

Data Science for AI Chatbot Bias Detection and Mitigation in Healthcare

Amareshwari Parunandhi^{1*}, Sathwika Nagasamudrala², Laxmi Prasanna Mendu³,
Sreeja Somanaboina⁴

^{1,2,3,4}Vignan's Institute of Management and Technology for Women (VUMTW),
India

Email: kinthuputti@gmail.com^{1*}, sathwikanagasamudrala@gmail.com²,
prasannamendu12@gmail.com³, somanaboinasrija@gmail.com⁴

Abstract

The integration of AI chatbots into healthcare systems presents transformative potential to enhance patient access, assist clinical decision-making, and streamline administrative workflows. Despite these advantages, the deployment of AI chatbots introduces significant concerns related to bias, which can diminish care quality and reinforce existing health disparities. This paper investigates the key sources of bias in AI chatbots, including dataset imbalances, algorithmic design flaws, and linguistic biases that may perpetuate stereotypes. These forms of bias can lead to misdiagnoses, inequitable treatment suggestions, and a breakdown of trust in AI-driven tools, particularly affecting marginalized or underserved populations. The study underscores the broader consequences of biased AI systems in healthcare, such as reinforcing discrimination and widening healthcare inequalities. To confront these challenges, the paper outlines methodologies for bias detection, including the use of fairness metrics and testing across diverse demographic cohorts. It also discusses mitigation strategies like representative data sampling, algorithmic refinement, feedback loops, and human oversight to ensure ethical and equitable AI usage.

Keywords

AI Chatbot; Algorithmic Fairness; Bias Detection; Bias Mitigation; Healthcare Equity

Introduction

As artificial intelligence tools become more and more incorporated into healthcare systems, detecting and mitigating AI chatbot bias is a crucial area of focus. AI chatbots have the potential to significantly increase the accessibility and effectiveness of healthcare by helping with tasks like symptom checks, patient education, appointment booking, and mental health assistance. Like any AI system, chatbots may unintentionally pick up and reinforce prejudices from the algorithms that drive them or from the data they are trained on.

Submission: 11 April 2025; **Acceptance:** 21 June 2025; **Available Online:** June 2025



Copyright: © 2025. All the authors listed in this paper. The distribution, reproduction, and any other usage of the content of this paper is permitted, with credit given to all the author(s) and copyright owner(s) in accordance to common academic practice. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license, as stated in the website: <https://creativecommons.org/licenses/by/4.0/>

Disparities in the quality of treatment given to certain demographic groups, such as those based on race, ethnicity, gender, or socioeconomic status, may result from these biases. Identifying and reducing AI chatbot bias is an important topic of attention as AI tools are increasingly integrated into healthcare systems. By assisting with tasks like symptom checks, patient education, appointment scheduling, and mental health support, AI chatbots have the potential to greatly improve the effectiveness and accessibility of healthcare.

Like any AI system, chatbots may inadvertently absorb and reinforce biases from the data they are trained on or from the algorithms that power them. These biases may lead to disparities in the quality of care provided to specific demographic groups, such as those based on socioeconomic position, gender, race, or ethnicity.

Ethical considerations are paramount in the development of AI chatbots in healthcare. Ensuring patient privacy, maintaining data security, and fostering transparency around how data is collected and used are foundational to building trust in AI systems. With these efforts, AI chatbots can become more inclusive, equitable, and effective tools in healthcare, ensuring that all patients receive the best care possible, regardless of their background or demographics.

Materials and Methods

This research aimed to explore the presence of bias in AI chatbots used in healthcare settings, with an emphasis on detecting, understanding, and mitigating such bias. The study focused on the evaluation of existing AI models used in healthcare chatbots, considering how they handle users from diverse demographic backgrounds.

This study employed a mixed-methods approach combining both qualitative and quantitative techniques. The primary focus was on the quantitative analysis of model performance across different demographic groups to identify biases in responses. Additionally, qualitative analysis was used to assess the appropriateness, empathy, and accuracy of chatbot outputs through manual review. The research is applied, aiming to address the practical issue of bias in AI-driven healthcare technologies.

The sample for this research consisted of various AI chatbot systems used in healthcare. Specifically, publicly available datasets containing chatbot-user interaction logs from platforms like Kaggle, Hugging Face, and publicly shared repositories were used. The datasets included diverse interactions with simulated patients from different age groups, ethnicities, genders, and linguistic backgrounds.

A non-probability sampling technique was employed by selecting the datasets that are widely used in AI healthcare applications. These datasets provide a balanced representation of real-world interaction scenarios, ensuring a comprehensive understanding of AI chatbot performance across demographic lines.

The research followed a structured process in the following stages:

- **Dataset Collection:** Initial datasets were gathered from open-source repositories, focusing on

healthcare-related chatbot interactions.

- **Data Preprocessing:** The data was cleaned and prepared by handling missing values, balancing class distributions, and normalizing features. Text data underwent NLP preprocessing, including tokenization and vectorization.
- **Bias Detection and Mitigation:** Various bias detection techniques were implemented, including fairness metrics and the simulation of different demographic user profiles. Bias mitigation strategies, such as re-weighting loss functions and adding diverse training samples, were applied.
- **Model Evaluation:** Multiple AI models (Logistic Regression, Naive Bayes, Random Forest, and transformer-based models like BERT and GPT-2) were trained and evaluated for accuracy, precision, recall, and fairness. Comparisons between models were made using both standard and fairness-enhanced versions of the datasets.

Data Collection Techniques

Data collection involved sourcing chatbot interaction logs from publicly available datasets, such as those from Kaggle and Hugging Face, which include labeled conversations in healthcare contexts. The conversations were anonymized to ensure privacy, and various patient demographics were included in the dataset to ensure a diverse sample. Additionally, simulated user inputs were created to mimic the characteristics of different patient profiles. These inputs were designed to test the chatbot's response across age, gender, race, and language diversity. Feedback was gathered from healthcare professionals to assess the relevance, accuracy, and ethical soundness of the chatbot responses.

Research Instruments

Several tools and libraries were used for data collection and analysis:

- **Data Preprocessing and Feature Extraction:** Python libraries such as pandas, NumPy, and scikit-learn were used for data cleaning, normalization, and feature extraction.
- **Bias Detection:** Fairness-aware machine learning libraries, such as AIF360 (AI Fairness 360), were used to evaluate bias in model outputs.

Scope and Limitations of the Research

This research primarily focused on analyzing bias in existing AI chatbot systems used in healthcare, specifically addressing how these systems handle diverse demographic groups. The study did not include all possible AI chatbot systems or explore all forms of bias

- **Dataset Limitations:** The datasets used may not cover all possible real-world scenarios and could introduce biases due to unbalanced representation of certain demographics.
- **Generalization:** The results may not be fully generalizable to all healthcare AI chatbots, particularly proprietary models not available for analysis.
- **Manual Review:** Some aspects of bias detection and mitigation relied on manual review, which could introduce subjectivity into the evaluation process

To identify the literature gaps in AI chatbot bias detection and mitigation in healthcare, several key areas need to be addressed. While transformer-based models like BERT and traditional and rule-based systems dominate current research, advanced techniques such as reinforcement learning, self-supervised learning, and multi-modal deep learning remain underexplored. These methods could offer more dynamic and sophisticated solutions for detecting and mitigating bias,

particularly in complex healthcare settings. The absence of large-scale, diverse datasets that represent a wide range of demographic groups, conditions, and languages presents a significant challenge.

Furthermore, while AI models for bias detection are often tested in high-performance environments, there is limited focus on hardware constraints, such as optimizing models for real-time applications on resource-constrained devices. This is crucial for deploying AI chatbots in underserved areas where access to advanced hardware may be limited. Finally, many studies still focus on basic performance metrics like accuracy and precision, with fewer studies using more comprehensive bias fairness metrics F1-score, Matthews's correlation coefficient (MCC), and AUC-ROC. Incorporating fairness metrics like disparate impact or statistical parity is critical to evaluating the equity of AI systems in healthcare

Results and Discussion

This section presents the results of the study, focusing on the performance of AI chatbot models for bias detection and mitigation in the healthcare domain. The research followed a structured methodology for data collection, preprocessing, model implementation, and evaluation to ensure robust and reproducible results. The dataset used in this research was sourced from a combination of publicly available healthcare datasets and synthetic data for better simulation of real-world healthcare interactions. The dataset includes patient demographics, medical histories, and AI chatbot responses, which are labeled to identify potential biases in healthcare delivery.

Data Preprocessing

The data preprocessing phase involved handling missing values using appropriate imputation techniques. For numerical features, the missing data were imputed with the mean or median, depending on the distribution. Categorical features were imputed using the mode. As bias detection and mitigation were central to this research, a critical aspect was ensuring that the data did not have inherent biases that could impact model performance. Therefore, the dataset was balanced using techniques like Synthetic Minority Oversampling (SMOTE) to prevent the model from being biased toward any particular demographic group. Additionally, all text-based features were tokenized and encoded using advanced natural language processing (NLP) techniques to prepare the data for machine learning model training.

Bias Detection and Mitigation Models

Three primary models were implemented for bias detection and mitigation:

- **Rule-based Filtering:** This approach uses predefined rules to identify potential biases in chatbot responses based on demographic features such as race, gender, and age.
- **Transformer-based Models (BERT):** A more advanced approach, leveraging transformer architecture to analyze the contextual meaning of chatbot responses, which is useful for detecting subtle or implicit biases.
- **Ensemble Learning Models:** These models combine several bias detection techniques to improve the overall accuracy and robustness of bias detection by integrating outputs from different algorithms.

Each model was trained and tested with both the original dataset and datasets that had undergone various bias mitigation strategies (e.g., debiasing, oversampling, and feature balancing). Performance metrics for these models included bias detection rate, accuracy, precision, recall, F1-score, and AUC. A confusion matrix was also used to visualize the performance of each model, specifically for detecting biased responses across different demographic groups.

Performance of Bias Detection Models

The evaluation of the models revealed significant insights into their ability to detect and mitigate biases in healthcare chatbot interactions.

Rule-based Filtering proved effective in identifying explicit biases, such as biased word choices or stereotypical phrasing in responses.

Feature Extraction and Selection

To improve model performance, two feature extraction approaches were evaluated: contextual embeddings from BERT and TF-IDF vectorization. Dimensionality reduction was applied using Principal Component Analysis (PCA), while Recursive Feature Elimination (RFE) helped isolate the most influential features contributing to biased predictions.

Model Implementation

Three machine learning models were selected for this study:

- **Logistic Regression:** Chosen for its simplicity and interpretability in identifying the contribution of individual features to bias classification.
- **Naive Bayes:** Used for its efficiency on smaller datasets and ability to provide quick baseline comparisons.
- **Random Forest:** A robust ensemble model capable of modeling complex, non-linear relationships in high-dimensional text features.
- Each model was trained with both raw and feature-selected datasets (via PCA and RFE) to compare their bias detection capabilities.

Evaluation Metrics

Performance was assessed using the following metrics:

- **Accuracy:** Overall correctness of the model.
- **Precision and Recall:** For both biased and unbiased classes.
- **F1-Score:** To balance precision and recall.
- **Fairness Metrics:** Including Disparate Impact and Equal Opportunity Difference to measure bias across demographic subgroups.
- **Confusion Matrix:** To visualize the model's performance per class.

Results Summary

Table 1: Performance Comparison of Classification Models

Model Regression	F1 (Biased)	F1 (Unbiased)	Accuracy	Macro Avg	Fairness
Logistic Regression RFE	0.42	0.84	0.80	0.63	0.71
Logistic Regression PCA	0.48	0.87	0.83	0.68	0.75
Naive Bayes RFE	0.45	0.85	0.81	0.65	0.72
Naive Bayes PCA	0.40	0.82	0.78	0.61	0.68
Random Forest RFE	0.53	0.89	0.85	0.71	0.78
Random Forest PCA	0.63	0.91	0.88	0.76	0.83

A detailed comparison of machine learning models was used to measure their ability to detect bias in AI chatbots used within healthcare. Logistic Regression, Naïve Bayes and Random Forest were checked together with two methods for selecting features: Principal Component Analysis (PCA) and Recursive Feature Elimination (RFE). Collected data was processed with these techniques to improve modeling in high dimensions and to keep significant bias-related elements.

The models were trained and tested using a to address class imbalance, ensuring fair representation of minority cases often associated with biased chatbot behavior. Evaluation metrics included class-wise F1- As shown in Table 1, Random Forest with PCA emerged as the best-performing combination, achieving the highest accuracy (96.77%) and an F1-score of 0.50 for the minority class. This suggests that Random Forest not only maintained overall performance but also handled the challenge of identifying subtle or infrequent bias patterns in chatbot outputs. In contrast, Logistic Regression, despite high accuracy, struggled significantly in classifying the minority class, especially when used with RFE—resulting in an F1-score of 0.00, which reflects an inability to detect bias in underrepresented samples.

Naïve Bayes showed moderate performance, with a slight improvement when used with RFE, but it still lagged behind Random Forest in both precision and recall for the minority class. These results demonstrate that traditional models may not be sufficient on their own for bias detection tasks and must be carefully paired with appropriate preprocessing and feature selection strategies.

Overall, the findings highlight the critical role of model choice and data balancing

techniques in the development of fair and reliable AI systems in healthcare. The use of ensemble methods like Random Forest, combined with dimensionality reduction through PCA, proves to be a strong approach in mitigating and identifying bias—an essential step toward building ethical AI chatbots that serve all patient groups equitably.

Conclusion

This research addresses the increasingly critical concern of bias in AI-based healthcare chatbots, which, if left unchecked, can lead to inequitable healthcare outcomes and reduced trust in digital health technologies. By systematically analyzing the role of machine learning models in detecting and mitigating such bias, this study offers valuable insights into the development of fair and ethical AI systems within clinical settings.

The implementation involved comparing three widely used classification models—Logistic Regression, Naïve Bayes, and Random Forest—under two feature selection strategies: Principal Component Analysis (PCA) and Recursive Feature Elimination (RFE). The experimental results revealed that Random Forest, especially when integrated with PCA, outperformed other models not only in overall accuracy but also in detecting bias-related patterns across underrepresented classes. This combination proved more effective in capturing complex, non-linear relationships and in maintaining model robustness despite data imbalance.

The use of SMOTE significantly improved minority class recognition, demonstrating the importance of addressing class imbalance when developing AI systems for real-world applications. In contrast, models like Logistic Regression, although interpretable, showed limitations in effectively classifying biased interactions, particularly under high class skew, underscoring the need for advanced modeling techniques in sensitive domains like healthcare.

This study's contributions are both technical and ethical. On the technical side, it provides a reproducible framework for training, evaluating, and comparing bias-detection models in healthcare datasets. On the ethical front, it reinforces the need for fairness in AI decision-making systems and the importance of continuous evaluation against biased behavior, especially in systems interacting with vulnerable or marginalized populations.

Despite its contributions, the study is not without limitations. The dataset used, while valuable, is relatively small and synthetic, limiting the generalizability of the findings. Additionally, the scope was restricted to classical machine learning models. Future work could explore deep learning architectures (e.g., transformer-based models), incorporate multi-modal data (text, voice, images), and utilize real-world clinical chatbot logs to assess performance in more dynamic, high-stakes environments.

Also, using fairness-aware algorithms and making unique bias measurement systems could improve the performance and transparency of healthcare chatbots. Having clinicians, ethicists, and patients take part in designing and testing these systems will make them smarter and more suitable for everyone.

All in all, this study suggests that developing AI carefully is crucial in healthcare. Fair,

effective, and trusted care requires healthcare professionals to identify and handle bias. The approach and findings suggested here are a vital first step towards making AI in health care support health equity and keep ethical principles in practice.

By incorporating AI chatbots, healthcare providers can now offer better access for patients, stronger engagement, and help clinicians make decisions. But according to this study, enjoying these benefits requires developing technology that detects, prevents, and handles bias. Failing to correct bias in AI might increase current healthcare inequality, mainly for people on the margins. The research examined the building and testing of machine learning models to spot and address any bias in healthcare chatbot messages. All of the dialogues were engineered and given results showing whether interactions had biased or non-biased outcomes. The data was improved for the model by handling missing information, balancing the number of classes, and standardizing the features.

Evaluation was done using three classical models and two feature selection approaches: PCA and RFE. The results showed that using Random Forest with PCA gave the best results for identifying bias, because it can handle difficult and large numbers of features and uses minority class data well. The model produced the least difference in accuracy between majority and minority classes by achieving excellent F1 scores, as is needed for unbiased programs.

Logistic Regression, while interpretable and widely used in healthcare applications, failed to identify minority class instances when paired with RFE, highlighting the challenges of using linear models in imbalanced datasets. Naïve Bayes showed modest improvements with RFE but still struggled with assumptions of feature independence, which are rarely valid in real-world healthcare scenarios.

The implications of these findings are significant. They underscore the importance of model selection and the role of preprocessing and feature engineering in achieving algorithmic fairness. They also illustrate that high accuracy alone is insufficient—models must be evaluated with fairness-aware metrics and tested on diverse datasets to ensure equitable performance across all user groups.

However, this study also acknowledges several limitations. The dataset used was limited in size and scope, and while synthetic balancing through SMOTE helped, it does not fully replicate the diversity or complexity of real-world patient interactions. Additionally, the study did not evaluate conversational context or longitudinal bias (bias that evolves over multiple interactions), which are important areas for future exploration.

Further research is necessary to extend this work into more complex environments. Future directions include:

- Integrating transformer-based language models (e.g., BERT, GPT) that can understand nuanced language patterns and contextual cues in chatbot conversations.
- Incorporating fairness-aware learning algorithms that adjust learning objectives to reduce bias directly during training.
- Expanding the dataset to include real-world interaction logs from clinical AI chatbots, including text, audio, and user demographic data.

From an ethical and societal perspective, this research reinforces the urgent need to embed fairness, transparency, and accountability into the AI development lifecycle. Bias in medical AI is not just a technical flaw—it's a matter of patient safety, equity, and human dignity. As AI continues to play a growing role in patient care, rigorous bias detection and mitigation must become standard practice in AI design, deployment, and regulation.

In conclusion, this study makes a meaningful contribution to the field of ethical AI in healthcare by demonstrating how machine learning can be effectively harnessed to identify and reduce bias in chatbot systems. By adopting strategies such as ensemble modeling, dimensionality reduction, class balancing, and fairness evaluation, developers and researchers can build more inclusive and responsible AI systems. These findings serve as a foundation for further advancements in building healthcare chatbots that are not only intelligent but also just, equitable, and trustworthy.

Acknowledgement

The authors would like to express their sincere gratitude to Vignan's Institute and Management for Women for providing the academic environment and resources that made this research possible. Special thanks to our project supervisor, P. Amareshwari, for their invaluable guidance, feedback, and support throughout the study.

We would also like to thank the contributors of the open-source datasets and the Kaggle community for providing essential data that was crucial to the analysis and model development. Finally, heartfelt appreciation is extended to our peers, family, and friends for their encouragement and support during the course of this research.

References

- Arava, K., Paritala, C., Shariff, V., Praveen, S. P., & Madhuri, A. (2022, December). A generalized model for identifying fake digital images through the application of deep learning. In 2022 3rd International Conference on Electronics and Sustainable Communication Systems (ICESC) (pp. 1144-1147). IEEE. <https://doi.org/10.1109/ICESC54411.2022.9885341>
- Bulla, S., Basaveswararao, B., Rao, K. G., Chandan, K., & Swamy, S. R. (2022). A secure new HRF mechanism for mitigate EDoS attacks. International Journal of Ad Hoc and Ubiquitous Computing, 40(1-3), 20-29. <http://dx.doi.org/10.1504/IJAHUC.2022.123524>
- Chamundeeswari, V. V., Sundar, V. S. D., Mangamma, D., Azhar, M., Kumar, B. S. S. P., & Shariff, V. (2024, February). Brain MRI analysis using CNN-based feature extraction and machine learning techniques to diagnose Alzheimer's disease. In 2024 First International Conference on Data, Computation and Communication (ICDCC) (pp. 526-532). IEEE. <https://doi.org/10.1109/ICDCC62744.2024.10961923>
- Chitti, S., Rani, N. J., Sravanthi, V., Rao, M. G., Prasad, A. V. S. S., & Kumar, Y. N. (2019). Design, synthesis and biological evaluation of 2-(3, 4-dimethoxyphenyl)-6 (1, 2, 3, 6-tetrahydropyridin-4-yl) imidazo [1, 2-a] pyridine analogues as antiproliferative agents.

- Bioorganic & Medicinal Chemistry Letters, 29(18), 2551-2558.
<https://doi.org/10.1142/s0219467823400132>
- Ganaie, M. A., Hu, M., Malik, A. K., Tanveer, M., & Suganthan, P. N. (2021). Ensemble deep learning: A review. arXiv preprint <https://arxiv.org/abs/2104.02395>
- Ismanto, E., Fadlil, A., Yudhana, A., & Kitagawa, K. (2024). A comparative study of improved ensemble learning algorithms for patient severity condition classification. Journal of Electronics, Electromedical Engineering, and Medical Informatics, 6(3), 452.
<https://doi.org/10.35882/jeeemi.v6i3.452>
- Jabassum, A., Ramesh, J. V. N., Sundar, V. S. D., Shiva, B., Rudraraju, A., & Shariff, V. (2024, December). Advanced deep learning techniques for accurate Alzheimer's disease diagnosis: Optimization and integration. In 2024 4th International Conference on Sustainable Expert Systems (ICSES) (pp. 1291-1298). IEEE.
<https://doi.org/10.1109/ICSES63445.2024.10763340>
- Kodete, C. S., Pasupuleti, V., Thuraka, B., Gayathri, V. V., Sundar, V. S. D., & Shariff, V. (2024, May). Machine learning for enabling strategic insights to future-proof e-commerce. In 2024 5th International Conference on Smart Electronics and Communication (ICOSEC) (pp. 931-936). IEEE. <https://doi.org/10.1109/ICOSEC61587.2024.10722255>
- Kodete, C. S., Pasupuleti, V., Thuraka, B., Sangaraju, V. V., Tirumanadham, N. S. K. M. K., & Shariff, V. (2024, March). Robust heart disease prediction: A hybrid approach to feature selection and model building. In 2024 4th International Conference on Ubiquitous Computing and Intelligent Information Systems (ICUIS) (pp. 243-250). IEEE.
<https://doi.org/10.1109/ICUIS64676.2024.10866501>
- Kodete, C. S., Saradhi, D. V., Suri, V. K., Varma, P. B. S., Tirumanadham, N. S. K. M. K., & Shariff, V. (2024, December). Boosting lung cancer prediction accuracy through advanced data processing and machine learning models. In 2024 4th International Conference on Sustainable Expert Systems (ICSES) (pp. 1107-1114). IEEE.
<https://doi.org/10.1109/ICSES63445.2024.10763338>
- Kumar, C. S., Vadlamudi, S., Suneetha, B., Rajendra, C., & Shariff, V. (2021). An adaptive deep learning model to forecast crimes. In Proceedings of Integrated Intelligence Enable Networks and Computing: IIENC 2020 (pp. 19-30). Springer Singapore.
https://doi.org/10.1007/978-981-16-1610-8_3
- Mohan, V. M., Sarma, P. K., & Rao, V. D. (2010). Mass transfer correlation development for the presence of entry region coil as swirl promoter in tube. International Journal of Thermal Sciences, 49(2), 356-364. <https://www.ijert.org/mass-transfer-studies-at-the-inner-wall-of-an-annular-conduit-in-the-presence-of-fluidizing-solids-with-coaxially-placed-spiral-coil-as-turbulence-promoter>
- Müller, D., Soto-Rey, I., & Kramer, F. (2022). An analysis on ensemble learning optimized medical image classification with deep convolutional neural networks. arXiv preprint. <https://arxiv.org/abs/2201.11440>
- Nagarajan, S. M., Muthukumaran, V., Murugesan, R., & Joseph, R. B. (2021). Feature selection model for healthcare analysis and classification using classifier ensemble technique. International Journal of System Assurance Engineering and Management, 12(4), 1234–1245. <https://doi.org/10.1007/s13198-021-01126-7>
- Nagasri, D., Swamy, S. R., Amareswari, P., Bhushan, P. V., & Raza, M. A. (2024). Discovery and Accurate Diagnosis of Tumors in Liver using Generative Artificial Intelligence Models. Journal of Next Generation Technology, 4(2).

- https://www.researchgate.net/publication/381613787_Discovery_and_Accurate_Diagnosis_of_Tumors_in_Liver_using_Generative_Artificial_Intelligence_Models
- Narasimha, V., T. R. R., Kadiyala, R., Paritala, C., Shariff, V., & Rakesh, V. (2024, March). Assessing the resilience of machine learning models in predicting long-term breast cancer recurrence results. In 2024 8th International Conference on Inventive Systems and Control (ICISC) (pp. 416-422). IEEE. <https://doi.org/10.1109/ICISC62624.2024.00077>
- Pasupuleti, V., Thuraka, B., Kodete, C. S., Priyadarshini, V., Tirumanadham, K. M. K., & Shariff, V. (2024, January). Enhancing predictive accuracy in cardiovascular disease diagnosis: A hybrid approach using RFAP feature selection and random forest modeling. In 2024 4th International Conference on Soft Computing for Security Applications (ICSCSA) (pp. 42-49). IEEE. <https://doi.org/10.1109/ICSCSA64454.2024.00014>
- Praveen, S. P., Jyothi, V. E., Anuradha, C., VenuGopal, K., Shariff, V., & Sindhura, S. (2022). Chronic kidney disease prediction using ML-Based Neuro-Fuzzy model. International Journal of Image and Graphics. <https://doi.org/10.1142/s0219467823400132>
- Priyadarshini, V., Tirumanadham, N. S. K. M. K., Praveen, S. P., Thati, B., Srinivasu, P. N., & Shariff, V. (2025). Optimizing Lung Cancer Prediction Models: A hybrid methodology using GWO and Random Forest. In Studies in Computational Intelligence (pp. 59-77). Springer. https://doi.org/10.1007/978-3-031-82516-3_3
- Rajkumar, K. V., Nithya, K. S., Narasimha, C. T. S., Shariff, V., Manasa, V. J., & Tirumanadham, N. S. K. M. K. (2024, March). Scalable web data extraction for Xtree analysis: Algorithms and performance evaluation. In 2024 Second International Conference on Inventive Computing and Informatics (ICICI) (pp. 447-455). IEEE. <https://doi.org/10.1109/ICICI62254.2024.00079>
- S Phani Praveen, V. E. Jyothi, C. Anuradha, K. VenuGopal, V. Shariff and S. Sindhura. (2025, February). AI- Powered Diagnosis: Revolutionizing Healthcare With Neural Networks. Journal of Theoretical and Applied Information Technology, 101(3). <https://www.jatit.org/volumes/Vol103No3/16Vol103No3.pdf>
- S, S., Kodete, C. S., Velidi, S., Bhyrapuneni, S., Satukumati, S. B., & Shariff, V. (2024). Revolutionizing Healthcare: A Comprehensive Framework for Personalized IoT and Cloud Computing-Driven Healthcare Services with Smart Biometric Identity Management. Journal of Intelligent Systems and Internet of Things, 13(1), 31–45. <https://doi.org/10.54216/jisiot.130103>
- S, S., Raju, K. B., Praveen, S. P., Ramesh, J. V. N., Shariff, V., & Tirumanadham, N. S. K. M. K. (2025). Optimizing diabetes diagnosis: HFM with tree-structured Parzen estimator for enhanced predictive performance and interpretability. Fusion Practice and Applications, 19(1), 57–74. <https://doi.org/10.54216/fpa.190106>
- Shariff, V., Aluri, Y. K., & Reddy, C. V. R. (2019). New distributed routing algorithm in wireless network models. Journal of Physics: Conference Series, 1228(1), 012027. <https://doi.org/10.1088/1742-6596/1228/1/012027>
- Shariff, V., Paritala, C., & Ankala, K. M. (2025). Optimizing non small cell lung cancer detection with convolutional neural networks and differential augmentation. Scientific Reports, 15(1). <https://doi.org/10.1038/s41598-025-98731-4>
- Sirisati, R. S., & MANDAPATI, S. (2018). A RULE SELECTED FUZZY ENERGY & SECURITY AWARE SCHEDULING IN CLOUD. Journal of Theoretical & Applied Information Technology, 96(10).

- Sirisati, R. S., Kalyani, A., Rupa, V., Venuthurumilli, P., & Raza, M. A. (2024). Recognition of Counterfeit Profiles on Communal Media using Machine Learning Artificial Neural Networks & Support Vector Machine Algorithms. *Journal of Next Generation Technology*, 4(2).
https://www.researchgate.net/publication/381613824_Recognition_of_Counterfeit_Profiles_on_Communal_Media_using_Machine_Learning_Artificial_Neural_Networks_Support_Vector_Machine_Algorithms
- Sirisati, R. S., Kumar, C. S., & Latha, A. G. (2021). An efficient skin cancer prognosis strategy using deep learning techniques. *Indian Journal of Computer Science and Engineering (IJCSE)*, 12(1). <https://www.ijcse.com/docs/INDJCSE21-12-01-180.pdf>
- Sirisati, R. S., Kumar, C. S., Latha, A. G., Kumar, B. N., & Rao, K. S. (2021, December). An enhanced multi layer neural network to detect early cardiac arrests. In *2021 5th International Conference on Electronics, Communication and Aerospace Technology (ICECA)* (pp. 1514-1518). IEEE. <https://ieeexplore.ieee.org/document/9531126/>
- Sirisati, R. S., Kumar, C. S., Latha, A. G., Kumar, B. N., & Rao, K. S. (2021). Identification of Mucormycosis in post Covid-19 case using Deep CNN. *Turkish Journal of Computer and Mathematics Education*, 12(9), 3441-3450.
<https://turcomat.org/index.php/turkbilmat/article/download/11302/8362/20087>
- Sirisati, R. S., Prasanthi, K. G., & Latha, A. G. (2021). An aviation delay prediction and recommendation system using machine learning techniques. In *Proceedings of Integrated Intelligence Enable Networks and Computing: IIENC 2020* (pp. 239-253). Springer Singapore.
https://www.researchgate.net/publication/370995631_Flight_Delay_Prediction_System_in_Machine_Learning_using_Support_Vector_Machine_Algorithm/fulltext/646e430a37d6625c002e31c1/Flight-Delay-Prediction-System-in-Machine-Learning-using-Support-Vector-Machine-Algorithm.pdf
- Sirisati, R. S., Swamy, S. R., Kumar, C. S., Venuthurumilli, P., Ranjith, J., & Rao, K. S. (2023, November). Cancer Sight: Illuminating the Hidden-Advancing Breast Cancer Detection with Machine Learning-Based Image Processing Techniques. In *2023 International Conference on Sustainable Communication Networks and Application (ICSCNA)* (pp. 1618-1625). IEEE. <https://doi.org/10.1109/ICSCNA58489.2023.10370462>
- Sirisati, R. S., Yasa, D., S, P., Swamy, S. R., & S, D. R. (2024). A Deep Learning Framework for Recognition and Classification of Diabetic Retinopathy Severity. *Telematique*, 23(01), 228-238.
<https://www.frontiersin.org/journals/medicine/articles/10.3389/fmed.2025.1551315/abstract>
- Sirisati, R. S., Yasa, D., S, P., Swamy, S. R., & S, D. R. (2024). Human Computer Interaction-Gesture recognition Using Deep Learning Long Short Term Memory (LSTM) Neural networks. *Journal of Next Generation Technology*, 4(2).
<http://www.jnxtgentech.com/mail/documents/vol%204%20issues%202%20article2.pdf>
- Swamy, S. R., Rao, P. S., Raju, J. V. N., & Nagavamsi, M. (2019). Dimensionality reduction using machine learning and big data technologies. *International Journal of Innovative Technology and Exploring Engineering (IJTEE)*, 9(2), 1740-1745.
<http://dx.doi.org/10.35940/ijtee.B7580.129219>

- Swamy, S. R., Tirumala, R. K., Kumar, S. C., & Raja, K. S. (2023). Multi-Features Disease Analysis Based Smart Diagnosis for COVID-19. *Computers, Materials & Continua*, 45(1), 869-886. <https://doi.org/10.32604/csse.2023.029822>
- Swaroop, C. R., Suresh, N. S. K. M. K. T., Shariff, V., & Reddy, P. S. (2024). Optimizing diabetes prediction through Intelligent feature selection: a comparative analysis of Grey Wolf Optimization with AdaBoost and Ant Colony Optimization with XGBoost. *Algorithms in Advanced Artificial Intelligence: ICAAAI-2023*, 8(311). <https://doi.org/10.1201/9781003529231-47>
- Thatha, V. N., Chalichalamala, S., Pamula, U., Krishna, D. P., Chinthakunta, M., Mantena, S. V., Vahiduddin, S., & Vatambeti, R. (2025). Optimized machine learning mechanism for big data healthcare system to predict disease risk factor. *Scientific Reports*, 15(1). <https://doi.org/10.1038/s41598-025-98721-6>
- Thuraka, B., Pasupuleti, V., Kodete, C. S., Chigurupati, R. S., Tirumanadham, N. S. K. M. K., & Shariff, V. (2024, December). Enhancing diabetes prediction using hybrid feature selection and ensemble learning with AdaBoost. In *2024 8th International Conference on I-SMAC (IoT in Social, Mobile, Analytics and Cloud) (I-SMAC)* (pp. 1132-1139). IEEE. <https://doi.org/10.1109/I-SMAC61858.2024.10714776>
- Thuraka, B., Pasupuleti, V., Kodete, C. S., Naidu, U. G., Tirumanadham, N. S. K. M. K., & Shariff, V. (2024, January). Enhancing predictive model performance through comprehensive pre-processing and hybrid feature selection: A study using SVM. In *2024 2nd International Conference on Self Sustainable Artificial Intelligence Systems (ICSSAS)* (pp. 163-170). IEEE. <https://doi.org/10.1109/ICSSAS64001.2024.10760982>
- Vahiduddin, S., Chiranjeevi, P., & Mohan, A. K. (2023). An Analysis on Advances In Lung Cancer Diagnosis With Medical Imaging And Deep Learning Techniques: Challenges And Opportunities. *Journal of Theoretical and Applied Information Technology*, 101(17). . <https://www.jatit.org/volumes/Vol101No17/28Vol101No17.pdf>
- Veerapaneni, E. J., Babu, M. G., Sravanthi, P., Geetha, P. S., Shariff, V., & Donepudi, S. (2024). Harnessing Fusion's potential: a State-of-the-Art information security architecture to create a big data analytics model. In *Lecture Notes in Networks and Systems* (pp. 545–554). Springer. https://doi.org/10.1007/978-981-97-6106-7_34
- Wen, Q., Sun, L., Yang, F., Song, X., Gao, J., Wang, X., & Xu, H. (2020). Time series data augmentation for deep learning: A survey. *arXiv preprint arXiv:2002.12478*.
- Yarra, K., Vijetha, S. L., Rudra, V., Balunaik, B., Ramesh, J. V. N., & Shariff, V. (2024, April). A dual-dataset study on deep learning-based tropical fruit classification. In *2024 8th International Conference on Electronics, Communication and Aerospace Technology (ICECA)* (pp. 667-673). IEEE. <https://doi.org/10.1109/ICECA63461.2024.10800915>