

Identification of Real and Fake Reviews Written in Roman Urdu

Asif Ahmed, Dr. Irfan Bacho, Dr. Shah Nawaz Talpur

Dept. of Computer Systems Engineering Mehran University of Engineering and Technology
Jamshoro, Pakistan

*Correspondence: aasif5056@gmail.com

Citation | Ahmed. A, Bacho. I, Talpur. S, “Identification of Real and Fake Reviews Written in Roman Urdu”, IJIST, Vol. 5 Issue. 4 pp 787-797, Dec 2023

DOI | <https://doi.org/10.33411/IJIST/603>

Received | Dec 04, 2023 **Revised** | Dec 17, 2023 **Accepted** | Dec 23, 2023 **Published** | Dec 28, 2023

The evolution of e-commerce has made reviews a crucial metric to judge the quality of online products or services. These reviews have a significant impact on the decision of the customer. Positive review catches more attraction while negative reviews impact sales of the product. Nowadays, deceptive reviews are being deliberately posted on e-commerce websites and social media stores to promote the product by illegal means. These reviews are sometimes posted in different local languages to build a fake virtual reputation among local customers. Thus, fake review detection is a wider area for ongoing research. This paper proposes several machine-learning approaches for the detection of fake reviews written in Roman Urdu. Furthermore, a comparative analysis of the performance of nine machine learning models on the given dataset is performed. The dataset is crawled from different e-commerce sites in Pakistan. The results show that the existing Support Vector Machine outperforms the rest of the models with an accuracy of 82%. This emphasized the critical role of robust detection methods in ensuring the authenticity of online reviews.

Keywords: Fake Reviews, Deep Learning Algorithms, E-Commerce, Support Vector Machine, Comparative Analysis.



Introduction:

E-commerce is based on the purchase and selling of products or e-services on the Internet. This mode of business has become gradually popular because of the convenience it offers to its customers, allowing them to shop from anywhere and at any time. One of the key aspects of e-commerce is customer reviews, which play a vital role in influencing purchasing decisions. A surge in customer reviews has been noticed to influence the buyers' selection [1][2]. Positive reviews can boost sales and increase customer loyalty, while negative reviews can have the opposite effect, leading to a loss of trust and potential revenue. As such, businesses must pay close attention to customer feedback and respond accordingly to ensure their e-commerce ventures are successful [3][4][5]. It is important to be aware that there is a risk of spammers posting fake reviews to promote a product through illegal means [6][7]. The reviews posted without encountering the items are known as fake reviews [8]. The one who posts a fake review is known as a spammer [8]. While a group of spammers having a similar target is referred to as a group of spammers [8].

In Pakistan, it is common for reviews to be written in Roman Urdu, a script that uses the Roman alphabet to write the Urdu language. This is because while Urdu is the national language of Pakistan, many Pakistanis are more comfortable writing in Roman script as it is easier to type and use on digital platforms. This trend is particularly noticeable on popular e-commerce websites in Pakistan, like Daraz.pk, where customers leave reviews in Roman Urdu to share feedback and insights on services and products. By using this script, reviews become more accessible to a wider audience and help promote greater understanding and communication among diverse groups of people [9]. There is always a possibility of fake reviews being posted across any platform, including e-commerce websites in Pakistan. This is a problem that is not limited to any specific language or country. In Pakistan, companies or e-commerce store owners may hire a spammer team to post reviews in Roman Urdu to misguide customers. However, it is important to note that e-commerce platforms must have measures in place to detect and remove fake reviews [2].

To sum up, detecting fraudulent reviews is a crucial and multifaceted concern that impacts both enterprises and customers. By devising techniques and algorithms to detect bogus reviews in Roman Urdu, this investigation has the potential to boost the authenticity and dependability of online reviews in this language. It could also help guarantee that people can access precise and reliable details.

This research gives a novel approach to detecting fake reviews written in Roman Urdu, a language widely spoken in South Asia and specifically in the context of Pakistani e-commerce. While existing studies have primarily focused on fake review detection in English, our work pioneers the application of advanced machine learning algorithms to address this issue in Roman Urdu. A key distinctive feature is the utilization of the newest and most comprehensive dataset gathered through meticulous crawling of e-commerce platforms. In the preprocessing phase, we applied TF-IDF on Bag-of-Words (BoW) representations, enhancing the purity and relevance of the dataset for subsequent machine learning algorithms. By venturing into this unexplored territory and employing state-of-the-art techniques, this research not only broadens the scope of fake review detection but also provides valuable insights for enhancing trust and reliability in online product reviews within the Roman Urdu domain.

Related Work:

In this research, two major features have been found; behavioral and textual features. The behavioral feature depends on the nonverbal attributes of the review. The main dependence of these types of features is behavior, such as emotional expressions, writing style, and the frequent times' the reviewer has written the review. While the textual feature is comprised, of the content written in the review [10]. However, investigating both features is crucial and challenging. Textual as well as behavioral reviews have a substantial influence on the

effectiveness of detecting fake reviews [11]. It seems challenging when these reviews are posted in local languages. Several languages or language translation tools can be used by spammers to post reviews in their local languages [12][13]. Research has been conducted to tackle these multilingual fake reviews as well. Textual and behavioral features have been found to play a significant role in the determination of multilingual reviews.

Li et al. [14] experimented with real-time reviews of 500 restaurants written in the Chinese language. The author used a model independent of language features. Heydari et al. [15] attempted to detect fake reviews by the analyses of linguistic features in the review. Three techniques are used by Ott et al. [16] to perform classification. The techniques are- text categorization, psycholinguistic deception, and genre identification. More linguistic features are also explored. Feng et al [17] used unlexicalized and lexicalized syntactic features. They constructed sentence parse trees to detect fake reviews. Their experiment showed that the accuracy of prediction can be improved by deep syntactic features. In [18], the author used five classifiers to detect the deception of review which are Naïve Bayes, SVM, KNN, decision tree, and K-star. The experiments have been performed on the labeled movie reviews dataset [19]. Also, in [20], the author used a Decision Tree, Naïve Bayes, SVM, Random Forest, and Maximum Entropy on the collected dataset. The dataset consists of over 10,000 negative tweets for Samsung services and products. In [21], the author utilized both Naïve Bayes classifiers and SVM. The dataset has been collected from Yelp which comprises 1600 reviews gathered from well-known hotels in Chicago.

Furthermore, several papers are based on behavioral features in the process of fake review detection. Such as, in [22], multiple behavioral features such as the ratio of the review count and average timing have been considered. In a different work [23], the author examined the influence of textual and behavioral features on the process of fake review detection. The author mainly focused on the hotel domain and restaurants.

Based on the previous discussion and information available to us, there has been no appropriate research conducted on detecting fake reviews in Roman Urdu, nor has a dataset been found containing reviews written in Roman Urdu. Our paper aimed to bridge this gap by crawling a dataset of 7710 Roman Urdu reviews from a popular e-commerce website, daraz. pk. Nine machine learning models were utilized to conduct experiments on the dataset, and text processing techniques were implemented to get the improved accuracy of the results. The text processing techniques included labeling, removal of characters and punctuation, transforming text to lowercase, elimination of stop words, stemming, lemmatizing, and applying TF-IDF on Bow. The use of these techniques helped in smoothing out the text and improving the accuracy of the results.

Objective:

This research aims to investigate and find effective methods for detection of the fake reviews written in Roman Urdu on e-commerce sites in Pakistan. The primary objectives include crawling and preprocessing a dataset of Roman Urdu reviews to establish a reliable model capable of distinguishing between fake and genuine reviews. Furthermore, the study involves the evaluation of several machine learning models on the preprocessed dataset to determine the most proficient approach. By achieving these objectives, this research contributes to the ongoing efforts to develop robust solutions for combating deceptive online reviews, particularly in the unique context of the Roman Urdu language.

Methodology:

In this research, a dataset has been crawled from an e-commerce site in Pakistan containing 7710 reviews written in Roman Urdu. Here we have each review containing a rating and a label i.e. OR (Human-generated original review) and F (Fake computer-generated review), category of product, and the text of the review. To achieve better results and find the outperforming model, the approach consists of four phases; review collection, pre-processing,

training classifier, and testing reviews to be classified. Figure 1 illustrates the proposed framework.

In our methodology, we employed a diverse set of nine machine learning algorithms to comprehensively analyze their effectiveness in detecting fake reviews. These algorithms include Support Vector Machine, Naïve Bayes, Logistic Regression, Decision Tree, K-Nearest Neighbors, Random Forest, Recurrent Neural Network (RNN), Convolutional Neural Networks (CNN), and Long Short-Term Memory (LSTM). The reason behind this selection was to ensure a robust evaluation, considering the varied strengths and capabilities of each algorithm. Naïve Bayes is known for its simplicity and efficiency, while Support Vector Machine excels in handling complex data. Logistic Regression provides a probabilistic approach, and Decision Tree and Random Forest offer ensemble learning. K-Nearest Neighbors utilizes proximity-based classification. The inclusion of deep learning models such as RNN, CNN, and LSTM allows for a nuanced exploration of sequential patterns in reviews. This comprehensive array of algorithms enables a thorough comparison, ensuring a nuanced understanding of their performance in the context of fake review detection.

Dataset Collection:

We collected a total number of 7710 reviews from Daraz. pk which is well well-known e-commerce site, widely used in Pakistan. The reviews that we collected contain Roman Urdu. Each review contains, the reviewer's username, date, and the category of the product. We have utilized 30% of the dataset for training purposes, while the remaining is reserved for testing. Out of these reviews, we have labeled 1710 reviews as fake out of which, 710 were deliberately written by us. The rest of the reviews are considered real.

Preprocessing:

The first phase in this approach starts with the preprocessing of the dataset [22] an essential step in the approach of machine learning. The step of data processing is so important because of the roughness of a real-world dataset like the usage of stop words, commas, numbers, etc. Sequential steps of preprocessing data have been applied to make raw data from the crawled dataset for several activities involving computation. The steps are given as:

Tokenization:

One of the frequently encountered common and important techniques in natural language processing is tokenization. It is commonly regarded as a fundamental step before applying any processing technique. The words in the text review are divided separately into "tokens". For example, if the review is "The product is amazing", after tokenization it will appear as ("The", "Product", "is", "amazing") [24].

Cleaning of Stop Words:

Stop words [23] are supposed to hold no value in the process of fake review detection. Some typical instances of stop words include; a, an, this, the, or. We cleaned all stop words to reduce the unnecessary word count and to make the process of fake review detection smooth.

Removal of Non-Alphabetic Words:

This step is used for removing punctuation marks, full stops, and digits from the reviews. The non-alphabetic words have no major impact on the processing of raw data.

Feature Extraction:

The unnecessary attributes can cause reduced accuracy of the model [25]. Therefore, feature extraction is performed to improve the efficiency of pattern recognition in machine learning systems. It is mainly used in the reduction phase of the dataset, which molds it into useful features that are fed into deep learning and machine learning models. Afterward, the data appears to be more meaningful.

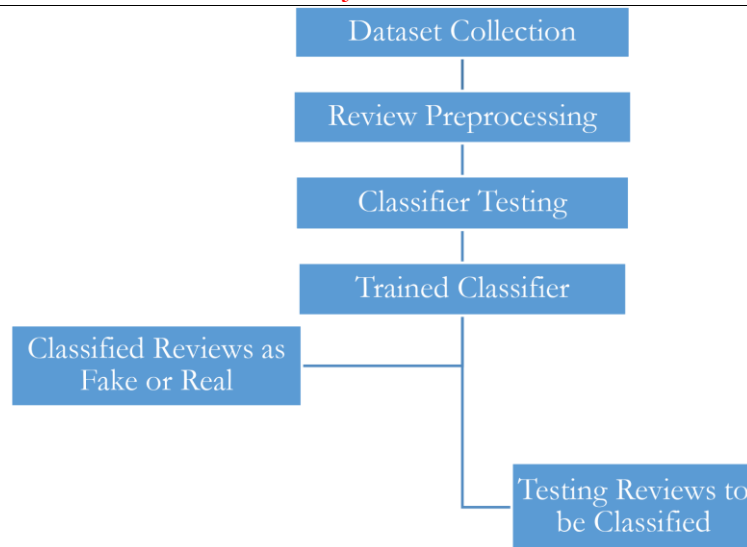


Figure 1: Proposed framework diagram

Several approaches in the literature have been developed to extract features from reviews. One of the popular approaches is textual features [26]. The textual feature contains sentimental classification [5] which depends upon negative and positive stances in the review e.g. “weak”, “good” or “strong”. Sometimes, the similarity of Cosine is used as a feature that refers to the cosine of the angle formed between two n-dimensional vectors and the result obtained by dividing the dot product of two vectors by the product of their magnitudes [27]. Another feature is true-false identification, also referred to as TF-IDF. This method counts the frequency of TF (True and false) and IDF (inverse document) scores for all words. The product of IDF and TF scores of the term is referred to as the weight of the TF-IDF term [28]. A confusion matrix is also used for the classification of reviews across four results; FP (Fake Positive), TP (True Positive), FN (Fake Negative), and TN (True Negative).

Equivalently, other features are behavioral and user personal profiles. These features can be used to identify spammers based on frequent time stamps or if the user is using redundant reviews and posting them everywhere regardless of the domain of the product. In this work, we used BoW (Bag of word features), also referred to as n-gram features, with TF-IDF. BoW features represent text into numerical representation i.e. the vector of numerical features that can be used to train the models. The features represent the frequency of each word appearing in the whole dataset of reviews. TF-IDF on the other hand is used to score each instance in the review text to make the detection more accurate.

Feature Engineering:

In this research, we combined TF-IDF with BoW to extract more accurate features from the vocabulary built by BoW and used these features to train multiple classifiers. In our crawled dataset, we have 7710 reviews and 7486 total vocabularies. TF calculates how often a word occurs across the document and the IDF is the measure of how important the word is. With the use of IDF, the unnecessary words in our dataset with 0 scores were reduced and we obtained the words with higher scores. A high score refers to more importance. Finally, as the product of IDF and TF, we obtained the TF-IDF score of the BoW. These features as a product of TF and IDF on BoW are taken into consideration to identify fraudulent reviews and conduct performance analysis on machine learning models.

Model Evaluation and Hyperparameter Tuning:

In this phase, following the preprocessing of our dataset through techniques e.g. the Bag-of-Words (BoW) and The Term Frequency-Inverse Document Frequency (TF-IDF), we subjected the refined dataset to a rigorous evaluation using nine distinct machine learning

models. These models, including Naïve Bayes, Support Vector Machine, Logistic Regression, Random Forest, Decision Tree, K-Nearest Neighbour, Recurrent Neural Network (RNN), Convolutional Neural Networks (CNN), and Long Short-Term Memory (LSTM), were chosen to comprehensively assess their efficacy in fake review detection. To optimize their performance, we systematically adjusted their hyperparameters, seeking the configurations that would yield the highest accuracy in our specific context. This meticulous process ensures a robust evaluation of each model's capabilities, facilitating an informed comparison of their effectiveness in identifying counterfeit reviews.

Naïve Bayes Machine Learning Probabilistic Algorithm:

Our initial trial involved implementing the Naïve Bayes machine learning probabilistic algorithm, which relies on principles of the Naïve Bayes theorem and is used widely in classification tasks. We have tuned the hyperparameter alpha to 0.7 to handle the problem of zero frequency. Tuning this hyperparameter greatly improved the results and we found an f1 score of 0.73.

The Decision Tree Classifier:

The tree-like structure efficiently portrays a collection of decisions and their potential outcomes, making it effective for resolving decision-based challenges. In our situation, we have chosen a depth of three strikes which is a good balance between model complexity and accuracy and is not too shallow or too deep. We observed that a decision tree with a shallow depth is more appropriate for smaller datasets and prior knowledge also suggests that a depth of three is sufficient to capture the relevant patterns and relationships in the data. It obtained an f1 score of 72.

Random Forest Classifier:

After the decision tree, we tested by increasing the number of trees, which is the case in a random forest classifier. We set the number of trees to 80 which is an approximately optimal value for our dataset. We also observed that the accuracy at 100 and 80 trees was the same, but the training time was reduced by 0.51 when we selected 80 trees. The model produces the f1 score of 0.75.

K-Nearest Neighbour:

The algorithm being described is a lazy learning and non-parametric approach that uses a simple distance metric to find the K-nearest neighbors to a given data point. We tried several K values for cross-validation and found that setting $K = 3$ balances the trade-off between bias and variance. This value made the model more flexible with low bias and high variance. With $k=3$, it obtained an f1 score of 0.75.

Support Vector Machine (SVM):

It finds the best possible decision boundary (or hyperplane) which is used to separate the data points among different classes. SVM can also handle dimensional data and is effective in dealing with nonlinear classification problems by using regularization parameters (C), and kernel tricks. We tuned C to 0.6, which is a good balance between bias and variance, and used Polynomial Kernel as a kernel trick. Tuning these hyperparameters got us an f1-score of 0.82.

Logistic Regression:

Logistic regression aims to predict a categorical outcome variable that is based on a minimum of one input variable or feature. We have L2 as the penalty and made 100 iterations for the solver. After tuning the value of C to 1.0, the model produces 0.77 as the f1 score.

Neural Networks:

In this phase, we examined three deep neural networks to find the most outperforming model. Deep learning methods, such as most representative neural networks, outperform traditional machine learning models in efficiently extracting valuable data features. The dataset undergoes testing using the following deep-learning models:

Convolution Neural Networks (CNN):

In the natural language processing task, CNN plays a momentous role in pulling in local features for classification. We have used 3 convolution layers with 32, 3 x 3 size filters, and a stride of 1. By tuning the hyperparameters to these given values and using ReLU (Rectified Linear Unit), which is an activation function, the model produces an f1-score of 0.75 on our dataset.

Recurrent Neural Network (RNN):

In comparison to conventional feed-forward neural networks, which process the data in one pass, RNNs can also keep a state or memory of past inputs, which makes them well-suited for sequential data. We used 32 hidden layers, a learning rate of 0.001, and tanh (hyperbolic tangent) as an activation function. The model obtained an f1-score of 0.72.

Long-short-term memory (LSTM):

It is an advancement of the RNN model. We used this model because of the long-term dependency problem in RNN. LSTM networks are designed to handle long-term dependencies and are capable of learning and remembering information over long periods. We have used 32 LSTM layers, 64 batch sizes, and a hyperbolic tangent (tanh) function as an activation function that produces an f1-score of 0.77, which is better than RNN.

Results and Performance Analysis:

We used six traditional machine learning algorithms; Support Vector Classifier, Logistic Regression, Decision Tree Classifier, K Nearest Neighbors, Multinomial Naive Bayes, Random Forests Classifier, and three neural networks; CNN, RNN, and LSTM, on our dataset. Two predefined labels are given to our dataset; F (Fake) which refers to fake reviews and OR (Original) which is real and written by a human. We have divided the dataset into 1710 F and 6000 OR reviews. Each review instance in the dataset contains a rating, Label (F or OR), and category of the product. Table 1 summarizes the statistics of the dataset.

Table 1: Statistics of the dataset

Total reviews	7710
Number of F reviews	6000
Number of OR reviews	1710
Maximum length of review	107
Minimum length of review	4
The average length of review	87

Before passing the dataset through the algorithms, we applied several text-pre-processing techniques to the dataset to make it free from roughness. After pre-processing, the code produces a new pre-processed XML file which is further utilized in the experiments. Counter Vectorizer takes place in the next step which converts the text to a number in terms of the sparse matrix. We then made BoW for our document to be used for further feature extraction from the dataset. To get better features, we further used TF-IDF with BoW. From TF-IDF we obtained the score of each instance in the review and it greatly improved the results.

We trained the classifiers using 30% of our dataset and 70% for testing. After training, we tested it with nine machine-learning models and generated the classification report for each which is further discussed in the results section.

Results:

Going forward in the results section, we present our thorough experiment results after running nine machine learning models on the dataset and analyzing the model performance by using the given metrics; recall, precision, and f1-score. For training purposes, 30% of the data has been utilized while 70% has been allocated for testing. The performance of all models is shown in Figure 2.

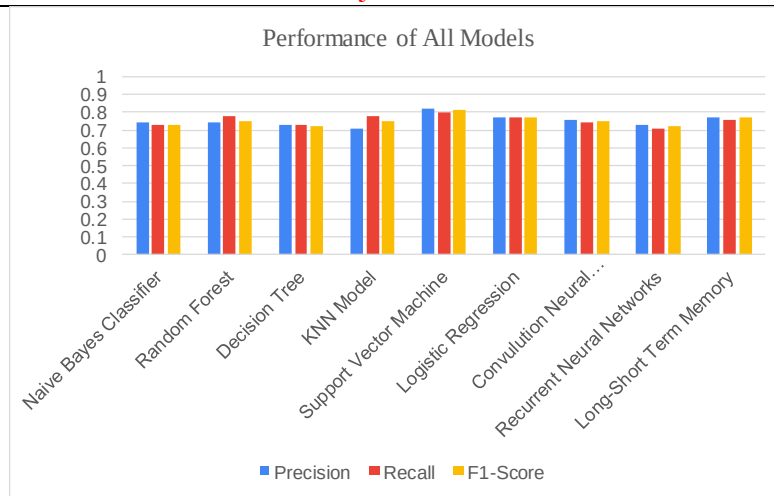


Figure 2: Performance of All Models

Performance Analysis and Discussion:

Neural networks are widely regarded as the most efficient and powerful models in the field of machine learning. On our dataset, neural networks did not outperform traditional machine learning algorithms as their performance directly relates to the quantity of the dataset. Producing an f1-score of 0.82, the SVM demonstrated superior performance compared to the other models. Figure 2 shows a comparison of the f1-scores of all the models. In our research, we have chosen the most processed dataset by performing TF-IDF on BoW to obtain well-related instances to minimize overfitting. By tuning preprocessing and applying TF-IDF on BoW, we were able to improve the accuracy and found the Support Vector Machine as the most outperforming model compared to the rest.

Table 2: Summarized classification reports of all the models.

Model	Labels	Precision	Recall	F1-Score	Support
Naive Bayes Classifier	F	0.71	0.79	0.74	421
	OR	0.77	0.7	0.73	629
	Accuracy			0.73	1050
	Macro Avg	0.74	0.74	0.73	1050
	Weighted Avg	0.74	0.73	0.73	1050
Random Forest	F	0.72	0.77	0.74	421
	OR	0.75	0.71	0.72	629
	Accuracy			0.73	1050
	Macro Avg	0.73	0.74	0.73	1050
	Weighted Avg	0.73	0.73	0.72	1050
Decision Tree	F	0.7	0.79	0.74	421
	OR	0.77	0.78	0.77	629
	Accuracy			0.75	1050
	Macro Avg	0.73	0.78	0.75	1050
	Weighted Avg	0.74	0.78	0.75	1050
KNN Model	F	0.63	0.77	0.7	421
	OR	0.76	0.78	0.77	629
	Accuracy			0.68	1050
	Macro Avg	0.7	0.77	0.74	1050
	Weighted Avg	0.71	0.78	0.75	1050
Support Vector Machine	F	0.81	0.79	0.8	421
	OR	0.83	0.81	0.82	629
	Accuracy			0.78	1050

		Macro Avg	0.82	0.8	0.81	1050
		Weighted Avg	0.83	0.81	0.82	1050
Logistic Regression		F	0.77	0.75	0.76	421
		OR	0.76	0.77	0.77	629
		Accuracy			0.71	1050
		Macro Avg	0.77	0.76	0.77	1050
		Weighted Avg	0.77	0.77	0.77	1050
Convolution Neural Network		F	0.78	0.71	0.75	421
		OR	0.73	0.76	0.75	629
		Accuracy			0.76	1050
		Macro Avg	0.76	0.73	0.75	1050
		Weighted Avg	0.76	0.74	0.75	1050
Recurrent Neural Networks		F	0.73	0.71	0.72	421
		OR	0.72	0.7	0.71	629
		Accuracy			0.72	1050
		Macro Avg	0.73	0.7	0.72	1050
		Weighted Avg	0.73	0.71	0.72	1050
Long-Short Term Memory		F	0.76	0.75	0.76	421
		OR	0.77	0.76	0.77	629
		Accuracy			0.75	1050
		Macro Avg	0.77	0.75	0.77	1050
		Weighted Avg	0.77	0.76	0.77	1050

Conclusion:

In this paper, we emphasized the significance of reviews and their potential to influence buyers' decisions. The role of product reviews in online purchases cannot be underestimated. Thus, this area of research is vivid and ongoing. This paper introduces a variety of machine-learning algorithms aimed at identifying counterfeit reviews composed in Roman Urdu. To accomplish this, we utilized a dataset comprising reviews written in Roman Urdu, gathered by crawling an e-commerce website based in Pakistan. We combined BoW and TF-IDF to extract more relatable features and observed better accuracy. Also, the hyperparameters of all models are tuned to get the maximum possible accuracy. It is revealed from the results that the Support Vector Machine outperforms the rest machine-learning models and has the highest f1-score of 0.82. The results are summarized in figure 2 and classification reports are summarized in table 2.

Future Work:

Given the foundational insights obtained from our research, future investigations should extend into the practical applications and industry implications of our findings. Exploring the integration of our optimized machine learning model, particularly the superior Support Vector Machine (SVM), into real-world e-commerce platforms could enhance the authenticity of online product reviews. Further research could also focus on adapting the model to various linguistic contexts beyond Roman Urdu, broadening its applicability. Continuous refinement of the model through ongoing monitoring and parameter tuning would ensure its sustained effectiveness in the dynamic landscape of fake review detection. Additionally, collaborative efforts with e-commerce platforms and regulatory bodies may facilitate the implementation of our model as a proactive tool in maintaining the integrity of customer feedback systems. Overall, future endeavors should align with the practical implementation and continuous improvement of our detection methodology in real-world scenarios.

References:

- [1] R. Xu, Y. Xia, K.-F. Wong, and W. Li, "Opinion Annotation in On-line Chinese Product

- Reviews.” 2008. Accessed: Dec. 25, 2023. [Online]. Available: http://www.lrec-conf.org/proceedings/lrec2008/pdf/415_paper.pdf
- [2] A. K. Samha, Y. Li, and J. Zhang, “Aspect-Based Opinion Extraction from Customer Reviews,” pp. 149–160, Apr. 2014, doi: 10.5121/csit.2014.4413.
- [3] I. Peñalver-Martinez et al., “Feature-based opinion mining through ontologies,” *Expert Syst. Appl.*, vol. 41, no. 13, pp. 5995–6008, Oct. 2014, doi: 10.1016/J.ESWA.2014.03.022.
- [4] A. M. Popescu and O. Etzioni, “Extracting product features and opinions from reviews,” *Nat. Lang. Process. Text Min.*, pp. 9–28, 2007, doi: 10.1007/978-1-84628-754-1_2/COVER.
- [5] M. Hu and B. Liu, “Mining and summarizing customer reviews,” *KDD-2004 - Proc. Tenth ACM SIGKDD Int. Conf. Knowl. Discov. Data Min.*, pp. 168–177, 2004, doi: 10.1145/1014052.1014073.
- [6] N. N. Ho-Dac, S. J. Carson, and W. L. Moore, “The Effects of Positive and Negative Online Customer Reviews: Do Brand Strength and Category Maturity Matter?,” <https://doi.org/10.1509/jm.11.0011>, vol. 77, no. 6, pp. 37–53, Nov. 2013, doi: 10.1509/JM.11.0011.
- [7] F. Zhu and X. (Michael) Zhang, “Impact of Online Consumer Reviews on Sales: The Moderating Role of Product and Consumer Characteristics,” <https://doi.org/10.1509/jm.74.2.133>, vol. 74, no. 2, pp. 133–148, Mar. 2010, doi: 10.1509/JM.74.2.133.
- [8] J. Ye and L. Akoglu, “Discovering opinion spammer groups by network footprints,” *Lect. Notes Comput. Sci. (including Subser. Lect. Notes Artif. Intell. Lect. Notes Bioinformatics)*, vol. 9284, pp. 267–282, 2015, doi: 10.1007/978-3-319-23528-8_17/COVER.
- [9] C. Jiang, X. Zhang, and A. Jin, “Detecting Online Fake Reviews via Hierarchical Neural Networks and Multivariate Features,” *Lect. Notes Comput. Sci. (including Subser. Lect. Notes Artif. Intell. Lect. Notes Bioinformatics)*, vol. 12532 LNCS, pp. 730–742, 2020, doi: 10.1007/978-3-030-63830-6_61/COVER.
- [10] Y. Li, F. Wang, S. Zhang, and X. Niu, “Detection of Fake Reviews Using Group Model,” *Mob. Networks Appl.*, vol. 26, no. 1, pp. 91–103, Feb. 2021, doi: 10.1007/S11036-020-01688-Z/METRICS.
- [11] R. Mohawesh, S. Tran, R. Ollington, and S. Xu, “Analysis of concept drift in fake reviews detection,” *Expert Syst. Appl.*, vol. 169, p. 114318, May 2021, doi: 10.1016/J.ESWA.2020.114318.
- [12] G. Satia Budhi, R. Chiong, Z. Wang, and S. Dhakal, “Using a hybrid content-based and behavior-based featuring approach in a parallel environment to detect fake reviews,” *Electron. Commer. Res. Appl.*, vol. 47, p. 101048, May 2021, doi: 10.1016/J.ELERAP.2021.101048.
- [13] F. Abri, L. F. Gutierrez, A. S. Namin, K. S. Jones, and D. R. W. Sears, “Linguistic Features for Detecting Fake Reviews,” *Proc. - 19th IEEE Int. Conf. Mach. Learn. Appl. ICMLA 2020*, pp. 352–359, Dec. 2020, doi: 10.1109/ICMLA51294.2020.00063.
- [14] H. Li, Z. Chen, B. Liu, X. Wei, and J. Shao, “Spotting Fake Reviews via Collective Positive-Unlabeled Learning,” *Proc. - IEEE Int. Conf. Data Mining, ICDM*, vol. 2015-January, no. January, pp. 899–904, Jan. 2014, doi: 10.1109/ICDM.2014.47.
- [15] A. Heydari, M. A. Tavakoli, N. Salim, and Z. Heydari, “Detection of review spam: A survey,” *Expert Syst. Appl.*, vol. 42, no. 7, pp. 3634–3642, May 2015, doi: 10.1016/J.ESWA.2014.12.029.
- [16] M. Ott, Y. Choi, C. Cardie, and J. T. Hancock, “Finding Deceptive Opinion Spam by Any Stretch of the Imagination,” *ACL-HLT 2011 - Proc. 49th Annu. Meet. Assoc.*

- Comput. Linguist. Hum. Lang. Technol., vol. 1, pp. 309–319, Jul. 2011, Accessed: Dec. 25, 2023. [Online]. Available: <https://arxiv.org/abs/1107.4557v1>
- [17] S. Feng, R. Banerjee, and Y. Choi, “Syntactic Stylometry for Deception Detection.” Association for Computational Linguistics, pp. 171–175, 2012. Accessed: Dec. 25, 2023. [Online]. Available: <https://aclanthology.org/P12-2034>
- [18] E. Elmurngi and A. Gherbi, “An empirical study on detecting fake reviews using machine learning techniques,” 7th Int. Conf. Innov. Comput. Technol. INTECH 2017, pp. 107–114, Nov. 2017, doi: 10.1109/INTECH.2017.8102442.
- [19] V. K. Singh, R. Piryani, A. Uddin, and P. Waila, “Sentiment analysis of Movie reviews and Blog posts,” Proc. 2013 3rd IEEE Int. Adv. Comput. Conf. IACC 2013, pp. 893–898, 2013, doi: 10.1109/IADCC.2013.6514345.
- [20] A. Molla, Y. Biadgie, and K. A. Sohn, “Detecting negative deceptive opinion from tweets,” Lect. Notes Electr. Eng., vol. 425, pp. 329–339, 2018, doi: 10.1007/978-981-10-5281-1_36/COVER.
- [21] S. Shojaei, M. A. A. Murad, A. Bin Azman, N. M. Sharef, and S. Nadali, “Detecting deceptive reviews using lexical and syntactic features,” Int. Conf. Intell. Syst. Des. Appl. ISDA, pp. 53–58, Oct. 2014, doi: 10.1109/ISDA.2013.6920707.
- [22] G. G. Chowdhury, “Natural language processing,” Annu. Rev. Inf. Sci. Technol., vol. 37, no. 1, pp. 51–89, 2003.
- [23] C. Silva and B. Ribeiro, “The Importance of Stop Word Removal on Recall Values in Text Categorization,” Proc. Int. Jt. Conf. Neural Networks, vol. 3, pp. 1661–1666, 2003, doi: 10.1109/IJCNN.2003.1223656.
- [24] J. Plisson, N. Lavrač, and D. Mladenović, “A Rule-based Approach to Word Lemmatization,” 2004.
- [25] C. Lee and D. A. Landgrebe, “Feature Extraction Based on Decision Boundaries,” IEEE Trans. Pattern Anal. Mach. Intell., vol. 15, no. 4, pp. 388–400, 1993, doi: 10.1109/34.206958.
- [26] N. Jindal and B. Liu, “Review spam detection,” 16th Int. World Wide Web Conf. WWW2007, pp. 1189–1190, 2007, doi: 10.1145/1242572.1242759.
- [27] R. Mihalcea, C. Corley, and C. Strapparava, “Corpus-based and Knowledge-based Measures of Text Semantic Similarity”, Accessed: Dec. 25, 2023. [Online]. Available: www.aaii.org
- [28] J. Ramos, “Using TF-IDF to Determine Word Relevance in Document Queries,” Proc. first Instr. Conf. Mach. Learn., vol. 242, no. 1, pp. 29–48.



Copyright © by authors and 50Sea. This work is licensed under Creative Commons Attribution 4.0 International License.