

RESEARCH ARTICLE



Deep Learning-Based Approach for Monitoring and Controlling Fake Reviews

Nilesh Sable¹, Parikshit Mahalle² , Kalyani Kadam³ , Bipin Sule⁴, Rahul Joshi⁵ and Mahendra Deore^{6,*}

¹Department of Computer Science and Engineering, BRAC's Vishwakarma Institute of Technology, India

²Department of Artificial Intelligence and Data Science, BRAC's Vishwakarma Institute of Technology, India

³Department of Artificial Intelligence and Machine Learning, Symbiosis Institute of Technology, India

⁴Department of Engineering, Sciences and Humanities, BRAC's Vishwakarma Institute of Technology, India

⁵Department of Computer Science and Engineering, Symbiosis Institute of Technology, India

⁶Department of Computer Engineering, MKSSS'S Cummins College of Engineering for Women, India

Abstract: In the last decade, E-commerce has developed into the world's biggest stage for shopping. It has allowed people around the world to directly communicate without any barriers to purchasing the products as per requirements. Internet technologies have reshaped E-commerce since product reviews have become a vital part of online shopping due to their rapid growth. But with widespread usage, it has also brought forth an influx in rates of fake reviews. Fake reviews, which are frequently used to influence public perception, are now a widespread occurrence due to the open nature of E-commerce. Using different learning techniques, many methods and techniques are implemented to spot false reviews and fake behavior. This research aims to use a recurrent neural network (RNN) to combine content and data to identify false product reviews. The proposed approach, which is related to spam indicators, makes use of both product reviews and reviewers' behavioral characteristics. The fine-grained burst pattern analysis is used to conduct a more thorough investigation of produced testimonials during "suspicious" periods in the proposed approach. Additionally, a customer's previous review data are utilized to determine their overall "authorship" reputation, which serves as a barometer for the authenticity of most recent reviews. For the proposed theory, we examined the real-world Amazon review dataset and produced more accurate findings than previous methodologies. In addition to this, our proposed deep learning-based model performance has been validated utilizing the benchmark Yelp Open dataset and IMDB dataset.

Keywords: E-commerce, recurrent neural network (RNN), authorship, suspicion, spam indicators

1. Introduction

During the 21st century, social media has become the world's biggest stage for communication of ideas. Specifically, websites and applications such as Twitter, Facebook, Instagram, LinkedIn, TikTok, Amazon, Flipkart, and WhatsApp are being used by the general public at an accelerated rate, giving each person the ability to freely voice their raw opinions on any topic [1]. The possibility to express one's opinions was provided through online social development. As a result, numerous companies, as well as solution providers, seem to be no longer able to supervise the elements of the parallel world. People who've been dissatisfied with a group's services or products post complaints on social media. This viewpoint might have an impact on other prospective clients, both good and bad. Before buying a product, prospective customers must learn more about an item [2–4]. An assessment of the emotion should be able to determine if the sensation is unfavorable or good almost

instantly. Sentiment analysis is a type of text classification that concentrates on people's feelings, moods, and views in their writing [5–7]. The basic concept behind sentiment classification is to classify the intensity of words and determine if they will be positive or negative. Speedy networking site expansion is frequently employed with opinion analytics. Society's views are becoming quite important in several locations. Gathering online test data has proven to be a challenge. As we know lots of review data are generated daily so using typical synthetic data streams and real-life data streams, Kokate et al. [8] examined a variety of data stream approaches and algorithms. A comparison of density-micro and density-grid-based clustering algorithms is presented and a variety of internal and external clustering assessment techniques are used to evaluate the algorithms. Distinct machine learning [9–13] and deep learning methods [14–17] have been explored for fake review detection and analysis.

Even though Naïve Bayes' approach is simplistic [18], accuracy is the crucial aspect we set out to achieve, which is given by our spam filtering system using RNN. Considering that the dataset increases the precision of Naïve Bayes falls significantly, whereas our implemented system doesn't fall prey to this aspect. Even in

*Corresponding author: Mahendra Deore, Department of Computer Engineering, MKSSS'S Cummins College of Engineering for Women, India. Email: Mahendra.deore@cumminscollege.in

K-nearest neighbors based approaches [19], accuracy didn't cross 78%, whereas our system performed with an accuracy of approx. 90%. Other ML-based algorithms such as support vector machine (SVM) also performed poorly in our proposed system. As surveyed by Zamil et al. [19], the achieved accuracy with an SVM-based spam filtering system is around 84%, whereas our approach is efficient and quick with an accuracy rate of 89%. This goes to show that our proposed system outperforms a lot of other machine learning-based approaches. Lately, a large number of item assessment sites have indeed been produced on the World Wide Web. It encourages researchers to do a sentiment analysis of customer reviews. This article looked at consumer feedback on brand assessments. The remainder of this paper is structured as presented: Section 2 continues the discussion of existing systems and relational work in detail. Section 3 covers proposed methods for fake reviews using deep learning. Section 4 discusses the results, and at last Section 5 presents a summary of the paper and related future work on the system.

2. Related Work

In Dematis et al. [20], a strategy that uses usage and content data to identify bogus product evaluations was provided. The method put forward makes use of both product reviews and reviewer behavioral characteristics, which are connected via specialized spam indicators. Granular bursting pattern recognition is implemented in this study to accurately evaluate the reviews written during "questionable" periods. The reviewer's prior review historical data are also used to assess the reviewer's "authorship" reputation, which provides an estimate of the genuineness of their most current evaluation. Using a big data analytical strategy, the research [21] explores the effect of digital consumer feedback on user flexibility and in return gives the performance of the product. The writers built a singular value decomposition-based linguistic keyword similarities technique to estimate user flexibility using a lot of consumer review texts and notes of product releases. By using a dataset from a smartphone application with over 30 lakh digital critiques, our empirical study shows that the reviews containing the dataset have a not-so-linear connection with consumer adaptability. Additionally, there is a curved link between customer adaptability and the performance metrics of the product. The influence of a business's ability to employ digital consumer evaluations of a product, for instance, is demonstrated in this study, which contributes to the pool of information on the subject of innovation. Moreover, this also helps to resolve discrepancies in the written literature about the relationships between the three aspects.

The quantity of reviews that reviewers post (in a given amount of time) follows an odd spatial variability that was never reported previously, according to Li et al. [22]. That means, their sharing patterns are likely multifactorial. Coursing is a process in which many fraudsters write reviews for a similar set of items in a brief period. In addition, the finding revealed some interesting patterns in individual reviewers' temporal dynamics as well as their co-bursting tendencies with other critiques. The writers propose a dual-mode system. Hidden Markov chains (HMC) are a type of Markov chain that depends on their findings and describes how to model spamming using just individual reviewers and their review submission timings. The writers then utilize the coupled hidden Markov Model to identify both the publishing actions of the reviewer and signals that are co-bursting. In Hazim et al. [23], the author contributed the above research recommend evaluating two multilingual datasets employing features based on statistics that are modeled using the supervised boosting approach such as the extreme gradient boost (XGBoost)

and the generalized boosted regression model (GBM) to facilitate the identification of sentiment spamming in the smartphone app marketplace (i.e., English and Malay language). The XGBoost is better for identifying sentiment spamming in the English dataset, whereas the GBM Gaussian is best for the Malay dataset, according to the assessment findings. In accordance to the comparison analysis, the suggested quantitative-based approaches had a diagnostic precision of 86.13 percent for the Malay dataset and 87.43 percent for the English dataset, respectively. In Kumar et al. [24], the researchers suggest a new hierarchical supervised learning technique to enhance the possibility of spotting irregularities by assessing various user characteristics to describe their cumulative actions homogeneously. To represent user traits and activities, the researcher uses univariate and multivariate distributions. The derivatives after that are stacked to produce robust meta-classifiers. This strategy is appealing to digital businesses because it may assist in reducing fake reviews and enhance customer satisfaction with the quality of their information digitally. This work participates in the knowledge body by integrating redistributive components of traits into machine-learning algorithms that can aid in the detection of phony reviews on electronic services. As per a study by Hussain et al. [25], several evaluation metrics are frequently utilized to calculate the efficacy of analyzing spam identification technologies. Finally, this study includes a suggested typology of spam identification algorithms, methods for assessment, and being accessible to the public datasets, and then an analysis of numerous extractions of feature approaches from datasets available for review. Gaps in the proposed study and upcoming implementation aim in the realm of faux identification of reviews that are also mentioned. As per this study, interconnections are one of the crucial elements for any method adopted to review false reviews. The attribute selection technique adopted determines the efficacy of review spam detection algorithms, and the extraction of features is dependent on the review dataset. In conclusion, these factors must be considered in tandem for effective putting the spamming prediction models into practice and with greater precision.

Judgments of consumers are highly impacted by internet reviews presently [26]. Everything from purchasing a clothing item on an E-commerce website to fine dining is now based on online ratings. Due to the sheer dependence on internet reviews, some persons and their businesses fabricate bogus positive comments to improve or destroy the standing of a person's service or business. As a consequence, determining what kind of review is harmful or spam is difficult, and manually categorizing all the reviews is similarly impossible. A spiral cuckoo search-based clustering technique is developed to identify fake comments. By adding an extra layer of cuckoo search with the Fermat spiral, the suggested solution addresses the cuckoo search method's convergence problem. The efficacy of the recommended approach was tested using quadruple datasets and a single spammer Twitter dataset. Following the rapid growth of the Worldwide Web [27], cyberspace organization is extremely prevalent. This is owing to the ease with which many items and services can be found on the internet. Therefore, evaluations of all the products and services by consumers as well as organizations are crucial. Regrettably, phone reviews are often used by scammers for the goal of profit or advertisement. Because of the bogus reviews created by fraudsters, customers and companies seem to be unable to make proper judgments about the items. To minimize fooling prospective buyers, these false reviews, also called spam comments, must be recognized and deleted. In this article, we employed a supervised learning technique to detect review spam. Models are built using a wide array of sentiment and value-based variable scores, and their progress is tracked with the help of

multiple classifications in the proposed study. The article by Barbado et al. [28] looks at fraudulent reviews in the end-user electronics space and proposes a functional technique for identifying them. The inputs are categorized into four groups: (i) create a dataset for identifying fraudulent reviews, in the end, user electronics market using scraping techniques in quadruple cities; (ii) establish a functional model for bogus comment recognition; (iii) create a false comment recognition approach related to the conceptual model; and (iv) assess, interpret the findings for each city under consideration. In a study by Ghai et al. [29], the author proposed that a review synthesis method be utilized. Certain dimensions are already present for gauging the efficacy of comments. These characteristics indicate how a comment stands out among the others, indicating that it is likely spam. This approach categorizes reviews as valuable or invaluable based on the score assigned to them.

Spam operations observed on prominent digital shopping websites (e.g., amazon.com) have piqued the interest of both industry and academics, where several digital commenters are employed to collectively produce false evaluations for select target items, according to Xu and Zhang's [30] study. The purpose is to alter the targets' presumed perceptions in their favor. The pair-wise characteristics are initially used for the identification of group quislings following spam campaigns in online product reviews, which can disclose spam campaign collisions in a more fine-grained manner. Reviews for products that are provided online have become a significant source of customer feedback, according to Fei et al. [31]. Scammers have been creating false or phony comments to advocate and/or reprimand some specific items or services for profit or recognition. Review spammers are also a type of impostor. Numerous techniques to get a hold of the problem have been offered in recent years. Take an alternative strategy in this paper, which takes advantage of the burrstone character of ratings to detect review spam.

Consumers will benefit much from reviews online of services and goods, according to Minnich et al.'s [32] study, but they must be secured from deception. Several researches to date have concentrated on examining internet evaluations from a unique hosting source. How may data from numerous comment-hosting sites be combined? This is the central issue in our research. As a result, create a consistent technique for combining, comparing, and evaluating evaluations from various hosting sites. Use more than 150 lakh reviews from more than 35 lakh individuals across three major travel sites to emphasize hotel website reviews. Customers are constantly relying on information gathered from the public, like comments on Amazon and Yelp, and popular comments and adverts on Facebook, according to Viswanath et al.'s [33] study. As a consequence, the marketplace for black hat promotional tactics such as spoofing (e.g., Sybil) and hacked accounts, as well as colluded networks, has grown. Most existing methods for detecting such behavior rely on supervised (or semi-supervised) gaining knowledge dependent on already known (or imagined) assaults. They can't identify attacks that the administrator overlooked when labeling or when the intruder switches tactics. Online customer reviews are an extremely crucial source for strategic planning and the launch of new products, according to the article by Li et al. [34]. Comment flooding, on the other hand, is a likely target of review systems. Although supervised learning has been used to identify review spam for years, the main truth of big data containing sets is still inaccessible, and most of the current supervised learning approaches are built on faux false comments instead of actual false reviews. We describe the first known work on false comment identification in Chinese, using comments which are filtered from Damping false review identification technique, in collaboration with Dianping1, the biggest Chinese review hosting site.

Customer reviews are increasingly one of the most essential pieces of knowledge for users on an array of products, according to Crawford et al. [35]. With their growing relevance, spammers and unscrupulous company owners have more opportunities to manufacture bogus feedback to fraudulently advertise their products and services or trash the services provided by their adversaries. Much research on the most efficient approaches to identifying spam using various ML algorithms has been conducted in reaction to current rising problems. Transformation of reviews to vectors containing words, which might generate thousands and thousands of attributes, is a consistent theme in most of these investigations. In Xue et al. [36], the author proposes a method for detecting review spammers that incorporates social relationships and is composed of two inferences: consumers are more inclined to believe reviews from people they know, and spammers have little to no probability of establishing a significant connection with regular consumers. This research makes two implications: (1) it explains a way interaction should be factored into comment rating forecast and put forth a faith-based prediction model about the rating that uses closeness as a key trust factor in rating and (2) it creates a rating variance-based belief detection model that recursively evaluates consumer-specific overall source credibility scores as a spam city predictor. The language and ranking properties from a review were used in Wahyuni and Djunaidy [37] to recognize false reviews for an item. In summary, the suggested system (ICF++) would assess the integrity of a review, the trustworthiness of the reviewers, and the item's dependability. Text mining techniques, as well as opinion mining techniques, will be used to determine a review's integrity rating. The experimental results reveal that the developed approach has greater precision than the iterative computation framework (ICF) method's outcome. Online social networks capture the dynamics of both person-to-person and person-to-technology interactions and are utilized for various purposes, such as commerce, academia, advertising, healthcare, and leisure [38]. However, they also enable illegal activities. Detecting anomalies is essential in this new perspective on social life, which represents and encapsulates offline connections, as these anomalies can signal significant issues or provide valuable insights for researchers.

Mangoes are classified into four categories [39, 40] in this study based on a machine learning approach: Green Mango, Yellow Mango, and Red Mango. This technique considers the Red, Green, and Blue color components, as well as the shape and texture of mangoes, to achieve a high probability of accurate classification. This, in turn, helps train the system to identify the correct ripeness of mangoes. The study employs two machine learning methods, Naive Bayes and SVM. In today's context, most implemented approaches [41, 42] are only marginally effective. As data become increasingly central, there is a growing need for more precise fake review monitoring and control systems to detect and mitigate the dangers posed by large volumes of fake data. Javed et al. [43] explore the classification of fake reviews using an ensemble of shallow convolutional neural networks. Their work focuses on employing advanced deep learning algorithms to enhance the accuracy of fake review detection. The study uses two Yelp datasets [44] to evaluate a deep learning-based method for automatically identifying fake reviews, analyzing both reviewing behavior and the textual content of reviews to improve detection accuracy. The findings suggest that high-level text representations are more predictive than traditional linguistic features and that feature learning significantly improves classification performance. Pal et al. [45] propose an approach for detecting fake reviews based on sentiment analysis that combines Word2vec, lexicons, and an attention mechanism.

A deep-transfer learning model is used to detect fake profiles in real time by analyzing data from various social media platforms, including posts, likes, comments, multimedia content, user activity, and login behaviors [46]. Asaad et al. [3] present a machine learning method that uses TF-IDF for feature extraction and preprocessing to identify fraudulent Yelp reviews. The method is tested on both balanced and unbalanced datasets and uses XGBoost, support vector classifier, and stochastic gradient descent for classification. The aim is to enhance the credibility of online reviews and protect businesses from the adverse effects of fake reviews. This work includes an extensive bibliometric analysis [47] of AI applications for detecting fake reviews, summarizing key ML-based research, highlighting major developments, and outlining trends and future directions. The findings indicate a growing interest in using AI for fake review detection in research.

To address the challenges associated with identifying fake reviews, the study by Pan and Xu [48] proposes a fully unsupervised method that integrates survey research, analysis of fake review features, calculation of a fake index, and selection of false reviews.

It also presents a recommendation-based performance score to assess detection techniques, which is especially helpful when analyzing review data that lacks objective authenticity classifications.

Zhang et al. [49] introduce a novel method called ImDetec-tor for identifying fraudulent reviewers. This approach addresses data imbalance using weighted latent Dirichlet allocation (LDA) and Kullback–Leibler (KL) divergence. Specifically, they develop a weighted LDA model to uncover latent topics among reviewers based on the characteristics of the reviews. Wu et al. [50] present 20 potential research questions and propose 18 hypotheses, noting that studies on fake reviews often suffer from a lack of high-quality datasets. Wang et al. [51] propose a new AUC-based extreme learning machine (AUC-ELM) for imbalanced binary classification, which is demonstrated to be effectively equivalent to an ELM applied to a transformed data space.

Song et al. [52] provide insights that help consumers make more informed decisions when faced with fake information, offer valuable strategies for marketers in promoting products within specific contexts, and contribute to the enhancement of online review monitoring mechanisms. Budhi et al. [53] propose two sampling techniques to improve the accuracy of detecting fake reviews in balanced datasets. By using random under-sampling and over-sampling with convolutional neural networks, they achieve an accuracy of up to 89%.

Rong et al. [54] present a framework that integrates a unified deep network, which models both visual and linguistic information. This framework encodes region-level and pixel-level visual features of natural scene images into spatial feature maps, which are subsequently decoded into saliency response maps for text instances. Nguyen and Hsu [55] suggest that E-commerce vendors consider three resource-matching dimensions to minimize excessive data collection while achieving sufficiently personalized recommendation results on their digital platforms.

Mohawesh et al. [56] explore the critical role customers play in determining the overall revenue of a company in the E-commerce sector. They note that consumers often check reviews before purchasing a product or service, and as a result, many E-commerce companies pay spammers to generate fake reviews for various products and services.

We propose a new model utilizing deep learning to address the limitations we previously encountered, employing a real-time dataset instead of a dummy dataset for more accurate evaluations.

Earlier algorithms, like the Naïve Bayes theorem, did not yield satisfactory results. Consequently, we developed a framework called SpamDup, which uses a recurrent neural network (RNN) for sentiment analysis and incorporates latent semantic analysis for semantic analysis.

3. Proposed Theory and Methods

In such an unstructured Amazon review dataset, the initial stage is to assess and quantify fraudster behavioral traits. This computation is performed on all dataset comments that used the behavioral characteristics technique for fake review identification.

Numerous research papers have concentrated on the topic of recognizing spammers and faux reviews in the recent decade. However, because the problem is complex and difficult, so not completely resolved. We can describe the results of our analysis. Previous studies are talked about in the three categories below. We express the idea as heterogeneity networking, with n nodes genuine sets of data constituents (like ratings, individuals, and things) or fake qualities, as earlier mentioned. We'll go over the major concepts and terminology in heterogeneous networking systems to help you understand the architecture.

3.1. Introduction

- 1) Unless there are $a (> 1)$ types of networks and $b (> 1)$ types of related linkages linking them, a heterogeneous information network is represented as a graph $G = (V, E)$. Each network V and link E is associated with one of the network and linkage varieties. The kinds of beginning and end nodes of two connections that belong to the same type are similar.
- 2) Gantt chart is a diagram that shows the relationship between two variables. $T = (a, r)$ is a metaphysical route involving entity type characterization using a heterogeneous network: V A and the connection modeling: E R , that is, a network built around entity type ' a ', with connections as connections via ' r '. G is the letter G . (V, E) . The schema explains a network's meta-structure (i.e., how several different sorts of nodes are there and where possible linkages exist).
- 3) There have been no connections linking two terminals or the same, but as previously indicated, there are paths. A biotic system P is stated by a group of relationships there in networking structure $T = (a, r)$, marked by the letters P provided a compound connection $P = R_1 o R_2 o \dots o (l_1)$ in both vertices, there is pattern $a_1(r_1) a_2(r_2) \dots a(l_1)al$, which establishes a polymeric connection $P = r_1 o r_2 o \dots o (l_1)$ among a collection of vertices, where o would be the compound provider on connections, a network of disparate data G is the letter G . (V, E) . A meta path exists only when no ambiguity is present = $a1a2 \dots al$, for example, may be expressed as a succession of node kinds. The meta path extends the linkage classes to route categories, a notion that depicts the various connections between node kinds via indirect links, i.e., routes, thus conveying a wide range of meanings.
- 4) Suppose that now in a heterogeneous system $G = (V, E)$, V is a portion of V that comprises either the declaration's or the declaration's connections (the network types that must be categorized).
- 5) Researchers have these great pre-endpoints in V associated with every category, like $c_1c_2c_3 \dots c_k-1c_k$, and we have several in V that are built on the preconfigured network with a satisfied client for every class, such as $c_1c_2c_3 \dots c_k-1c_k$. The categorization challenge aims to estimate all of V 's unmarked networks' labels.

3.2. Featured types

- 1) User-Linguistic: These traits are drawn from consumers' own words and indicate in a way how they communicate their beliefs or viewpoints of what they've seen. I've experienced the following experiences as a client of a specific company. This is what we use to determine a troll's communication style; search for the traits listed below. Average content similarity and maximum content similarity (MCS) are dual frameworks in this area (MCS). These two features emphasize the differences between two evaluations posted by two different people. The wording of spammers there looks to be very much like reviews based on pre-written text templates.
- 2) Review-Linguistic: The characteristics of the group are made based on the content of the evaluation. Two essential criteria in the RL category that we apply in research are the proportion of first-person diminutives (PP1) to exclaim expressions including "I" (RES).
- 3) Behavioral Review: Instead of the real review content, content is used in the production of the attribute style. There are two RB category features: an early time frame and deviation of review for threshold rating.
- 4) User Behavioral: Because these attributes are unique to each user individually and are determined per person, we may take advantage of them to sum up just about all the opinions posted by that user. The Overabundance of Testimonials submitted by a person in accordance with the meaning of a user's bad ratio offered to multiple companies are the two primary elements in the set.

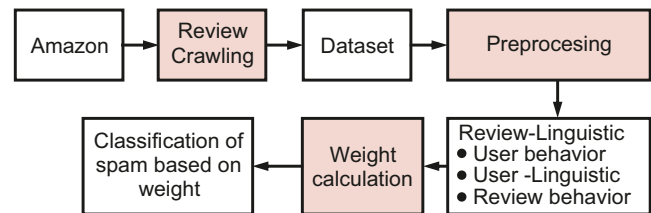
Based on the features and survey done we made utilization of the RNN recurrent network for emotion evaluation in our proposed system.

3.3. System model

Figure 1 shows the planned research methodology. Spam Dup reviewing systems are modeled as heterogeneity interconnection in this paradigm, which is a unique network-based method. A novel fake feature ranking approach is presented to estimate the relative relevance of each characteristic and to demonstrate how efficient each characteristic is in distinguishing junk from legitimate ratings. In consideration of computational cost, the system beats the state of the art, which is highly impacted by the number of characteristics utilized to recognize a malicious review.

The flowchart below depicts how to identify and classify fake reviews using supervised and unsupervised machine learning methodologies. A variable dataset (Amazon) was used to illustrate the usefulness of this method. Using a variable dataset makes the technique more conclusive. The provided dataset is first modeled in such a way into a review dataset collection. Consequently, the dataset is transformed into a HIN. Tokenization is one of the strategies to be used to preprocess the evaluation data. Tokens can be considered as pieces. Stop word [a stop list is created to omit often used words such as (a, an, the, and for) so they don't take up unneeded database space] with inclusion to stemming. The preprocessed dataset is then connected to distinct terminals that take the form of opinions and are labeled on grounds of importance. Term frequency-inverse document frequency (TF-IDF) is a term for identifying them on the grounds of their frequency distribution. As data labeling is done for distinct nodes on the grounds of their feature

Figure 1
Research methodology



significance, they are tagged having a special intent weightage taken into consideration during the final classification of bogus testimonials. For classification, the Network Spam Novel approach is used wherein the transformed data are filtered, and the fraudulent evaluations are highlighted to illustrate the provided review dataset in HIN. The basic idea behind our proposed architecture is to define a database consisting of reviews as a hierarchical network system and then move the fake identification problem to an FIN job. A design rating database, specifically, wherein evaluations are connected to use multiple network types.

3.4. Recurrent neural network (RNN)

The very first neurological approach for imposter class detection, HIS-RNN, a heterogeneity information system (HIS) compliant recurrent neural network that employs mutual information and does not require hand-crafted features, is described in this paper. The HIS-RNN provides a uniform structure for each author's classification tasks, with the beginning vector including the total number of word vectors of all evaluation texts supplied by the source author, combined with the portion of bad comments. With a co-review understanding of the interaction inspectors who have reviewed the very same objects with equivalent evaluations and the examiner's vectors description, the HIS-RNN learning records a cooperative grid.

The general method that our platform does is to identify a particular dataset and model it into a heterogeneity information system. Heterogeneous information network helps us to model real-time information like Amazon data is given due regard in this particular case and make a graphical model and inculcate it with real-world information. The faux review identification is modeled and mapped into a collection of disparate information classification problems.

In summary, distinct node types are connected as a type of opinions to the heterogeneous information network where our database is structured like in HIN.

After classification into distinct nodes, each node is then taken into consideration on the grounds of the feature's significance of the node. For this task, a weighting algorithm is implanted to describe and analyze each feature's importance. Following the analysis of the weightage of each feature, every node is provided with a final identification label for the reviews, which are characterized using both supervised and unsupervised machine learning methodologies.

So, we propose the Spam Dup novel technique where our provided dataset is modeled and structured networks as HIN. The process performs uniquely where this segregation uses multiple classes of meta paths, which are very ingenious in the anti-spam domain.

3.5. Feature selection

Algorithm: Identification of fake reviews using behavioral features approach

```

Input review Ri.  $\tau = 0.5$ . 0.55. 0.6 //threshold value for labeling the review
Output: Spam or Not-Spam
for each review Ri in the review dataset do
//behavior features (F1, F2, F3... F13)
for each behavior feature Fi calculate normalize value
//variable V is calculating normalized value of F
Vi = calculate normalized value Fi
Sum += Vi
end for
//calculating average score
Average Score = Sum/13
for each Value Vi do
//calculating drop score
Drop Score = (Sum - Vi) / 12
if |Average Score - Drop Score| >= 0.05 then
assign weight Wi      2      ◀
Total weight += 2
else
assign weight Wi      1      ◀
Total Weight += 1
end if
end for
for each value Vi do
//calculating total spam score
Score += Wi * Vi
end for
Spam Score = Score / Total Weight
if Spam Score >  $\tau$  then
label Ri      ◀ Spam
else
label Ri      ◀ Not-Spam
end if
end for

```

According to Figure 2 RNN, suppose x represents a series with length T , where $x = x_1, x_2, \dots, x_T$, and x_t represents a feature representation for each item. The output nodes state y_t , the current hidden layer state h_t , and the preceding hidden layer state h_{t-1} may all be computed at time step t by Equations (1) and (2),

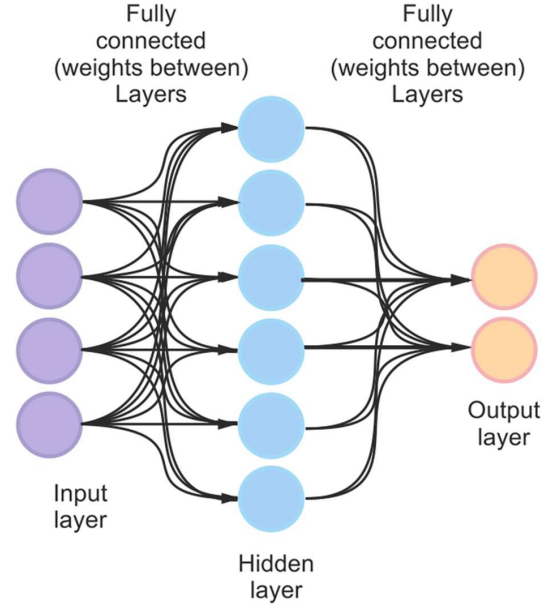
$$h_t = \sigma_h(w_h x_t + U_h h_{t-1} + b_h) \quad (1)$$

$$y_t = \sigma_y(w_y h_t + b_y) \quad (2)$$

where U_h is the matrix containing the recurrent weights between the hidden layer and itself at two consecutive time steps, b_h and b_y are the biases, and h and y signify the convolution layers. W_h and W_y are the insight input and concealed output weighted matrix, respectively.

The information is transmitted in a typical feed-forward neural manner at every time step, and then a training rule is implemented. Because of the back linkages, the context units always keep a copy of the hidden units' prior values. To do tasks like sequence learning that go beyond the capabilities of conventional multilayer

Figure 2
Recurrent neural network (RNN)



perceptrons, the network must be able to retain a state. Equation (3) provides the formula for computing the present state.

$$h_t = \int_t^{t-1} (h_{t-1}, x_t) \quad (3)$$

Where:

h_t -> Current state

h_{t-1} -> Previous state

x_t -> Input state

The formula for applying the activation function is given by Equation (4)

$$h_t = \text{activation}(W_{hh} h_{t-1} + w_{xh} x_t) \quad (4)$$

Where:

W_{hh} -> Weight at the recurrent neuron

w_{xh} -> Weight at input neuron

The formula for calculating output is given by Equation (5):

$$y_t = w_{hy} h_t \quad (5)$$

Where:

y_t -> Output

w_{hy} -> Weight at the output layer

In the next section, we implemented our proposed system and discussed the result.

4. Result and Discussion

We describe the significant parameters utilized in model training for fake review recognition and categorization to promise reproducibility and provide insight into the methodology. The specifications are compiled in the following Table 1. Testing was carried out on a laptop with the following specifications: AMD Ryzen 7 4800 H with Radeon Graphics 2.90 GHz, 16.0 GB (15.4 GB usable), Windows 10, MySQL 5.1 backend database, and Java platform 1.8. The website software works on the Tomcat web server and is used to design code in Eclipse. The learning rate, batch size,

dropout rate, optimizer, activation function, and sequence length are 0.01, 64, 0.5, Adam, ReLU, and 50, respectively.

Table 1
The implementation details of the proposed method

Parameters	Specification
Learning rate	0.01
Batch size	64
Dropout rate	0.5
Optimizer	Adam
Activation function	ReLU
Length of sequence	50
Processor	AMD Ryzen 7 4800 H
Window	10

Figure 3(a) and (b) depict the performance of the existing and proposed system, using performance measures such as precision, recall, F1-measure, and accuracy. The evaluation metrics' formulas are described as follows:

$$Accuracy = \frac{TP + TN}{TP + TN + FP + FN} \quad (6)$$

$$Precision = \frac{TP}{TP + FP} \quad (7)$$

$$Recall = \frac{TP}{TP + FN} \quad (8)$$

Figure 3

(a) Accuracy improvement between existing and proposed system and (b) analysis of performance between proposed and existing systems

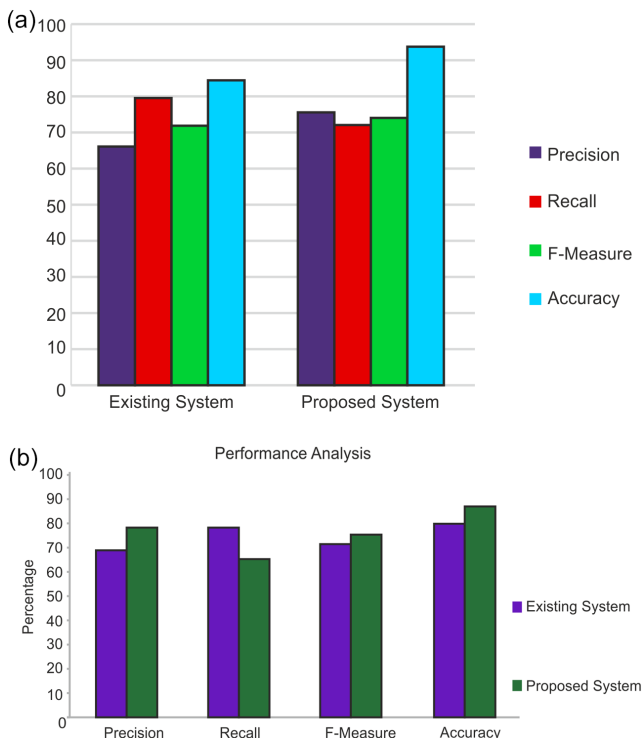


Table 2
Factor comparison between Naïve Bayes and RNN

Factors	Naïve Bayes	RNN
Precision	68.05	80.19
Recall	81.49	64.64
F-Measure	74.14	80.32
Accuracy	80.12	90.14

Figure 4
Distinguishing precision, recall, and F1-score of different thresholds of similar reviews

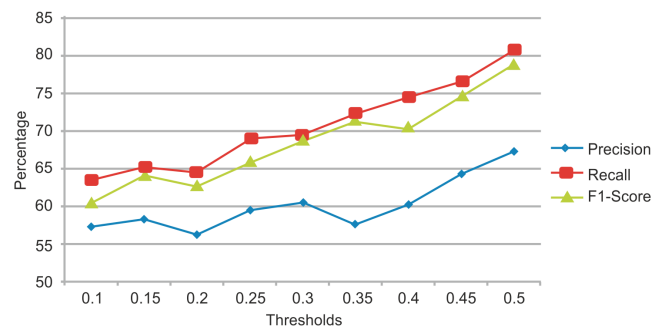


Table 3
Distinguishing precision, recall, and F1-score of different thresholds of similar reviews

Threshold	Precision	Recall	F1-score
0.1	57.22	63.44	60.44
0.15	58.22	65.14	64.24
0.2	56.15	64.34	62.77
0.25	59.33	68.99	65.88
0.3	60.44	69.23	68.74
0.35	57.48	72.35	71.34
0.4	60.15	74.44	70.18
0.45	64.33	76.49	74.66
0.5	67.18	80.66	78.99

$$F1 \text{ Score} = \frac{2 \times \text{Recall} \times \text{Precision}}{\text{Recall} + \text{Precision}} \quad (9)$$

where the TP is the true positive, TN is the true negative, FP is the false positive, and FN is the false negative.

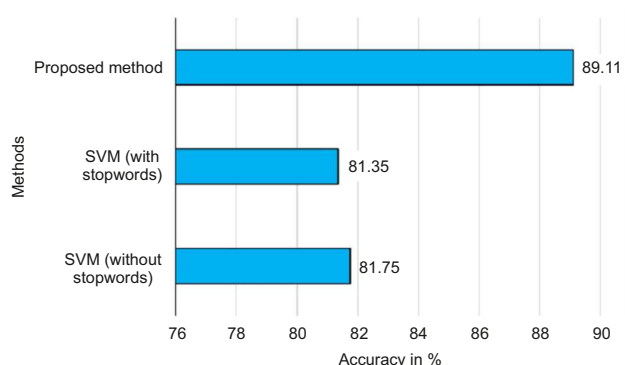
The precision, recall, F-measure, and accuracy between Naive Bayes and RNN in Table 2 states that RNN has given us better performance concerning the Naive Bayes algorithm, we see our accuracy has risen to 89.11. We have also seen different thresholds in comparison in Table 3 of similar reviews. Our current system produces better results and performance than discussed above with different approaches.

Results in Figure 4 show the comparisons of precision, F-measure, and recall of different thresholds from similar reviews and, as per observation at different threshold values such as 0.3, the precision is as high as 60.44 and the F-measure is 68.74 as well as recall with 69.23. The values differentiate based on the threshold as seen in Table 2.

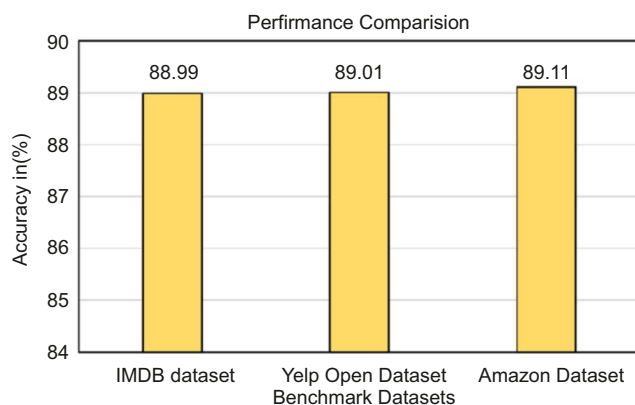
Figure 5 [57] depicts the comparative analysis of the proposed method and the state-of-the-art method. The proposed fake review monitoring approach obtains an accuracy of 89.11%. However, Elmurungi and Gherbi state-of-the-art methods namely SVM

Figure 5

The comparative analysis of the proposed method with state-of-the-art fake review detection methods, i.e., SVM (with stop words) and SVM (without stop words)

**Figure 6**

The comparative analysis of measured accuracy on different benchmark datasets used for model validation



(with stopwords) and SVM (without stopwords) obtained the accuracy of 81.35% and 81.75%, respectively. Thereby, the proposed fake review detection method outperforms the state-of-the-art techniques.

Figure 6 illustrates the comparative analysis of measured accuracy on different benchmark datasets used for model validation. This proposed model performance was evaluated on distinct benchmark datasets namely the IMDB, Yelp Open datasets, and Amazon dataset. The measured accuracy on these datasets namely the IMDB, Yelp Open dataset, and Amazon dataset are 88.99%, 89.01%, and 89.11%, respectively. The measured accuracy is enhanced on all datasets; however, the accuracy is observed best on the Amazon benchmark dataset.

5. Conclusion

Spam reviews may make or break a purchase for any consumer, but when a person purchases a product that includes fake reviews, the end user suffers. This study offers an innovative technique for identifying fake comments. We employ a range of bogus indicators on an item and customer level to collect and utilize the majority of the critical data. The proposed approach uses advanced analytical variables based on bursting clustering and classification to identify abnormal periods and ratings. The study investigates the reputations of writers by quickly scanning the prior ratings and interactions to successfully assess the authenticity of one of their most current ones.

We tested the suggested method using the Amazon real-time dataset and found it useful for identifying negative feedback. The suggested method, which has an accuracy of 89.11% and an F-measure of 79.31%, employs sentiment analysis utilizing RNN, latent semantic analysis, and the Spam Dup framework to detect spam. A Netspam extension, known as the Spam Dup framework, is utilized to profit from superior finishing high point methods and can also take execution to lead to improved performance. Latent semantic analysis helps to reduce similar reviews in a dataset and produce datasets with distinct nodes i.e., reviews, which help to classify reviews more accurately. The suggested model may become more optimized in future work as we transition to a data-centric culture, which will speed up data extraction and result in a more accurate model.

Acknowledgment

We authors are thankful to the research support provided by Symbiosis Institute of Technology, Pune, Vishwakarma Institute of Technology, Pune and MKSSS'S Cummins College of Engineering for Women, Pune, India.

Ethical Statement

This study does not contain any studies with human or animal subjects performed by any of the authors.

Conflicts of Interest

The authors declare that they have no conflicts of interest to this work.

Data Availability Statement

Data are available on request from the corresponding author upon reasonable request.

Author Contribution Statement

Nilesh Sable: Conceptualization, Methodology, Software, Validation, Data curation, Writing – original draft, Writing – review & editing. **Parikshit Mahalle:** Conceptualization, Software, Resources, Supervision. **Kalyani Kadam:** Formal analysis, Writing – review & editing, Visualization. **Bipin Sule:** Methodology, Writing – original draft. **Rahul Joshi:** Investigation, Visualization. **Mahendra Deore:** Software, Validation, Writing – original draft, Writing – review & editing.

References

- [1] Kahn, W. J. (2021). *Social media in the 21st century: Perspectives, influences and effects on well-being*. USA: Nova Science Publishers.
- [2] Gayathri, M., Siva Teja, Y. S. N., & Ajay Sharma, K. (2023). Fake review detection using machine learning. *Computational Intelligence and Machine Learning*, 4(2), 10–12. <https://doi.org/10.36647/ciml/04.02.a002>
- [3] Asaad, W. H., Allami, R., & Ali, Y. H. (2023). Fake review detection using machine learning. *Revue d'Intelligence Artificielle*, 37(5), 1159–1166. <https://doi.org/10.18280/ria.370507>
- [4] Ram, N. C. S., Vakati, G., Nadimpall, J. V., Sah, Y., & Datla, S. K. (2022). Fake reviews detection using supervised machine learning. *International Journal for Research in Applied Science*

- & Engineering Technology, 10(5), 3718–3727. <https://doi.org/10.22214/ijraset.2022.43202>
- [5] Kaddoura, S., Chandrasekaran, G., Popescu, D. E., & Duraisamy, J. H. (2022). A systematic literature review on spam content detection and classification. *PeerJ Computer Science*, 8, e830. <https://doi.org/10.7717/peerj-cs.830>
 - [6] Kerrysa, N. G., & Utami, I. Q. (2023). Fake account detection in social media using machine learning methods: Literature review. *Bulletin of Electrical Engineering and Informatics*, 12(6), 3790–3797. <https://doi.org/10.11591/eei.v12i6.5334>
 - [7] Sharma, D. K., Singh, B., Agarwal, S., Garg, L., Kim, C., & Jung, K.-H. (2023). A survey of detection and mitigation for fake images on social media platforms. *Applied Sciences*, 13(19), 10980. <https://doi.org/10.3390/app131910980>
 - [8] Kokate, U., Deshpande, A., Mahalle, P., & Patil, P. (2018). Data stream clustering techniques, applications, and models: Comparative analysis and discussion. *Big Data and Cognitive Computing*, 2(4), 32. <https://doi.org/10.3390/bdcc2040032>
 - [9] Ennaouri, M., & Zellou, A. (2023). Machine learning approaches for fake reviews detection: A systematic literature review. *Journal of Web Engineering*, 22(5), 821–848. <https://doi.org/10.13052/jwe1540-9589.2254>
 - [10] Capuano, N., Fenza, G., Loia, V., & Nota, F. D. (2023). Content-based fake news detection with machine and deep learning: A systematic review. *Neurocomputing*, 530, 91–103. <https://doi.org/10.1016/j.neucom.2023.02.005>
 - [11] Lu, J., Zhan, X., Liu, G., Zhan, X., & Deng, X. (2023). BSTC: A fake review detection model based on a pre-trained language model and convolutional neural network. *Electronics*, 12(10), 2165. <https://doi.org/10.3390/electronics12102165>
 - [12] Han, S., Wang, H., Li, W., Zhang, H., & Zhuang, L. (2023). Explainable knowledge integrated sequence model for detecting fake online reviews. *Applied Intelligence*, 53(6), 6953–6965. <https://doi.org/10.1007/s10489-022-03822-8>
 - [13] Liu, Y., Wang, L., Shi, T., & Li, J. (2022). Detection of spam reviews through a hierarchical attention architecture with N-gram CNN and Bi-LSTM. *Information Systems*, 103, 101865. <https://doi.org/10.1016/j.is.2021.101865>
 - [14] Jia, Z., Cai, X., & Jiao, Z. (2022). Multi-modal physiological signals based squeeze-and-excitation network with domain adversarial learning for sleep staging. *IEEE Sensors Journal*, 22(4), 3464–3471. <https://doi.org/10.1109/JSEN.2022.3140383>
 - [15] Khairnar, N. V., Mankar, S. L., Pandav, M. R., Kotecha, H., & Ranjanikar, M. (2022). A survey paper on fake review detection system. In R. Srivastava & A. K. S. Pundir (Eds.), *New frontiers in communication and intelligent systems* (pp. 625–634). SCRS Publications. <https://doi.org/10.52458/978-81-95502-00-4-64>
 - [16] Athira, A. B., Madhu Kumar, S. D., & Chacko, A. M. (2023). A systematic survey on explainable AI applied to fake news detection. *Engineering Applications of Artificial Intelligence*, 122, 106087. <https://doi.org/10.1016/j.engappai.2023.106087>
 - [17] Zhang, X., Guo, F., Chen, T., Pan, L., Beliaikov, G., & Wu, J. (2023). A brief survey of machine learning and deep learning techniques for e-commerce research. *Journal of Theoretical and Applied Electronic Commerce Research*, 18(4), 2188–2216. <https://doi.org/10.3390/jtaer18040110>
 - [18] Metsis, V., Androutsopoulos, I., & Paliouras, G. (2006). Spam filtering with Naive Bayes – Which Naive Bayes? In *Third Conference on Email and Anti-Spam*, 1–9.
 - [19] Zamil, Y. K., Ali, S. A., & Naser, M. A. (2019). Spam image email filtering using K-NN and SVM. *International Journal of Electrical and Computer Engineering*, 9(1), 245–254. <https://doi.org/10.11591/ijece.v9i1.pp245-254>
 - [20] Dematis, I., Karapistoli, E., & Vakali, A. (2018). Fake review detection via exploitation of spam indicators and reviewer behavior characteristics. In *SOFSEM 2018: Theory and Practice of Computer Science: 44th International Conference on Current Trends in Theory and Practice of Computer Science*, 581–595. https://doi.org/10.1007/978-3-319-73117-9_41
 - [21] Zhou, S., Qiao, Z., Du, Q., Wang, G. A., Fan, W., & Yan, X. (2018). Measuring customer agility from online reviews using big data text analytics. *Journal of Management Information Systems*, 35(2), 510–539. <https://doi.org/10.1080/07421222.2018.1451956>
 - [22] Li, H., Fei, G., Wang, S., Liu, B., Shao, W., Mukherjee, A., & Shao, J. (2017). Bimodal distribution and co-bursting in review spam detection. In *Proceedings of the 26th International Conference on World Wide Web*, 1063–1072. <https://doi.org/10.1145/3038912.3052582>
 - [23] Hazim, M., Anuar, N. B., Ab Razak, M. F., & Abdullah, N. A. (2018). Detecting opinion spams through supervised boosting approach. *PLOS ONE*, 13(6), e0198884. <https://doi.org/10.1371/journal.pone.0198884>
 - [24] Kumar, N., Venugopal, D., Qiu, L., & Kumar, S. (2018). Detecting review manipulation on online platforms with hierarchical supervised learning. *Journal of Management Information Systems*, 35(1), 350–380. <https://doi.org/10.1080/07421222.2018.1440758>
 - [25] Hussain, N., Turab Mirza, H., Rasool, G., Hussain, I., & Kaleem, M. (2019). Spam review detection techniques: A systematic literature review. *Applied Sciences*, 9(5), 987. <https://doi.org/10.3390/app9050987>
 - [26] Pandey, A. C., & Rajpoot, D. S. (2019). Spam review detection using spiral cuckoo search clustering method. *Evolutionary Intelligence*, 12(2), 147–164. <https://doi.org/10.1007/s12065-019-00204-x>
 - [27] Narayan, R., Rout, J. K., & Jena, S. K. (2018). Review spam detection using opinion mining. In *Progress in Intelligent Computing Techniques: Theory, Practice, and Applications: Proceedings of ICACNI 2016*, 2, 273–279. https://doi.org/10.1007/978-981-10-3376-6_30
 - [28] Barbado, R., Araque, O., & Iglesias, C. A. (2019). A framework for fake review detection in online consumer electronics retailers. *Information Processing & Management*, 56(4), 1234–1244. <https://doi.org/10.1016/j.ipm.2019.03.002>
 - [29] Ghai, R., Kumar, S., & Pandey, A. C. (2019). Spam detection using rating and review processing method. In *Smart Innovations in Communication and Computational Sciences: Proceedings of ICSICCS 2017*, 2, 189–198. https://doi.org/10.1007/978-981-10-8971-8_18
 - [30] Xu, C., & Zhang, J. (2015). Combating product review spam campaigns via multiple heterogeneous pairwise features. In *Proceedings of the 2015 SIAM International Conference on Data Mining*, 172–180. <https://doi.org/10.1137/1.9781611974010.20>
 - [31] Fei, G., Mukherjee, A., Liu, B., Hsu, M., Castellanos, M., & Ghosh, R. (2013). Exploiting burstiness in reviews for review spammer detection. *Proceedings of the International AAAI Conference on Web and Social Media*, 7(1), 175–184. <https://doi.org/10.1609/icwsm.v7i1.14400>
 - [32] Minnich, A. J., Chavoshi, N., Mueen, A., Luan, S., & Faloutsos, M. (2015). TrueView: Harnessing the power of multiple review sites. In *Proceedings of the 24th International*

- Conference on World Wide Web, 787–797. <https://doi.org/10.1145/2736277.2741655>
- [33] Viswanath, B., Bashir, M. A., Crovella, M., Guha, S., Gum-madi, K. P., Krishnamurthy, B., & Mislove, A. (2014). Towards detecting anomalous user behavior in online social networks. In *Proceedings of the 23rd USENIX Security Symposium*, 223–238.
- [34] Li, H., Chen, Z., Liu, B., Wei, X., & Shao, J. (2014). Spotting fake reviews via collective positive-unlabeled learning. In *IEEE International Conference on Data Mining*, 899–904. <https://doi.org/10.1109/ICDM.2014.47>
- [35] Crawford, M., Khoshgoftaar, T. M., & Prusa, J. D. (2016). Reducing feature set explosion to facilitate real-world review spam detection. In *Proceedings of the Twenty-Ninth International Florida Artificial Intelligence Research Society Conference*, 304–309.
- [36] Xue, H., Li, F., Seo, H., & Pluretti, R. (2015). Trust-aware review spam detection. In *IEEE Trustcom/BigDataSE/ISPA*, 1, 726–733. <https://doi.org/10.1109/Trustcom.2015.440>
- [37] Wahyuni, E. D., & Djunaidy, A. (2016). Fake review detection from a product review using modified method of iterative computation framework. In *3rd Bali International Seminar on Science & Technology*, 58, 03003. <https://doi.org/10.1051/mateconf/20165803003>
- [38] Hassanzadeh, R. (2014). *Anomaly detection in online social networks using data mining techniques and fuzzy logic*. PhD Thesis, Queensland University of Technology.
- [39] Pise, D., & Upadhye, G. D. (2018). Grading of harvested mangoes quality and maturity based on machine learning techniques. In *International Conference on Smart City and Emerging Technology*, 1–6. <https://doi.org/10.1109/ICSCET.2018.8537342>
- [40] Papat, S. N., & Singh, Y. P. (2017). Efficient research on the relationship standard mining calculations in data mining. *Journal of Advances in Science and Technology*, 14(2), 42–48.
- [41] Sable, N. P., Rathod, V. U., Mahalle, P. N., & Birari, D. R. (2022). A multiple stage deep learning model for NID in MANETs. In *International Conference on Emerging Smart Computing and Informatics*, 1–6. <https://doi.org/10.1109/ESCIS3509.2022.9758191>
- [42] Sable, N. P., Powar, S. R., Fernandes, Q., Gade, N. A., & Shingade, A. B. (2022). Pragmatic approach for online document verification using block-chain technology. In *International Conference on Automation, Computing and Communication*, 44, 03001. <https://doi.org/10.1051/itmconf/20224403001>
- [43] Javed, M. S., Majeed, H., Mujtaba, H., & Beg, M. O. (2021). Fake reviews classification using deep learning ensemble of shallow convolutions. *Journal of Computational Social Science*, 4, 883–902. <https://doi.org/10.1007/s42001-021-00114-y>
- [44] Zhang, D., Li, W., Niu, B., & Wu, C. (2023). A deep learning approach for detecting fake reviewers: Exploiting reviewing behavior and textual information. *Decision Support Systems*, 166, 113911. <https://doi.org/10.1016/j.dss.2022.113911>
- [45] Pal, K., Poddar, S., Jayalakshmi, S. L., Choudhury, M., Saif Ahmed, S. K., & Halder, S. (2023). Opinion mining-based fake review detection using deep learning technique. In *Proceedings of the 4th International Conference on Data Science, Machine Learning and Applications*, 13–20. https://doi.org/10.1007/978-981-99-2058-7_2
- [46] Aditya, B. L. V. S., & Mohanty, S. N. (2023). Heterogenous social media analysis for efficient deep learning fake-profile identification. *IEEE Access*, 11, 99339–99351. <https://doi.org/10.1109/ACCESS.2023.3313169>
- [47] Jabeur, S. B., Ballouk, H., Arfi, W. B., & Sahut, J.-M. (2023). Artificial intelligence applications in fake review detection: Bibliometric analysis and future avenues for research. *Journal of Business Research*, 158, 113631. <https://doi.org/10.1016/j.jbusres.2022.113631>
- [48] Pan, Y., & Xu, L. (2024). Detecting fake online reviews: An unsupervised detection method with a novel performance evaluation. *International Journal of Electronic Commerce*, 28(1), 84–107. <https://doi.org/10.1080/10864415.2023.2295067>
- [49] Zhang, W., Xie, R., Wang, Q., Yang, Y., & Li, J. (2022). A novel approach for fraudulent reviewer detection based on weighted topic modelling and nearest neighbors with asymmetric Kullback–Leibler divergence. *Decision Support Systems*, 157, 113765. <https://doi.org/10.1016/j.dss.2022.113765>
- [50] Wu, Y., Ngai, E. W., Wu, P., & Wu, C. (2020). Fake online reviews: Literature review, synthesis, and directions for future research. *Decision Support Systems*, 132, 113280. <https://doi.org/10.1016/j.dss.2020.113280>
- [51] Wang, G., Wong, K. W., & Lu, J. (2021). AUC-based extreme learning machines for supervised and semi-supervised imbalanced classification. *IEEE Transactions on Systems, Man, and Cybernetics: Systems*, 51(12), 7919–7930. <https://doi.org/10.1109/TSMC.2020.2982226>
- [52] Song, Y., Wang, L., Zhang, Z., & Hikkerova, L. (2023). Do fake reviews promote consumers' purchase intention? *Journal of Business Research*, 164, 113971. <https://doi.org/10.1016/j.jbusres.2023.113971>
- [53] Budhi, G. S., Chiong, R., Wang, Z., & Dhakal, S. (2021). Using a hybrid content-based and behaviour-based featuring approach in a parallel environment to detect fake reviews. *Electronic Commerce Research and Applications*, 47, 101048. <https://doi.org/10.1016/j.elerap.2021.101048>
- [54] Rong, X., Yi, C., & Tian, Y. (2020). Unambiguous scene text segmentation with referring expression comprehension. *IEEE Transactions on Image Processing*, 29, 591–601. <https://doi.org/10.1109/TIP.2019.2930176>
- [55] Nguyen, T., & Hsu, P.-F. (2022). More personalized, more useful? Reinvestigating recommendation mechanisms in e-commerce. *International Journal of Electronic Commerce*, 26(1), 90–122. <https://doi.org/10.1080/10864415.2021.2010006>
- [56] Mohawesh, R., Xu, S., Tran, S. N., Ollington, R., Springer, M., Jararweh, Y., & Maqsood, S. (2021). Fake reviews detection: A survey. *IEEE Access*, 9, 65771–65802. <https://doi.org/10.1109/ACCESS.2021.3075573>
- [57] Elmurngi, E., & Gherbi, A. (2017). An empirical study on detecting fake reviews using machine learning techniques. In *Seventh International Conference on Innovative Computing Technology*, 107–114. <https://doi.org/10.1109/INTECH.2017.8102442>

How to Cite: Sable, N., Mahalle, P., Kadam, K., Sule, B., Joshi, R., & Deore, M. (2025). Deep Learning-Based Approach for Monitoring and Controlling Fake Reviews. *Journal of Computational and Cognitive Engineering*, 4(3), 377–386. <https://doi.org/10.47852/bonviewJCCE42023602>