

Cross-platform Fake Review Detection: A Comparative Analysis of Supervised and Deep Learning Models

Faryad Bigdeli

School of Computing and Creative Technologies, University of the West of England, Bristol, UK
E-mail: Faryad2.Bigdeli@live.uwe.ac.uk
ORCID iD: <https://orcid.org/0009-0005-9675-146X>

Received: 04 September 2024; Revised: 10 January 2025; Accepted: 28 March 2025; Published: 08 June 2025

Abstract: This project addresses the growing issue of fake reviews by developing models capable of detecting them across different platforms. By merging five distinct datasets, a comprehensive dataset was created, and various features were added to improve accuracy. The study compared traditional supervised models like Logistic Regression and SVM with deep learning models. Notably, simpler supervised models consistently outperformed deep learning approaches in identifying fake reviews. The findings highlight the importance of choosing the right model and feature engineering approach, with results showing that additional features don't always improve model performance.

Index Terms: Fake Review Detection, Supervised Learning, Deep Learning, Feature Engineering, Cross-platform Analysis

1. Introduction

The advent of e-commerce has revolutionized how consumers shop, offering convenience and a wide array of choices. However, this transformation has been accompanied by the proliferation of fake reviews, posing significant challenges to the integrity of online marketplaces [1, 2]. These fraudulent practices not only mislead consumers but also tarnish the reputation of platforms and genuine sellers, undermine consumer trust and distort purchasing decisions. Current literature demonstrates the effectiveness of various machine learning models, including Logistic Regression, Random Forest, Support Vector Machine, and Naïve Bayes, in identifying fake reviews [3-6]. However, significant challenges remain, such as the need for large-scale, diverse datasets and the limitations of current methods in handling unbalanced data [7, 8].

A critical gap identified in the literature is the need for cross-platform detection capabilities. Many existing models are trained on datasets from a single e-commerce platform, limiting their generalizability. Alsubari et al. [1] explored the detection of fake hotel reviews using supervised machine learning techniques, finding that Random Forest outperformed other models. However, the study noted the scarcity of large labeled datasets, highlighting a significant limitation in current research. Similarly, Soldner et al. [9] found that detecting fake reviews is more challenging with cross-platform datasets due to additional variability and complexity. Their findings suggest that product ownership and data origin confound fake review detection, leading to overestimation of model performance when data is sourced from different platforms.

Feature engineering plays a crucial role in enhancing the performance of fake review detection models. Nguyen [4] proposed a machine learning framework for detecting fake reviews using text features and found that SVC was the most reliable. However, the study suggested exploring other feature extraction methods, like Bag of Words and Word2Vec, and considering deep learning approaches in future work.

Recently, deep learning methods have been increasingly adopted in fake review detection research, with studies reporting their effectiveness [10, 11]. Vyas et al. [12] proposed a deep learning-based model using features such as word count, sentence length, sentiments, and N-grams, achieving promising results. Their preliminary results indicated that the LSTM-based approach has the potential to outperform existing machine learning methods.

This study aims to address the identified gaps by developing and evaluating both supervised and deep learning models for fake review detection across multiple e-commerce platforms. The project specifically focuses on three key objectives: (1) assessing the performance of the models on a cross-platform dataset to evaluate their generalizability, (2)

analyzing the impact of feature engineering on model accuracy and reliability, and (3) comparing the efficacy of traditional machine learning approaches with deep learning techniques. By addressing these goals, this study seeks to contribute to the development of more robust and adaptable fake review detection systems, ultimately enhancing the integrity and trustworthiness of online marketplaces.

2. Dataset Overview

Initially, around 20 user review datasets were sourced from platforms like GitHub and Kaggle. After careful review, five datasets were selected for their compatibility in structure and essential features like labels, allowing them to be merged seamlessly. The selected datasets include three Amazon review datasets, one Yelp dataset, and one hotel review dataset.

The Amazon datasets contain reviews across different product categories, with approximately 21,000, 18,000, and 40,000 reviews respectively. The Yelp dataset includes around 360,000 reviews, primarily focused on restaurants. The hotel dataset comprises 1,600 reviews from multiple platforms, such as TripAdvisor.

These datasets were chosen because they provided raw, unprocessed data that maintained the originality of the analysis and allowed for cross-platform applicability. The diversity of review sources and product categories helps ensure that the dataset is representative and suitable for developing models with broad applicability.

Subsequently, a distinctive dataset was created by merging the five chosen datasets to reflect a diverse array of consumer feedback across various e-commerce platforms. The final dataset, post-merging, comprises 8 attributes with 7960 observations, evenly split with 1592 instances from each of the five datasets, containing 796 genuine and 796 fake reviews per dataset. After feature engineering, the modified dataset expanded to 16 attributes, which were utilized for the modeling phase.

A new test sample consisting of 800 reviews, evenly split between 400 fake and 400 real reviews, was used to evaluate the performance of the supervised models. This dataset includes products from different categories than those in the dataset used to create the models. The aim was to assess how well the models, developed using a cross-platform dataset, perform with new and unseen data.

3. Methodology

3.1. Data Collection and Initial Exploration

Figure 1 presents a schematic flowchart of the research steps undertaken in this study. The initial step was collecting relevant datasets through web scraping. Datasets were sourced from platforms including Amazon, Yelp, and specific domains like hotel and restaurant reviews. The initial exploration involved assessing the structure, content, and review authenticity indicators within these datasets. This step was critical in planning the data cleaning and integration strategies to ensure the datasets' compatibility and relevance to the study's objectives.

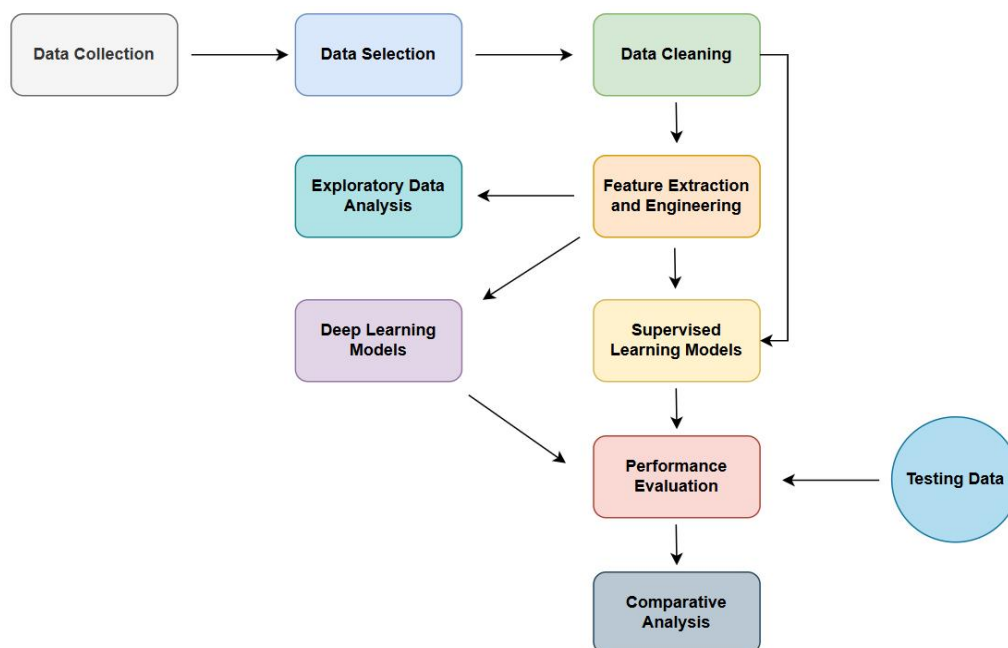


Fig.1. Framework for the proposed methodology

3.2. Data Cleaning Process

The data cleaning phase streamlined the datasets by pruning irrelevant or redundant columns, renaming columns for consistency, and standardizing review labels to 'F' for fake and 'R' for real. It also involved discarding rows without review text, essential for authenticity analysis, and removing duplicates to ensure data uniqueness.

3.3. Data Sampling and Merging

A balanced approach was applied in sampling real and fake reviews from each dataset, ensuring diversity in the final dataset. These sampled subsets were then merged, resulting in a dataset that represents various user-generated content across e-commerce platforms.

3.4. Feature Engineering

The initial step involved calculating straightforward yet informative attributes such as review text length, sentiment polarity, Flesch-Kincaid score, stopword, capitalization and punctuation count directly from the review texts. A custom function as "Product Name Mention" also was developed to detect the presence of the product's name within the review text. This aimed to assess the review's relevance and specificity towards the product.

The feature engineering process employed a Random Forest classifier with CountVectorizer to identify keywords critical for predicting review authenticity. The top ten influential words were visualized for relevance, as shown in Figure 2. From these, six words (great, good, like, love, really, quality) were chosen for their significant generic relevance across different datasets, leading to the creation of the IMPORTANT_WORDS_COUNT feature. This feature aggregates the frequency of these words within each review, providing a concise, predictive attribute.

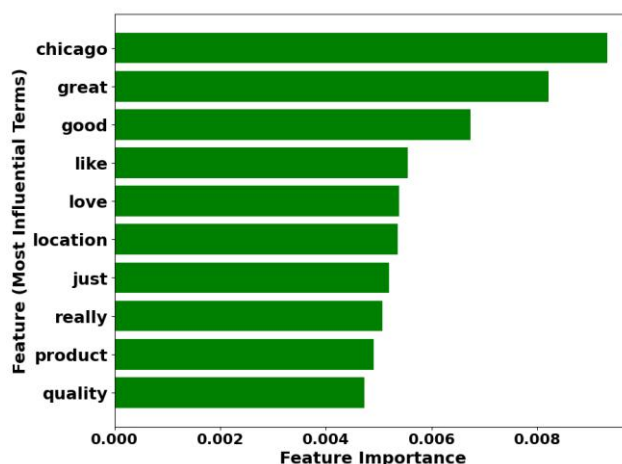


Fig.2. Top 10 most influential terms identified using a random forest classifier

Another novel feature, `RATIO_REPEATING_WORDS`, was engineered to quantify redundancy within review texts. This metric is calculated using the formula:

$$\text{Ratio of Repeating Words} = \frac{\text{Count of repeating words}}{\text{Total word count in a review}}$$

This proportion of words repeated more than once provides insight into the originality and potential spamminess of the content. By normalizing for the length of the review, this ratio offers a measure of repetitiveness that is unbiased by review length.

3.5. Dataset Preparation

The preprocessing steps applied to the dataset were crucial to ensure uniformity and enhance the model's ability to detect fake reviews effectively. These steps included removing non-word characters, converting the text to lowercase, and applying lemmatization and stop-word filtering. Lemmatization was chosen to reduce words to their base or root form, ensuring consistency across different variations of a word, which is essential for improving model learning. Stop-word filtering was applied to eliminate common but insignificant words that do not contribute meaningful information to the classification task. These choices were made based on standard practices in NLP, where they help reduce noise and enhance the focus on critical words that can affect model performance.

3.6. Model Training and Hyperparameter Tuning

For model training, a pipeline was established, incorporating preprocessing steps for numeric, categorical, and textual features (Figure 3). Building on existing literature that explores various machine learning approaches for detecting fake reviews, this study selected four models including KNN, SVM, Random Forest (RF), and Logistic Regression (LR) based on their strengths in different aspects of classification tasks. KNN was chosen for its simplicity and ability to classify based on proximity to known data points, while SVM was selected for its robust performance in high-dimensional spaces, making it effective for complex text data. Random Forest was included for its ensemble learning capability, which helps reduce overfitting by averaging multiple decision trees. Logistic Regression, a widely-used and interpretable model, was chosen for its efficiency in binary classification problems. These models were compared across a cross-platform dataset to determine which one offers the best balance between accuracy, adaptability, and computational efficiency for detecting fake reviews. [4-6].

The study employed GridSearchCV along with 5-fold cross-validation for several key purposes including optimal hyperparameter selection, model validation, performance stability, resource efficiency and model benchmarking.

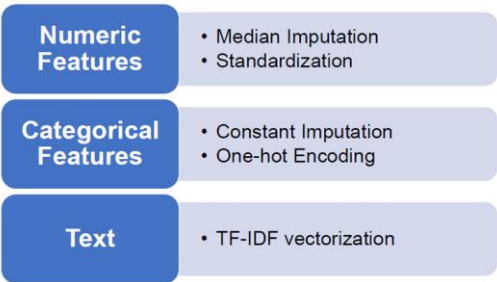


Fig.3. Preprocessing steps applied to different feature types in the pipeline

3.7. Performance Evaluation

After training, each model's performance was thoroughly evaluated using several key metrics: accuracy, precision, recall, and F1 score. Confusion matrices were also generated to provide detailed insights into each model's predictive accuracy across different classes, highlighting their specific strengths and weaknesses in identifying fake reviews.

To further validate the models, a separate test set comprising 400 fake and 400 real reviews was utilized to verify the models' performance on unseen data. This step ensured that the models could generalize well beyond the training data.

Additionally, a version of the dataset without the engineered features was subjected to the same preprocessing and modelling procedures. This allowed for a direct comparison of performance metrics to determine the efficacy of the feature engineering process.

3.8. Deep Learning Model Development

As the final step of this project, three primary neural network architectures were developed and evaluated: Long Short-Term Memory (LSTM), Convolutional Neural Network combined with LSTM (CNN-LSTM), and a hybrid LSTM-Recurrent Neural Network (LSTM-RNN). The choice of these methods was based on recent literature, as researchers have successfully employed these architectures for fake review detection [11, 12].

These deep learning methods were selected as basic approaches to compare their performance with the previously developed supervised models. This comparison aims to evaluate the effectiveness of deep learning techniques in detecting fake reviews and to determine if they offer significant advantages over traditional supervised learning methods.

4. Results and Discussion

The histograms in Figure 4 show similar rating distributions for fake and real reviews, highlighting the challenge of detecting fake reviews based solely on ratings, as fake reviews are crafted to closely mirror genuine ones.

Figure 5 presents word clouds for both fake and real reviews, illustrating the frequent use of similar terms in both categories. This visual similarity highlights the challenge in distinguishing fake reviews based purely on word frequency or common terms. The overlap suggests that fake reviews are crafted to resemble authentic ones closely, further complicating detection efforts. This finding underscores the need for more advanced techniques beyond basic term analysis, such as feature engineering and contextual evaluation, to improve the accuracy of fake review detection models.

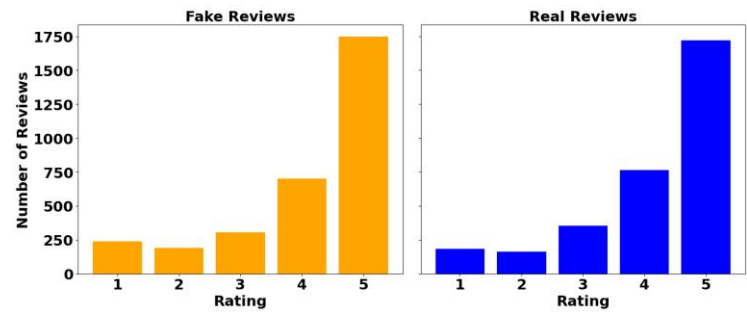


Fig.4. Distribution of rating for fake (left) and real (right) reviews



Fig.5. Word clouds for fake (left) and real (right) reviews

Figure 6 demonstrates the boxplots for stopwords, punctuation, and capitalization features, revealing higher occurrences in real reviews compared to fake ones. This pattern indicates that real reviews tend to exhibit a more conversational tone, characterized by a higher use of stopwords—common in natural speech—and greater punctuation and capitalization, suggesting expressive and less structured language. These linguistic nuances, more prevalent in real reviews, offer clues for distinguishing them from fake reviews, which tend to be more polished and formulaic. The differences in these features underscore the value of analyzing multiple linguistic characteristics when detecting fake reviews.

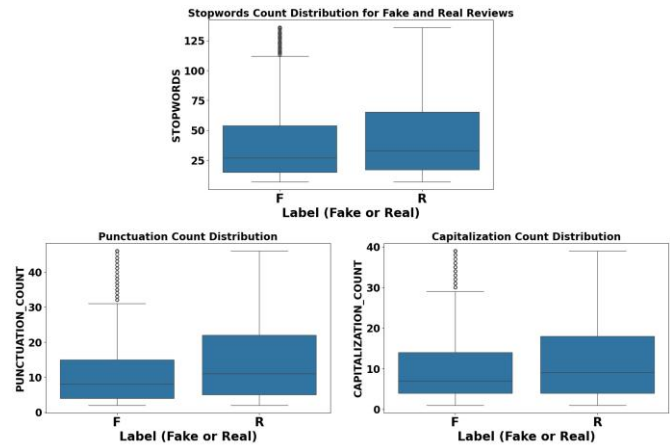


Fig.6. Boxplots for the distribution of stopwords, punctuation, and capitalization counts in fake and real reviews

The boxplots in Figure 7 illustrate that fake reviews tend to show higher sentiment polarity and a greater ratio of repeating words compared to real reviews. This trend suggests that fake reviews often exhibit exaggerated positivity to enhance their perceived authenticity and appeal. The increased repetition of certain words or phrases in fake reviews points to an intentional strategy to reinforce key terms, making the review appear more substantial. These tactics are often employed to give fake reviews a sense of credibility, but they can serve as markers for detecting inauthentic content.

Figure 8 depicts boxplots for review length and FK score (Flesch-Kincaid readability score), revealing that real reviews are generally longer and have higher FK scores compared to fake reviews. This suggests that real reviews tend to be more complex and varied in readability, often reflecting authentic experiences that require more detailed expression. In contrast, fake reviews are typically shorter, with simpler language and less variety in sentence structure, indicating a more formulaic or automated approach to review generation. The greater complexity and length of real reviews can serve as important differentiators when identifying fake content.

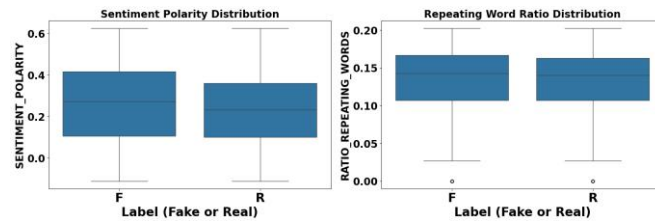


Fig.7. Boxplots for the distribution of sentiment polarity and the ratio of repeating words in fake and real reviews

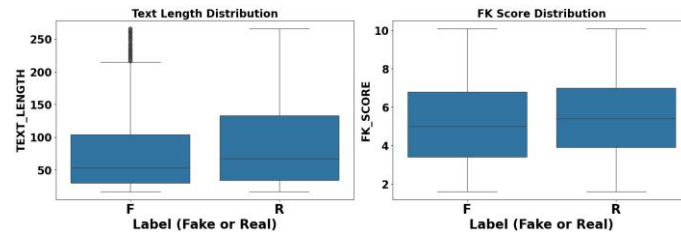


Fig.8. Boxplots for the distribution of text length and the FK score in fake and real reviews

The bar charts in Figure 9 reveal that fake reviews are more likely to mention the product name and use a higher frequency of important words. The higher occurrence of product name mentions suggests a deliberate tactic in fake reviews to mimic authentic review behavior by appearing more specific and detailed. Similarly, the increased count of important words in fake reviews indicates an effort to make the reviews sound convincing and credible, likely by strategically using persuasive or impactful terms. These features highlight how fake reviews often rely on specific strategies to enhance their legitimacy, providing valuable indicators for detection.

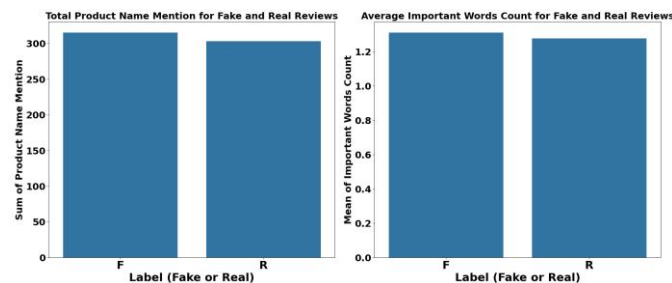


Fig.9. Bar charts for total product name mention and average important words count in fake and real reviews

The selection of the most appropriate models was conducted through GridSearchCV, employing a 5-fold cross-validation strategy to optimize the models' hyperparameters and assess their effectiveness across multiple data partitions. Table 1 shows the cross-validation accuracy and optimal hyperparameters for each of the supervised models.

Table 1. Cross-validation accuracy and optimal hyperparameters for supervised models

Model	Cross-Validation Accuracy	Optimal Hyperparameters
KNN	0.7637	7 Neighbors
SVM	0.7995	RBF Kernel, C = 1
RF	0.7952	200 Estimators, No Maximum Depth
LR	0.7894	Regularization Strength C = 0.1

Table 2 presents the initial performance metrics of the four supervised models. The SVM model emerged as the top performer, achieving the highest scores across all performance metrics. This highlights SVM's superior ability to effectively balance the trade-offs between different types of errors, which is crucial in real-world applications.

Figure 10 shows the confusion matrix results. SVM not only showed fewer false positives compared to the other models but also maintained a high number of true positives and true negatives, supporting its high-performance scores in Table 2. This suggests that SVM is particularly effective at managing non-linear patterns, which likely contributed to its superior performance on the test set. Logistic Regression (LR) also performed well, with a significant number of true positives and the fewest false positives, reflecting its high recall and precision scores. Both RF and KNN were competitive, though their performance was slightly below that of SVM and LR. RF excelled in identifying fake reviews with the highest number of true positives but had a higher rate of false negatives compared to LR and SVM.

Table 2. Initial model performance metrics

Model	Accuracy	Precision	Recall	F1 Score
KNN	0.7714	0.7724	0.7714	0.7714
SVM	0.8216	0.8235	0.8216	0.8216
RF	0.8122	0.8122	0.8122	0.8121
LR	0.8172	0.8179	0.8172	0.8173

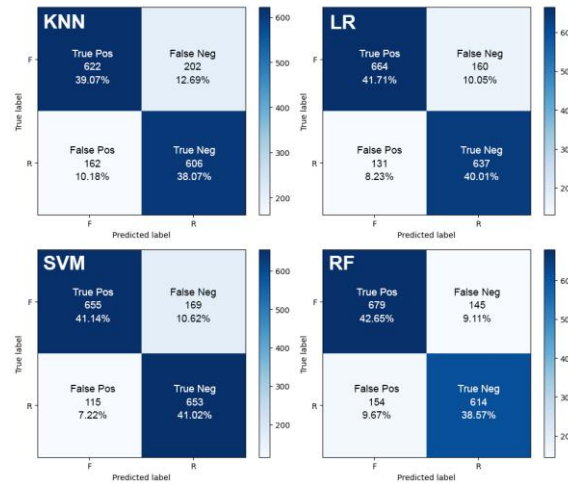


Fig.10. Confusion matrix results for supervised models

SVM's superior performance can be attributed to its ability to handle high-dimensional feature spaces, which makes it well-suited for text classification tasks where input data often involves numerous features derived from natural language processing. The model's use of kernel functions such as the radial basis function in this case allows it to find non-linear decision boundaries, which may be particularly useful when the separation between fake and real reviews is not easily defined in the input space.

Furthermore, SVM excels at balancing the trade-off between margin maximization and error minimization. By adjusting its regularization parameters (C), SVM effectively controls overfitting and underfitting, resulting in better generalization to unseen data. The other models may lack SVM's capacity to define complex decision boundaries or, in the case of KNN, struggle with the curse of dimensionality, particularly when dealing with a large number of features.

While Random Forest (RF) performed well, its ensemble structure may not be as efficient in capturing subtle distinctions between fake and real reviews as SVM. Logistic Regression (LR), while simpler and faster to train, relies on linear decision boundaries, which may not always be appropriate for the non-linear characteristics of the dataset. KNN, being sensitive to the choice of the number of neighbors (k) and distance metrics, may also struggle with noise and sparsity in the dataset.

The analysis of model performance before and after feature engineering provides key insights into its varied impact across different supervised models (Table 3). Feature engineering improved metrics for KNN, SVM, and Logistic Regression, with SVM showing the most significant gains of around 3%, suggesting enhanced discriminative power from the new features. However, the Random Forest model saw a slight decline in performance, likely due to noise or redundancy introduced by the new features, highlighting that feature engineering's effectiveness depends on the specific algorithm. These results underscore the importance of a tailored approach to feature engineering, as more features do not necessarily improve model performance.

Table 3. Relative percentage change in model performance metrics after feature engineering

Model	Accuracy (%)	Precision (%)	Recall (%)	F1 Score (%)
KNN	0.4926	0.1813	0.4926	0.6871
SVM	3.0852	3.2829	3.0852	3.0635
RF	-0.8426	-0.8426	-0.8426	-0.8426
LR	0.4672	0.4547	0.4672	0.4672

The supervised models were tested on a new dataset, with results showing a slight decline in performance compared to initial metrics, as expected with the introduction of a new and unseen dataset featuring different product types. However, Figure 11 illustrates the models' strong ability to generalize to new data, suggesting that they can be effectively applied across different e-commerce platforms.

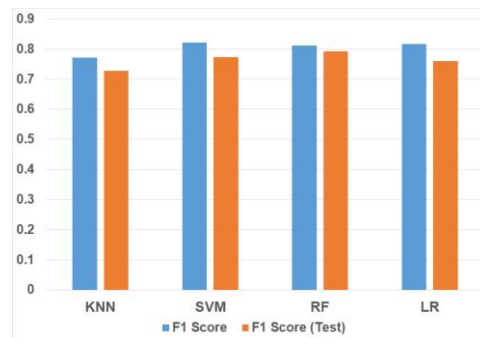


Fig.11. Comparison of initial and test set F1 scores for the four supervised models

Figure 12 shows that deep learning models (LSTM, CNN-LSTM, and LSTM-RNN) did not outperform the SVM model in this application. The LSTM model achieved an F1 Score of 0.6589, the CNN-LSTM improved slightly to 0.6849, and the LSTM-RNN reached 0.6025. These scores fall short compared to SVM's superior performance.

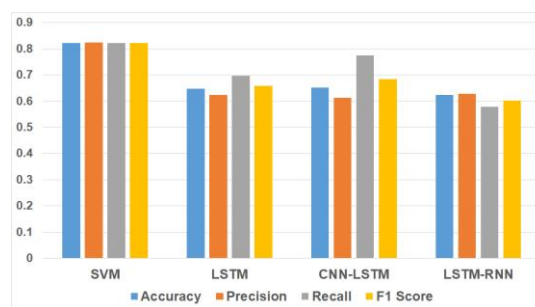


Fig.12. Comparison of performance metrics for SVM and deep learning models

The performance disparity between deep learning models and simpler supervised models can be attributed to several key factors. While deep learning models have powerful feature extraction capabilities, they generally require significantly larger and more diverse datasets to fully realize their potential, particularly when working with raw text data without the aid of explicit feature engineering [11, 13]. In this study, the dataset, though diverse in its sources, was relatively small and homogenous compared to the vast volumes typically required for deep learning to excel. This limitation likely impacted the ability of the deep learning models to effectively learn complex patterns from the data.

Supervised models, by contrast, were better aligned with the dataset's size and characteristics. The feature engineering process provided well-defined, relevant inputs that improved the performance of these models. Supervised approaches excel in leveraging explicitly defined features tailored to the problem at hand, allowing them to generalize more effectively on the test data compared to deep learning models. This advantage became evident in the context of this project, where simpler models were able to strike a balance between performance and resource efficiency.

Additionally, deep learning models are prone to overfitting, particularly when extensive hyperparameter tuning and regularization are not applied. Without sufficient tuning, these models can memorize patterns in the training data rather than generalize to unseen data, leading to poorer test performance [14]. This contrasts with the supervised models, which, benefiting from feature engineering, avoided overfitting and achieved better generalization on the test dataset.

These findings underscore the ongoing relevance and efficiency of supervised learning models in tasks such as fake review detection. While deep learning models hold promise, particularly for large-scale and unstructured datasets, their deployment in smaller, more focused datasets requires careful consideration of factors such as dataset size, hyperparameter tuning, and computational costs. This study highlights the advantages of supervised models in terms of both performance and interpretability, suggesting that for certain tasks, simpler approaches may offer a more effective and efficient solution. Looking ahead, this research provides a foundation for future explorations aimed at optimizing the balance between model complexity and real-world applicability in the fight against fake reviews.

5. Conclusions

This study highlights the strengths of machine learning models in detecting fake reviews, emphasizing the importance of cross-platform datasets, feature engineering, and model adaptability. The SVM model consistently outperformed others, showing superior precision and flexibility, particularly when combined with optimized kernel functions and regularization parameters. While deep learning shows potential, simpler supervised models like SVM demonstrated greater efficacy with smaller, more focused datasets. The role of feature engineering varied, offering significant advantages for some models (e.g., SVM) but adding unnecessary complexity for others like RF.

The findings from this study have practical implications for e-commerce platforms in their efforts to combat fake reviews. The high accuracy of models like SVM suggests they could be effectively deployed for real-time review monitoring, flagging suspicious reviews as they are posted. By integrating such models into existing moderation systems, platforms could enhance their ability to detect fake reviews alongside human oversight.

User feedback could be incorporated to continuously retrain and refine the models, improving detection accuracy over time. The deployment of these models not only helps maintain the credibility of product reviews but also mitigates the economic impact of fake reviews, preserving consumer trust.

For future research, several opportunities exist to advance fake review detection. Expanding dataset diversity to include various platforms and product categories can improve model generalizability. Incorporating metadata such as reviewer history, product characteristics, and temporal review patterns could enhance accuracy. Additionally, exploring hybrid models that combine supervised learning with deep learning may boost performance by leveraging the strengths of both approaches.

References

- [1] Alsubari, S.N., Deshmukh, S.N., Alqarni, A.A., Alsharif, N., Aldhyani, T.H.H., Alsaade, F.W., & Khalaf, O.I. (2021). Data analytics for the identification of fake reviews using supervised learning. *Computers, Materials & Continua*, 64, 1-19.
- [2] Abi Priya, M., Hema, S., & Dhivya Praba, R. (2020). Fake product review monitoring and removal for genuine online product review using IP address tracking. *International Journal of Advanced Research in Electrical, Electronics and Instrumentation Engineering*, 9, 1-10.
- [3] Salminen, J., Kandpal, C., Kamel, A.M., Jung, S.-g., & Jansen, B.J. (2022). Creating and detecting fake reviews of online products. *Journal of Retailing and Consumer Services*, 64, 102771.
- [4] Nguyen, N. (2023). An empirical study on fake review detection. *Travinh University Journal of Science*, 13(Special Issue), 1-10.
- [5] Sajid, T., Jamshed, W., Kumar Goyal, N., Keswani, B., Cieza Altamirano, G., Goyal, D., & Keswani, P. (2023). Analysis and challenges in detecting the fake reviews of products using Naïve Bayes and Random Forest techniques. 1-20.
- [6] Mohawesh, R., Xu, S., Tran, S.N., Ollington, R., Springer, M., Jararweh, Y., & Maqsood, S. (2021). Fake reviews detection: A survey. *IEEE Access*, 9, 1-20.
- [7] Mewada, A., & Dewang, R.K. (2022). Research on false review detection methods: A state-of-the-art review. *Journal of King Saud University – Computer and Information Sciences*, 34, 7530-7546.
- [8] Zaki, N., Krishnan, A., Turaev, S., Rustamov, Z., Rustamov, J., Almusalami, A., Ayyad, F., Regasa, T., & Iriho, B.B. (2023). Node embedding approach for accurate detection of fake reviews: A graph-based machine learning approach with explainable AI. 1-15.
- [9] Soldner, F., Kleinberg, B., & Johnson, S.D. (2022). Confounds and overestimations in fake review detection: Experimentally controlling for product-ownership and data-origin. *PLoS ONE*, 17(12), 1-15.
- [10] Qayyum, H., Ali, F., Nawaz, M., & Nazir, T. (2023). FRD-LSTM: a novel technique for fake reviews detection using DCWR with the Bi-LSTM method. *Multimedia Tools and Applications*, 82, 31505-31519.
- [11] Mohawesh, R., Al-Hawawreh, M., Maqsood, S., & Alqudah, O. (2023). Factitious or fact? Learning textual representations for fake online review detection. *Cluster Computing*, 1-15.
- [12] Vyas, P., Liu, J., & El-Gayar, O. (2021). Fake news detection on the web: An LSTM-based approach. *Americas Conference on Information Systems (AMCIS)*, 1-10.
- [13] Dean, J., Corrado, G., Monga, R., Chen, K., Devin, M., Le, Q.V., Mao, M.Z., Ranzato, M., Senior, A., Tucker, P., Yang, K., & Ng, A.Y. (2012). Large scale distributed deep networks. *Advances in neural information processing systems*, 25, 1223-1231.
- [14] Srivastava, N., Hinton, G., Krizhevsky, A., Sutskever, I., & Salakhutdinov, R. (2014). Dropout: A simple way to prevent neural networks from overfitting. *The journal of machine learning research*, 15(1), 1929-1958.

Authors' Profiles



Faryad Bigdeli is a recent graduate with a Master's degree in Data Science, specializing in machine learning and natural language processing. With a solid foundation in data analysis and model development, he has focused on addressing the challenges of fake review detection in e-commerce platforms. In his research, he explored the effectiveness of various machine learning models in detecting fake reviews across different datasets. His work highlights the practical implications of integrating machine learning models into real-time review monitoring systems to enhance the detection of fraudulent reviews while preserving consumer trust. Faryad Bigdeli is passionate about leveraging data science to create impactful solutions in the digital landscape and continues to pursue opportunities for research and application in this evolving field.

How to cite this paper: Faryad Bigdeli, "Cross-platform Fake Review Detection: A Comparative Analysis of Supervised and Deep Learning Models", *International Journal of Information Technology and Computer Science(IJITCS)*, Vol.17, No.3, pp.52-60, 2025. DOI:10.5815/ijitcs.2025.03.04