

Article

Employee Turnover Prediction and Interpretability Analysis Based on SMOTE and SHAP

Ruoqin Liu^{1,*}¹ School of Economics and Management, Southwest Petroleum University, Chengdu, China

* Correspondence: Ruoqin Liu, School of Economics and Management, Southwest Petroleum University, Chengdu, China

Abstract: Employee turnover imposes substantial financial and operational burdens on organizations by increasing recruitment and training expenditures while undermining institutional stability and knowledge continuity. In the era of data-driven decision-making, accurately identifying employees at risk of departure has emerged as a critical challenge in human resource management. This study leverages the IBM HR Analytics Employee Attrition dataset to develop and evaluate predictive models for employee turnover. Guided by prior literature and the imperative of model interpretability, ten key features are systematically selected, encompassing demographic, financial, and organizational variables such as age, monthly income, stock option level, years with the current manager, and overtime status. To mitigate the inherent class imbalance in the dataset, the Synthetic Minority Over-sampling Technique (SMOTE) is employed to generate balanced training samples. Three predictive models—Logistic Regression, Random Forest, and XGBoost—are constructed and rigorously compared. Model outputs are further interpreted using SHapley Additive exPlanations (SHAP) to elucidate feature-level contributions. Empirical results demonstrate that XGBoost achieves the highest overall predictive accuracy at 82.31%, followed by Random Forest at 80.73% and Logistic Regression at 74.38%. Regarding the area under the receiver operating characteristic curve (AUC), Logistic Regression attains 0.770, marginally surpassing XGBoost (0.759) and Random Forest (0.756). Notably, Logistic Regression yields the highest recall of 0.732 in identifying departing employees, substantially outperforming Random Forest (0.296) and XGBoost (0.282). SHAP-based interpretability analysis reveals that overtime significantly elevates turnover risk, whereas higher monthly income and stock option levels exert a protective effect. By constructing a parsimonious yet effective feature system, this study successfully balances predictive performance with interpretability, offering actionable insights for evidence-based human resource management.

Keywords: employee turnover; machine learning; data resampling; model interpretability; ensemble learning; human resource management

Received: 13 May 2026

Revised: 22 June 2026

Accepted: 06 July 2026

Published: 12 July 2026



Copyright: © 2026 by the authors. Submitted for possible open access publication under the terms and conditions of the Creative Commons Attribution (CC BY) license (<https://creativecommons.org/licenses/by/4.0/>).

1. Introduction

In the era of data-driven management, HR analytics has emerged as a crucial tool for managing employee turnover. High turnover rates incur substantial human resource costs and negatively impact organizational performance. Consequently, proactively identifying turnover risks and developing targeted retention strategies have become critical priorities in modern corporate management [1].

The IBM HR Analytics Employee Attrition dataset is a widely used benchmark dataset for turnover prediction. While previous studies have developed models such as logistic regression, random forest, and XGBoost using this dataset, notable limitations persist. Existing literature exhibits two primary deficiencies [2]. First, class imbalance is frequently inadequately addressed, leading to model bias and suboptimal performance. Second, many studies overemphasize overall accuracy metrics while neglecting

evaluation of the underrepresented class and the interpretability of features driving predictive outcomes.

To address these gaps, this study applies the Synthetic Minority Over-sampling Technique (SMOTE) to resolve class imbalance and compares three predictive models: logistic regression, random forest, and XGBoost [3]. Furthermore, the Shapley Additive exPlanations (SHAP) method is incorporated to enhance model interpretability, thereby identifying key determinants of turnover and the direction of their impact.

The contributions of this study are threefold. First, methodologically, the integration of SMOTE and SHAP alleviates class imbalance while simultaneously boosting model interpretability. Second, regarding model comparison, this study systematically evaluates the performance disparities among logistic regression, random forest, and XGBoost under imbalanced conditions, interpreting these findings through a managerial lens [4]. Finally, critical driving factors are identified via SHAP analysis, providing an actionable foundation for enterprises to design targeted employee retention strategies.

2. Literature Review

2.1. Machine Learning in Employee Turnover Prediction

Employee turnover prediction is an important research area within human resource management. Traditional approaches rely on survey questionnaires and statistical regression; however, these methods are often constrained by inherent subjectivity and limited sample sizes. In recent years, advancements in big data and artificial intelligence have facilitated the integration of machine learning techniques into turnover prediction. These techniques demonstrate the capability to automatically capture complex interactions among variables within high-dimensional datasets [5].

Regarding algorithm selection, logistic regression is frequently employed as a baseline model due to its robust interpretability. Random forest, an ensemble of multiple decision trees, is recognized for managing nonlinear relationships and feature interactions, thereby exhibiting high robustness [5]. Meanwhile, XGBoost—an optimized gradient boosting implementation—typically yields superior predictive accuracy on structured data. Empirical evidence further supports these applications; for instance, a study predicting the turnover of nurses in South Korea demonstrated that a random forest model achieved 97% accuracy, identifying compensation as the most critical determinant of attrition. Similarly, using the IBM HR Analytics dataset, a comparative analysis of nine machine learning algorithms found that logistic regression and random forest yielded the highest accuracies, at 87.7% and 85.7%, respectively.

2.2. Class Imbalance and Model Interpretability

Employee turnover datasets are notoriously characterized by class imbalance [2]. If left unaddressed, predictive models tend to favor the majority class, thereby compromising the detection of employees who leave. To address this, the Synthetic Oversampling Technique (SMOTE) has become a widely adopted method; it balances datasets by generating synthetic samples for underrepresented classes. Empirical evidence supports its effectiveness: both random forest and learning machine algorithms have demonstrated improved predictive performance on the IBM HR dataset after applying SMOTE. Similarly, validation studies in nurse turnover forecasting confirm that enhanced random forest models exhibit the strongest predictive capability when SMOTE is employed.

Regarding model transparency, the Shapley Additive exPlanations (SHAP) framework—grounded in Shapley values—has emerged as the industry standard for model interpretation. It quantifies the marginal contribution of individual features to predictive outcomes [6]. Recent research has leveraged this capability to bridge the gap between complex algorithms and practical decision-making. For instance, a weighted LightGBM-XGBoost ensemble integrated with SHAP analysis identified compensation, position, and tenure as critical determinants of turnover, achieving an AUC of 0.9504.

Furthermore, bibliometric analyses indicate that explainable artificial intelligence (XAI) currently represents the most prominent frontier in turnover prediction research.

2.3. Research Gaps and the Present Study

A synthesis of the existing literature indicates that while employee turnover prediction has achieved significant progress, notable limitations persist. Most studies prioritize model accuracy comparisons while providing limited analysis of interpretability. Moreover, research using the IBM HR dataset that simultaneously integrates SMOTE, multi-model comparative analysis, and SHAP-based interpretation remains scarce [7].

To address these gaps, this study applies SMOTE to mitigate class imbalance and develops logistic regression, random forest, and XGBoost models for comparative evaluation [8]. SHAP is employed to enhance model transparency. The objective is to identify key determinants of employee turnover and clarify their directional effects, thereby supporting evidence-informed HR decision-making.

3. Research Methodology

3.1. Dataset and Features

This study employs the IBM HR Analytics Employee Attrition dataset, which contains 1,470 employee records [9]. Employees who left the organization constitute approximately 16.1% of the total sample, indicating a notable class imbalance. The dataset includes diverse attributes—such as demographic characteristics, compensation levels, work history, satisfaction metrics, and organizational behavior indicators—making it well-suited for predictive modeling of employee turnover.

Variables with no theoretical or empirical link to turnover—namely employee number, employee count, over-18 indicators, and standard hours—are excluded to streamline the feature space. Based on established human resource management principles and the need for operational clarity, ten variables are selected as predictors: Age, MonthlyIncome, StockOptionLevel, YearsWithCurrManager, NumCompaniesWorked, RelationshipSatisfaction, EnvironmentSatisfaction, JobSatisfaction, WorkLifeBalance, and OverTime [10].

Feature contribution analysis is conducted using the XGBoost model (parameter settings specified in Section 3.3) combined with SHAP for interpretability. Figure 1 presents the ranking of feature importance based on mean absolute SHAP values. OverTime, YearsWithCurrManager, StockOptionLevel, Age, and MonthlyIncome emerge as the top predictors, exhibiting strong discriminative power. Although satisfaction-related features show comparatively lower importance scores, they are retained due to their conceptual relevance and utility in guiding practical organizational actions [11]. The target variable Attrition is binary-coded: 1 denotes employees who departed, and 0 denotes those retained.

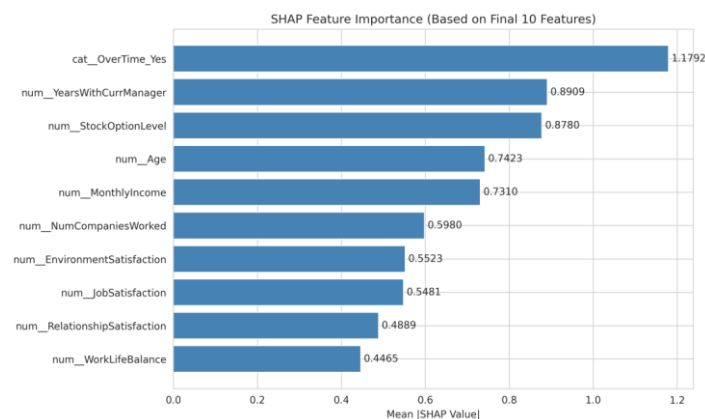


Figure 1. SHAP Feature Importance Ranking (Mean Absolute SHAP Value)

3.2. Data Preprocessing and Smote

Data preprocessing is performed to enhance model training effectiveness [12]. The target variable is binary-encoded (attrition = 1, non-attrition = 0); the categorical variable OverTime is transformed using one-hot encoding; and Z-score standardization is applied to the nine numerical features to eliminate scale disparities. Stratified sampling is then used to partition the dataset into training and testing sets at a 7:3 ratio, preserving the proportional representation of attrition cases in both subsets.

In real-world organizational turnover prediction, instances of employee attrition are typically underrepresented. To improve classifier performance for this underrepresented class, synthetic minority oversampling technique (SMOTE) is applied exclusively to the training set, generating synthetic samples to balance the class distribution. The testing set retains its original composition to ensure unbiased model evaluation [13].

3.3. Models and Evaluation Metrics

Three machine learning models are constructed for this analysis [5]. Logistic regression serves as the linear classification baseline, with maximum iterations set to 5,000. The random forest model consists of 200 decision trees with a maximum depth of 8. The XGBoost model comprises 300 trees with a maximum depth of 5, a learning rate of 0.1, a subsample ratio of 0.8, and a column subsample ratio of 0.8.

To enhance model reproducibility and practical deployment feasibility, these hyperparameters are selected using heuristic guidelines and conservative configurations, avoiding extensive grid search tuning. This approach supports model generalizability and mitigates overfitting to the training sample distribution, thereby improving reliability when applied to new organizational data. Model performance is evaluated using five metrics: accuracy, precision, recall, F1-score, and the Area Under the ROC Curve (AUC).

3.4. SHAP Method for Interpretability

The SHAP TreeExplainer is applied to interpret the trained XGBoost model. Rooted in Shapley values from cooperative game theory, this method quantifies each feature's marginal contribution to prediction outcomes [14]. SHAP summary plots visualize both the relative importance ranking of features and their directional influence—positive or negative—on turnover probability. Analyses are conducted using the final set of 10 retained features to identify key drivers of employee attrition.

4. Results and Analysis

4.1. Turnover Distribution and Overtime Analysis

As illustrated in Fig [9]. 2, the original dataset comprises 1,233 retained employees (83.9%) and 237 turnover employees (16.1%), indicating a pronounced class imbalance. The synthetic oversampling technique is applied exclusively to the training set, effectively balancing the proportion of positive and negative samples. Fig. 3 demonstrates that the turnover rate among employees with OverTime = Yes is significantly higher than that of their non-overtime counterparts. Specifically, 127 overtime employees departed (approximately 30.5% of the total overtime workforce), whereas only 110 non-overtime employees departed (roughly 10.4% of the non-overtime workforce). These findings suggest a significant association between overtime behavior and turnover propensity, implying that imbalances in work and personal life management may represent a critical factor driving employee attrition (As shown in Figure 2).

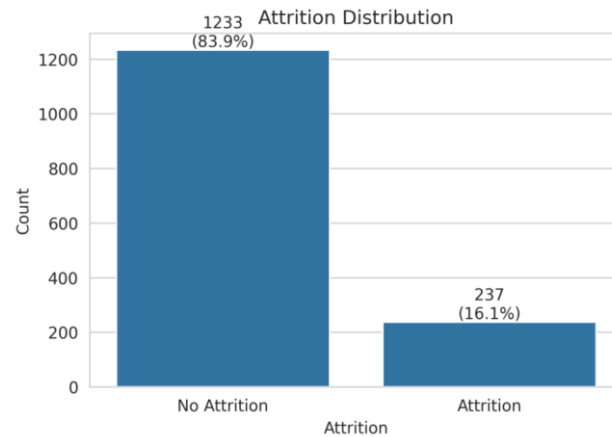


Figure 2. Employee Attrition Distribution

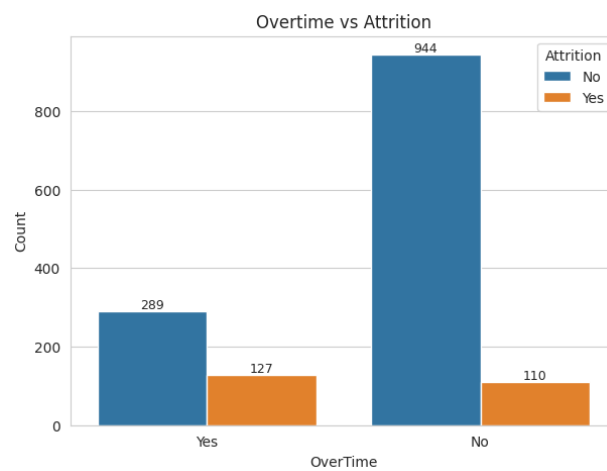


Figure 3. Cross-Analysis of Overtime and Attrition

4.2. Overall Model Performance Comparison

Table 1 presents the performance metrics of the three predictive models on the testing set, which comprises 441 samples, including 71 employees who subsequently left the organization.

Table 1. Overall Model Performance Comparison

Model	Accuracy	Precision	Recall	F1-Score	AUC
Logistic Regression	0.7438	0.3562	0.7324	0.4793	0.7701
Random Forest	0.8073	0.3750	0.2958	0.3307	0.7564
XGBoost	0.8231	0.4255	0.2817	0.3390	0.7586

Regarding classification performance, XGBoost achieves the highest accuracy (82.31%), followed by random forest (80.73%) and logistic regression (74.38%). In terms of AUC, logistic regression performs slightly better (0.770) than XGBoost (0.759) and random forest (0.756), with all three models substantially outperforming random classification [15].

For identifying employees who subsequently left the organization, logistic regression yields the highest recall (0.7324), significantly outperforming random forest (0.2958) and XGBoost (0.2817). This indicates that logistic regression is more sensitive to this outcome class; however, its relatively low precision suggests a higher rate of false-positive predictions. Conversely, while ensemble models prioritize overall classification accuracy, they exhibit limitations in detecting this outcome class.

The suitability of these models depends on distinct managerial objectives. If the primary goal is overall predictive accuracy, XGBoost offers a clear advantage. Conversely,

if the objective is to screen employees with potential turnover risks, logistic regression holds greater practical value. In operational settings, classification thresholds may be adjusted according to organizational risk appetite to achieve an optimal trade-off between recall and precision.

4.3. Confusion Matrix and ROC Curves

Fig. 4 presents the confusion matrix for the XGBoost model [9]. On the testing set, XGBoost correctly predicts 343 retained employees, accounting for 92.7% of the actual retained cohort. Conversely, 27 retained employees are misclassified as departed. Regarding turnover employees, the model correctly identifies 20 individuals, representing 28.2% of actual turnover cases, while yielding 51 false negatives. Although the model demonstrates robust identification capabilities for the majority class, performance in detecting turnover cases requires further optimization.

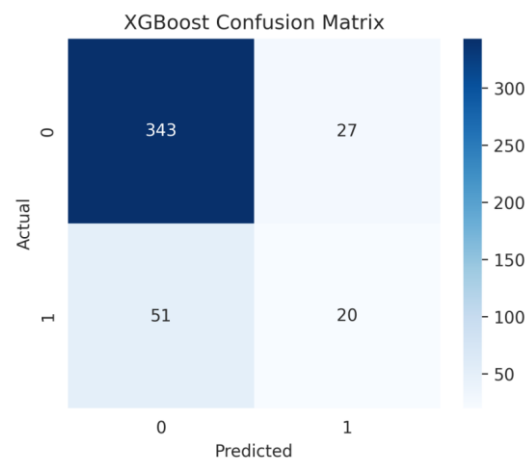


Figure 4. XGBoost Confusion Matrix

The Receiver Operating Characteristic (ROC) curves in Fig. 5 illustrate that the AUC values are 0.770 for logistic regression, 0.759 for XGBoost, and 0.756 for random forest [3]. All three curves lie clearly above the diagonal reference line, indicating that the evaluated models possess sound classification capabilities.

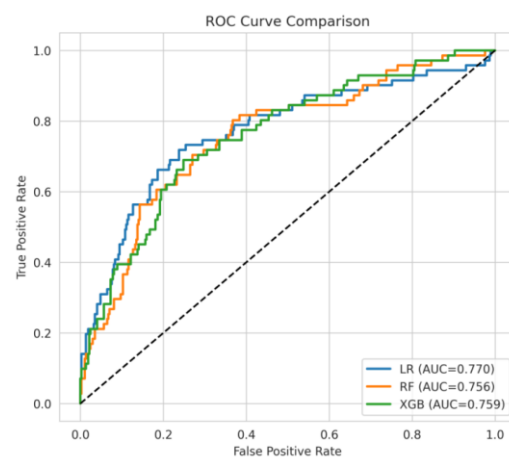


Figure 5. ROC Curve Comparison

4.4. Feature Correlation Analysis

To explore the linear relationships between individual features and employee turnover, Pearson correlation coefficients are calculated for the selected features and the target variable, Attrition (Fig. 6). The variable OverTime (encoded as 1 = Yes) exhibits a positive correlation with Attrition ($r = 0.25$), indicating that overtime work is associated

with a higher propensity to leave. Conversely, MonthlyIncome ($r = -0.16$), Age ($r = -0.16$), and YearsWithCurrManager ($r = -0.16$) demonstrate negative correlations with attrition; higher income, older age, and longer supervisory tenure correspond to a reduced probability of turnover. EnvironmentSatisfaction ($r = -0.10$) and JobSatisfaction ($r = -0.10$) also exhibit weak negative associations. In contrast, WorkLifeBalance ($r = -0.06$) and NumCompaniesWorked ($r = 0.04$) display negligible correlations with attrition [14].

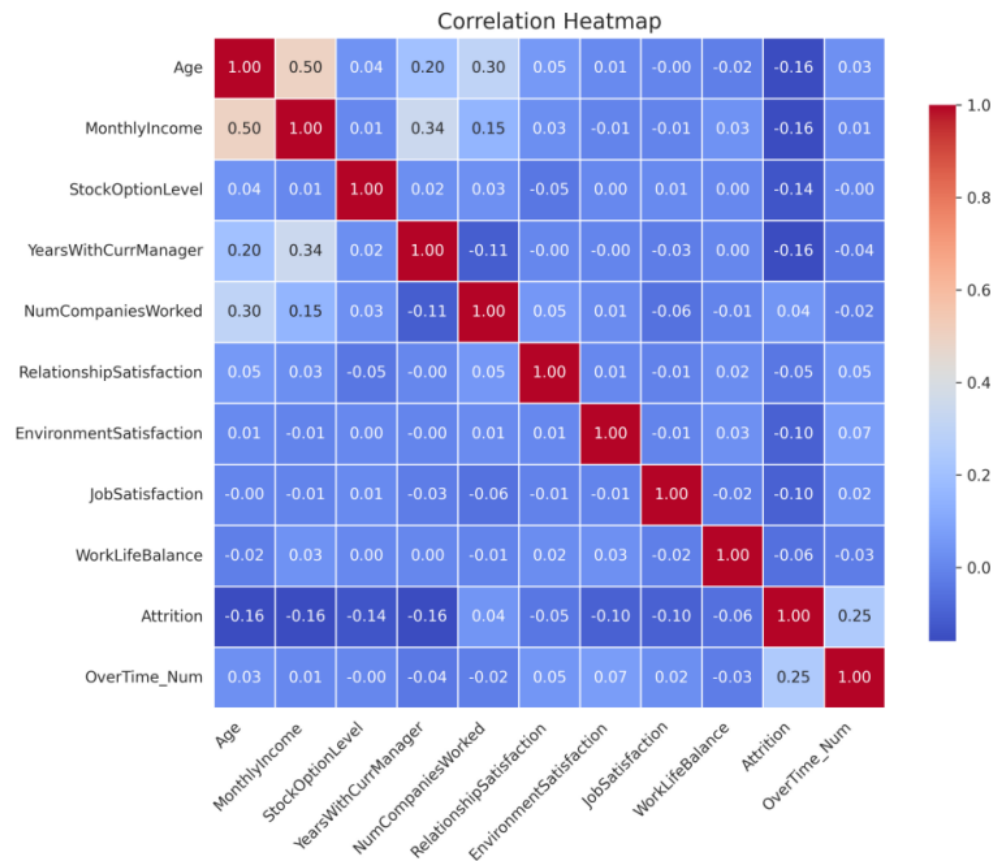


Figure 6. Feature Correlation Heatmap

These directional relationships are consistent with findings reported in prior empirical studies, supporting the relevance of the selected features for turnover prediction [12]. As correlation coefficients reflect only linear associations, the comprehensive impact of these features on turnover propensity is further explored via SHAP analysis in the subsequent section.

4.5. SHAP Interpretability Analysis

To reveal the directional impacts of individual features on turnover probability, the SHAP method is applied to interpret the XGBoost model. Figure 7 presents the SHAP summary plot. The horizontal axis represents the SHAP value, where positive values indicate an increased probability of turnover, while negative values denote a decreased probability. Feature value magnitudes are represented by color intensity, with red and blue signifying high and low values, respectively.

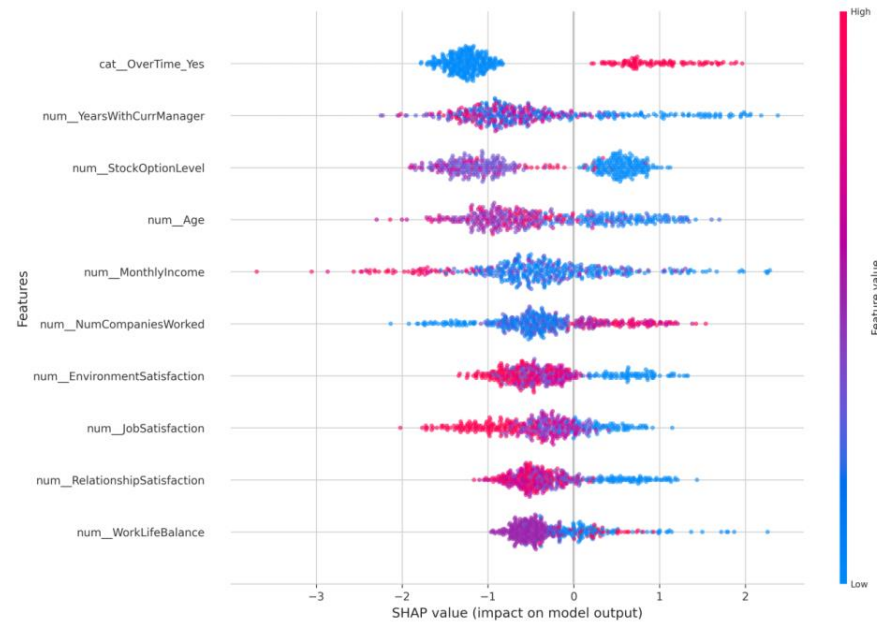


Figure 7. SHAP Summary Plot

OverTime: Samples with overtime status are concentrated on the positive side of the SHAP axis with a broad distribution, identifying overtime as the most influential positive risk factor for turnover [13, 15].

YearsWithCurrManager: Longer tenure corresponds to negative SHAP values and a reduced probability of turnover, suggesting that stable superior-subordinate relationships strengthen retention intentions [5].

StockOptionLevel: High stock option levels are concentrated on the negative side, whereas low levels appear on the positive side, indicating that equity-based incentives effectively mitigate turnover risk [3].

MonthlyIncome: Higher monthly income is predominantly located on the negative side, confirming an inverse relationship between compensation level and turnover propensity.

Age: Older age corresponds more distinctly to negative SHAP values, indicating greater inclination toward organizational stability among senior employees.

NumCompaniesWorked: The SHAP values for this feature exhibit a broad, interspersed range of positive and negative impacts, implying an inconsistent directional effect likely dependent on individual circumstances [4].

Satisfaction-related features: High levels of EnvironmentSatisfaction, JobSatisfaction, RelationshipSatisfaction, and WorkLifeBalance skew toward the negative side, whereas low levels skew toward the positive. This indicates that enhancing employee satisfaction contributes to reducing turnover, although the magnitude of this effect is notably weaker compared to compensation and overtime factors.

5. Discussion

5.1. Primary Findings on Model Performance

The experimental comparison of logistic regression, random forest, and XGBoost reveals distinct performance trade-offs across different classification models. Consistent with expectations, ensemble methods (XGBoost and random forest) demonstrate superior overall predictive accuracy, whereas the linear baseline (logistic regression) exhibits enhanced sensitivity toward the underrepresented class, as evidenced by its significantly higher recall [2].

These findings suggest that the choice of model should be contingent upon the specific organizational objective: XGBoost is preferable for high-precision accuracy requirements, while logistic regression offers practical advantages for high-recall

screening of employees at elevated risk [3, 13]. This disparity highlights the fundamental difference between ensemble methods, which aim to capture complex nonlinear relationships to improve overall accuracy, and linear classifiers, which provide a more balanced trade-off in the presence of severe class imbalance.

5.2. Comparison with Existing Studies

Predictive accuracy observed in this research is somewhat lower than that reported in certain prior studies using the same dataset. This discrepancy arises primarily from three methodological distinctions.

First, the feature set differs in scale and composition [11]. The present study employs a streamlined set of 10 features—yielding approximately 12 input variables after encoding—whereas prior work often adopts higher-dimensional feature sets to maximize predictive precision. This streamlining prioritizes business interpretability and reduced deployment costs over marginal improvements in accuracy.

Second, hyperparameter tuning follows a conservative strategy to enhance model reproducibility and practical deployment feasibility, rather than pursuing peak predictive performance [4]. This approach reflects the baseline behavior of lightweight models under near-default configurations, supporting rapid deployment scenarios.

Third, model evaluation adheres to stricter protocol: stratified sampling partitions the data into training and testing subsets, and synthetic oversampling is applied exclusively to the training subset. Maintaining the original class distribution in the testing subset prevents data leakage, ensuring evaluation metrics better reflect real-world operational conditions [15].

5.3. Discussion of Key Influencing Factors

Empirical findings are interpretable through established organizational behavior and human resource management frameworks. SHAP analysis identifies *OverTime*, *YearsWithCurrManager*, *StockOptionLevel*, *MonthlyIncome*, and *Age* as the five most influential predictors.

Overtime emerges as the most potent positive risk factor, with its SHAP contribution substantially exceeding that of other features. Prolonged overtime depletes psychological and physiological resources, inducing emotional exhaustion and intensifying turnover intentions. This finding aligns with both work-family interface theory and the conservation of resources theory.

Tenure with the current supervisor, stock option levels, monthly income, and age exert significant negative impacts on attrition. Stable superior-subordinate relationships facilitate the accumulation of social capital and organizational trust, thereby reducing perceived uncertainty. Concurrently, high stock option levels and monthly income serve as economic incentives, reinforcing retention intentions through perceived organizational rewards. Furthermore, the negative impact of age reflects the higher organizational embeddedness and increased switching costs typically associated with more senior employees [15].

Finally, while satisfaction-related features—including environment satisfaction, job satisfaction, relationship satisfaction, and work-life balance—exhibit relatively weaker contributions to turnover, their directional impacts align with theoretical expectations [11]. These features remain valuable as reference variables for supplementary managerial support measures.

5.4. Managerial Implications

Explainable AI (XAI) plays an important role in managerial decision support by translating predictive results into actionable human resource insights.

First, enterprises should optimize overtime management through streamlined workflows and rational scheduling to mitigate turnover propensity associated with chronic resource depletion [11].

Second, organizations should enhance salary competitiveness and refine long-term incentive structures, such as stock options, to strengthen employees' perceptions of organizational rewards and reinforce retention intentions.

Third, organizations must prioritize the stability of superior-subordinate relationships by intensifying communication and leadership training for frontline managers to minimize turnover risks arising from managerial instability [12].

Finally, although satisfaction-related variables exert relatively weaker influence, sustained improvement in employee satisfaction remains essential; enhancing the work environment, optimizing promotion mechanisms, and providing career development support are key strategies for long-term turnover reduction.

5.5. Limitations and Future Research

This research is subject to several limitations that suggest avenues for future inquiry [3]. First, the dataset scale is relatively limited and drawn from a single source, which may constrain the generalizability of the findings. Second, the absence of time-dynamic features precludes survival analysis or modeling of turnover evolution over time. Third, while synthetic oversampling effectively mitigates class imbalance, generated samples may not fully reflect the underlying data distribution. Future research could explore advanced data augmentation methods, such as Generative Adversarial Networks (GANs) or CTGAN, to improve representational fidelity. Finally, the current feature set may omit important determinants, including psychological states and perceived organizational fairness.

Future studies should aim to integrate larger-scale, cross-industry datasets while incorporating deep learning architectures or causal inference models. Such advancements would further enhance both predictive accuracy and explanatory depth in turnover research [12].

6. Conclusion

This study evaluated logistic regression, random forest, and XGBoost for predicting employee turnover using the IBM HR Analytics dataset. To mitigate class imbalance, SMOTE was applied; SHAP was used to enhance model interpretability. Empirical results show that XGBoost attains the highest overall accuracy (82.31%), whereas logistic regression yields superior recall for departing employees (0.732), making it especially valuable for high-risk screening. SHAP analysis identifies overtime as the strongest positive predictor of turnover, while higher income, greater stock options, longer tenure with the current supervisor (YearsWithCurrManager), and older age emerge as key protective factors.

A streamlined set of ten features achieves an optimal trade-off between predictive performance and interpretability. These findings support the development of low-cost, practical early-warning systems for employee turnover. Future work could extend this approach by integrating longitudinal dynamic data and causal inference methodologies.

References

1. İ. T. Baydili and B. Tasci, "Predicting employee attrition: XAI-powered models for managerial decision-making," *Systems*, vol. 13, no. 7, p. 583, 2025.
2. P. Ajit, "Prediction of employee turnover in organizations using machine learning algorithms," *Algorithms*, vol. 4, no. 5, p. C5, 2016.
3. D. Avrahami, D. Pessach, G. Singer, and H. Chalutz Ben-Gal, "A human resources analytics and machine-learning examination of turnover: implications for theory and practice," *Int. J. Manpower*, vol. 43, no. 6, pp. 1405–1424, 2022.
4. L. Breiman, "Random forests," *Mach. Learn.*, vol. 45, no. 1, pp. 5–32, 2001.
5. N. V. Chawla, K. W. Bowyer, L. O. Hall, and W. P. Kegelmeyer, "SMOTE: synthetic minority over-sampling technique," *J. Artif. Intell. Res.*, vol. 16, pp. 321–357, 2002.
6. C. H. Chang, H. W. Lin, W. H. Tsai, W. L. Wang, and C. T. Huang, "Employee satisfaction, corporate social responsibility and financial performance," *Sustainability*, vol. 13, no. 18, p. 9996, 2021.
7. R. C. Chen et al., "An end to end of scalable tree boosting system," *Sylwan*, vol. 165, no. 1, pp. 1–11, 2020.

8. S. P. Das and J. Samal, "People analytics for employee turnover prediction using Extreme Learning Machine," *J. Bus. Analytics*, vol. 9, no. 2, pp. 108–123, 2026.
9. Y. Fang and Z. Zhang, "Employee Turnover Prediction Model Based on Feature Selection and Imbalanced Data Handling," *IEEE Access*, 2025.
10. A. Bujold, I. Roberge-Maltais, X. Parent-Rochelleau, J. Boasen, S. Sénécal, and P. M. Léger, "Responsible artificial intelligence in human resources management: a review of the empirical literature," *AI and Ethics*, vol. 4, no. 4, pp. 1185–1200, 2024.
11. L. Kohoutová et al., "Toward a unified framework for interpreting machine-learning models in neuroimaging," *Nat. Protoc.*, vol. 15, no. 4, pp. 1399–1435, 2020.
12. D. Yin, X. Lu, and T. Mei, "Explainable deep learning for healthcare workforce attrition: a methodological study on the Watson healthcare synthetic benchmark," *Front. Public Health*, vol. 14, p. 1861450, 2026.
13. A. Dimri, P. Kumar, and D. Garg, "Exploring job embeddedness in the era of artificial intelligence (AI) and machine learning: a bibliometric analysis and future trends," *Benchmarking: An Int. J.*, vol. 33, no. 7, pp. 1991–2018, 2026.
14. S. Chowdhury, S. Joel-Edgar, P. K. Dey, S. Bhattacharya, and A. Kharlamov, "Embedding transparency in artificial intelligence machine learning models: managerial implications on predicting and explaining employee turnover," *Int. J. Hum. Resour. Manag.*, vol. 34, no. 14, pp. 2732–2764, 2023.
15. W. Wu and S. Fukui, "Using human resources data to predict turnover of community mental health employees: Prediction and interpretation of machine learning methods," *Int. J. Ment. Health Nurs.*, vol. 33, no. 6, pp. 2180–2192, 2024.

Disclaimer/Publisher's Note: The statements, opinions and data contained in all publications are solely those of the individual author(s) and contributor(s) and not of Publisher and/or the editor(s). Publisher and/or the editor(s) disclaim responsibility for any injury to people or property resulting from any ideas, methods, instructions or products referred to in the content.