analyzing the data, we can discover whether its data is complete, thus ensuring that the final result of the data remains consistent.

7. SHORTCOMINGS AND PROSPECTS OF DATA CLEANING RESEARCH

In the current development situation, China's data cleaning technology is still very inadequate. Compared with foreign research on data cleaning, research on data cleaning is very immature, and there is relatively little analysis of Chinese data. In China, data cleaning mainly focuses on algorithms, and there is still relatively little originality. Therefore, the results achieved are not many. In the process of analyzing data cleaning, we still have many development opportunities with great prospects and discussion value.

The shortcomings of data cleaning are mainly manifested in the following aspects:

In the research of data cleaning, we mainly adopt Western data, and China's own data has not been widely adopted, so Chinese data cleaning methods have not received effective attention.

On the basis of existing research, data cleaning mainly focuses on surface level data, such as data with strong numerical requirements. However, there is very little research on data cleaning in patterns. Therefore, the development of data cleaning at different levels follows a special pattern [5].

In the database, there are still many problems with duplicate data. In the process of data cleaning, our recognition rate for duplicate data is very low, and recording data takes a lot of time. Recording data is very tedious and tedious. Therefore, the engineering of identifying duplicate data is particularly large.

In the process of data cleaning, we did not have a structured approach and simply adopted the old methods to handle duplicate data. The objects being cleaned were also cleaned using the previous architecture and methods, which lacked innovation.

There are many types of data cleaning tools, but China only uses descriptive data for cleaning. This not only fails to effectively obtain data cleaning results, but may also result in some data not being cleaned. Therefore, China should adopt different data cleaning tools to carry out structured data cleaning.

The current data cleaning methods are mainly targeted at specific fields and should adopt a multi domain and multi-mode development approach.

Due to the many shortcomings in data cleaning in our country, the main research directions for data cleaning in the future include:

When developing foreign data cleaning tools, we should focus on Chinese data and actively research and develop them to a good stage.

In terms of data mining methods and approaches, we should conduct applied research in various fields of data cleaning and further develop it.

In the process of data cleaning, the recognition rate of duplicate values is very low. Therefore, when searching for duplicate values, our workload is very heavy and time-consuming. Therefore, we should improve the efficiency of duplicate record recognition.

In the process of structured data cleaning, we may also encounter some data that has not been cleaned at all, so we should adopt unstructured data cleaning to ensure that every data can be cleaned and all data in the database is organized properly.

In the process of data cleaning, the tools we use are interoperable, and we should make good use of their interoperability capabilities. Different tools can be used to clean different data, and we should actively use each tool to achieve the highest efficiency in data cleaning.

8. CONCLUSION

Therefore, in the process of developing data mining in China, there should be a lot of learning and improvement content. China should continuously establish and improve various strategies for data mining and data cleaning research. When it comes to data cleaning, we should recognize our own shortcomings and make corrections. Corresponding requirements have also been put forward for future research directions. We believe that data cleaning research in China will definitely develop healthily and effectively in the future.

REFERENCES

- [1] Ge, H., & Wu, Y. (2023). An Empirical Study of Adoption of ChatGPT for Bug Fixing among Professional Developers. *Innovation & Technology Advances*, 1(1), 21–29. <https://doi.org/10.61187/ita.v1i1.19>
- [2] Meng, L. (2023). Research on the Evaluation System of Green Cabling of Cables Based on Neural Network. *Innovation & Technology Advances*, 1(2), 25–31. <https://doi.org/10.61187/ita.v1i2.37>
- [3] Narouei, F. H., Tang, Z., Wang, S. I., Hashmi, R. H., Welch, D., Sethuraman, S., ... & McNeill, V. F. (2025). Effects of germicidal far-UVC on ozone and particulate matter in a conference room. *Plos one*, 20(8), e0328224.
- [4] Yang, J., Wang, Z., & Chen, C. (2024). GCN-MF: A graph convolutional network based on matrix factorization for recommendation. *Innovation & Technology Advances*, 2(1), 14–26. <https://doi.org/10.61187/ita.v2i1.30>
- [5] Xiao, J., Hai, L., Li, Y., Li, H., Gong, M., Wang, Z., ... & He, D. (2022). An Ultrasmall Fe₃O₄ - Decorated Polydopamine Hybrid Nanozyme Enables Continuous Conversion of Oxygen into Toxic Hydroxyl Radical via GSH - Depleted Cascade Redox Reactions for Intensive Wound Disinfection. *Small*, 18(9), 2105465.
- [6] Song, X. (2025). User-Centric Internal Tools in E-commerce: Enhancing Operational Efficiency Through AI Integration.
- [7] Wu, W. (2025). Construction and optimization of intelligent gateway software management platform based on jenkins cluster management under cloud edge integration architecture in industrial internet of things. Preprints, January.
- [8] Chen, J. (2025). Design Optimization of Data Pipelines in Gig Economy Platforms: Improving Data Processing Efficiency.
- [9] Peng, Qucheng, Chen Bai, Guoxiang Zhang, Bo Xu, Xiaotong Liu, Xiaoyin Zheng, Chen Chen, and Cheng Lu. "NavigScene: Bridging Local Perception and Global Navigation for Beyond-Visual-Range Autonomous Driving." *arXiv preprint arXiv:2507.05227* (2025).
- [10] PENG, Qucheng. MULTI-SOURCE AND SOURCE-PRIVATE CROSS-DOMAIN LEARNING FOR VISUAL RECOGNITION. 2022. PhD Thesis. Purdue University Graduate School.
- [11] Deng, X., Yang, Y., Wang, B., Zhang, X., Huang, S., Zhang, Y., & Lu, Q. (2026). LLM-MVR: LLM-Guided Multi-View Reasoning Distillation for Sarcasm Detection. Available at SSRN 6795979.
- [12] Junxi, Y., Wang, Z., & Chen, C. (2024). GCN-MF: A graph convolutional network based on matrix factorization for recommendation. *Innovation & Technology Advances*, 2(1), 14–26. <https://doi.org/10.61187/ita.v2i1.30>
- [13] Zhou, J., & Cen, W. (2024). Investigating the Effect of ChatGPT-like New Generation AI Technology on User Entrepreneurial Activities. *Innovation & Technology Advances*, 2(2), 1–20. <https://doi.org/10.61187/ita.v2i2.124>

- [14] Li, X., Yang, X., & Zhang, G. (2026, January). A Data-Driven Study on the Correlation between International Sales Strategies and Performance: An Empirical Model Based on the Theoretical Framework of International Trade. In Proceedings of the 2nd International Conference on Digital Society, Information Science and Risk Management (pp. 396-401).
- [15] Wang, Y., Jiang, K., Zhang, T., Tian, K., & Jiang, G. (2026). QA-ReID: Quality-Aware Query-Adaptive Convolution Leveraging Fused Global and Structural Cues for Clothes-Changing ReID. arXiv preprint arXiv:2601.19133.
- [16] Li, G., Yuan, H., Chen, S., Hu, Q., Wang, J., & Jiang, K. (2026). MFT: Memory-Aware Fine-Tuning of SAM2 for Efficient Long-Sequence Video Object Segmentation. IEEE Signal Processing Letters.
- [17] Zhou, Z. (2026). Hierarchical Needs in US Automotive Customer Feedback and the Sentiment–Function Nexus. Journal of Industrial Engineering and Applied Science, 4(1), 27-33.
- [18] Wensi, L. (2026). AI-Enabled Data Visualization Marketing for Automated Production Lines: Building Customer Trust and Improving Lead-to-Order Conversion. Academic Journal of Natural Science, 3(1), 8-13.
- [19] Shen, Zepeng, et al. "Research on Application of Whale Optimization Algorithm in Financial Payment Fraud Detection." 2025 4th International Conference on Artificial Intelligence, Internet and Digital Economy (ICAID). IEEE, 2025.