



Simulation-Based Environmental Risk Prioritization for Industrial Waste Management

Najwa H Samrgandi

Data Science Department, College of Computing, Umm Al-Qura University, Saudi Arabia

(Received: 14th February 2026; Accepted: 27th April 2026)

Abstract: Industrial processes produce various types of waste which, when poorly handled, pose public health risks and compromise environmental integrity. There is an increase in environmental regulatory data and tools; however, most waste management practices focus on established regulatory provisions and retrospective management, which limits their ability to anticipate emerging environmental risks and support early intervention. In industrial contexts, this limitation is particularly serious, as delays in risk identification contribute to cumulative environmental degradation in soil and water systems. This study proposes a data-driven intelligent framework for industrial waste management that integrates machine-learning-based classification of regulatory-informed proxy risk levels, waste data, operational standards, and environmental indicators for waste management. A representative synthetic (simulation-based) dataset was generated based on national technical guidelines and regulatory waste classification schemes to enable controlled evaluation of the proposed framework under well-defined regulatory and environmental conditions. An integrated Environmental Risk Index (ERI) was developed by combining hazard intensity, normalized waste quantity, and environmental sensitivity. A Random Forest model, capable of handling uncertainty in regulatory and environmental datasets, was used to classify proxy environmental risk levels (rather than directly modeling real-world contamination processes). The analysis showed 89% classification accuracy and identified high-risk waste streams with high precision, while maintaining consistent differentiation between medium- and low-risk categories. The proposed framework supports structured environmental risk prioritization and encourages evidence-based regulatory decision-making that takes environmental, social, and governance (ESG) factors into consideration, within a simulation-based proof-of-concept setting.

Keywords: industrial waste management; environmental risk assessment; machine learning; decision support systems; AI-based modeling.



(*) Corresponding Author:

Najwa H Samrgandi

Data Science Department, College of Computing, Umm Al-Qura University, Saudi Arabia

E-mail: nhsamrgandi@uqu.edu.sa

1. Introduction

Industrial activities generate a wide range of waste materials, including sludge, tailings, spent chemicals, and contaminated packaging. What makes the problem worse is that in many cases, reliance on delayed monitoring and incomplete datasets affects handling and disposal decisions, which leads to risks to environmental integrity and public health (Chen, 2025). Across many industrial sectors, waste management still depends on fixed regulatory standards and reactive inspection methods that are not effective for early risk detection. Although regulatory instruments, including Extended Producer Responsibility (EPR), exist, limitations in enforcement capacity, as well as the practice of post-event supervision affect environmental protection outcomes (Tran et al., 2025) and leave early warning signals of environmental risk unnoticed.

Reliable information on the generated waste is central to waste management. Therefore, it is important to design frameworks integrating waste characteristics with operational and environmental conditions in order to support consistent risk identification and management actions (Dawar et al., 2025; Xu & Zheng, 2026). Recent studies emphasize the need for more integrated and adaptive environmental risk assessment frameworks to address increasing system complexity and uncertainty (Geng et al., 2024). Moreover, integrating waste characterization procedures with environmental risk assessment tools (Bisciotti et al., 2025), as well as introducing automated analysis that would not only reduce the burden of analysis and increase efficiency of monitoring (Rao et al., 2025). Such approaches also help align waste management strategies with environmental priorities in light of competing economic and social constraints (Shahab et al., 2026).

Accordingly, this study proposes a data-driven intelligent framework for industrial waste management that integrates waste characterization data, operational parameters, and environmental indicators to: (1) classify industrial waste streams according to proxy environmental risk levels; (2) support the structured prioritization of environmental risk under data-limited conditions; and (3) provide recommendations to support decision-making for treatment, recycling, controlled storage, and final disposal.

1.1 Problem Statement

Although there is an increasing number of industrial waste systems, many of them still rely on reactive, rule-based frameworks which are not efficient in identifying environmental risks early. Moreover, there is no platform that would offer an integrated and data-driven approach that would link waste characteristics with environmental sensitivity indicators.

1.2 Objectives of the Study

To address the identified research gap, this study pursues the following objectives:

- To develop and evaluate a data-driven analytical framework that uses publicly available data to classify regulatory-informed proxy environmental risk levels linked to industrial waste types.
- To implement supervised machine learning methods through a simulation-based case study to support risk-informed decision-making in industrial waste management as a methodological proof-of-concept under data-limited conditions.

1.3 Research Contributions

This study presents three main contributions.

- First, it brings together regulatory logic, environmental indicators, and machine learning into a single, integrated framework for decision support in industrial waste management. Instead of treating these elements separately, the framework shows how they can work together in a more practical and structured way.
- Second, it examines how machine learning models behave when they operate under predefined rules. In such settings, models tend to follow the logic already built into the data rather than discovering entirely new patterns.
- Third, it evaluates how classification performance is affected when the labels are not perfectly reliable. By simulating uncertain and imperfect data conditions, the study shows how model performance responds to noise and highlights some limitations of relying on data-driven approaches in real-world scenarios.

1.4 Research Questions

This study seeks to address the following research questions:

- Q1:** To what extent can publicly available waste and environmental data be used to classify proxy environmental risk levels linked to types of industrial waste?
- Q2:** How can data-driven risk classification approaches enhance industrial waste management decision-making in light of limited data?

2. Related works

AI-assisted predictive models are more appropriate than systems that use conventional statistics in certain situations (Abdeljaber et al., 2025; Chau et al., 2025). For example, AI-based sorting systems produce accurate results that can distinguish between material types, aiding

better recycling opportunities (Ahmad et al., 2025). Such schemes are useful for industry since it produces a vast assortment of waste and information (Islam et al., 2025). Other studies require biowaste identification using YOLOv8-SPP architectures (Balasubramaniam et al., 2025) and risk assessments for different forms of waste (Bisciotti et al., 2025). These analyses combine modeling with environmental and biochemical testing.

AI-based prediction of waste generation has been identified by Alsabt et al. (2024) and (Alhathloul et al., 2025). However, issues relating to the rigor of such systems and their validation have raised concerns. Mullane et al. (2025) claimed that inadequate validation practices produce false outcomes, which produce poor decisions. In addition, AI-based real-time optimization has been shown to be possible (Rao et al., 2025; Anitha & Parthiban, 2025).

AI in decision-support systems for smart cities is now available (Nesmachnow et al., 2025). Nevertheless, limitations include poor data quality and the transparency of such models (Zou et al., 2025). Forecasting, classification, decision-making, and risk assessments are usually treated as isolated subjects. Although a combination of some of these topics is present in the Literature, such research usually examines a single stage in the process, which is unsuitable for industrial-scale use (Dao et al., 2025). Recent work has also emphasized the integration of emerging technologies within waste management systems, although such efforts often remain fragmented and limited to specific operational contexts (Martínez-Ramón et al., 2025).

2.1 Research gap

Overall, the literature documents substantial advances in applying AI and machine learning to waste forecasting, classification, and environmental risk estimation. Systematic reviews confirm the growing use and diversity of computational methods and decision-support systems deployed across the waste management lifecycle, particularly within smart city and IoT-enabled contexts (Nesmachnow et al., 2025), but challenges related to the quality of data, transparency of the used models, and analytical integration are commonly found (Zou et al., 2025).

Despite the progress, most studies treat forecasting, classification, decision support and risk assessment as isolated analytical components. Although some multilayer approaches link sensing, sorting, planning, and treatment processes, their focus is usually on specific lifecycle stages and there is still a lack of integrated structures suitable for industrial-scale implementation (Dao et al., 2025). The lack of such analytical frameworks restricts the transferability of research outcomes to regulatory and

industrial decision-making, which further demonstrates the need for such an integrated approach that would combine waste characterization, environmental sensitivity assessment, risk modeling, and decision-support.

3. Methodology

This study develops a structured framework to prioritize environmental risks associated with industrial waste. It integrates regulatory logic, environmental indicators, and supervised machine learning within a unified analytical framework, rather than treating them as separate components.

Due to the limited availability of detailed facility-level data and observed contamination outcomes, the study adopts a simulation-based proof-of-concept design. This approach is commonly used to evaluate operational systems under controlled conditions (Dawar et al., 2025). Instead of attempting to directly model real-world contamination processes, the framework operates within a controlled synthetic environment that reflects regulatory principles and realistic industrial scenarios.

The primary objective of this study is methodological rather than predictive: to examine the interaction between rule-based classification, composite risk indices, and machine learning models, and to assess their consistency and robustness under uncertainty. Accordingly, the study focuses on risk prioritization and decision-support design, rather than aiming to predict real-world environmental outcomes

3.1 Data Sources and Variables

The selected variables capture key dimensions of environmental risk in industrial waste management, including hazard severity, scale of waste generation, and environmental exposure. These dimensions are consistent with those commonly used in environmental risk assessment frameworks. The selection of variables is informed by established regulatory approaches, while the ISIC Rev.4 system (United Nations, 2008) is used to classify industrial activities within the dataset. This classification system is also implemented within the Saudi statistical framework (GASTAT, 2019). To balance realism with analytical flexibility, the dataset integrates regulatory-based categorical variables with simulated continuous variables. This hybrid structure preserves the integrity of regulatory classifications while introducing controlled variability in operational and environmental conditions. Specifically, variables such as hazard classification, waste type, and storage conditions follow established regulatory categories, whereas variables including waste quantity and distance to water sources are generated to reflect plausible variation under controlled assumptions.

It is important to clarify the nature of the outcome variable. The target variable, referred to as `risk_proxy`, does not represent an observed environmental outcome. Instead, it is a regulatory-informed proxy classification introduced strictly for methodological purposes. Accordingly, the framework does not aim to model real-world contamination processes; rather, it seeks

to examine how risk categorization can be structured and evaluated within a controlled, regulation-informed setting. Importantly, the proxy labels are derived using rule-based logic based on the same input variables included in the model. Therefore, the model is expected to reproduce the underlying rule-based structure rather than identify independent causal relationships.

Table 1. Variables, roles, and sources within the simulation-based framework

Variable	Description	Type	Role in Framework	Source
waste-type	Type of industrial waste	Categorical	Contextual variable (not used in final model)	NCM (2025)
hazard	Hazard classification	Binary	Core predictor (H_n)	NCM (2025)
quantity	Generated waste quantity (kg)	Continuous	Core predictor (Q_n)	Simulated
storage	Storage method	Categorical	Contextual variable	NCM (2025)
industry	Industrial sector	Categorical	Contextual variable	ISIC Rev.4 (GASTAT, 2019)
dist_water	Distance to nearest water source (km)	Continuous	Environmental variable (used to derive S)	GIS / Simulated
ERI	Environmental Risk Index	Continuous	Prioritization index (not used in model)	Computed
risk-proxy	Proxy environmental risk class	Categorical	Target variable (rule-based + noise)	Synthetic

Table 1 summarizes the variables included in the framework, along with their descriptions, types, roles, and sources. It distinguishes between variables used directly in the predictive model and those included for contextual or decision-support purposes.

To further illustrate how these variables are operationalized, the analytical workflow can be conceptualized as a sequence of interconnected steps. The process begins with the core variables hazard, waste quantity, and distance to water sources, which are transformed into derived features (H_n , Q_n , and S). These features form the foundation for both index construction and predictive modeling.

These components are subsequently combined to compute the Environmental Risk Index (ERI), which serves as a composite measure for prioritizing scenarios.

Although ERI is not used as an input to the predictive model, it provides a useful benchmark for comparing relative risk levels across cases.

In the subsequent step, proxy risk labels (`risk_proxy`) are generated using a rule-based approach informed by regulatory logic. To better reflect real-world uncertainty, a controlled level of noise is introduced into these labels, simulating potential inconsistencies in environmental classification.

Finally, supervised machine learning models are trained using the core predictors (H_n , Q_n , and S). The resulting predictions are interpreted alongside the ERI within a conceptual decision-support framework, enabling a comparative perspective between rule-based and data-driven approaches to risk prioritization.

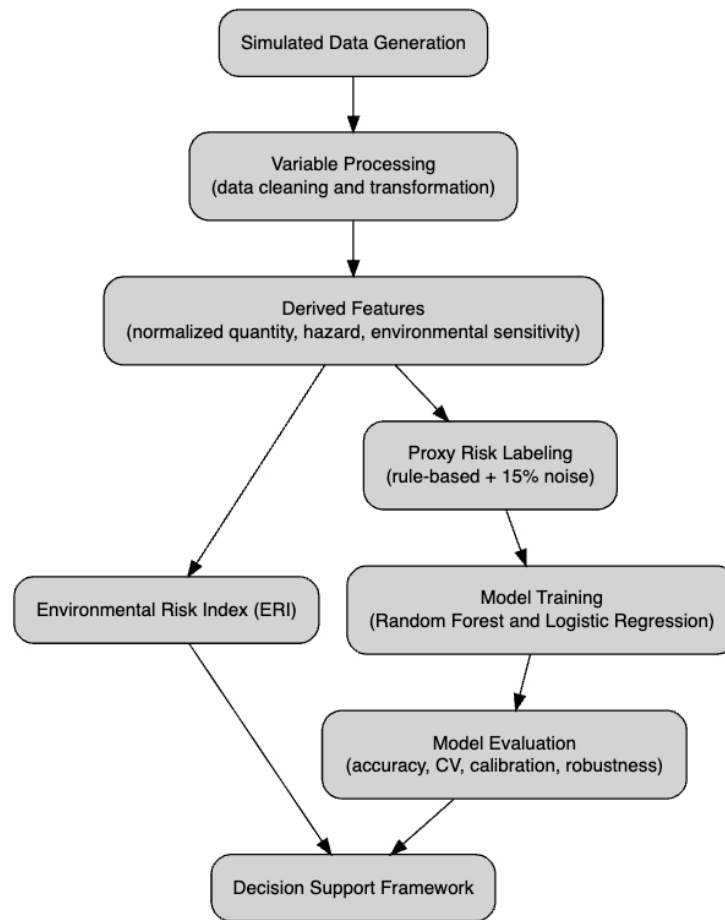


Figure 1. Workflow of the simulation-based environmental risk classification framework

Figure 1 presents the overall workflow of the simulation-based environmental risk classification framework, illustrating how data generation, feature engineering, rule-based labeling, and machine learning components interact within the proposed system.

3.2 Dataset Generation

To support controlled experimentation, a synthetic dataset consisting of 500 observations was constructed. The objective was not to generate arbitrary random data, but to create a structured dataset that reflects plausible industrial waste scenarios under regulatory constraints.

To ensure transparency and reproducibility, the data generation process was defined using explicit distributions and parameter ranges for each variable. The hazard variable was simulated as a binary outcome using a Bernoulli distribution ($p = 0.6$ for non-hazardous waste), reflecting typical industrial patterns.

Waste quantity was generated as a continuous variable within a predefined operational range (e.g., 10–2000 kg), representing realistic variability in industrial production volumes. Distance to water sources was simulated within a realistic spatial range (e.g., 0.2–12 km), serving as a proxy for environmental exposure.

Categorical variables, including waste type, storage method, and industry sector, were sampled from predefined sets aligned with regulatory classifications. These categories were treated as fixed inputs to maintain consistency with existing waste management frameworks.

All simulations were conducted using a fixed random seed to ensure full reproducibility of the dataset. The generation process was intentionally designed to balance realism and simplicity, allowing controlled evaluation of the proposed framework under clearly defined conditions. A simplified description of the data generation procedure is provided to facilitate replication.

3.3 Data Preprocessing and Feature Engineering

Before building the model, the dataset was cleaned up and aligned so that all variables were consistent.

Waste quantity was standardized using z-score normalization to produce (Q_n). Hazard levels were converted into numerical values (H_n) to reflect their relative impact, with higher values assigned to more hazardous waste. Environmental sensitivity (S) was calculated based on how close each case is to water sources, the closer the distance, the higher the sensitivity.

These elements were then combined to form the Environmental Risk Index (ERI), which serves as a simple way to compare different scenarios. It's important to note that ERI is not meant to represent actual environmental risk, but rather to give a relative sense of priority across cases.

In that sense, the index is used for interpretation rather than prediction. This type of index-based approach has also been used in structuring environmental risk assessments, especially in hazardous waste contexts (Zheng et al., 2025).

Since Q_n is standardized, ERI values can sometimes be negative. This doesn't indicate anything problematic, it just reflects where a case sits relative to others in the dataset.

Finally, testing different weighting schemes showed that ERI remains fairly stable. This suggests that the

index captures a consistent ranking pattern, rather than depending too heavily on a specific set of weights.

Importantly, ERI was not included as an input to the predictive model. This decision was made to avoid redundancy and prevent information leakage, given that the index is derived from the same underlying variables.

3.4 Risk Labeling and Uncertainty Modeling

The classification labels used in this study are intentionally designed as synthetic constructs. Rather than relying on observed environmental outcomes, the labels are generated through a simplified rule-based structure informed by general regulatory principles in hazardous waste management.

To ensure transparency and reproducibility, the deterministic rules used to assign proxy risk levels are explicitly defined in Table 2.

Table 2. Rule-based structure for proxy risk classification

Hazard Indicator	Quantity Threshold	Rule Description	Assigned Risk Class
1 (Hazardous)	> 1500 kg	Exceeds high-risk threshold	High
1 (Hazardous)	$\leq$ 1500 kg	Within moderate range	Medium
0 (Non-hazardous)	Any	No hazardous properties	Low

The rules presented in Table 2 define the deterministic logic used to generate the target variable (risk_proxy). These rules are intentionally simplified and are based on a predefined threshold ($T_n = 1500$ kg), which is introduced for simulation purposes rather than derived from a specific regulatory standard. This abstraction allows the study to operate within a controlled, rule-constrained environment while preserving key aspects of regulatory reasoning.

While simplified, this structure is consistent with general regulatory approaches that consider both hazard characteristics and waste generation scale in risk prioritization. It therefore serves as a conceptual approximation of screening-level decision logic, rather than a direct representation of real-world classification systems.

Importantly, the proxy labels are derived from the same underlying variables used as model inputs. As a result, the learning task is inherently constrained, and the model is expected to reproduce the predefined rule structure rather than identify independent relationships. This design is intentional and aligns with the methodological objective of evaluating model behavior within a controlled analytical setting.

To introduce realism and avoid a fully deterministic system, a controlled level of uncertainty is incorporated into the labeling process. Specifically, 15% of the class

labels are randomly reassigned. This simulates potential inconsistencies in environmental classification and reflects the imperfect nature of real-world data (Lee & Shin, 2025). It also aligns with noise-aware learning approaches that assess model robustness under uncertain labeling conditions.

3.5 Model Development

Within this framework, machine learning is not intended to uncover new environmental relationships or to predict real-world risk outcomes. Rather, it is explicitly used as a controlled analytical instrument to evaluate how predefined rule-based classifications behave under varying levels of uncertainty. This shifts the role of the model from predictive inference to methodological evaluation, with a focus on consistency, robustness, and behavior under controlled conditions.

To operationalize this evaluation, a Random Forest classifier was selected as the primary modeling approach due to its flexibility and robustness when applied to structured datasets. The dataset was divided into training (80%) and testing (20%) subsets to allow for out-of-sample evaluation. The model was trained using 300 trees, and feature importance measures were extracted to support interpretability.

The model was intentionally restricted to three core predictors: normalized waste quantity (Q_n), encoded

hazard level (H_n), and environmental sensitivity (S). This constraint ensures alignment between the model inputs and the underlying rule-based structure used to generate the proxy labels, while maintaining a transparent and interpretable modeling setup.

In addition to the Random Forest model, a multinomial logistic regression model was implemented as a baseline. This provides a point of comparison between a flexible machine learning approach and a simpler statistical model under the same controlled conditions.

Given that the target variable is derived from the same underlying predictors, the modeling task is inherently constrained. Accordingly, the models are not intended to identify independent relationships, but rather to assess the extent to which the predefined classification structure can be consistently reproduced, particularly under conditions of label uncertainty.

3.6 Model Evaluation

Model performance was evaluated using a set of standard classification metrics, including accuracy, precision, recall, and F1-score. Confusion matrices were used to examine class-level performance and to identify potential imbalances in classification behavior.

However, given the rule-constrained nature of the modeling task, evaluation focused not only on overall predictive performance, but also on model stability and robustness under varying levels of uncertainty.

To assess this, performance was systematically evaluated across different levels of label noise (0%, 15%, and 30%). This allows for a controlled examination of how sensitive the models are to inconsistencies in the target labels, and whether they can maintain stable classification behavior under degraded conditions.

In addition to classification metrics, probabilistic performance was assessed using Brier scores and calibration analysis. These measures provide insight into how well predicted probabilities align with the underlying label distribution, offering a more nuanced view of model reliability beyond discrete predictions.

A hold-out validation approach was combined with five-fold cross-validation to ensure that the results are not dependent on a single data split. This provides a more robust assessment of model consistency across different subsets of the data.

Importantly, given that the target labels are derived from predefined rules, the evaluation is not intended to demonstrate predictive superiority, but rather to assess the consistency, reliability, and robustness of model behavior within a controlled simulation environment. This evaluation design allows for distinguishing between

models that merely replicate deterministic rules and those that maintain stable behavior under uncertainty.

3.7 Decision-Support Framework

To bridge analytical outputs with practical application, a conceptual decision-support layer was developed. This component translates model outputs and index-based assessments into structured risk management actions.

Within this framework, risk categories derived from the classification model are linked to differentiated response strategies. High-risk cases are associated with immediate intervention and stricter control measures, while medium-risk cases require targeted monitoring and mitigation. Low-risk cases are managed through routine operational procedures. Rather than prescribing fixed actions, these categories provide a structured basis for prioritization under constrained resources.

Importantly, the Environmental Risk Index (ERI) is evaluated alongside model predictions to offer a complementary perspective on risk prioritization. While the model reflects rule-constrained classification behavior under uncertainty, the ERI provides a continuous measure of relative risk across scenarios. The combined use of these components enables a more nuanced interpretation of risk, particularly in cases where classification boundaries may be sensitive to noise.

This decision-support layer is intentionally presented as a conceptual proof-of-concept. Its purpose is not to define an operational system, but to illustrate how rule-based logic, index-based prioritization, and model outputs can be integrated within a unified analytical framework to support structured decision-making.

3.8 Reproducibility

All analyses were conducted in the R programming environment using open-source tools. Reproducibility was ensured through the use of fixed random seeds and explicitly defined procedures for data generation, preprocessing, and model training.

The full analytical workflow, from synthetic data construction to model evaluation, was designed to be fully transparent and replicable. Each stage of the process follows clearly specified assumptions and parameter settings, allowing independent reproduction of the results under identical conditions.

Given the simulation-based nature of the study, reproducibility is a central component of the methodological design. The controlled structure of the dataset and the deterministic generation of proxy labels provide a consistent baseline for evaluating model behavior across repeated experiments.

4. Results and Discussion

4.1 Dataset Characteristics and Class Distribution

The dataset consists of 500 synthetically generated observations designed to represent a range of industrial waste scenarios under controlled simulation conditions.

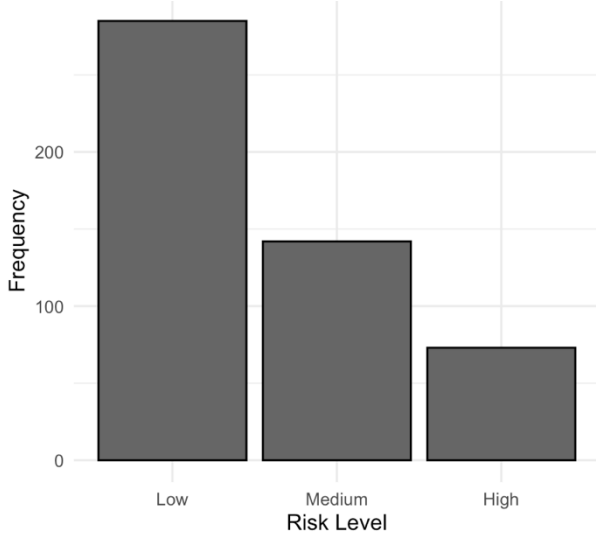


Figure 2. Class Distribution

As shown in Figure 2, the initial class distribution is imbalanced, with Low-risk cases representing the majority (62.6%), followed by Medium (26.4%) and High (11.0%). This imbalance reflects the underlying rule-based structure, where non-hazardous conditions are more frequently represented in the simulated dataset. After introducing 15% label noise, the distribution becomes slightly more dispersed (Low: 57.0%, Medium: 28.4%, High: 14.6%), increasing class overlap and introducing controlled ambiguity into the classification structure. This shift is methodologically significant, as it moves the dataset beyond a purely deterministic structure and introduces a level of uncertainty that more closely resembles classification behavior under imperfect conditions. By introducing controlled noise, the dataset no longer reflects a perfectly rule-driven system, but instead provides a more appropriate testbed for evaluating model robustness and stability under uncertainty.

4.2 Environmental Risk Index (ERI)

The Environmental Risk Index (ERI) is constructed as a composite measure integrating normalized waste quantity (Q_n), hazard level (H_n), and environmental sensitivity (S), providing a continuous representation of relative risk across simulated scenarios.

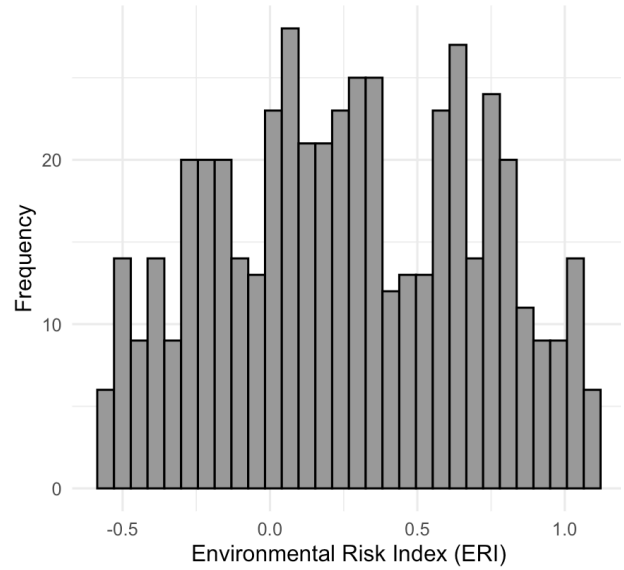


Figure 3. ERI Distribution

As illustrated in Figure 3, ERI values are reasonably well distributed across the dataset, with a mean of 0.264 and a median of 0.255, indicating a relatively balanced spread of risk values rather than concentration within a narrow range. The observed range (−0.558 to 1.093) captures variability in both waste characteristics and environmental exposure, reflecting the combined influence of hazard, quantity, and spatial context.

Negative values arise from the normalization process and do not affect interpretability, as ERI is designed for relative comparison rather than absolute classification thresholds.

Table 3. Sensitivity analysis of ERI weighting schemes

Weights (Q, H, S)	Mean	Min	Max
(0.4, 0.4, 0.2)	0.264	-0.558	1.093
(0.5, 0.3, 0.2)	0.208	-0.761	1.156
(0.3, 0.5, 0.2)	0.320	-0.354	1.033

The sensitivity of the index to weighting choices is summarized in Table 3, where alternative weight configurations produce only minor variations in overall behavior. This indicates that the index is not overly dependent on a specific parameterization, suggesting that its structure is primarily driven by the underlying variables rather than by arbitrary weighting choices.

While the selected weights (0.4, 0.4, 0.2) are not derived from a formal optimization procedure, they represent a balanced and interpretable configuration aligned with screening-level prioritization approaches, where hazard and quantity are typically treated as primary drivers of risk.

Importantly, the stability of ERI supports its role as a complementary analytical tool for relative prioritization, rather than a standalone predictive measure. ERI was therefore excluded from the predictive model to avoid redundancy and potential information leakage, as the index is directly derived from the same underlying variables used as model inputs.

4.3 Model Performance

The Random Forest model achieved an overall accuracy of 0.89, with a Kappa value of 0.806. The associated 95% confidence interval (0.812–0.944) indicates that the observed performance is consistent across the evaluated sample.

While the model outperforms the No Information Rate (0.55), this comparison should be interpreted with caution given the structured nature of the dataset. This apparent improvement primarily reflects the model's alignment with the underlying rule-based structure embedded in the data.

Accordingly, the observed performance should not be interpreted as evidence of predictive capability, but rather as an indication of how effectively the model reproduces the predefined rule-based classification under conditions of uncertainty. In this context, model performance is structurally bounded by the rule-constrained design, and therefore reflects consistency with the label generation process rather than the discovery of independent relationships within the data.

4.4 Classification Behavior

The classification behavior is summarized in Table 4 and further illustrated in Figure 4, providing a detailed view of class-level prediction patterns.

Table 4. Confusion matrix for the test dataset

Predicted \ Actual	Low	Medium	High
Low	55	6	3
Medium	0	18	2
High	0	0	16

As shown in Table 4, the majority of predictions align with the true class labels, indicating consistent reproduction of the underlying classification structure.

Notably, no Low-risk cases are misclassified as High-risk, and no High-risk cases are misclassified as Low-risk. This absence of extreme misclassification is particularly important from a decision-making perspective, as it preserves critical risk boundaries and reduces the likelihood of severe misallocation of management actions.

Misclassifications are primarily concentrated between adjacent categories (Low–Medium and

Medium–High), reflecting the transitional nature of the classification boundaries rather than random error. This pattern is consistent with the introduction of label noise and suggests that classification uncertainty is localized near boundary conditions, where small variations in input features can influence class assignment.

This behavior is further reflected in Figure 4, where off-diagonal values remain limited and are largely confined to neighboring classes.

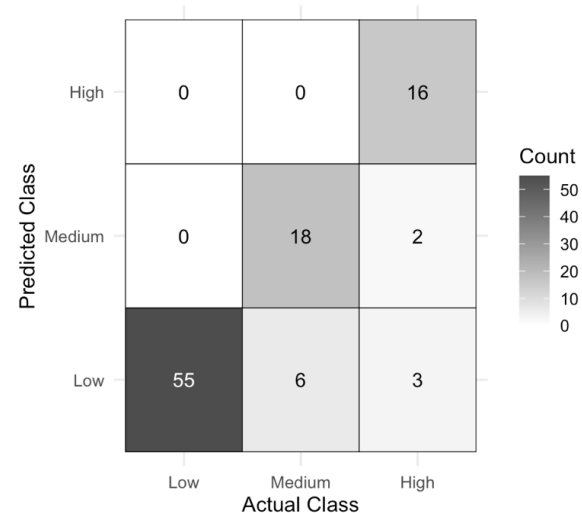


Figure 4. Confusion Matrix

4.5 Class-Level Performance

Table 5. Class-wise performance metrics

Class	Precision	Recall	F1-score
Low	0.859	1.000	0.924
Medium	0.900	0.750	0.818
High	1.000	0.762	0.865

Class-wise performance metrics are summarized in Table 5, providing a more detailed view of model behavior across different risk categories.

The results indicate generally consistent performance across classes, although interpretation should be considered in light of the rule-constrained nature of the labeling process.

Low-risk cases achieve perfect recall, indicating that all such cases are correctly identified. This behavior is consistent with the underlying rule structure, where non-hazardous conditions are more distinctly separated from higher-risk categories.

High-risk cases achieve perfect precision, indicating that all cases classified as high risk correspond to true high-risk instances. From a decision-making perspective, this reduces the likelihood of false alarms in critical scenarios, although it should be interpreted within the constraints of the predefined labeling structure.

The Medium-risk class exhibits lower recall, reflecting its intermediate position between low and high-risk categories. This reduction is consistent with increased overlap at classification boundaries, where small variations in input features and introduced label noise can influence class assignment.

4.6 Model Stability

Model stability was assessed using five-fold cross-validation to evaluate the consistency of performance across different data partitions.

The resulting average accuracy (0.874) and Kappa value (0.772) are consistent with the test results, indicating stable performance across different splits rather than sensitivity to a specific partition.

This consistency further supports the interpretation that model behavior is driven by the underlying rule-constrained structure, with limited variability introduced by data partitioning.

4.7 Feature Importance

Feature importance results are presented in Figure 5, providing insight into how the model weights the relative contribution of input variables.

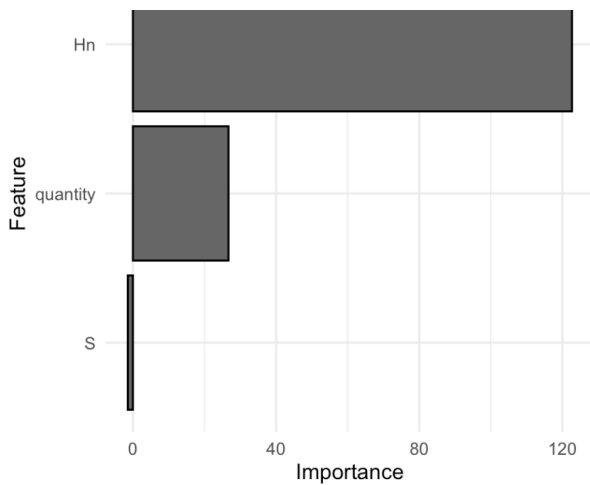


Figure 5. Feature Importance

As shown in Figure 5, hazard level (H_n) appears as the most influential predictor, followed by waste quantity, while environmental sensitivity (S) contributes less prominently.

This hierarchy is structurally expected, as hazard and quantity are the primary drivers embedded in the rule-based labeling process. As such, feature importance

4.9 Comparison with Logistic Regression

A comparison between the Random Forest model and multinomial logistic regression is presented in Table 7.

Table 7. Comparison of Random Forest and Multinomial Logistic Regression

Model	Accuracy	CI Lower	CI Upper	Kappa
Random Forest	0.89	0.812	0.944	0.806
Multinomial Logistic Regression	0.86	0.776	0.921	0.752

should not be interpreted as the discovery of independent relationships, but rather as a reflection of the predefined classification logic.

This alignment reflects consistency with general risk assessment principles, where hazard severity and scale of exposure are central factors, while also reinforcing the rule-constrained nature of the modeling framework.

4.8 Comparison with Rule-Based Classification

The comparison with a rule-based baseline reveals a central methodological result

Table 6. Comparison between rule-based classification and Random Forest

Model	Accuracy	Kappa
Rule-based classification	0.89	0.806
Random Forest	0.89	0.806

As shown in Table 6, both approaches achieve identical performance (accuracy = 0.89, Kappa = 0.806). This equivalence is not coincidental, but a direct consequence of the rule-constrained design, where the learning task is structurally aligned with the label generation process. Because the target labels are derived from predefined rule-based criteria, the machine learning model is inherently constrained to learn and reproduce the same underlying structure, even in the presence of moderate uncertainty introduced through label noise.

Rather than weakening the study, this result provides a clear delineation of the role of machine learning within rule-constrained environments. Under the current setup, machine learning does not provide additional predictive capability beyond the predefined rule structure, but instead demonstrates its ability to reproduce and maintain that structure under uncertainty.

In practical applications, the value of machine learning is expected to increase when data are not strictly governed by predefined rules, and when variability, missing values, and nonlinear relationships become more prominent. In such contexts, data-driven models can extend beyond fixed classification schemes and support more adaptive and scalable decision-making processes.

Taken together, these findings define the boundary conditions under which machine learning is most effective, positioning machine learning as a complementary analytical tool rather than a replacement for established rule-based regulatory frameworks.

The Random Forest model shows a modest but consistent improvement in both accuracy and Kappa. The relatively small performance gap is expected, given that both models operate under the same rule-constrained structure.

This improvement likely reflects the flexibility of the Random Forest model in handling boundary conditions and label noise, rather than the identification of fundamentally new relationships.

4.10 Probability Calibration

Probability calibration results are summarized in Table 8 and further illustrated in Figure 6, providing insight into the reliability of predicted probabilities across different risk classes.

Table 8. Brier Scores by Class and Model

Model	Class	Brier Score
Random Forest	Low	0.0887
Random Forest	Medium	0.0913
Random Forest	High	0.0657
Multinomial Logistic Regression	Low	0.0772
Multinomial Logistic Regression	Medium	0.1077
Multinomial Logistic Regression	High	0.0901

As shown in Table 8, the Random Forest model achieves lower Brier scores for the high-risk class compared to logistic regression, indicating improved probabilistic calibration, particularly in scenarios where accurate confidence estimation is critical for decision-making.

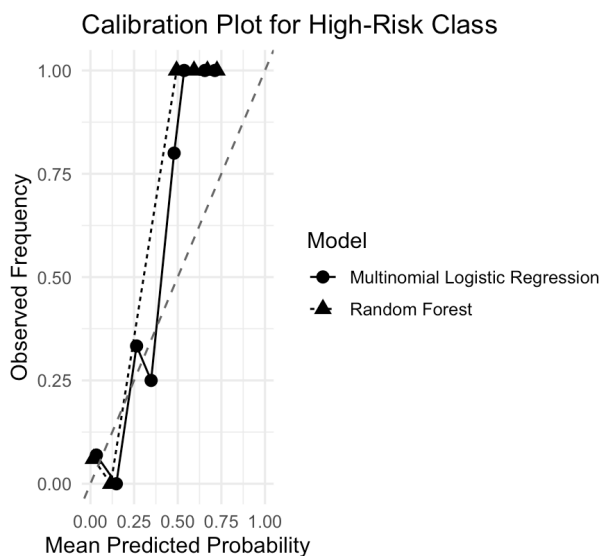


Figure 6. Calibration Plot for High-Risk Class

This is further supported by Figure 6, where predicted probabilities show a generally consistent alignment with observed frequencies, indicating acceptable calibration performance across probability ranges.

Given the rule-constrained nature of the labeling process, calibration results should be interpreted as reflecting the model's ability to maintain probabilistic consistency under uncertainty, rather than capturing true underlying probability distributions.

4.11 Robustness to Uncertainty

The impact of label noise on model performance is evaluated in Table 9 and further illustrated in Figure 7, providing a controlled assessment of robustness under increasing levels of uncertainty.

Table 9. Model performance under different noise levels

Noise Level	Accuracy	95% CI
0%	1.00	(0.964–1.000)
15%	0.89	(0.812–0.944)
30%	0.84	(0.753–0.906)

As shown in Table 9, accuracy declines progressively as noise increases, dropping from near-perfect performance at 0% noise to 0.84 at 30% noise. This trend reflects the model's sensitivity to increasing uncertainty in the target labels.

Performance remains relatively stable, indicating that the model retains a consistent ability to reproduce the underlying classification structure under degraded conditions.

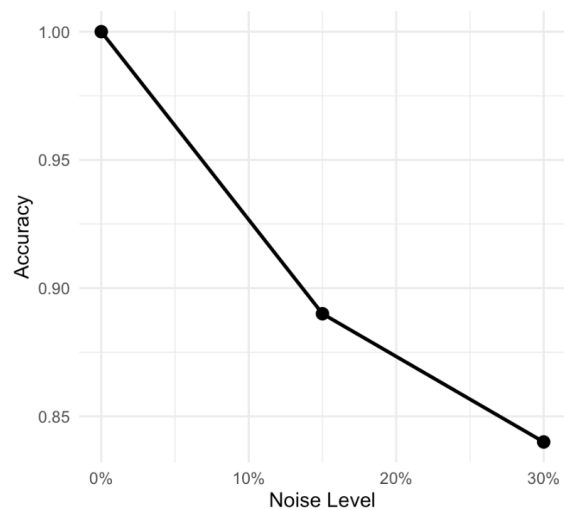


Figure 7. Noise Robustness

Figure 7 illustrates a smooth degradation pattern rather than a sharp drop, suggesting that performance deterioration occurs gradually as uncertainty increases. This behavior indicates that the model is not overly sensitive to moderate levels of label noise, and that classification performance degrades in a controlled and predictable manner.

Given the rule-constrained nature of the problem, this robustness should be interpreted as the model's ability to maintain structural consistency under uncertainty, rather than as resilience in discovering new patterns within noisy data.

4.12 Decision-Support Implications

The decision-support component of the framework establishes a clear linkage between predicted risk levels and corresponding waste management actions.

Table 10. Decision-support rules for waste management actions

Predicted Risk Level	Recommended Management Action
High	Immediate intervention, strict containment, priority inspection, and specialized treatment or disposal
Medium	Controlled handling, regular monitoring, and reinforced storage and compliance checks
Low	Routine waste management procedures and standard regulatory compliance

As illustrated in Table 10, each risk category is associated with a distinct set of recommended interventions. High-risk cases necessitate immediate action and stringent control measures. Medium-risk cases require continuous monitoring and targeted mitigation strategies, whereas low-risk cases can be effectively managed through routine procedures and standard compliance practices.

This mapping serves as a practical bridge between analytical outputs and operational decision-making. Although the framework is presented as a proof of concept, it highlights the potential of translating risk classification results into actionable priorities within real-world waste management systems.

5. Limitations and Future Directions

Although the proposed system demonstrates strong performance within a simulated environment, its results are inherently influenced by the assumptions underlying synthetic data generation and proxy labeling. Consequently, it does not fully capture the complexity of real-world industrial waste systems, where risk is shaped by context-specific interactions among contaminants, environmental conditions, and operational factors (He et al.,2024).

Moreover, the model is trained on proxy labels derived from predefined regulatory rules based on the same input variables. This dependency limits its capacity to provide substantial improvements over traditional rule-based approaches. Additionally, the incorporation of controlled label noise (15%) represents a simplifying assumption and may not adequately reflect the uncertainty present in real-world scenarios.

Furthermore, the absence of real operational data constrains the generalizability of the findings beyond the simulated context.

Future research should therefore investigate hybrid modeling approaches that integrate rule-based logic with data-driven techniques, particularly in data-scarce environments. Combining domain expertise with machine learning has the potential to mitigate the limitations associated with rule-derived proxy labels and enhance model adaptability and robustness in real-world applications.

6. Conclusion

A data-driven system is presented for industrial waste management, which accounts for national guidelines on environmental risks. Combining AI-based modeling with environmental sensitivity indicators, this system provides a structured and rigorous approach for classifying risk levels. An environmental risk index and a random forest classifier were used. The system achieved an overall accuracy of 89% and, most importantly, identified higher-risk waste with precision and consistency. Importantly, this performance should be interpreted within the constraints of the rule-based classification structure, where model outcomes primarily reflect the reproduction of predefined regulatory logic under controlled uncertainty.

An important finding of this study is that, within a rule-constrained setting, the machine learning model reproduces the underlying regulatory logic rather than outperforming it, clarifying its role as a tool for supporting structured decision-making rather than replacing rule-based systems.

Despite successes, a limitation arises relating to its use of simulated data, which affects transferability of the findings to other contexts. Although this enables controlled and reproducible analysis, it does not fully capture the complexity of real-world waste management systems, including variability in operational practices and data quality. Nonetheless, the findings offer a methodological basis for further exploration and subsequent field-based validation.

The suggestion for future studies is to use real-time monitoring data, actual operational records from industry, and geospatial information to improve external validity and increase applicability of the model to diverse settings. In addition, the integration of XAI techniques may increase transparency and support regulatory acceptance. Future work may also examine the framework under less constrained conditions, where variability and complexity extend beyond predefined rule structures. The proposed framework is intended to contribute to the development of more sustainable and data-driven approaches to industrial waste management.

7. References

- Abdeljaber, A., Al-Smadi, S., Abu-Talib, M., & Abdallah, M. (2025). Comparative analysis of machine learning and conventional methods for waste generation forecasting. *Cleaner Engineering and Technology*, 4, 100992. <https://doi.org/10.1016/j.clet.2025.100992>
- Ahmad, G., Aleem, F. M., Alyas, T., Abbas, Q., Nawaz, W., Ghazal, T. M., Aziz, A., Aleem, S., Tabassum, N., & Ibrahim, A. M. (2025). Intelligent waste sorting for urban sustainability using deep learning. *Scientific Reports*, 15, 27078. <https://doi.org/10.1038/s41598-025-08461-w>
- Alhathloul, N., Lakhout, A., Abdalla, G. M. T., Alghamdi, A., Shaban, M., Alshahir, A., Alshahr, S., Alali, I., & Alshammari, F. M. (2025). Assessing waste management using machine learning forecasting. *Sustainability*, 17(19), 8654. <https://doi.org/10.3390/su17198654>
- Alsabt, R., Alkhaldi, W., Adenle, Y. A., & Alshuwaikhat, H. M. (2024). Optimizing waste management strategies through artificial intelligence and machine learning: An economic and environmental impact study. *Cleaner Waste Systems*, 8, 100158. <https://doi.org/10.1016/j.clwas.2024.100158>
- Anitha, R., & Parthiban, A. (2025). Smart waste ecosystems under Industry 5.0: A next-generation framework. *Results in Engineering*, 9, 107988. <https://doi.org/10.1016/j.rineng.2025.107988>
- Balasubramaniam, S., Nagu, B., Bansal, S., Faruque, M. R. I., & Al-mugren, K. S. (2025). Advanced predictive analytics for bio-waste management using YOLOv8-SPP to enhance waste prediction and sustainability in smart cities. *Scientific Reports*, 15, 23583. <https://doi.org/10.1038/s41598-025-09433-w>
- Biscioti, A., Brombin, V., Song, Y., Bianchini, G., & Cruciani, G. (2025). Classification and predictive leaching risk assessment of construction and demolition waste using multivariate statistical and machine learning analyses. *Waste Management*, 196, 60–70. <https://doi.org/10.1016/j.wasman.2025.02.033>
- Chau, K. Y., Moslehpour, M., Pan, S.-H., & Akhundzada, A. (2025). Advanced predictive modeling of municipal solid waste management using robust machine learning. *Scientific Reports*, 15, 43247. <https://doi.org/10.1038/s41598-025-27237-w>
- Chen, B. (2025). Analysis of industrial solid waste and the possibility of improving management practices. *International Journal of Environmental Technology and Management*, 28(4), 345–362. <https://doi.org/10.1080/17518253.2025.2493155>
- Dao, S. V. T., Le, T. M., Tran, H. M., Pham, H. V., Vu, M. T., & Chu, T. (2025). Integrating artificial intelligence for sustainable waste management: Insights from machine learning and deep learning. *Waste Engineering*, 8, 100158. <https://doi.org/10.1016/j.wsee.2025.07.001>
- Dawar, I., Srivastava, A., Singal, M., Dhyani, N., & Rastogi, S. (2025). A systematic literature review on municipal solid waste management using machine learning and deep learning. *Artificial Intelligence Review*, 58, 183. <https://doi.org/10.1007/s10462-025-11196-9>
- General Authority for Statistics (GASTAT). (2019). Government entities begin implementing the national classification for economic activities (ISIC4). Kingdom of Saudi Arabia. <https://stats.gov.sa/en/classification-news>
- Geng, J., Zhang, Y., Wang, L., Liu, H., & Chen, X. (2024). Advances and future directions of environmental risk assessment: A comprehensive review. *Science of the Total Environment*, 908, 167878. <https://doi.org/10.1016/j.scitotenv.2023.167878>
- He, C., Li, Y., Zhang, Q., Wang, X., & Liu, J. (2024). Risk assessment of e-waste: Liquid crystal monomers re-suspension caused by coastal dredging operations. *Science of the Total Environment*, 933, 173176. <https://doi.org/10.1016/j.scitotenv.2024.173176>
- Islam, M. M., Hasan, S. M. M., Hossain, M. R., Uddin, M. P., & Mamun, M. A. (2025). Towards sustainable solutions: Effective waste classification framework via enhanced deep convolutional neural networks. *PLOS ONE*, 20(6), e0324294. <https://doi.org/10.1371/journal.pone.0324294>
- Lee, J.-S., & Shin, D.-C. (2025). Prediction of waste generation using machine learning: A regional study in Korea. *Urban Science*, 9(8), 297. <https://doi.org/10.3390/urbansci9080297>
- Martínez-Ramón, N., Istrate, R., Iribarren, D., & Dufour, J. (2025). Unlocking advanced waste management models: Machine learning integration of emerging technologies into regional systems. *Resources, Conservation and Recycling Advances*, 11, 200253. <https://doi.org/10.1016/j.rcradv.2025.200253>
- Mullane, P., Fitzpatrick, C., & Grua, E. M. (2025). Rethinking machine learning evaluation in waste management research. *Sustainability*, 17(23), 10707. <https://doi.org/10.3390/su172310707>

- National Center for Waste Management (NCM). (2025). *Standards and technical guidelines for sorting at source – industrial waste*. Riyadh, Saudi Arabia. <https://istitlaa.ncc.gov.sa/en/civil/ncwm/sortingatsourceindustrialwaste/Documents/Standards%20&%20Technical%20Guidelines%20Sorting%20at%20Source%20Industrial%20Waste.pdf>
- Nesmachnow, S., Rossit, D., & Moreno-Bernal, P. (2025). A literature review of recent advances on innovative computational tools for waste management in smart cities. *Urban Science*, 9(1), 16. <https://doi.org/10.3390/urbansci9010016>
- Rao, R., Singh, S., Salas, M., Sarker, A., Kumar, R., Wang, Y., Lucia, L., Mittal, A., Yarbrough, J., Barlaz, M. A., Singh, A., & Pal, L. (2025). AI-powered municipal solid waste management: A comprehensive review from generation to utilization. *Frontiers in Energy Research*, 13. <https://doi.org/10.3389/fenrg.2025.1670679>
- Shahab, S., Simic, V., Dutta, A. K., Pamucar, D., & Anjum, M. (2026). Evaluating blockchain-based waste management investments in smart cities using a multi-criteria decision support framework. *Scientific Reports*, 16, 3179. <https://doi.org/10.1038/s41598-025-33085-5>
- Tran, T. Y. A., Herat, S., & Kaparaju, P. (2025). Current status and compliance management of Extended Producer Responsibility (EPR) regulations for packaging waste in Vietnam. *Circular Economy and Sustainability*, 5, 2921–2958. <https://doi.org/10.1007/s43615-025-00568-6>
- United Nations. (2008). *International standard industrial classification of all economic activities (ISIC), Rev.4*. New York: United Nations.
- Xu, P., & Zheng, H. (2026). A multi-AI approach to predicting municipal solid waste generation and recycling demand in Hong Kong. *Resources, Conservation and Recycling*, 225, 108590. <https://doi.org/10.1016/j.resconrec.2025.108590>
- Zheng, F., Jin, X., Huo, H., Tian, Y., & Li, H. (2025). A new method for constructing an environmental risk analysis and assessment index system for hazardous waste. *Results in Engineering*, 106986. <https://doi.org/10.1016/j.rineng.2025.106986>
- Zou, Q., Duan, H., Yang, Z., Wang, X., Tian, Y., Hou, H., & Yang, J. (2025). Machine learning-assisted solid waste life-cycle management: Applications, constraints, and future opportunities. *Resources, Conservation and Recycling*, 219, 108320. <https://doi.org/10.1016/j.resconrec.2025.108320>