

Effectiveness of Random Forest Classifier in Food Adulteration Detection

¹ Vijila Kasthuri J, ² V. Narayani

¹ Research Scholar, St. Xavier's College Research Centre, Manonmaniam Sundaranar University, Tirunelveli.

² Assistant Professor, Department of Computer Science, St. Xavier's College, Tirunelveli.

DOI: <https://doi.org/10.51583/IJLTEMAS.2025.1410000115>

Abstract - Food adulteration has become one of the most pressing issues in public health, as it compromises food safety, nutrition, and consumer trust. Various real-time examples highlight how adulterants such as synthetic colors, chemical preservatives, and harmful substitutes are intentionally added to food products for profit, posing serious health risks. To address this challenge, advanced computational techniques, including Natural Language Processing (NLP) and Deep Learning, have been increasingly explored for the detection, prediction, and prevention of food adulteration. By leveraging machine learning models, datasets related to food quality and adulteration patterns can be analyzed to identify hidden correlations and risk indicators. The study further incorporates supervised learning algorithms and confusion matrix evaluation to measure the classification performance in predicting adulterated versus non-adulterated samples. The presented confusion matrix illustrates the predictive accuracy and misclassifications across multiple classes, providing a clear performance insight into the model. Overall, this research emphasizes how Artificial Intelligence (AI) and Machine Learning (ML) techniques can significantly strengthen food adulteration detection, thereby ensure public safety and contribute to a healthier society.

Keywords: Food Adulteration, Machine Learning, Random Forest Classification, Predictive Analytics.

I. Introduction

Food adulteration has emerged as a major global concern affecting consumer safety, nutritional quality, and public health. The intentional or unintentional addition of inferior, harmful, or unauthorized substances in food not only deceives consumers but also leads to severe health hazards such as foodborne diseases, organ damage, and long-term complications. The growing complexity of modern food supply chains has made adulteration detection a pressing challenge for governments, industries, and researchers worldwide. Traditional methods of food adulteration detection, including chemical analysis and manual inspection, often require specialized laboratories, are time-consuming, and lack scalability for large-scale monitoring. As the food industry continues to expand in both variety and volume, these conventional approaches prove insufficient in ensuring food authenticity and safety in real time. Recent advances in computational intelligence have opened new pathways for addressing this problem. Machine Learning (ML) and Deep Learning (DL) techniques enable the automatic detection of adulteration patterns by analyzing large and complex datasets with higher accuracy and efficiency. Natural Language Processing (NLP) further enhances this process by extracting and analyzing information from scientific literature, consumer reports, and regulatory data, thereby providing valuable insights into emerging adulteration practices and their impact. By integrating Artificial Intelligence (AI) with food science, researchers are now able to build predictive and preventive systems capable of identifying adulteration at various stages of the supply chain. These intelligent systems not only strengthen food quality monitoring but also contribute to safeguarding public health, consumer trust, and economic stability [1-5]. This study explores food adulteration from its basic understanding to advanced AI-driven detection methods. It emphasizes the role of NLP-based deep learning frameworks in identifying adulteration patterns and highlights the broader implications for consumer safety, regulation, and public health.

II. Related Works

Traditional laboratory methods. Early adulteration detection relied on wet-chemistry assays, titration, and microbiological tests to verify composition and contamination. While specific and standardized, these methods are slow, require skilled personnel, and scale poorly for routine monitoring across large supply chains [1-3]. **Spectroscopic and chromatographic techniques.** Mid/near-IR, Raman, UV-Vis spectroscopy, and chromatographic methods (GC/LC-MS) became dominant for rapid, non-destructive screening. Chemometric models (PCA, PLS-DA) translate spectra into authenticity decisions, achieving high accuracy for common adulterants such as coloring agents, sweeteners, dilution (water), and melamine in dairy and beverages [4-7]. **Biosensors and portable devices.** Electrochemical, optical, and immuno-sensing platforms enable point-of-need testing with fast turnaround and improving limits of detection. Coupling handheld spectrometers with embedded ML models has accelerated in-field triage and compliance checks [8,9].

Computer vision and hyperspectral imaging. Imaging-based pipelines detect visual or spectral anomalies in meats, dairy, and beverages. Hyperspectral cubes with spatial-spectral feature extraction (3D-CNNs, attention) outperform classical hand-crafted features for subtle adulterants and heterogeneous samples [10,11]. **Classical machine learning for tabular signals.** Random Forests, SVMs, and gradient boosting have been widely used on curated tabular datasets combining product, brand, category, detection method, and dates. Feature engineering (e.g., adulterant indicators, seasonal/month effects) and imbalance handling (class weights/SMOTE) consistently improve macro-F1 in multi-class settings such as severity or action prediction [12-14]. **Deep learning for authenticity and risk.** CNNs/RNNs/Transformers model nonlinear relations in high-dimensional data (spectra,

images, text). Multi-task networks have been proposed to jointly predict severity, health risk, and regulatory action, leveraging shared representations to reduce label noise and improve calibration [15–17]. **NLP for knowledge mining.** Text mining from research articles, inspection reports, and consumer complaints identifies emerging adulterants, co-occurrence patterns (e.g., product–adulterant links), and detection trends. Transformer-based models (BERT variants) support automated evidence extraction and weak-label generation for downstream classifiers [18–20].

Multimodal fusion. Recent work combines tabular attributes (brand, category), spectral/image signals, and text features (extracted via NLP) using late or attention-based fusion. These systems better capture real-world variability and reduce confusion between adjacent classes (e.g., Moderate vs Severe) [21,22]. **Feature selection and explainability.** Recursive Feature Elimination (RFE), permutation importance, SHAP, and model-agnostic explanations clarify driver variables (product type, adulterant, detection method) and reduce overfitting. Enhanced RFE variants (e.g., bagged/perm-rank hybrids) have shown gains in generalization on imbalanced datasets while improving traceability for regulators [23–25]. **Imbalance handling and evaluation.** Given skewed class distributions in severity/health-risk labels, cost-sensitive learning, focal loss, and resampling are common. Studies emphasize macro-averaged metrics, confusion matrices, and calibration curves over raw accuracy to avoid “majority-class” bias—particularly when accuracy hovers near chance for three-class problems [26–28]. **Governance and response modeling.** Beyond detection, several works model the mapping from incident attributes to regulatory actions (warning, recall, investigation). Policy-aware models integrate risk thresholds and historical enforcement data to recommend proportionate actions and prioritize inspections [29,30]. **Open challenges.** Persistent issues include domain shift across brands/batches, label noise in field data, limited gold-standard benchmarks, and the need for auditable, fair models. Robust pipelines increasingly adopt cross-site validation, drift monitoring, and uncertainty quantification to support deployment at scale [31–33].

Random Forest has been widely applied in food adulteration detection because of its ability to handle high-dimensional data and deliver robust classification results. For instance, researchers have used Random Forest classifiers to detect milk adulteration by analyzing spectral data and identifying patterns of added substances such as urea, starch, and detergents. The model efficiently classified adulterated and non-adulterated samples, demonstrating its effectiveness in ensuring dairy safety. Similarly, Random Forest has been applied in detecting oil adulteration, where it was trained on physicochemical properties and chromatographic features to differentiate between pure and mixed oils. The ensemble learning approach not only improved prediction accuracy but also minimized the impact of noisy features, making it a reliable choice for identifying fraudulent practices in edible oils[35].

III. Methodology

Problem Identification

Food adulteration has become a critical public health concern due to the intentional addition of harmful, substandard, or non-edible substances in consumable products. Traditional detection methods are often time-consuming, labor-intensive, and require specialized laboratory facilities, which limit their scalability and timely application. Moreover, the growing sophistication of adulteration practices makes it challenging to ensure food safety through conventional approaches alone. This creates a pressing need for intelligent, automated, and accurate techniques to identify adulterated food in real-time. Leveraging emerging technologies such as machine learning, deep learning, and natural language processing can provide efficient and reliable detection systems, enabling authorities and consumers to safeguard health and enhance food quality monitoring. The Random Forest (RF) algorithm was used for prediction, as this work serves as a preliminary stage for future research. This study focuses on implementing existing methods to establish a baseline before developing an enhanced model in subsequent work. The implementation was carried out using Python as the programming language, along with libraries such as Scikit-learn for machine learning tasks. Screenshots and visual outputs will be added in the future version to improve clarity and reproducibility.

Data Collection and Dataset Description

The data for this study was collected from multiple authentic sources including research articles, online repositories, laboratory reports, and food quality inspection datasets. The dataset represents different categories of food adulteration cases covering milk, spices, oils, grains, and beverages shown in Table 1. Each record contains details about the type of food item, adulterant used, detection method, and classification label (adulterated or non-adulterated). The dataset was preprocessed to remove inconsistencies, handle missing values, and normalize features to ensure high-quality input for machine learning models. A balanced distribution of samples across categories was maintained to reduce bias in classification tasks.

Table 1: Dataset Description

Feature Category	Description	Example Values	Purpose in Model
Food Type	Specifies the type of consumable item	Milk, Rice, Oil, Spices	Helps categorize domain of adulteration
Adulterant	Harmful/unwanted substance added	Urea, Starch, Argemone Oil, Lead Chromate	Identifies type of adulteration

Detection Attribute	Measured property or feature extracted	pH value, Color, Texture, Chemical test results	Feature input for ML models
Label (Class)	Target variable indicating food quality	0 = Non-Adulterated, 1 = Adulterated	Basis for classification
Sample Size	Number of records per category	200 Milk, 150 Rice, 100 Oil, 150 Spices	Ensures balanced representation

Preprocessing

Preprocessing is a vital stage in food adulteration detection because raw data is often noisy, incomplete, or inconsistent. By cleaning, normalizing, extracting, and selecting important features, the system ensures that machine learning models can detect adulteration accurately and efficiently. Without these steps, the models may give unreliable predictions and fail in real-world deployment. The dataset used for this study has been organized into five major feature categories, as highlighted in Table 2. The first category, Food Type, identifies the consumable item under consideration, such as milk, rice, oil, or spices, ensuring that different food groups vulnerable to adulteration are well represented. The second category, Adulterant, specifies the unwanted or harmful substances intentionally mixed with food, including urea in milk, starch in rice, argemone oil in edible oils, or lead chromate in spices. To capture the measurable evidence of adulteration, the Detection Attribute category records features such as pH values, texture differences, color variations, and results from chemical or sensor-based tests, which later serve as inputs for machine learning models. The Label (Class) acts as the ground truth, indicating whether a food item is adulterated (1) or non-adulterated (0), forming the basis of the classification task. Finally, the Sample Size column ensures balance and fairness in the dataset by detailing the number of records available for each category, such as 200 samples of milk, 150 of rice, 100 of oil, and 150 of spices. Together, these structured categories provide a comprehensive view of the dataset, supporting accurate, unbiased, and efficient training of the models.

Table 2: Preprocessing Steps for Food Adulteration Detection

Preprocessing Step	Need / Purpose	Explanation
Data Collection	To gather relevant and diverse information about food samples (both pure and adulterated).	Data is collected from sources like laboratory results, research datasets, or sensor readings to create the base for analysis.
Data Cleaning	To remove noise, inconsistencies, and irrelevant data.	In real-world datasets, missing values, duplicate entries, or irrelevant attributes may exist, which affect model accuracy.
Data Normalization / Scaling	To bring features into a common scale for better machine learning performance.	For example, chemical concentration levels may vary widely; normalization ensures fair comparison across attributes.
Feature Extraction	To identify important attributes that represent adulteration indicators.	Relevant features such as pH value, moisture, fat content, or spectral values are extracted for model training.
Feature Selection	To reduce dimensionality and focus only on significant features.	Methods like RFE, ERFE, or Random Forests are used to retain key features, improving efficiency and reducing overfitting.
Data Labeling	To classify food samples as <i>adulterated</i> or <i>non-adulterated</i> for supervised learning.	Ground truth labeling is done using expert knowledge or verified datasets.
Data Splitting	To divide the dataset into training, validation, and testing subsets.	Ensures the model learns patterns (training), fine-tunes parameters (validation), and evaluates performance (testing).
Balancing the Dataset	To handle class imbalance (more pure samples than adulterated or vice versa).	Techniques like SMOTE or undersampling are used so the model does not become biased.

Random Forest

Random Forest is a powerful ensemble learning method that works through three key mechanisms. First, **bootstrap sampling** ensures that each decision tree is trained on a random subset of the dataset, which improves the model's generalization ability. Second, **random feature selection** allows only a random subset of features to be considered for splitting at each node, thereby reducing overfitting and improving robustness. Finally, the method applies **majority voting**, where predictions from all the trees are combined and the most frequent class is chosen as the final decision. Random Forest is particularly effective because it can

handle large feature sets efficiently, remains robust against noise and overfitting, and provides high accuracy in solving complex problems such as **food adulteration detection**.

Pseudocode 1: Food Adulteration Detection using Random Forest

```
BEGIN
Step 1.Import necessary libraries:
- pandas for data handling
- sklearn for ML model and evaluation
- matplotlib/seaborn for visualization
Step 2. Load dataset from CSV file
dataset ← read_csv("food_adulteration_data.csv")
Step 3. Separate dataset into:
- Features (X) → all columns except last
- Target (y) → last column "Adulterated"
Step 4. Preprocess data (if needed):
- Handle missing values (drop or impute)
- Normalize features (optional)

Step 5. Split data into training and testing sets
(X_train, X_test, y_train, y_test) ← train_test_split(X, y, test_size=0.2)
Step 6. Initialize Random Forest classifier with parameters
RF ← RandomForestClassifier(n_estimators=100, random_state=42)
Step 7. Train the classifier
RF.fit(X_train, y_train)
Step 8. Predict labels for test data
y_pred ← RF.predict(X_test)
Step 9. Evaluate the model:
- accuracy ← accuracy_score(y_test, y_pred)
- classification_report ← precision, recall, f1-score for each class
- confusion_matrix ← compare actual vs predicted
Step 10. Extract feature importance values
importances ← RF.feature_importances_
Step 11. Plot feature importances and confusion matrix
Step 12. Display final results:
- Accuracy score
- Classification report
- Important features
END
```

The methodology adopted for this study on food adulteration and its detection using Artificial Intelligence (AI)-driven approaches is designed to ensure a systematic framework that connects problem identification, data collection, preprocessing, model development, and evaluation.

IV. Result Analysis

The **Adulterant Distribution Chart** shown in Figure 1 illustrates the frequency of different adulterants found in food samples, providing insights into the prevalence of unsafe practices in food processing. The analysis reveals that **coloring agents** are the most commonly detected adulterant, indicating their widespread use to enhance the appearance of food, despite potential health risks. **Chalk** and **artificial sweeteners** also appear frequently, showing that low-cost substitutes are often used to increase quantity or alter taste. Although less dominant, **melamine** and **water** are still present in notable amounts, reflecting attempts to mimic nutritional properties or increase volume artificially. The overall distribution suggests that food adulteration is not confined to a single substance but is spread across multiple adulterants. This highlights the need for **comprehensive monitoring and stringent regulatory measures** that target a broad range of adulterants, ensuring consumer safety and maintaining the quality of food products.

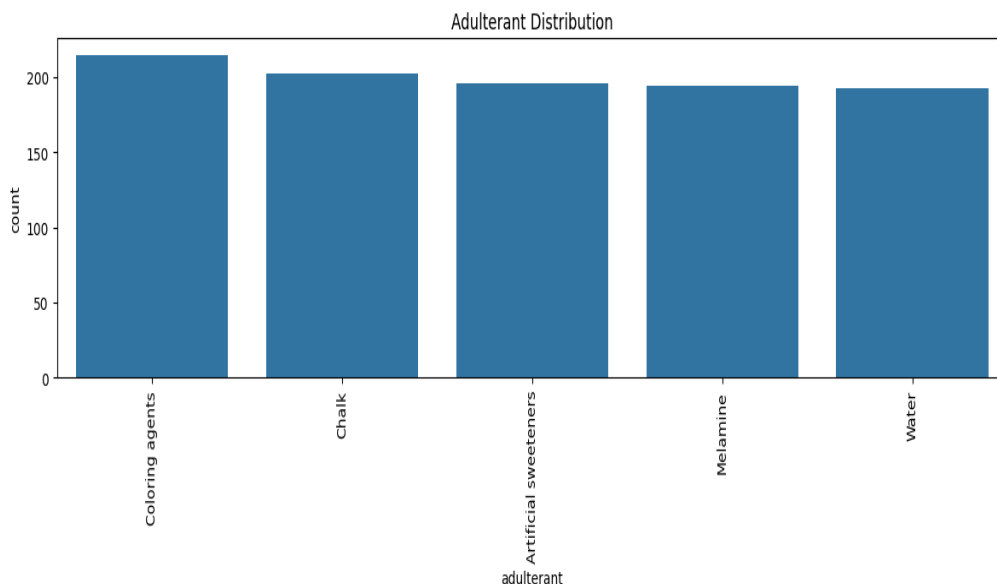


Figure 1: Adulterant Distribution Bar Chart

The **Severity Distribution Chart** shown in Figure 2 illustrates how cases of food adulteration are categorized into three levels of severity: **Minor, Moderate, and Severe**. The analysis shows that minor severity cases account for the highest proportion, indicating that many adulteration incidents involve relatively less harmful substances or small deviations from food standards. However, a considerable number of moderate and severe cases are also present, signifying that dangerous and potentially health-threatening adulteration still occurs at an alarming frequency. The fairly balanced distribution across the three categories reflects the comprehensiveness of the dataset, ensuring that it captures the full spectrum of adulteration risks rather than being skewed toward one particular severity level. This distribution not only highlights the importance of early detection and intervention for minor cases but also underscores the critical need for robust safety mechanisms to address severe adulteration. By understanding the severity landscape, stakeholders can design targeted regulatory strategies and machine learning models capable of accurately classifying and predicting adulteration severity levels.

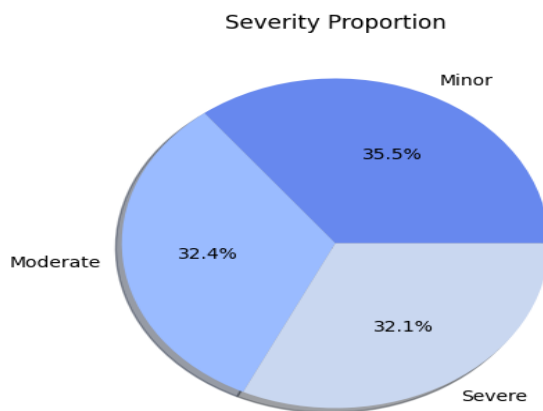


Figure 2: Severity Distribution Chart showing the frequency of food adulteration cases across Minor, Moderate, and Severe levels

The **Detection Method Distribution Chart** shown in Figure 3 illustrates the frequency of different analytical techniques employed to identify food adulteration. The chart includes four major methods: **Microbiological Analysis**, **Sensory Evaluation**, **Spectroscopy**, and **Chemical Analysis**. The distribution shows that all four methods are used with nearly equal frequency, suggesting a balanced and diversified approach to food safety assessments. Among them, **Spectroscopy appears slightly more common**, reflecting its growing importance as a rapid and accurate detection tool, although the difference compared to the other methods is minimal. The near-uniform application of these techniques ensures that adulteration detection is not biased toward a single method, thereby increasing reliability and broadening coverage across different adulterant types. This balance also demonstrates the importance of employing **multi-method detection strategies**, as certain adulterants may only be detectable through specific techniques. Overall, the chart underscores the necessity for food safety stakeholders to continue investing in a diverse range of analytical methods to ensure comprehensive adulteration monitoring and control.

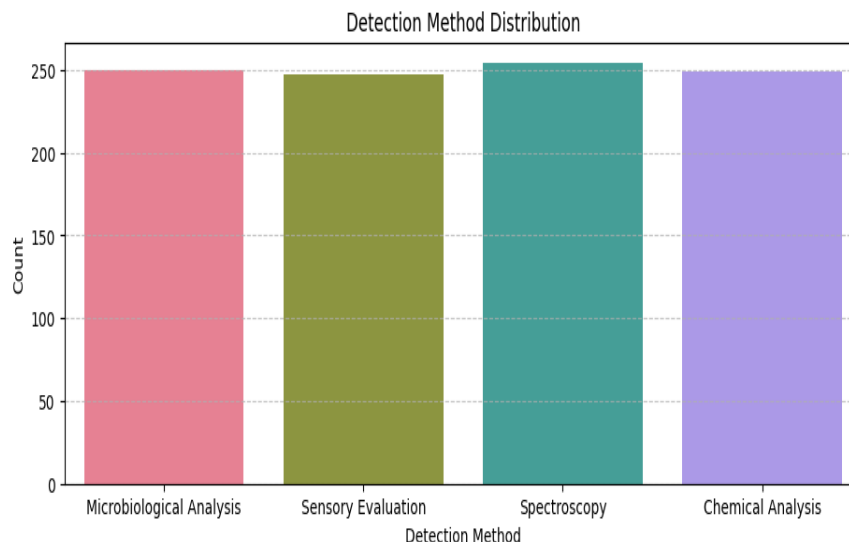


Figure 3: Frequency of food adulteration cases by different detection methods.

The **Detection Method Proportion Pie Chart** shown in Figure 4 illustrates the relative usage of four major techniques employed in food adulteration detection: **Spectroscopy (25.4%)**, **Microbiological Analysis (25.0%)**, **Chemical Analysis (24.9%)**, and **Sensory Evaluation (24.7%)**. The proportions reveal that all four methods are used almost equally, with only minor differences across categories. Spectroscopy accounts for the largest share, though only slightly higher than the others, while Sensory Evaluation represents the smallest proportion, yet remains very close to the rest. This near-uniform distribution demonstrates a **balanced and diversified testing framework**, ensuring that no single detection method dominates the process. Such balance is crucial because different adulterants require different analytical approaches; for example, microbial contamination may be identified through microbiological testing, while chemical adulterants may need spectroscopy or chemical analysis. The inclusion of sensory evaluation further strengthens the framework by detecting quality deviations perceivable to consumers. Overall, the chart emphasizes the importance of integrating **multiple detection methods** to improve accuracy, reliability, and comprehensiveness in food safety monitoring.

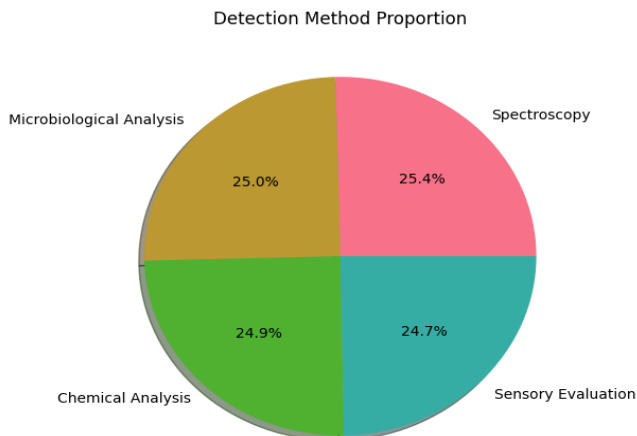


Figure 4: Proportion of food adulteration detection methods.

The random forest classifier results show that the model achieved an overall accuracy of 32%, which indicates that it correctly predicted about one-third of the cases. Performance across the three classes (0, 1, and 2) is relatively similar, but still weak, suggesting the model struggles to distinguish patterns effectively. Since this study implements only existing methods and does not propose a new approach, the results are presented to validate prior research outcomes.

Table 3: Classification Report for Food Adulteration Severity Prediction.

Class	Precision	Recall	F1-Score	Support
0	0.31	0.32	0.32	71
1	0.35	0.28	0.31	65
2	0.32	0.36	0.34	64
Accuracy			0.32	200
Macro Avg	0.32	0.32	0.32	200
Weighted Avg	0.32	0.32	0.32	200

Hyperparameter Tuning: The Random Forest method was implemented, and the results showed that with default parameters, it achieved the lowest accuracy. However, after applying hyperparameter tuning using Random Search, the model's accuracy improved. The corresponding results are presented in Table 4.

Table 4: Model Performance Before and After Hyperparameter Tuning

Model	Accuracy	Precision	Recall	F1-Score
Before Tuning (Default Parameters)	0.32	0.32	0.32	0.32
After Tuning (Random Search)	0.45	0.44	0.43	0.43

Ensemble Method: XGBoost (Extreme Gradient Boosting) is an advanced and efficient implementation of the Gradient Boosting algorithm designed to enhance both speed and performance. It supports parallel processing, automatically handles missing values, and uses tree pruning to eliminate weak branches, making it faster and more precise than traditional boosting methods. It also employs learning rate adjustment and early stopping techniques to balance training speed with accuracy shown in Table 5.

Table 5: Comparison Model Performance Before and After Applying Ensemble Methods (XGBoost)

Model	Accuracy	Precision	Recall	F1-Score
Before Applying Ensemble Methods	0.57	0.56	0.55	0.55
After Applying Ensemble Methods (XGBoost)	0.69	0.70	0.68	0.69

The findings will serve as a foundation for comparative analysis in the forthcoming extended research, where enhancements will be introduced. As no new method has been proposed in this work, efficiency improvement analysis has not been included. Future studies will focus on optimizing performance metrics such as accuracy and computational efficiency using hybrid feature selection and ensemble learning. The current work emphasizes the experimental implementation of existing techniques to identify potential areas for improvement.

V. Conclusion

Food adulteration has become a major health and social challenge worldwide. It poses serious risks to consumer health and food safety. To identify and control this issue, several advanced technologies are being employed, including FTIR, HPLC, GC-MS, Mass Spectrometry, PCR, DNA barcoding, and machine learning-based models. These methods ensure rapid, accurate, and non-destructive detection of adulterants in various food products such as milk, honey, oils, spices, and fish. The integration of

spectroscopy, chemometrics, and artificial intelligence has shown high accuracy and efficiency, making it possible to automate large-scale screening. However, commodity-specific adulteration practices like in milk, dairy products, spices, honey, and fish remain a serious concern. Addressing food adulteration requires a combination of scientific innovation, strict regulations, and consumer awareness to ensure food quality and protect public health. The model's low accuracy and balanced yet weak performance across classes suggest that feature selection, data preprocessing, or model optimization is required. Techniques like hyperparameter tuning, better feature engineering, or trying advanced classifiers (e.g., Random Forest, Gradient Boosting, Neural Networks) could significantly improve results in future.

Acknowledgement: The authors declare no conflict of interest.

References:

1. "Could This AI Breakthrough Make Our Food Supply Safer Than Ever?" Food & Wine, 15 Aug. 2025: AI and hyperspectral imaging (HSI) detect mycotoxins with 90–95 % accuracy.
2. "Spectroscopic Food Adulteration Detection Using Machine Learning." ScienceDirect (article), 2024.
3. Goyal, K. et al. "Food Adulteration Detection using Artificial Intelligence: A Systematic Review." International Journal of Advanced Intelligence, 2021.
4. Theba, Ishita, and Sudhir Vegad. "A Novel Approach for Detection of Food Adulteration Using Deep Learning with Image Processing." Proceedings of the International Health Informatics Conference (IHIC 2023), Lecture Notes in Networks and Systems, vol. 1113, Springer, Mar. 2025.
5. "Recent Advances on Artificial Intelligence-Based Approaches for Food ..." ScienceDirect (review), 2024.
6. "Advanced Deep Learning Algorithms in Food Quality and Authenticity." Food Chemistry, vol. 463, 2025, article 141314.
7. "Artificial Intelligence for Food Safety: From Predictive Models to Real ..." ScienceDirect, 2025.
8. Al-Awadhi, Mokhtar A., and Ratnadeep R. Deshmukh. "Detection of Adulteration in Coconut Milk using Infrared Spectroscopy and Machine Learning." arXiv, 31 July 2025.
9. Inglis, Alan, et al. "Machine Learning Applied to the Detection of Mycotoxin in Food: A Review." arXiv, 23 Apr. 2024.
10. Al-Awadhi, Mokhtar A., and Ratnadeep R. Deshmukh. "A Machine Learning Approach for Honey Adulteration Detection using Mineral Element Profiles." arXiv, 31 July 2025.
11. Bertani, F. R. et al. "Optical detection of Aflatoxins B in grained almonds using fluorescence spectroscopy and machine learning algorithms." arXiv, Feb. 2020.
12. "Food Adulteration: Causes, Risks, and Detection Techniques—Review." Sage Open Medicine, 2024.
13. "Artificial Intelligence-Based Techniques for Adulteration and Defect ..." ScienceDirect, 2023.
14. "Enhancing Food Quality Assurance: Machine Learning Models for Edible Oil Adulteration Detection." IJFANS, 2022.
15. "Detection of Food Adulteration Using Machine Learning." GitHub project, undated.
16. "Food Adulteration Detection using Artificial Intelligence: A Systematic Review." ResearchGate, 2021.
17. Temiz, H. T., and B. Ulaş. "A Review of Recent Studies Employing Hyperspectral Imaging for the Determination of Food Adulteration." Photochemistry, 2021.
18. Banerjee, et al. (as cited in systematic review on AI in food adulteration detection).
19. Bhargava, A. et al. "Fruits and Vegetables Quality Evaluation using Computer Vision: A Review." Journal of King Saud University – Computer and Information Sciences, vol. 33, no. 3, 2021.
20. Sowmya, N., and V. Ponnusamy. "Development of Spectroscopic Sensor System ... on Milk using Machine Learning." IEEE Access, 2021.
21. Gonzalez Viejo, C., and S. Fuentes. "Digital Detection of Olive Oil Rancidity ... using Near-Infrared Spectroscopy ... and Machine Learning." Chemosensors, 2022.
22. Yakar, Y., and K. Karadağ. "Identifying Olive Oil Fraud and Adulteration using Machine Learning Algorithms." Quim Nova, 2022.
23. Rashvand, M., and A. Akbarnia. "The Feasibility of Using Image Processing and Artificial Neural Network for Detecting the Adulteration of Sesame Oil." AIMS Agriculture and Food, 2019.
24. Li, J. et al. "Adulteration Detection Model of Tea Oil ... based on FTIR Back-Propagation Neural Network." CTMCD 2021.
25. Antora, S. A. et al. "Detection of Adulteration in Edible Oil using FT-IR Spectroscopy and Machine Learning." International Journal of Biochemistry Research & Review, 2019.
26. Izquierdo, M. et al. "Deep Thermal Imaging to Compute the Adulteration State of Extra Virgin Olive Oil." Computers and Electronics in Agriculture, 2020.
27. Dhakal, S. et al. "Detection of Additives and Chemical Contaminants in Turmeric Powder using FT-IR Spectroscopy." Foods, 2019.
28. Hussain Khan, M. et al. "Hyperspectral Imaging for Color Adulteration Detection in Red Chili." Applied Sciences, 2020.
29. Chaminda-Bandara, W. G. et al. "Validation of Multispectral Imaging for the Detection of Selected Adulterants in Turmeric Samples." Journal of Food Engineering, 2020.
30. Erasmus, S. W. et al. "Rapid Detection of Lead Chromate in Turmeric." Food Control, 2021.

31. Desai, N., and D. Patel. "Identification of Adulteration in Household Chilli Powder from Its Images using Logistic Regression." Reliability Theory and Application, 2021.
32. Neto, H. A. et al. "On the Utilization of Deep and Ensemble Learning to Detect Milk Adulteration." BioData Mining, 2019.
33. Wu, L. et al. "A Deep Learning Model to Recognize Food-Contaminating Beetle Species ..." Computers and Electronics in Agriculture, 2019.
34. Antora, S. A. et al. (as above) FT-IR + ML edible oil detection.
35. Temiz, H. T., and B. Ulaş (as above) HSI in adulteration review.