



International Journal of Artificial Intelligence and Machine Learning

Publisher's Home Page: <https://www.svedbergopen.com/>



Research Paper

Open Access

Fairness-Aware Learning Framework For Bias Mitigation In Algorithmic Decision Systems

Uma Maheswari G¹, Liu Taiyu², Manisha Sagar Pawar³, Keerthika K⁴, Abhinav Rathour⁵, Dr. S. T. Santhanalakshmi⁶, Ashish Kumar Sharma⁷, Ubaydulloyeva Sabrina Kodirovna⁸

¹ Assistant Professor, Department of Mathematics, Meenakshi College of Arts and Science, Meenakshi Academy of Higher Education and Research, India. Email: umamahes@maher.ac.in

² Research Scholar, School of Education, Lincoln University College, Malaysia. Email: taiyu.phdscholar@lincoln.edu.my

³ Department of Engineering, Science & Humanities, Vishwakarma Institute of Technology, Pune, Maharashtra 411037, India. Email: manisha.pawar@vit.edu

⁴ Assistant Professor, Department of Computer Science, Meenakshi College of Arts and Science, Meenakshi Academy of Higher Education and Research, India. Email: keerthikak@maher.ac.in

⁵ Centre for Research Impact & Outcome, Chitkara University Institute of Engineering and Technology, Chitkara University, Rajpura, Punjab 140401, India. Email: abhinav.rathour.orp@chitkara.edu.in ORCID: 0009-0008-4434-1073

⁶ Assistant Professor Grade I, Department of Computer Science and Engineering, Panimalar Engineering College, Chennai, Tamil Nadu, India. Email: santhanalakshmi.pecse2024@gmail.com ORCID: 0000-0002-1377-1408

⁷ Assistant Professor, Department of Computer Application, Dev Bhoomi Uttarakhand University, Dehradun, Uttarakhand, India. Email: socse.asishkumar@dbuu.ac.in ORCID: 0009-0008-1487-8662

⁸ Assistant, Department of Social and Humanitarian Sciences, Faculty of Pedagogy and Psychology, Samarkand State Medical University, Uzbekistan. Email: ubajdullaevsabrina4@gmail.com ORCID: 0009-0002-2873-1860

Abstract

The increased reliance on Machine Learning (ML) based algorithmic decisions has given rise to some critical issues such as fairness and bias because the machine learning models, if trained on biased data, tend to make discriminatory predictions. Therefore, to counter these problems, this study puts forth an innovative Fairness-Aware Grouping Learning Framework (FAGLF). To improve the quality of the datasets, missing value imputation is utilized to replace incomplete data points, while the Synthetic Minority Over-sampling Technique (SMOTE) addresses the problem of class imbalance. Furthermore, Independent Component Analysis (ICA) based feature extraction is employed to enhance data quality and enable more equitable predictive learning. The approach embeds data instances into a feature space and organizes them into homogeneous subgroups based on similarity patterns. For each subgroup, a specialized predictive model is trained, allowing more precise learning while maintaining overall consistency. The framework is formulated as a three-level optimization problem, solved using Dynamic Transient Search Optimized Extreme Gradient Boosting (DTSO-XGBoost), which balances exploration and exploitation during optimization. For better performance on the smaller and class-imbalanced datasets, techniques for domain adaptation have been used. Fairness is ensured by keeping equal distribution among sensitive variables such as gender and socio-economic status. The developed model is tested on the dataset "Fairness Decision Records" using fairness metrics implemented in Python. It achieves a precision of 0.972, a recall of 0.948, and an F1 score of 0.960. Results demonstrate significant improvement in fairness-accuracy trade-offs while maintaining strong predictive performance, highlighting the effectiveness of the proposed framework.

Keywords: Fairness-aware, Machine Learning, Algorithmic decision systems, Group-based learning, Discriminatory bias detection, Domain adaptation

1. Introduction

Big data and ML at academic and institutional level, have become an important component for educational systems supporting decision-making, individualized learning, performance assessment, and administration [1]. The use of enhanced ML algorithms within AI and proper methodologies within software engineering practices have lead to a ubiquitous application of ML in several industries assisting with decision-making processes [2]. AI is changing a variety of businesses with its classification and forecasting capabilities from given data. For instance, banks leverage AI to improve their credit scoring decision-making process [3]. Digitally accessible collective data have enabled algorithmic decision-making but they may raise a concern on fairness, transparency, human interference, and social damage [4]. Edge-Based Social Network (EBSN) recommendation system, suffers from failure to balance the users' interest since individual preferences compete, cold-start issue further aggravates suggestion dependability, user engagement, and it causes a drop in the accuracy of the recommendations [5]. Some specific challenges involve lacking methods to realize fairness for sequential recommendation algorithms, address interaction fairness issue, filter bubble and fairness issue across user attribute groups [6]. FL approach to realize fairness are implemented through optimizing global objectives and by sharing client data bias thus compromising for poor scalability and threatening users' privacy and the cause for lack of efficiency is due to non-IID data distribution [7]. The traditional methods are credit scoring processes using logistic regression, scorecards, and rule-based models constructed on historical financial data. Such methods are usually opaque, have a historical bias and do not hold a sense of fairness [8].

Research Aim: The objective of the research is to create the FAGLF framework. It can be useful in combating the problems of discrimination caused by the decision-making algorithms in the process. The proposed methodology leverages the benefits of dynamic subgroup organizing, representation learning, and proper DTSO-XGBoost model training techniques. Besides, it ensures equitable results and treatment for all genders and economic backgrounds without any compromise on accuracy in prediction.

Research Organization: Background and rationale for bias mitigation in algorithmic decision-making systems are provided in Section 1. Literature survey related to techniques of group-based learning, approaches of bias mitigation, and fairness-aware ML is carried out in Section 2. The proposed FAGLF is introduced in Section 3 concerning optimization, subgroup learning, and dynamic grouping. Section 4 discusses results and experimental findings. Conclusions of this research are discussed in Section 5.

2. Related Works

Earlier research that looks at the fairness of AI algorithms within finance and decision-making processes has been compared in Table 1. In this case, factors like the objective of the research, use of approaches including XAI, the measurements used for fairness, and ways of reducing bias, and others, have been considered. However, some limitations remain even with many advances made.

Table 1: Comparison of Fairness-Aware AI Methods and Limitations

Ref	Objective	Methods	Result	Limitation
[9]	Transparent and fair AI-driven trading decisions	Explainable AI (XAI) techniques with fairness-aware ML	Improved interpretability, compliance, and decision reliability	High complexity and limited real-time scalability
[10]	Ensure fairness and ethics in financial AI	Audits, inclusive data, regulations, and ethical design	Improved transparency, fairness, and trust in AI	Complex implementation and evolving regulatory challenges
[11]	Develop a fair and transparent AI banking framework	Bias-aware data, fairness models, XAI techniques	Improved fairness with minimal accuracy loss	Complex integration and limited real-world scalability
[12]	Ensure ethical, fair, and transparent financial AI	XAI method, bias analysis, and regulatory frameworks	Improved interpretability, fairness, and regulatory compliance	Complex models and evolving global regulatory challenges
[13]	Analyze bias in AI-based credit risk models	Fairness metrics, debiasing model, interdisciplinary approach	Improved fairness awareness and ethical decision-making	Data bias persistence and regulatory implementation challenges

[14]	Evaluate bias across multiple protected classes	Fairness metrics, threshold comparison, quantitative evaluation heuristic	Reduced bias versus traditional Fair Isaac Corporation (FICO) scoring	Limited generalization beyond specific case study data
[15]	Improve fairness in neural network image classification	Fairness-aware loss prioritizing the worst-performing sensitive groups	Enhanced fairness with minimal performance reduction	Slight decrease in overall model accuracy
[16]	Assess fairness in credit scoring systems	ML fairness method and bias mitigation approaches	Reduced bias with minimal accuracy trade-offs	Performance loss, complexity, and limited real-world generalization

Research gap: Previous research, including Fair Banking Framework [11] and Credit Scoring Fairness Model [16], exhibits improvements in terms of fairness through bias-aware algorithms; but they suffer from problems like difficulty in integration, poor performance, and lack of scalability in reality. Also, current approaches employ global models, which cannot identify the discrepancy of subgroups. Therefore, the introduced FAGLF utilizes dynamic grouping along with DTSO-XGBoost optimization to address the above-mentioned problems.

3. Proposed Fairness-Aware Grouping Learning Framework (FAGLF)

The procedures involved in creating the FAGLF framework are depicted in Figure 1. Data about the fairness decision-making process is collected in the first step, which is followed by SMOTE balance and missing value imputation. To make the procedure more scalable and compatible with parallel computing, dynamic problem decomposition is incorporated. Predictive accuracy is improved by hyperparameter adjustment with DTSO-XGBoost optimization. When it comes to addressing the problems caused by group heterogeneity in decision-making, group-based predictive learning is crucial.

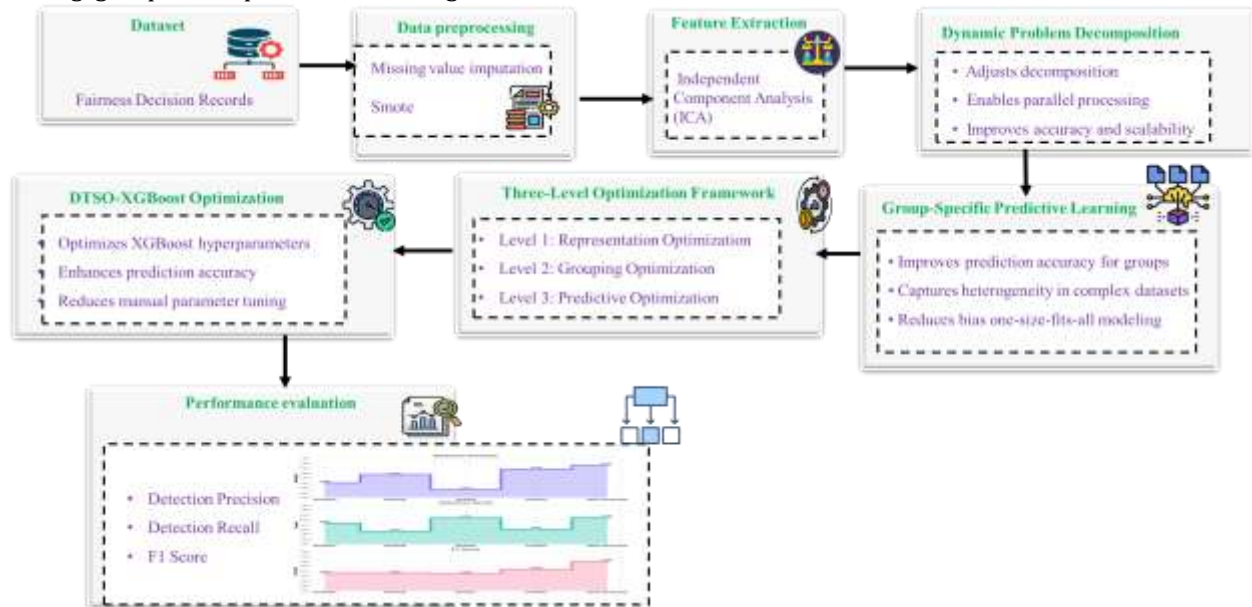


Figure 1: Intelligent Bias Reduction Framework for Predictive Learning

3.1 Data collection

The Fairness Decision Records dataset would be a valuable resource to analyze bias and fairness in the algorithmic decision system since it provides fairly predictable results. This dataset consists of 11,000 observations having properties related to demographics, financial, behavioral, socioeconomic and fairness, income dynamics, credit behavior, debt level, geographical factors, and risk parameters conditional upon the considered populations are amongst the key variables explored with the dataset while including groups such as socio-economic status and gender. The dataset's goal variable is Decision Outcome, which has three classes: Approved, Manual Review, and Rejected. The dataset is divided into 25% for testing and 75% for training.

Kaggle Source: <https://www.kaggle.com/datasets/colabsss/fairness-decision-records>

3.2 Data Preprocessing

The preprocessing of data through taking care of missing values, normalization of data, and balancing of the data ensures that the data for the decision of the fairness is dependable, and accuracy enables fair decision-making.

Missing value imputation: Data imputation techniques are used in data preprocessing to improve the accuracy and reliability of the datasets for fairness-conscious learning. Missing values in the data are estimated through different methods that include statistical and model-based approaches. This process is aimed at increasing data quality and minimizing the influence of bias on the data analysis.

$$CSS = \frac{(z_j \cdot \hat{z}_j)}{\|z_j\| * \|\hat{z}_j\|} \quad (1)$$

In Equation (1), Cosine Similarity Score (CSS) evaluates the degree to which the two feature vectors resemble each other in terms of their orientation. z_j original feature vector of the j th data record. This is the true vector of the dataset before any process of filling in the missing information. $\hat{z}_j$ imputed feature vector of the same data record. $z_j, \hat{z}_j$ inner product of the two vectors. Indicates the closeness between the two in terms of the directions they take. $\|z_j\|$ L2 norm of the original feature vector. $\|\hat{z}_j\|$ L2 norm of the imputed vector. $\|z_j\| * \|\hat{z}_j\|$ normalization factor that ensures the similarity score is independent of vector magnitude and depends only on direction.

SMOTE: Generating artificial data through SMOTE in the minority group, rather than repeating existing ones, the SMOTE technique is a data generation strategy and can be employed to handle class imbalance problems. This ensures that the model is generalized, the balance of data is improved, and there is minimal bias toward the majority group. SMOTE has been applied in the present research to ensure fairness in learning despite class disparity.

$$o_{ji} = w_j + rand(0,1) \times (w_{ji} - w_j) \quad (2)$$

In this Equation (2), the optimized feature o_{ji} improves representation of data quality and prediction fairness using imputed and observed dataset values. Here, i denotes the feature index in the dataset, and j indicates the sample index. w_j represents the current feature weight, while w_{ji} reflects alternative feature values generated through imputation or neighboring data patterns. *rand* refers to the randomization factor used to enhance exploration of feature space.

3.3 Feature extraction: Independent Component Analysis (ICA)

ICA is a feature extraction method which utilizes statistical independence and non-Gaussian nature of the multivariate data to identify statistically independent components. ICA aids in revealing complex relationships between variables, which cannot be revealed using only correlation analysis. In the context of this study, ICA will be used to remove noise and redundancies in order to improve feature quality.

$$w = At \quad (3)$$

Equation (3), w is the output vector that is the set of independent components, making it the noise-free version of the original data. A is the transformation matrix that is learnt to transform the observed vector into independent components. t is the observed vector that includes different financial features in the form of mixed and correlated financial data.

$$v = Xw = XAt \quad (4)$$

In Equation (4), v stands for the predictive signal generated after transforming the output signals of ICA decomposition through the weight matrix of the prediction model. X represents the projection matrix of the financial predictive model, such as a classifier.

$$e_v(v) = \prod_j e_{v_j}(v_j) \quad (5)$$

In Equation (5), $e_v(v)$ refers to the joint probability of the vector v . In the expression, v_j refers to the j th independent component of vector v , while $e_{v_j}(v_j)$ refers to the marginal probability of that particular independent component. The joint distribution of v is equal to the product of its marginal distributions represented by $\prod_j$.

3.4 Dynamic Problem Decomposition (DPD)

DPD is an adaptive method that decomposes a large and complicated problem in optimization or ML into many small, dynamically changing problems according to distribution or similarity. This method is used in this study to mitigate data heterogeneity, decrease bias, and promote fairness among different classes of sensitive attributes in algorithmic decision-making systems.

$$S_{ij} = \|Z_i - Z_j\|^2 \quad (6)$$

In Equation (6), S_{ij} is the Euclidean distance (L2 norm), which determines the degree of dissimilarity between two instances in the embedding space. $\|$ denotes the norm in mathematics, which is used to measure the magnitude or distance of a vector Z_i and Z_j are embeddings of two data samples i and j . The lower the value of S_{ij} the more similar the two instances, while the higher the value, the more dissimilar the two instances, which helps in clustering.

3.5 Group-Specific Predictive Learning

Group-Specific Models is a learning method whereby a predictive model is learned in an individualized or adaptive manner with respect to specific subgroups in a given set of data, as opposed to having one model for the entire population. It is utilized to solve concerns related to bias and unfairness associated with algorithmic decision-making systems.

$$\hat{y}_k = g_k(Z_k) \quad (7)$$

In Equation (7), $\hat{y}_k$ is the prediction of output corresponding to k -th subgroup, the resultant output from model prediction for this particular subgroup. g_k subgroup-dependent model function that is independently trained to predict output for k -th subgroup. Z_k input feature space belonging to k -th subgroup, obtained from a specific distribution of this subgroup's data.

3.6 Three-Level Optimization Framework

The three-level optimization approach is a hierarchical learning methodology that works towards enhancing the performance, fairness, and feature selection capabilities at each level of the process, is shown in Table 2. This includes global optimization, group optimization, and local optimization processes.

Table 2: Fairness-Aware Multi-Level Optimization Framework

Levels	Title	Description
Level 1	Representation Optimization	Transforms raw data into meaningful features by reducing noise, bias, and redundancy
Level 2	Grouping Optimization	Clusters data into fair groups to enhance learning and reduce bias
Level 3	Predictive Optimization	Builds accurate models using optimized features, ensuring fairness and generalization

$$\min L_{total} = \alpha L_{pred} + \beta L_{fair} + \gamma L_{group} \quad (8)$$

In Equation (8), $\min$ is the minimum that, L_{total} loss function to optimize, which comprises model accuracy, fairness, and clustering loss. L_{pred} prediction loss that captures the difference between predictions and true labels. L_{fair} fairness loss that penalizes any bias in the model predictions across various sensitive subgroups. L_{group} grouping loss that ensures samples from the same cluster/subgroup are treated similarly by the model. α importance of prediction loss in the overall optimization process. β importance of fairness loss. γ weight controlling the importance of grouping consistency.

3.7 DTSO-XGBoost to improve fairness and accuracy

The use of DTSO-XGBoost optimization involves applying the XGBoost algorithm of ML along with that of DTSO Optimization with the view to gaining high levels of accuracy and bias in the process of algorithmic decision-making. To enhance predictive performance, reduce bias, improve generalization, and ensure robust, fair, and efficient decision-making across diverse data distributions. Using the technique of dynamic solution space exploration, the dynamic two-stage optimization is employed in the first stage with regard to feature selection

and improvement. This implies the deletion of redundant and noisy features along with the selection of only significant features.

- **XGBoost Optimization for classification and regression**

The XGBoost model is a supervised ML algorithm that can be employed for solving both classification and regression problems. Within the context of this study, it serves as the main prediction model for socio-economic and financial data while considering issues of fairness. As an extension of gradient boosting, it builds an ensemble of decision trees where subsequent trees learn from mistakes made by their predecessors.

$$M^{(r)} = \sum_{j=1}^u d(p_j, \check{p}_j^{(r-1)} + f_r(x_j)) + \Omega(f_r) \quad (9)$$

In Equation (9), $M^{(r)}$ denotes objective value at iteration, $\sum$ represents accumulation of total losses, j indicates index of data sample, u denotes total training samples, $d(\cdot)$ represents prediction error measurement, p_j denotes actual target output, $\check{p}_j^{(r-1)}$ indicates previous prediction output, f_r represents current tree prediction, x_j denotes input feature vector, r indicates boosting iteration number, Ω denotes model complexity regularization, $+$ represents addition operation.

- **Dynamic Transient Search Optimization (DTSO) for adaptive optimization strategy**

DTSO is a recently developed population-based metaheuristic designed to solve complex, high-dimensional optimization problems. It effectively avoids local optima by dynamically balancing exploration and exploitation through adaptive tuning of its search characteristics. In this study, DTSO is applied for both feature selection and hyperparameter optimization of the predictive model. A multi-objective fitness function is employed, incorporating both prediction accuracy and fairness (bias reduction) to evaluate candidate solutions. Due to its transient search strategy, DTSO is highly effective in handling noisy, high-dimensional, and biased datasets. Its capacity for avoiding premature convergence while retaining the ability for global search makes it a useful tool in optimizing for robustness and fairness-conscious decisions.

Algorithm 1: Fairness-Aware Grouping Learning Framework (FAGLF)

Input:

$D \rightarrow$ Fairness-aware dataset,
 $MaxIter \rightarrow$ Maximum DTSO iterations,
 $P \rightarrow$ Population size,
 $T \rightarrow$ Total boosting rounds

Output:

$M^* \rightarrow$ Optimized fairness-aware predictive model

Begin

1. Load dataset D
2. Data Preprocessing using Missing Value Imputation:
3. For each data sample z_j :
 Compute similarity score: $CSS(z_j, \hat{z}_j)$
 Estimate imputation error: $MSE(z_j, \hat{z}_j)$
 Replace missing values and normalize data
3. Data Balancing using SMOTE:
 - Identify minority class samples
 - Generate synthetic samples: $o_{ji}(w_j, w_{ji})$
- Construct a balanced dataset D_B
4. Feature Extraction using ICA:
 - Transform observed data: $w = At$
- Select optimized feature subset F_O
5. Dynamic Problem Decomposition (DPD):
 - Partition the dataset into homogeneous groups G_k
6. Group-Specific Predictive Learning:
7. For each subgroup G_k :
 - Train subgroup predictor g_k
 - Predict subgroup output $\hat{y}_k$
8. Three-Level Optimization Framework:
 - Perform representation optimization

- Execute grouping optimization
 - Conduct predictive optimization
9. DTSO-XGBoost Optimization:
 - Initialize DTSO population
 - Optimize feature selection and XGBoost parameters
 - Evaluate fairness and prediction accuracy
 10. Predictive Modeling using XGBoost:
 11. For $t = 1$ to T :
 - Train weak learner $F_t(x)$
 - Update ensemble prediction
 12. Return optimized predictive model M^*

End

Algorithm 1 represents FAGLF, which achieves improved fairness-aware predictions using techniques such as missing value imputation, SMOTE, ICA feature extraction, dynamic subgroup decomposition, and learning with subgroups. The optimization technique of DTSO-XGBoost provides better feature selection, fairness, and improved predictions.

4. Result and Discussion

FAGLF is an effective way to guarantee fairness reduction in algorithmic decision-making processes by combining dynamic grouping, fairness-aware learning, and optimization using DTSO-XGBoost. It maintains outstanding classification performance while achieving enhanced fairness features, decision stability, and accuracy. Python on Windows is used to address this issue, by using Scikit-learn and XGBoost.

Data feature exploration analysis: The density distributions of bias risk scores with regard to the decision output of approval, review, and rejection are the primary focus of the comparison of bias risk and fairness-adjusted scores across numerous categories. The Group Bias Risk distributions, where approved decisions peak at 1-4, manual review at 5-8, and rejected at 10-13, extending to 25, indicate increasing bias influence of Figure 2(a). Figure 2(b) shows the distribution of density for fairness-adjusted score in terms of regions, where urban regions have their peak values at 65-70, semi-urban at 60-65, and rural at 55-60, with ranges of 30-90.

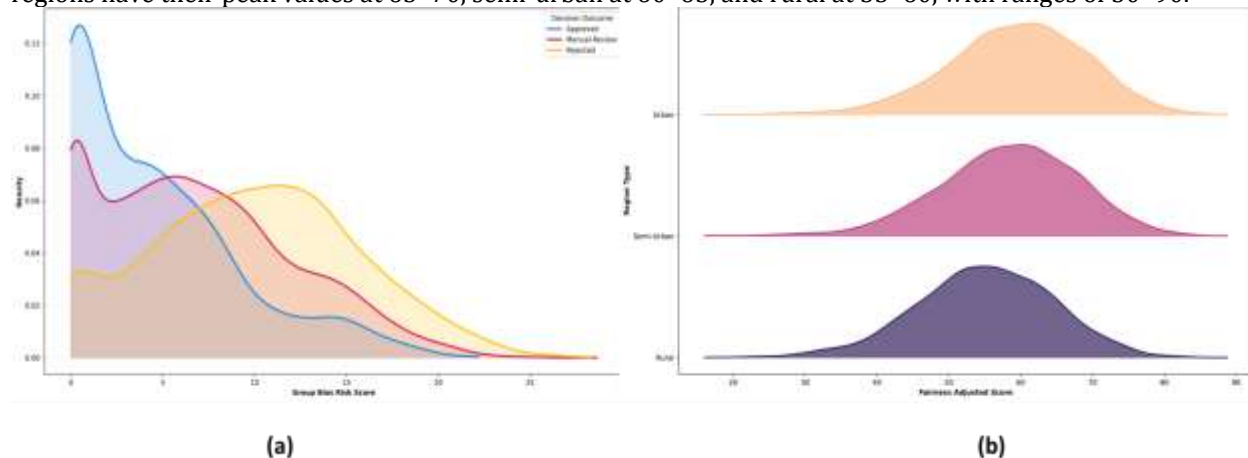


Figure 2: Comparative Analysis of Bias Risk and Fairness Scores (a) Distribution of Group Bias Risk Across Decisions (b) Fairness Adjusted Score Distribution Across Region Types

Metrics explanation: **Detection Precision:** Determines the proportion of accurate identification of the fair decisions among all those that have been identified as positive, representing the extent of correctness of the algorithm while identifying unbiased decisions without any errors. **Detection Recall:** The ability of the algorithm to detect all fair decisions in reality, indicating how well the algorithm is working at recognizing all unbiased decisions without any omissions. **F1 Score:** The harmonic mean of both these parameters, acting as an indicator of how effectively the algorithm detects fair decisions.

Performance evaluation using the credit scoring dataset [16]: The proposed FAGLF were trained using credit scoring dataset [16], along with some baselines, including Fair Detect (Adversarial Networks), Bias Audit (Ensemble Methods), Fair Sight (Representation Learning), and Credit Fair (Causal Modeling) [16]. As a result, it has been found that this suggested model has shown better results with respect to the Detection Precision, Recall, and F1 score metrics represented in Table 3. This means that the FAGLF model is more effective in determining fairness and unbiasedness than others.

Table 3: Performance Analysis of Fairness-Aware Detection Models using the existing dataset

Framework	Detection Precision	Detection Recall	F1 Score
Fair Detect [16]	0.884	0.912	0.898
Bias Audit [16]	0.921	0.876	0.898
Fair Sight [16]	0.856	0.934	0.893
Credit Fair [16]	0.945	0.887	0.915
FAGLF [proposed]	0.949	0.924	0.935

Performance evaluation using the fairness decision records dataset: Both existing and proposed trained models using the trained using decision records dataset, the detection performance of cutting-edge techniques, such as Fair Detect, Bias Audit, Fair Sight, and Credit Fair, works well; nevertheless, FAGLF outperforms the others in terms of precision, recall, and F1-score, as shown in Table 4 and Figure 3. This implies that the suggested FAGLF provides improved forecast fairness and accuracy. On the suggested dataset, it can be inferred that the FAGLF performs better than the most advanced techniques.

Table 4: Comparison of Detection Metrics Across Frameworks using the fairness decision records dataset

Framework	Detection Precision	Detection Recall	F1 Score
Fair Detect	0.901	0.926	0.913
Bias Audit	0.934	0.892	0.913
Fair Sight	0.873	0.947	0.909
Credit Fair	0.956	0.901	0.928
FAGLF [proposed]	0.972	0.948	0.960

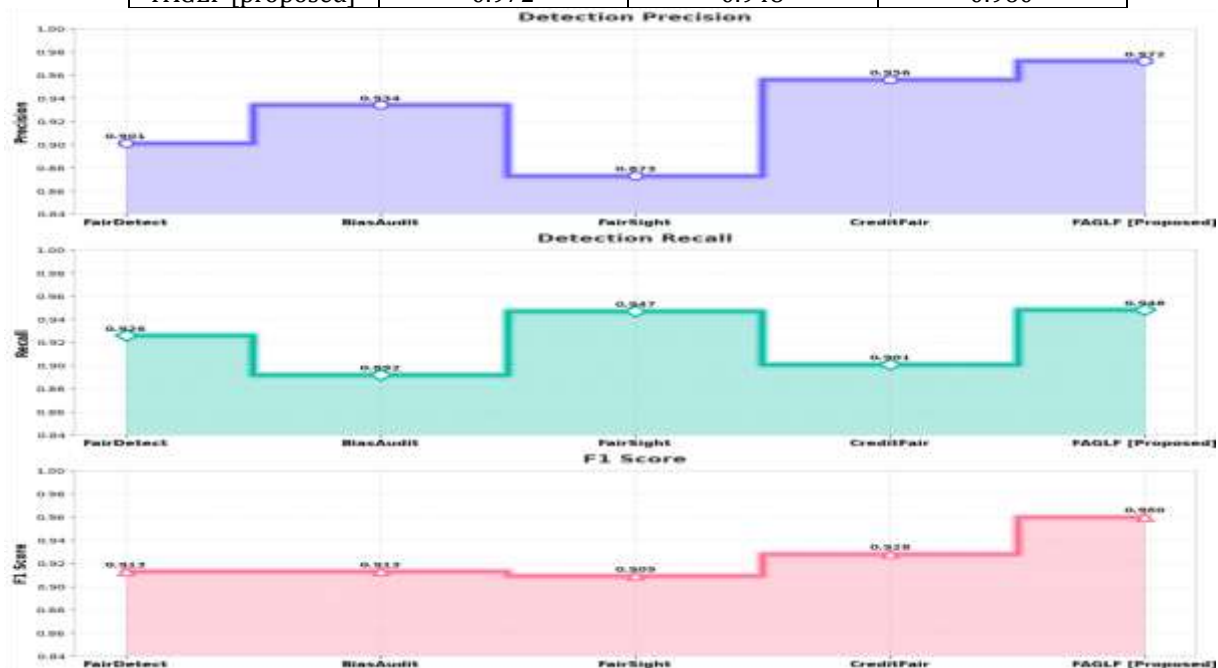


Figure 3: Presentation of the proposed and existing models trained using the fairness decision records dataset

Existing research efforts on a fair banking framework [11] and a credit scoring fairness model [16] attempted to solve the problem of bias through bias-aware approaches; however, they encountered challenges related to scalability and subgroup-specific biases. These problems can be solved using the proposed FAGLF framework by introducing dynamic grouping, allowing for detailed identification of subgroup bias problems. In addition, DTSO-XGBoost can optimize the model and further improve its flexibility and computing capability.

5. Conclusion

By imposing constraints of fairness, creating group equality, performing better than traditional techniques, and using reliable and scalable techniques for decision-making using AI algorithms, the FAGLF is capable of solving biases to make accurate predictions. Although there are methods that attempt to solve issues with bias in their predictions, they fail to incorporate solutions related to fairness, scale poorly, and negatively affect their accuracy. To achieve high levels of generalization without losing accuracy, the FAGLF approach tries to balance learning and scalability along with fairness constraints and adaptability to reduce bias. The FAGLF model has the best Precision of 0.972, recall of 0.948, and F1 score of 0.960 compared to other conventional models that serve as baselines. However, in most cases, such an approach is hard to implement owing to significant resource consumption and moderate adjustments of parameters. When expanding this method to the Deep Learning (DL) algorithms, implementing present fairness assessment techniques, studying dynamic environments, solving the scale issues, and designing reliable bias handling systems become the objectives for future research.

References

1. Bhat, J., 2024. Responsible Machine Learning in Student-Facing Applications: Bias Mitigation & Fairness Frameworks. *American International Journal of Computer Science and Technology*, 6(1), pp.38-49. <https://doi.org/10.63282/3117-5481/AIJCS-T-V6I1P104>
2. Ferrara, C., Sellitto, G., Ferrucci, F., Palomba, F. and De Lucia, A., 2024. Fairness-aware machine learning engineering: how far are we?. *Empirical software engineering*, 29(1), p.9. <https://doi.org/10.1007/s10664-023-10402-y>
3. Sargeant, H., 2023. Algorithmic decision-making in financial services: economic and normative outcomes in consumer credit. *AI and Ethics*, 3(4), pp.1295-1311. <https://doi.org/10.1007/s43681-022-00236-7>
4. Rinta-Kahila, T., Someh, I., Gillespie, N., Indulska, M. and Gregor, S., 2022. Algorithmic decision-making and system destructiveness: A case of automatic debt recovery. *European Journal of Information Systems*, 31(3), pp.313-338. <https://doi.org/10.1080/0960085X.2021.1960905>
5. Liang, Y., 2023. DFGR: Diversity and fairness awareness of group recommendation in an event-based social network. *Neural Processing Letters*, 55(8), pp.11293-11312. <https://doi.org/10.1007/s11063-023-11376-0>
6. Li, C.T., Hsu, C. and Zhang, Y., 2022. Fairsr: Fairness-aware sequential recommendation through multi-task learning with preference graph embeddings. *ACM Transactions on Intelligent Systems and Technology (TIST)*, 13(1), pp.1-21. <https://doi.org/10.1145/3495163>
7. Zhang, J., Li, Y., Wu, D., Zhao, Y. and Palaiahnakote, S., 2025. SFFL: Self-aware fairness federated learning framework for heterogeneous data distributions. *Expert Systems with Applications*, 269, p.126418. <https://doi.org/10.1016/j.eswa.2025.126418>
8. Kisten, M. and Khosa, M., 2024. Enhancing fairness in credit assessment: Mitigation strategies and implementation. *IEEE Access*, 12, pp.177277-177284. <https://doi.org/10.1109/ACCESS.2024.3505836>
9. Umeaduma, C.M.G., 2025. Explainable AI in algorithmic trading: mitigating bias and improving regulatory compliance in finance. *Int J Comput Appl Technol Res*, 14(4), pp.64-79. <https://doi.org/10.7753/IJCATR1404.1006>
10. <https://doi.org/10.7753/IJCATR1404.1006>
11. Agu, E.E., Abbulimen, A.O., Obiki-Osafiele, A.N., Osundare, O.S., Adeniran, I.A. and Efunniyi, C.P., 2024. Discussing ethical considerations and solutions for ensuring fairness in AI-driven financial services. *International Journal of Frontier Research in Science*, 3(2), pp.001-009. <https://doi.org/10.56355/ijfrms.2024.3.2.0024>
12. Selvam, M., 2025. Ethical AI for Personalized Banking: Addressing Bias and Fairness Challenges. *LatIA*, (3), p.361. <https://doi.org/10.62486/latia2025361>

13. Oyasiji, O., Okesiji, A., Imediegwu, C.C., Elebe, O. and Filani, O.M., 2023. Ethical AI in financial decision-making: Transparency, bias, and regulation. *International Journal of Scientific Research in Computer Science, Engineering and Information Technology*, 9(5), pp.453-471. <https://doi.org/10.32628/IJSRCSEIT>
14. Akhamere, G.D., 2023. Fairness in credit risk modeling: Evaluating bias and discrimination in AI-based credit decision systems. *International Journal of Advanced Multidisciplinary Research and Studies*, 3(6), pp.2061-2070. <https://doi.org/10.62225/2583049X.2023.3.6.4716>
15. Ford, J., 2025. Fairness in focus: quantitative insights into bias within machine learning risk evaluations and established credit models. *Management System Engineering*, 4(1), p.8. <https://doi.org/10.1007/s44176-025-00043-4>
16. Trinh, T.K. and Zhang, D., 2024. Algorithmic fairness in financial decision-making: Detection and mitigation of bias in credit scoring applications. *Journal of Advanced Computing Systems*, 4(2), pp.36-49. <https://doi.org/10.69987/JACS.2024.40204>