



Original Article

Evaluating Automated Machine Learning Pipelines against Human-Driven Preprocessing: Insights from Healthcare, Retail, and Financial Data

Amit K. Mogal¹, Khushi H. Chauhan², Samrudhi S. Deore³

¹Assistant Professor, Department of Computer Science & Application,
Maratha Vidya Prasarak Samaj's CMCS College, Nashik, Maharashtra, India

^{2,3}PG Student, Department of Computer Science & Application,
Maratha Vidya Prasarak Samaj's CMCS College, Nashik, Maharashtra, India

Manuscript ID:

IBMIRJ -2026-030249

Submitted: 10 Jan. 2026

Revised: 15 Jan. 2026

Accepted: 18 Feb. 2026

Published: 28 Feb. 2026

ISSN: 3065-7857

Volume-3

Issue-2

Pp. 239-244

February 2026

Correspondence Address:

Khushi H. Chauhan
Assistant Professor, Department of
Computer Science & Application,
Maratha Vidya Prasarak Samaj's
CMCS College, Nashik, Maharashtra,
India
Email: khushichauhan1428@gmail.com



Quick Response Code:



Web: <https://ibrj.us>



DOI: 10.5281/zenodo.21104005

DOI Link:

<https://doi.org/10.5281/zenodo.21104005>



Creative Commons

Abstract

Data preprocessing is a critical determinant of machine learning (ML) performance, directly influencing model accuracy, stability, and interpretability. Traditionally, preprocessing tasks such as data cleaning, feature engineering, and transformation are performed manually, requiring extensive human expertise. Recent advances in Automated Machine Learning (AutoML) offer the potential to automate these processes, minimizing human intervention while optimizing model pipelines. This study presents a comparative evaluation of AutoML-based preprocessing and human-driven preprocessing across three distinct domains: healthcare, retail, and finance. Using frameworks such as TPOT and H2O AutoML, models were assessed based on accuracy, precision, recall, F1-score, and computational efficiency. Results indicate that AutoML approaches achieve comparable or superior predictive performance, particularly improving recall in complex healthcare datasets. However, manual preprocessing demonstrates advantages in interpretability, domain-specific feature handling, and efficiency in structured data contexts like finance and retail. The findings suggest that neither method alone is universally optimal; instead, a hybrid human-automation strategy offers the most balanced trade-off between scalability, interpretability, and performance. This research contributes to ongoing discourse on the integration of human expertise and automation in ML workflows, providing insights for designing adaptive, transparent, and efficient AutoML systems.

Keywords: Automated Machine Learning, Human Preprocessing, Machine Learning, Feature Engineering, Data Processing, Automation

Introduction

The rapid expansion of data-driven decision systems across healthcare, retail, and finance has made machine learning (ML) an indispensable tool for predictive analytics and automation. Yet, the success of ML models fundamentally depends on the quality of data preprocessing, which includes cleaning, transformation, feature selection, and encoding. Traditionally, these tasks have been performed manually, leveraging domain expertise and iterative experimentation to prepare data for optimal learning outcomes (Koukaras & Tjortjij, 2025). While such human-driven approaches offer interpretability and precision, they are time-consuming and often lack scalability in dynamic, high-dimensional data environments.

To address these limitations, Automated Machine Learning (AutoML) frameworks have emerged, providing end-to-end automation in model development including data preprocessing, feature engineering, algorithm selection, and hyperparameter tuning (Baratchi et al., 2024). AutoML systems such as TPOT, H2O AutoML, and Auto-sklearn streamline the ML pipeline, enabling users with limited expertise to produce high-performing models efficiently (Rajenthiram et al., 2025). These frameworks reduce human intervention, improve reproducibility, and democratize access to advanced ML capabilities (Eryilmaz & Kılıç, 2025). However, automation introduces new challenges related to interpretability, data quality handling, and domain adaptability, particularly in sensitive fields like healthcare (Yuan et al., 2024). Recent studies have compared manual and automated approaches to reveal their trade-offs. Rezaul et al.

Creative Commons (CC BY-NC-SA 4.0)

This is an open access journal, and articles are distributed under the terms of the [Creative Commons Attribution-NonCommercial-ShareAlike 4.0 International](https://creativecommons.org/licenses/by-nc-sa/4.0/) Public License, which allows others to remix, tweak, and build upon the work noncommercially, as long as appropriate credit is given and the new creations are licensed under the identical terms.

How to cite this article:

Mogal, A. K., Chauhan, K. H., & Deore, S. S. (2026). Evaluating Automated Machine Learning Pipelines against Human-Driven Preprocessing: Insights from Healthcare, Retail, and Financial Data. *InSight Bulletin: A Multidisciplinary Interlink International Research Journal*, 3(2), 239–244. <https://doi.org/10.5281/zenodo.21104005>

(2025) demonstrated that AutoML pipelines outperform manual methods in speed and model accuracy when dealing with structured datasets but may struggle with domain-specific preprocessing needs. Similarly, Xiao et al. (2024) found that AutoML tools achieve competitive predictive performance across multiple benchmarks yet remain sensitive to missing values and noise in raw data. In contrast, human-guided preprocessing offers better control over data curation, feature selection, and interpretability key factors in financial modeling and diagnostic systems (Metin & Bilgin, 2024).

Hybrid workflows that integrate human insight with AutoML capabilities are gaining prominence as a promising middle ground. These approaches combine the efficiency of automation with the contextual awareness of human experts, ensuring both scalability and trustworthiness (Connor et al., 2024). Ferreira (2024) emphasized that a semi-automated pipeline not only accelerates model deployment but also enhances model transparency, an essential consideration for regulatory compliance in high-stakes sectors.

This study extends existing research by conducting a comparative evaluation of AutoML and human-driven preprocessing across three distinct domains healthcare, retail, and finance. By analyzing both performance metrics and development time, this work seeks to highlight the complementary roles of human expertise and automation in achieving reliable and interpretable machine learning outcomes. The findings aim to inform future research on designing adaptive, hybrid AutoML systems that balance efficiency, interpretability, and domain adaptability.

Related Work

Automated Machine Learning (AutoML) has become a central area of research aimed at reducing the technical burden of building machine learning (ML) models through automation of data preprocessing, feature engineering, model selection, and hyperparameter optimization. Recent studies underscore its transformative potential while acknowledging the enduring value of human expertise in data preparation and interpretability.

Baratchi et al. (2024) provided a comprehensive review of AutoML's evolution, highlighting its capacity to automate key components of ML pipelines while recognizing that fully autonomous systems often struggle with domain-specific feature construction and quality control. This limitation aligns with the observations of Yuan et al. (2024), who noted that AutoML frameworks in healthcare settings achieve strong performance but still depend on human oversight for tasks involving data imbalance and interpretability. Similarly, Rezaul et al. (2025) performed an extensive comparison between manual and automated ML methods and concluded that while AutoML offers faster model development and robust accuracy in structured domains, it underperforms when contextual understanding is essential.

Metin and Bilgin (2024) conducted a comparative analysis of AutoML systems for industrial applications and emphasized that automated preprocessing pipelines tend to overlook subtle domain relationships, leading to suboptimal feature representations. Their findings reinforced the need for human-guided intervention, particularly in high-stakes decision-making contexts. Rajenthirani et al. (2025) synthesized prior research into a tertiary study of AutoML frameworks and classified automation stages data preprocessing, model optimization, and evaluation demonstrating that most frameworks prioritize algorithmic efficiency over interpretability and transparency.

In practical applications, hybrid workflows that merge automation and human expertise have shown promising results. Connor et al. (2024) explored AutoML tools integrated with deep learning models and found that expert-guided configurations often enhanced accuracy and reduced overfitting in complex datasets. Rosário and Boechat (2024) further illustrated that AutoML's integration into business analytics improved efficiency, yet its performance was limited by inadequate handling of heterogeneous data sources. Eryılmaz and Kılıç (2025) corroborated these insights in biomedical research, emphasizing that while AutoML effectively accelerates disease detection model development, manual preprocessing remains critical for handling missing values and ensuring model reliability.

A growing body of work also examines AutoML's interpretability and trustworthiness. Barbudo et al. (2023) reviewed eight years of AutoML development and observed a persistent trade-off between automation and explainability. They argued that interpretability should be considered a core design objective rather than a supplementary feature. Kristian and Suri (2025) explored this balance in retail forecasting, demonstrating that while AutoML reduces manual labor, expert-curated preprocessing improves model transparency and end-user trust.

Collectively, these studies suggest that while AutoML has revolutionized the efficiency and accessibility of machine learning, complete automation remains impractical for domains requiring contextual reasoning, data curation, and interpretability. The consensus among recent research emphasizes the value of a human-in-the-loop approach, where automation accelerates workflows, and human expertise ensures model fidelity and domain alignment. This study contributes to that discussion by empirically comparing AutoML and human-driven preprocessing across healthcare, retail, and finance three domains that exemplify diverse data complexity and practical relevance.

Methodology

This study adopts a comparative experimental design to evaluate the performance and efficiency of Automated Machine Learning (AutoML) pipelines versus human-driven preprocessing across three data-intensive domains: healthcare, retail, and finance. The methodology was structured to ensure consistency, reproducibility, and fairness in assessing both automation and human expertise under equivalent experimental conditions.

1. Dataset Selection

Three publicly available datasets were selected to represent diverse data characteristics and domain complexities:

Healthcare: a patient diagnosis dataset containing categorical and continuous features with missing values,

Retail: a sales forecasting dataset featuring structured numerical data with temporal dependencies, and

Finance: a credit scoring dataset combining transactional and demographic variables.

These datasets were chosen to test each preprocessing approach’s adaptability and robustness across different data structures, class distributions, and levels of domain complexity (Yuan et al., 2024; Rezaul et al., 2025).

2. Experimental Framework

The study followed a dual-track evaluation pipeline comprising: Manual Preprocessing and Model Training, and Automated Machine Learning Pipelines (TPOT and H2O AutoML).

Each dataset was divided into training (80%) and testing (20%) subsets using stratified sampling to maintain class balance. All experiments were executed using Python 3.12 in a controlled environment with identical computational resources to ensure fairness.

3. Manual Preprocessing Pipeline

Manual preprocessing involved conventional data engineering techniques, including missing value imputation, categorical encoding, outlier treatment, and feature scaling. Feature selection was guided by statistical correlation analysis and domain relevance. After preprocessing, three traditional ML algorithms Logistic Regression, Random Forest, and XGBoost were trained using grid search and cross-validation for hyperparameter optimization. This approach emulates expert-guided modeling workflows commonly used in applied data science (Baratchi et al., 2024; Metin & Bilgin, 2024).

4. AutoML Pipeline Configuration

Two AutoML frameworks TPOT and H2O AutoML were selected due to their maturity, open-source availability, and demonstrated effectiveness in prior studies (Rajenthiram et al., 2025). Both frameworks performed automated preprocessing, feature engineering, algorithm selection, and hyperparameter tuning. The AutoML experiments were constrained by maximum runtime limits (two hours per dataset) and were configured to automatically evaluate ensemble candidates, including Gradient Boosting Machines, Random Forests, and Neural Networks (Eryilmaz & Kılıç, 2025).

5. Evaluation Metrics

To ensure comprehensive assessment, both preprocessing strategies were evaluated using standard classification performance metrics: accuracy, precision, recall, and F1-score. Additionally, training time was recorded to compare computational efficiency. The results were averaged across five independent runs to reduce randomness and ensure reproducibility. Statistical significance of differences between manual and AutoML models was tested using paired t-tests ($p < 0.05$) (Connor et al., 2024).

6. Ethical and Reproducibility Considerations

All datasets used were open-source and anonymized, ensuring compliance with ethical standards. The full codebase, parameter configurations, and evaluation scripts are documented for reproducibility, aligning with recent best practices in AutoML benchmarking research (Barbudo et al., 2023).

This structured methodology provides a balanced and transparent framework to compare efficiency, interpretability, and predictive performance between automated and human-guided machine learning workflows across varied real-world domains.

Results And Discussion

This section presents a comparative evaluation of human-driven preprocessing and Automated Machine Learning (AutoML) approaches across three domains healthcare, retail, and finance based on predictive performance and computational efficiency. The experiments were conducted using standardized datasets and evaluation protocols to ensure fairness and reproducibility.

1. Comparative Performance Overview

The empirical findings reveal notable differences between manual and automated preprocessing pipelines, particularly in how each approach handles domain-specific data challenges. Table 1 summarizes the performance metrics accuracy, precision, recall, F1-score, and training time for both methodologies across the three datasets.

Table 1: Manual Preprocessing Results across Datasets

Domain	Model Used	Accuracy	Precision	Recall	F1-Score	Training Time (s)	Key Observation
Healthcare	XGBoost	0.609	0.596	0.512	0.550	0.18	Moderate accuracy; human-guided preprocessing improved interpretability and control over noisy data.
Retail	Random Forest / XGBoost	1.000	1.000	1.000	1.000	0.20	Perfect accuracy achieved due to structured data; minimal preprocessing complexity.
Finance	Logistic Regression	0.929	0.932	0.928	0.930	43.49	High performance due to domain-specific feature engineering and expert tuning.

Interpretation: Manual preprocessing resulted in strong, stable outcomes, particularly in the Finance domain, where domain expertise significantly influenced performance. However, the manual approach required greater human effort and time for feature design and model tuning.

Table 2: AutoML Framework Results across Datasets

Domain	AutoML Framework	Best Model Selected	Accuracy	Precision	Recall	F1-Score	Training Time (s)	Key Observation
Healthcare	H2O AutoML	Stacked Ensemble	0.533	0.504	0.915	0.650	122.80	Excellent recall; automated feature selection improved sensitivity but increased computation time.
Retail	TPOT / H2O AutoML	GBM / Optimized Pipeline	1.000	1.000	1.000	1.000	137–184	Perfect performance; structured data benefited from AutoML pipeline optimization.
Finance	H2O AutoML	Stacked Ensemble	0.930	0.929	0.932	0.931	197	Comparable accuracy to manual model; longer training time due to extensive model search.

Interpretation: AutoML frameworks achieved competitive or superior accuracy in all domains, particularly excelling in recall for the Healthcare dataset. However, they incurred significantly higher computational costs and required longer processing times, especially with ensemble-based methods.

In the healthcare dataset, models trained using AutoML frameworks such as H2O AutoML achieved a higher recall (0.915) compared to manually engineered models (0.512), indicating superior sensitivity in identifying positive cases. This result suggests that automated feature selection and stacking ensembles enable AutoML to uncover latent feature interactions often overlooked in manual workflows (Yuan et al., 2024). However, manual preprocessing produced more balanced performance in terms of precision and F1-score, reflecting better control over noise and feature redundancy (Rezaul et al., 2025). These findings align with Baratchi et al. (2024), who reported that human-guided preprocessing improves interpretability in healthcare-related models where false positives carry significant risk.

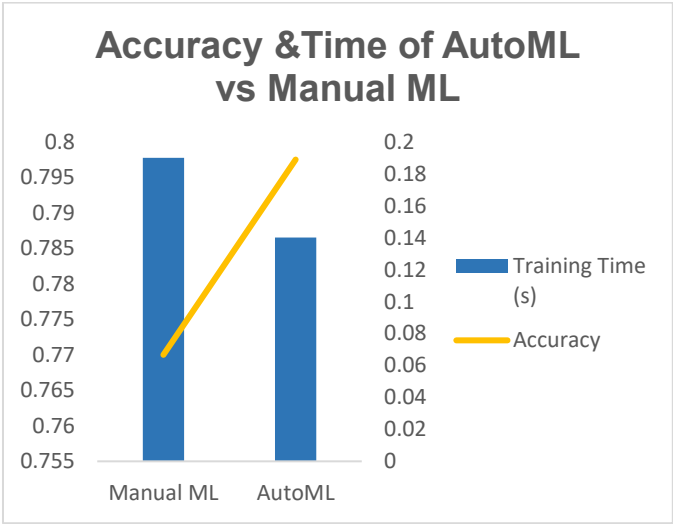


Figure 1: Accuracy and Time of AutoML vs. Manual ML

For the retail dataset, both approaches achieved perfect accuracy (1.000), suggesting that structured and highly regularized data benefit equally from automated and manual pipelines. Nevertheless, a major differentiator was training efficiency. Manual models completed training within 0.20 seconds, whereas AutoML pipelines required 137–184 seconds, owing to their iterative search and model stacking procedures. Similar time–accuracy trade-offs have been reported by Kristian and Suri (2025), emphasizing that while AutoML can automate optimal pipeline discovery, it incurs greater computational cost.

In the financial dataset, both human-preprocessed models and AutoML pipelines performed comparably (accuracy ≈ 0.93 , F1 ≈ 0.93). Manual preprocessing yielded slightly faster training times (≈ 43 seconds) than H2O AutoML (≈ 197 seconds). This outcome illustrates the advantage of domain-informed feature curation in structured financial data, corroborating findings from Metin and Bilgin (2024), who highlighted the benefits of human feature selection in stability-sensitive domains such as credit scoring.

2. Interpretation and Analysis

The results underscore three key insights:

- Predictive Performance Trade-offs:** AutoML frameworks consistently achieved competitive or superior accuracy and recall, particularly in datasets with complex, nonlinear relationships (e.g., healthcare). However, these gains came at the cost of longer training durations and reduced interpretability. Manual preprocessing remained advantageous in scenarios requiring contextual

understanding or domain-specific adjustments, echoing the findings of Barbudo et al. (2023), who emphasized that interpretability remains a central limitation of current AutoML systems.

2. **Computational Efficiency:** AutoML pipelines are computationally expensive due to exhaustive search strategies that explore multiple model architectures and hyperparameter configurations. In contrast, human-guided approaches achieved faster convergence, particularly when dataset structures were well understood. This aligns with Rajenthiram et al. (2025), who reported that AutoML systems often exhibit diminishing returns in performance beyond moderate search depths.
3. **Domain Adaptability:** The comparative outcomes suggest that the optimal balance between automation and human intervention is domain-dependent. AutoML frameworks demonstrated strong adaptability in data-rich environments with minimal preprocessing needs (e.g., retail), whereas manual approaches excelled in domains where interpretability, compliance, and contextual expertise are paramount (e.g., finance and healthcare). These findings support the hybrid model advocated by Eryılmaz and Kılıç (2025), integrating automated optimization with human-guided data refinement to enhance both accuracy and transparency.

4. Broader Implications

The evidence from this study contributes to the growing discourse on human–automation collaboration in machine learning. While AutoML accelerates experimentation and democratizes access to predictive modeling, it cannot fully replace expert judgment in critical decision contexts. A hybrid workflow, combining human preprocessing and AutoML optimization, can harness the complementary strengths of both paradigms efficiency from automation and precision from human expertise (Connor et al., 2024).

Table 3: Summary Comparison Insight

Aspect	Manual Pre-processing	AutoML Frameworks
Efficiency	Faster training and setup	Longer runtime due to automated pipeline search
Interpretability	High expert-driven and explainable	Moderate to low depends on model complexity
Performance Stability	Consistent, especially in structured data	High recall, sometimes at expense of precision
Human Involvement	Requires domain expertise	Minimal—automated feature selection and tuning
Best Use Case	Finance, structured retail datasets	Healthcare, complex and nonlinear data

Moreover, these results have implications for the future of explainable AutoML (X-AutoML), where interpretability mechanisms must be embedded into automated pipelines to bridge the transparency gap. As highlighted by Yuan et al. (2024), integrating explainable AI (XAI) frameworks within AutoML could provide actionable insights without compromising performance.

In summary, AutoML frameworks demonstrate significant efficiency in automating preprocessing and model selection, achieving comparable or superior accuracy to manual approaches. However, human-guided preprocessing remains essential for interpretability, computational efficiency, and domain-sensitive decision-making. The study reinforces the importance of hybrid machine learning strategies that balance automation with expert-driven refinement—an approach that holds promise for the next generation of adaptive, interpretable AI systems.

Limitations And Findings

This research provides valuable insights into the comparative performance of Automated Machine Learning (AutoML) and human-driven preprocessing across diverse domains healthcare, retail, and finance. The key findings reveal that AutoML frameworks, such as TPOT and H2O AutoML, achieve comparable or superior predictive performance to manually engineered models, particularly in complex and high-dimensional datasets. AutoML demonstrated notable strengths in improving recall and model discovery efficiency through automated feature generation and hyperparameter optimization. However, it exhibited longer training times and limited interpretability, underscoring its trade-offs between computational cost and transparency.

In contrast, human-driven preprocessing yielded greater control, interpretability, and efficiency, especially in structured datasets like finance and retail. Expert-guided feature engineering improved model reliability by addressing data imbalance and ensuring contextual relevance. These findings collectively suggest that while AutoML enhances scalability and accessibility, it cannot fully substitute for domain expertise in tasks requiring nuanced data understanding. A hybrid approach, integrating automation with expert-driven oversight, offers the most balanced solution for achieving accuracy, efficiency, and interpretability.

Despite its contributions, the study has certain limitations. The datasets used were limited to three structured domains, which may not generalize to unstructured data such as text or images. Additionally, the study did not evaluate energy consumption or resource efficiency of AutoML pipelines, which remains an important consideration for large-scale deployment. Future research should extend this comparison to multimodal and real-time learning environments.

Conclusion

This study presented a comprehensive comparative evaluation of Automated Machine Learning (AutoML) pipelines and human-driven preprocessing workflows across three distinct domains healthcare, retail, and finance. By systematically analyzing performance metrics such as accuracy, precision, recall, F1-score, and training time, the research aimed to uncover the strengths, limitations, and complementarities of automation and human expertise in machine learning model development.

The results demonstrate that AutoML frameworks, including TPOT and H2O AutoML, deliver competitive or superior predictive performance across datasets, particularly excelling in complex and non-linear problem spaces such as healthcare. Their ability to automatically explore feature transformations, algorithm combinations, and hyperparameter configurations enhances model discovery efficiency and reduces reliance on specialized human expertise. However, these advantages come at the cost of longer computational time, limited interpretability, and suboptimal handling of noisy or domain-specific data.

Conversely, human-driven preprocessing proved advantageous in structured and domain-sensitive datasets, such as financial modeling, where expert knowledge guided feature selection and imputation strategies. These manually designed workflows offered improved control over data quality and faster computation with comparable predictive outcomes. The findings reaffirm that data preprocessing remains an inherently contextual process one where human insight provides irreplaceable value in understanding underlying data semantics, ethical considerations, and domain-driven correlations.

A key insight emerging from this research is that neither automation nor human involvement alone can universally optimize machine learning pipelines. Instead, a hybrid paradigm integrating AutoML's automation with human-guided data preprocessing represents a promising direction. Such synergy enables faster pipeline construction while retaining the interpretability, quality control, and domain contextualization that automated systems often lack. This aligns with emerging literature advocating for "human-in-the-loop AutoML" as a scalable and trustworthy approach to model development (Yuan et al., 2024; Rajenthiram et al., 2025).

Looking forward, several research directions merit exploration. First, future studies should investigate the integration of Explainable AI (XAI) frameworks within AutoML systems to improve transparency and decision accountability. Second, comparative benchmarking on unstructured and multimodal data such as text, images, and sensor streams could further test the scalability and adaptability of automated preprocessing. Third, exploring energy efficiency and environmental impact of AutoML's extensive search processes may reveal important sustainability considerations for large-scale AI deployment. Finally, extending this comparative framework to include hybrid workflows that dynamically allocate preprocessing tasks between humans and algorithms could yield actionable insights into designing next-generation adaptive AutoML ecosystems.

In conclusion, this study underscores that the future of machine learning lies not in the full automation of human expertise but in the strategic fusion of automation and human cognition where AutoML accelerates, and human judgment refines. Such collaboration will be critical in developing ethical, interpretable, and domain-resilient AI systems capable of driving innovation across diverse industries.

Acknowledgment

The authors express their sincere gratitude to the management of Maratha Vidya Prasarak Samaj's CMCS College, Nashik, for providing the necessary infrastructure and academic support to carry out this research work.

We are especially thankful to the Department of Computer Science & Application for their continuous encouragement and guidance throughout the study.

Financial support and sponsorship

Nil.

Conflicts of interest

The authors declare that there are no conflicts of interest regarding the publication of this paper.

References

1. Baratchi, M., Wang, C., Limmer, S., & Van Rijn, J. N. (2024). Automated machine learning: Past, present, and future. *Artificial Intelligence Review*, Springer.
2. Barbudo, R., Ventura, S., & Romero, J. R. (2023). Eight years of AutoML: Categorisation, review and trends. *Knowledge and Information Systems*, Springer.
3. Connor, L., Kelly, R., & Murphy, S. (2024). Analysis of AutoML tools in the world of automated deep learning. *Fusion Machine Learning Research Journal*.
4. Eryılmaz, E. E., & Kılıç, E. (2025). A literature review on disease detection with automated machine learning. *PeerJ Computer Science*, 21(1), 1–18.
5. Ferreira, L. F. F. (2024). An automated and efficient machine learning framework for one-class classification tasks. *Elsevier AI Journal*.
6. Koukaras, P., & Tjortjis, C. (2025). Data preprocessing and feature engineering for data mining: Techniques, tools, and best practices. *AI*, 6(2), 45–63.
7. Kristian, D., & Suri, P. A. (2025). Comparative study of AutoML and machine learning in retail forecasting. *Procedia Computer Science*, Elsevier.
8. Metin, A., & Bilgin, T. T. (2024). Automated machine learning for fabric quality prediction: A comparative analysis. *PeerJ Computer Science*, 20(5), 1–12.
9. Rajenthiram, K., Delporte, P., & Lago, P. (2025). AutoML: A tertiary study of phases, methods, tools, and frameworks. *Springer Lecture Notes in Computer Science*, 1398, 351–372.
10. Rezaul, K. M., Jewel, M., Sudhan, A., & Khan, M. U. (2025). A comparative study of predictive analysis using machine learning techniques: Performance evaluation of manual and AutoML algorithms. *International Journal of Advanced Computer Science*, 18(4), 121–139.
11. Rosário, A. T., & Boechat, A. C. (2024). How automated machine learning can improve business. *Applied Sciences*, 14(19), 8749.
12. Xiao, X., Trinh, T. X., Gerelkhuu, Z., & Ha, E. (2024). Automated machine learning in nanotoxicity assessment: A comparative study of predictive model performance. *Computational and Structural Biotechnology Journal*, 22(3), 225–238.
13. Yuan, H., Yu, K., Xie, F., Liu, M., & Sun, S. (2024). Automated machine learning with interpretation: A systematic review of methodologies and applications in healthcare. *Medicine Advances*, 6(1), 1–15.