

Original Paper

Research on Safety Management of University Laboratories Based on Isolation Forest Algorithm

Junjie Xue¹

¹ Experimental and Practical Training Center, Shanxi University of Finance and Economics, Taiyuan, Shanxi, China

Received: August 24, 2024 Accepted: October 11, 2024 Online Published: October 22, 2024

doi:10.22158/wjer.v11n5p185 URL: <http://dx.doi.org/10.22158/wjer.v11n5p185>

Abstract

This paper mainly discusses the application of the Isolation Forest Algorithm in the safety management of university laboratories. Through the analysis of the current situation of safety management in university laboratories, the deficiencies of traditional management methods are expounded. Then, the principle and advantages of the Isolation Forest Algorithm are introduced in detail, and the effectiveness of this algorithm in the abnormal detection of laboratory safety is verified through experiments in a simulated laboratory environment. The experimental results show that the Isolation Forest Algorithm can quickly and accurately detect abnormal situations in the laboratory, providing a new and effective means for the safety management of university laboratories.

Keywords

Isolation Forest Algorithm, University Laboratories, Safety Anomaly Detection, Safety Management

1. Introduction

University laboratories, as crucial positions for teaching, scientific research, and talent cultivation, are of self-evident importance in terms of safety management. They are the cradles of knowledge exploration and innovation, carrying the core mission of experimental teaching tasks in universities. However, the laboratory environment is complex and harbors many potential safety hazards. Fires may be triggered by improper storage of chemical reagents or malfunctions of electrical equipment; explosion risks may stem from improper handling of flammable and explosive substances; chemical leaks may occur due to experimental operation errors or equipment aging and damage; electrical accidents may happen because of aging circuits or overloaded use of electrical appliances. Once a safety accident occurs, the consequences are unimaginable. Casualties will bring great pain to families, property losses may cause severe damage to the scientific research and teaching resources of

universities, and the normal order of teaching and scientific research work will be completely disrupted, delaying teaching progress and affecting the quality and efficiency of talent cultivation (Zheng, Zhu, Xu et al., 2024).

The traditional laboratory safety management model mainly relies on manual inspections and empirical judgments. Manual inspections require a large amount of manpower and time, and it is difficult for inspectors to maintain a high level of concentration for a long time, making it easy to overlook certain aspects. Empirical judgments often have subjectivity and limitations, and different people may have different judgment criteria for the same safety hazard. This method cannot meet the management requirements of large-scale laboratories in terms of efficiency, and its accuracy is difficult to guarantee. It may not be able to detect some potential and unnoticeable safety hazards in a timely manner. At the same time, it is also difficult to achieve real-time monitoring of the laboratory safety status and cannot capture abnormal situations and respond in a timely manner (Wang, Zhang, You et al., 2024).

With the rapid development of information technology, laboratory safety management has encountered new opportunities. Sensor technology can perceive various physical and chemical parameters in the laboratory environment in real time, such as temperature, humidity, gas concentration, etc.; Internet of Things technology can connect various devices and sensors to achieve real-time data transmission and sharing; data analysis technology can process and analyze a large amount of monitoring data and mine potential laws and abnormal information. Among these technologies, machine learning algorithms, as a powerful data analysis tool, provide new ideas and method for laboratory safety management.

The Isolation Forest Algorithm, as an abnormal detection algorithm based on random forests, has unique advantages in the field of laboratory safety management. It has high computational efficiency and can quickly process a large amount of monitoring data, meeting the needs of laboratory real-time monitoring. It has good accuracy and can accurately identify abnormal patterns in the data, timely detecting potential safety hazards. The characteristic of not requiring prior knowledge enables it to be applicable to various complex laboratory environments without the need to have a deep understanding of the data distribution and abnormal situations in advance. Applying the Isolation Forest Algorithm to the safety management of university laboratories is expected to break through the limitations of traditional management methods and improve the efficiency and accuracy of management. Through the real-time analysis of laboratory environment data and equipment operation data, abnormal situations can be quickly located, providing a solid guarantee for the safe operation of university laboratories. This not only helps to prevent the occurrence of safety accidents, protect the lives of teachers and students and the property of schools, but also ensures the smooth progress of teaching and scientific research work and promotes the continuous development of academic undertakings in universities.

2. Literature Review

2.1 The Current Situation of Safety Management in University Laboratories

(A) Imperfect Safety Management System

At present, although many universities have formulated some laboratory safety management systems, these systems are often not perfect and have loopholes and deficiencies. For example, some systems lack specific operation procedures and standards, resulting in difficulties in implementation during the actual execution process; some systems lack effective supervision and assessment mechanisms, resulting in poor implementation effects (Yang, Li, & Hong, 2022).

(B) Weak Safety Awareness

Some teachers and students have insufficient understanding of the importance of laboratory safety and weak safety awareness. During the experiment process, there are problems such as illegal operations and negligence. For example, not following the laboratory safety operation procedures, randomly placing experimental instruments and drugs; not wearing personal protective equipment such as gloves and goggles; eating and smoking in the laboratory (Zhang & Lai, 2023).

(C) Insufficient Safety Facilities

The safety facilities in some university laboratories are insufficient and cannot meet the needs of laboratory safety management. For example, lacking necessary fire-fighting facilities, ventilation facilities, protection facilities, etc.; some facilities are aging and damaged and cannot be used normally.

(D) Low Level of Safety Management Informatization

At present, the safety management of many university laboratories still remains in the traditional manual management mode, with a low level of informatization. There is a lack of an effective laboratory safety management information system and it is impossible to achieve real-time monitoring and early warning of laboratory safety.

2.2 Principle and Advantages of the Isolation Forest Algorithm

(A) Principle of the Isolation Forest Algorithm

The Isolation Forest Algorithm is an abnormal detection algorithm based on Random Forests. Its design concept is unique and efficient. The core idea of this algorithm is to isolate each sample in the data set by randomly selecting features and partition points to form Isolation Tree.

For a given data set, the specific processing process is as follows: First, randomly select a feature and randomly select a partition point between the minimum and maximum value of this feature. According to this partition point, the data set is divided into two subsets. Then, repeat the above steps in each subset, that is, randomly select features and partition points again for data division until specific stopping conditions are met. These stopping conditions can be that each sample is isolated or the preset tree height is reached (Wang, He, & Yuan, 2024).

To understand the principle of the Isolation Forest Algorithm more in-depth, we introduce some mathematical concepts and formulas for explanation.

In an Isolation Forest, the construction of each Isolation Tree can be regarded as a recursive process. Suppose we have a data set D with a size of n and we want to build an Isolation Tree T .

From a mathematical perspective, each partition can be represented as selecting a feature j and

determining a partition point p such that the samples in the data set with a value of feature j less than p are divided into the left subtree, and those with a value greater than or equal to p are divided into the right subtree. That is:

Left subtree data set: $D_{left} = \{x \in D \mid x_j < p\}$

Right subtree data set: $D_{right} = \{x \in D \mid x_j \geq p\}$

This partitioning process is repeated in the subtrees until the stopping conditions are met.

The stopping conditions for the construction of an Isolation Tree usually have the following two common cases:

- a. The tree reaches the specified height or depth h .
- b. The data cannot be further divided, that is, there is only one sample left in the subset or all the selected feature values of the samples are the same.

Theoretically, the key to the Isolation Forest Algorithm lies in quickly isolating abnormal points through random partitioning. Intuitively, clusters with high density require multiple partitions to isolate each point, while abnormal points, due to their sparse distribution, can often be isolated more quickly, that is, form a shorter path in the tree.

To determine whether a sample is an abnormal value, it is necessary to calculate the path length of the sample in the Isolation Tree. The calculation of the path length is the number of edges passed from the root node to the leaf node where the sample is located (Zhao & Ma, 2024).

(B) Advantages of the Isolation Forest Algorithm

- a. High computational efficiency: The Isolation Forest Algorithm constructs isolated trees by randomly partitioning and does not require complex matrix operations and optimization solutions, so the computational efficiency is very high. Its time complexity is close to linear and it can quickly process large-scale data sets, which gives it a significant advantage in scenarios such as real-time monitoring and early warning that require high time efficiency.

From a mathematical analysis perspective, each partition only requires simple comparison and data allocation operations, with a low complexity. Moreover, since there is no need to calculate complex distances, densities, etc., the amount of calculation is greatly reduced. Specifically, for a data set containing n samples, the time complexity of building an isolated tree is approximately $O(\log n)$. In an Isolation Forest, usually multiple isolated trees (assumed to be t trees) are built, but since the construction of each tree is independent and can be performed in parallel, the overall time complexity remains at a relatively low level, approximately $O(t \log n)$, which enables the algorithm to efficiently process large-scale data.

- b. Good accuracy: The Isolation Forest Algorithm can effectively avoid overfitting problems by randomly selecting features and partition points. The process of randomly selecting features and partition points each time makes the model not overly dependent on certain specific features or

partitioning methods, thus enhancing the model's generalization ability.

At the same time, this algorithm can handle high-dimensional data and can accurately detect abnormal values for complex data sets. In high-dimensional data space, the definition and detection of abnormal points are often challenging, but the Isolation Forest can search in different dimensions and positions through random partitioning, making it more likely to capture those data points that exhibit abnormal behavior in local regions.

It can be proved mathematically that through multiple random divisions and integrating the results of multiple isolated trees, the variance can be effectively reduced and the detection accuracy can be improved. Specifically, for a given abnormal point, it has a high probability of forming a shorter path length in multiple isolated trees, and thus can be accurately identified as abnormal during comprehensive evaluation.

c. No need for prior knowledge: The Isolation Forest Algorithm does not need to know the distribution of data and the proportion of abnormal values in advance. Only a given data set is required, and it can automatically detect abnormal values. This makes the algorithm very versatile and adaptable.

Mathematically, the randomness of the algorithm and the judgment method based on path length make it not dependent on specific data distribution assumptions. Regardless of the distribution of the data, abnormal points usually exhibit a relative short path length and can thus be detected by the algorithm. This characteristic of not requiring prior knowledge is very important in practical applications because in many cases we may not know the specific distribution of the data, or the proportion of abnormal values may be unknown or dynamically changing.

d. Easy to explain: The results of the Isolation Forest Algorithm are easy to explain. It can be determined whether a sample is an abnormal value by observing the path length of the sample in the isolated tree. The shorter the path length, the more likely the sample is to be an abnormal value.

At the same time, the algorithm can also output the score of abnormal values, facilitating users to evaluate and handle abnormal situations. From a mathematical perspective, the calculation of the abnormal score is based on the path length of the sample in multiple isolated trees. Specifically, by calculating the path length $h(x)$ of the sample in each tree and combining it with the average value $c(\psi)$ of the path length when a given sample number ψ is provided for standardization processing, an abnormal score is obtained. For example, a commonly used abnormal score calculation formula is:

$s(x) = 2^{-\frac{E(h(x))}{c(\psi)}}$, where $E(h(x))$ is the average value of $h(x)$. Such a score has a clear physical

meaning, and users can intuitively judge the degree of abnormality of the sample according to the size of the score and can set a threshold according to specific requirements to determine the judgment standard of abnormality. This interpretability makes the results of the algorithm easier to understand and accept and also facilitates users to conduct further analysis and decision-making according to the actual situation.

The Isolation Forest Algorithm, with its unique principle and significant advantages, shows strong

capabilities and wide application prospects in the field of abnormal detection. Its high computational performance, good accuracy, characteristic of not requiring prior knowledge, and easy interpretability make it a powerful tool for processing complex data and real-time detection of abnormalities, applicable to many fields such as network security, financial transactions, industrial production, etc. However, like any algorithm, the Isolation Forest Algorithm is not perfect. It may be sensitive to data noise in some cases and may have certain limitations in dealing with local abnormalities. In practical applications, it is necessary to combine specific problems and data characteristics, fully consider the advantages and disadvantages of the algorithm to achieve the best abnormal detection effect. At the same time, continuous research and improvement are also further expanding and improving the application range and performance of the Isolation Forest Algorithm (Dong, 2023).

3. Design Strategies and Implementation for Experiments

3.1 Construction of Experimental Environment

In order to accurately simulate the university laboratory environment and effectively verify the application effect of the Isolation Forest Algorithm in laboratory safety management, we have carefully constructed the experimental environment.

We selected a room with a specific configuration as a simulated laboratory. There are 40 computers placed in this room. The configuration and layout of these computers are set according to the common usage situation of computers in university laboratories. The computer models are selected as mainstream models widely used in universities to ensure the authenticity and representativeness of the experimental environment.

In terms of sensor installation, we have carried out a comprehensive and detailed layout. Temperature sensors are evenly distributed in every corner of the room and around computer equipment to accurately obtain environmental temperature data. Humidity sensors are installed in positions where air circulation is relatively stable to avoid being interfered by local water vapor and ensure the reliability of humidity data. Power monitoring devices are connected to the power lines of each computer and can monitor the power consumption of computers in real time. At the same time, in order to ensure the accuracy and stability of data acquisition, we have selected high-precision sensor devices whose measurement accuracy can meet the experimental requirements.

In addition, we have established a data acquisition and transmission system. This system can collect data obtained by various sensors and computer monitoring devices in real time and transmit them to the central data processing unit. The data acquisition frequency is set to once every 5 minutes. This can not only ensure that enough data is obtained for analysis but also avoid data redundancy caused by too high an acquisition frequency. Through such an experimental environment construction, we can obtain real-time environmental data and equipment operation state data in the simulated laboratory, providing a solid data basis for subsequent experiments.

3.2 Source of Experimental Data

The experimental data mainly comes from various sensors in the simulated laboratory and the strict monitoring records of computer equipment operation state.

For environmental parameter data, temperature sensors can feed back the temperature values at different positions in the room in real time. The acquisition range of these temperature values covers the heat generated by the operation of computer equipment and the influence of the natural environment temperature. Humidity sensors accurately measure the water vapor content in the air, providing data support for analyzing the influence of environmental humidity on laboratory equipment and experimental processes. Power monitoring devices record in detail the power consumption of each computer, including the power used during startup, standby, and different running tasks, which is crucial for understanding the energy utilization efficiency of computer equipment and possible power abnormal situations.

For equipment parameter data, by installing professional monitoring software in the computer operating system, we can obtain key indicators such as the running time, CPU usage, and memory usage of the computer. The running time data can reflect the frequency and duration of the use of computer equipment, helping to analyze the wear and tear of the equipment. The CPU usage and memory usage data can intuitively present the resource utilization situation of the computer under different tasks, which is of great significance for judging whether the equipment has an abnormal running state.

By integrating and analyzing the data from these different data sources, we can comprehensively understand the running state of the simulated laboratory, providing rich data resources for the application and verification of the Isolation Forest Algorithm.

3.3 Parameter Settings of the Isolation Forest Algorithm

The parameter settings of the Isolation Forest Algorithm are crucial for the performance of the algorithm and the accuracy of the experimental results. After considering various factors such as the scale, dimension, distribution characteristics of the experimental data and the performance requirements of the algorithm, we finally determined the following parameter settings, aiming to achieve the best balance between the accuracy and computational efficiency of the algorithm and ensure the effectiveness and reliability of the Isolation Forest Algorithm in the laboratory safety management experiment.

The number of Isolation Tree is set to 100. The number of Isolation Tree is a key factor affecting the performance of the Isolation Forest Algorithm. A larger number of Isolation Tree can improve the accuracy of the algorithm, but it will also increase the computational cost and time. In our experimental environment, considering the data scale and complexity of the simulated laboratory, 100 Isolation Trees can ensure that the algorithm has sufficient detection ability for abnormal data and can complete the calculation within an acceptable time range. Through pre-experiments, we found that when the number of trees is less than 80, the algorithm may not be able to accurately detect some complex abnormal

patterns; when the number of trees exceeds 120, although the accuracy is slightly improved, the calculation time is significantly increased, which is not conducive to real-time monitoring and analysis. The height of the Isolation Tree is set to 8. The height of the tree determines the construction depth of the isolated tree and further affects the algorithm's ability to divide data and detect abnormalities. A smaller tree height may lead to insufficiently detailed data division and inability to effectively separate abnormal points; a larger tree height will increase the computational complexity and time. In our experiment, according to the dimension and distribution characteristics of the data, the height of 8 can ensure the effectiveness of data division while avoiding excessive calculation. Through tests on different heights, we found that when the tree height is less than 6, some abnormal data cannot be effectively isolated; when the tree height is greater than 10, the computational efficiency is significantly decreased and the improvement in the accuracy of abnormality detection is not significant.

The sub-sample size is set to 256. The sub-sample size refers to the number of samples selected each time when building an isolated tree. A suitable sub-sample size can balance the accuracy and computational efficiency of the algorithm. In our experimental data, the sub-sample size of 256 can make the algorithm neither lose information due to too small a sample nor increase unnecessary computational volume due to too large a sample when processing data. Through comparison of the experimental results of different sub-sample sizes, we found that when the sub-sample size is less than 200, the accuracy of the algorithm will be affected and it may not be able to accurately identify some abnormal data; when the sub-sample size is greater than 300, the computational efficiency will be significantly reduced and it is not conducive to fast data processing.

3.4 Experimental Design

(A) Data Preprocessing

Data Cleaning: After obtaining various data from the simulated laboratory, the data cleaning operation is carried out first. Since the data collected by sensors may be affected by environmental interference and some factors of the equipment itself, there are noise and outliers. For temperature data, there may be instantaneous spikes or troughs, which may be caused by a short-term failure of the sensor or a sudden change in the local environment. For computer equipment parameters, such as CPU usage, there may be an unreasonably high value due to a short-term abnormality of some background programs. We adopted multiple methods for data cleaning. For data points that significantly deviate from the normal range, they are removed by setting a reasonable threshold. For example, according to historical data and the normal fluctuation range of the laboratory environment, data points of temperature data outside the range of $[-10^{\circ}\text{C}, 50^{\circ}\text{C}]$ are determined as outliers and deleted. For continuous data sequences, a sliding window method is used for smoothing processing to eliminate local fluctuation noise.

Data Normalization: After data cleaning, data normalization processing is carried out. The data collected by different sensors and equipment have different dimensions and value ranges. For example, temperature data may be between $[-10^{\circ}\text{C}, 50^{\circ}\text{C}]$, while power consumption data may be between $[0,$

1000] watts. Such data differences will affect the processing effect and comparison accuracy of the algorithm. We adopted a linear normalization method to map the data to the $[0, 1]$ interval. For a certain data feature x , its normalization formula is $x_{new} = \frac{x - x_{\min}}{x_{\max} - x_{\min}}$, where $x_{\min}$ and $x_{\max}$ are the minimum and maximum values of the feature data respectively. Through data normalization, all data are in the same dimension, facilitating the processing and comparison of the Isolation Forest Algorithm.

(B) Experimental Comparison

In order to comprehensively verify the effectiveness of the Isolation Forest Algorithm, we compared it with traditional anomaly detection algorithms in an experiment. The distance-based anomaly detection algorithm and the density-based anomaly detection algorithm were selected as comparison objects.

Principle of the Distance-based Anomaly Detection Algorithm: This algorithm judges outliers based on the distance between data points. Usually, the distance between each data point and other data points is calculated, such as the Euclidean distance. If the average distance between a data point and other data points exceeds a certain threshold, it is judged as an outlier.

Principle of the Density-based Anomaly Detection Algorithm: This algorithm identifies anomalies according to the density around data points. If the density of the area where a data point is located is significantly lower than that of the surrounding area, the data point is considered an outlier.

Comparison Indicators: We set detection accuracy, detection time, and false alarm rate as comparison indicators. Detection accuracy refers to the proportion of outliers correctly detected by the algorithm; detection time refers to the time taken by the algorithm to process data and obtain results; false alarm rate refers to the proportion of normal data misjudged as outliers by the algorithm.

4. Experimental Process and Result Analysis

4.1 Experimental Process

We conducted multiple experiments to ensure the reliability of the results. In each experiment, data from different time periods were randomly selected from the data collected in the simulated laboratory for anomaly detection. Specifically, the amount of data in each experiment was about 20% of the total data volume and covered different operating states of the laboratory, such as high-load and low-load periods of computer equipment, as well as different environmental conditions, such as periods with higher temperature and higher humidity.

4.2 Result Analysis

(A) Detection Accuracy

Isolation Forest Algorithm: In different experiments, the detection accuracy of the Isolation Forest Algorithm can reach more than 90%. This benefits from its unique algorithm principle. By randomly selecting features and partition points, it can effectively capture abnormal patterns in the data. For example, in an experiment for computer equipment anomaly detection, the Isolation Forest Algorithm accurately detected the abnormal increase in CPU usage caused by heat dissipation problems and the

abnormal fluctuation in memory usage caused by software conflicts.

Traditional Algorithms: The detection accuracy of the distance-based anomaly detection algorithm and the density-based anomaly detection algorithm is around 80%. When dealing with high-dimensional data, the distance-based algorithm may ignore some local abnormal patterns due to the complexity of distance calculation. In the case of uneven data distribution, the density-based algorithm may misjudge some normal data in low-density areas as abnormal values.

(B) Detection Time

Isolation Forest Algorithm: Due to its unique random partitioning method, the Isolation Forest Algorithm can quickly process large-scale data sets. In the experiment, the detection time of the Isolation Forest Algorithm is only about 1/3 of that of traditional anomaly detection algorithms. For example, when processing a data set containing 10,000 data points, the Isolation Forest Algorithm only needs about 10 seconds to complete the detection, while the distance-based anomaly detection algorithm needs about 30 seconds and the density-based anomaly detection algorithm needs about 25 seconds.

Traditional Algorithms: The distance-based anomaly detection algorithm and the density-based anomaly detection algorithm have relatively low computational efficiency because they need to perform complex distance calculations or density estimations.

(C) False Alarm Rate

Isolation Forest Algorithm: Through careful analysis of experimental data and parameter adjustment, the false alarm rate of the Isolation Forest Algorithm is only about 1/2 of that of traditional anomaly detection algorithms. This indicates that the Isolation Forest Algorithm can more accurately identify real abnormal situations and reduce the interference of false alarms on laboratory safety management. For example, in the detection of environmental temperature and humidity data, the Isolation Forest Algorithm can accurately distinguish between normal temperature fluctuations caused by seasonal changes and abnormal temperature changes caused by equipment failures, while traditional algorithms may produce more false alarms due to inaccurate grasping of data characteristics.

Traditional Algorithms: The distance-based anomaly detection algorithm and the density-based anomaly detection algorithm are prone to false alarm situations in some complex data environments.

To present the experimental results more intuitively, we have organized the following analysis Table 1.

Table 1. The Performance of Different Algorithms in Laboratory Anomaly Detection

Algorithm	Detection Accuracy	Detection Time (for processing 10,000 data points)	False Alarm Rate
Isolation Forest Algorithm	More than 90%	About 10 seconds	1/2 of traditional algorithms

Distance-based Detection Algorithm	Anomaly	About 80%	About 30 seconds	High
Density-based Detection Algorithm	Anomaly	About 80%	About 25 seconds	High

Through the above experimental process and result analysis, it can be clearly seen that the application of the Isolation Forest Algorithm in laboratory safety management has significant advantages. It is superior to traditional anomaly detection algorithms in terms of detection accuracy, detection time, and false alarm rate, and can more accurately and quickly detect abnormal situations in the laboratory, providing a more reliable guarantee for laboratory safety management.

5. Conclusion

Based on the analysis of the current situation of safety management in university laboratories, this paper proposes a method of applying the Isolation Forest Algorithm to laboratory safety management. The effectiveness of the Isolation Forest Algorithm in laboratory safety anomaly detection has been verified through experiments in a simulated laboratory environment. The experimental results show that the Isolation Forest Algorithm has the advantages of high detection accuracy, short detection time, low false alarm rate, no need for prior knowledge, and easy interpretation. It can quickly and accurately detect abnormal situations in the laboratory, providing a new and effective means for the safety management of university laboratories.

In future research, the parameter settings of the Isolation Forest Algorithm can be further optimized to improve the performance and stability of the algorithm. At the same time, other machine learning algorithms and data analysis techniques can be combined to construct a more complete laboratory safety management system, realizing comprehensive monitoring and early warning of laboratory safety and providing a more powerful guarantee for the safe operation of university laboratories.

References

- Dong, Y. Q. (2023). Research on the Detection Method of Potential Attacks in Railway Communication Networks Based on Isolation Forest Algorithm. *Railway Signalling & Communication*, 59(07), 54-59.
- Wang, F., Zhang, L. C., You, T. T. et al. (2024). Analysis and Prospect of the Current Situation of Safety Management in University Teaching Laboratories in the Digital Intelligence Era. *Laboratory Science*, 27(04), 204-209.
- Wang, X., He, Y., Yuan, M. W. (2024). Defense Strategy Against False Data Injection Attacks Based on Stacking and Isolation Forest. *Intelligent Computer and Application*, 14(07), 222-226.
- Yang, F. Q., Li, X., & Hong, Y. D. (2022). Safety Management of University Laboratories Based on the Meta-Model of Behavioral Safety Management. *Research and Exploration in Laboratory*, 195

41(07), 307-312.

Zhang, K., & Lai, Y. N. (2023). Potential Electrical Safety Hazards and Countermeasures in University Laboratories. *Popular Standardization*, 2023(18), 124-126.

Zhao, D., Ma, T. R. (2024). Application of Power Monitoring Platform Based on Isolation Forest Algorithm. *Electrical Engineering*, 2024(10), 126-129.

Zheng, J. W., Zhu, Z. L., Xu, W. B. et al. (2024). Research on the Current Situation and Improvement Strategies of Safety Management in University Laboratories. *Henan Science and Technology*, 51(14), 80-83.