

A Systematic Literature Review on Fake Review Detection Using Machine Learning and Deep Learning Techniques

Ruman Maharjan¹, Dr. Bhoj Raj Ghimire²

¹Student, Faculty of Science Health & Technology, Nepal Open University

²Associate Professor, Faculty of Science Health & Technology, Nepal Open University

Abstract

The rapid expansion of e-commerce has intensified consumer dependence on online reviews for making informed purchasing decisions. However, the widespread occurrence of fabricated reviews created to artificially inflate or deflate product and service ratings has emerged as a significant threat to digital trust and market fairness. This systematic literature review synthesizes contemporary research on fake review detection utilizing machine learning and deep learning methodologies. Through rigorous analysis of 31 peer-reviewed publications spanning 2019–2024, this study examines feature engineering approaches, algorithmic frameworks, validation protocols, and performance benchmarks. The findings reveal that hybrid models integrating behavioral and linguistic characteristics consistently outperform unimodal approaches. Ensemble learning techniques and deep neural architectures demonstrate superior detection capabilities compared to conventional classifiers, with Support Vector Machines and ensemble methods such as XGBoost and AdaBoost exhibiting robust performance across heterogeneous datasets. However, these advancements, persistent challenges including concept drift, data imbalance, domain generalization limitations, and adversarial vulnerabilities remain unresolved. This review outlines critical research gaps and proposes actionable directions for future inquiry, emphasizing the imperative for real-time detection architectures, cross-domain adaptability, and resilient countermeasures against evolving deceptive strategies.

Keywords: Fake Review Detection, Machine Learning, Deep Learning, Natural Language Processing, Feature Engineering, Opinion Spam Detection

1. Introduction

1.1 Background

The digital revolution has fundamentally transformed commercial transactions, with online review platforms emerging as pivotal decision-making instruments for contemporary consumers. High-speed internet connectivity and ubiquitous digital access have enabled consumers to evaluate products and services through peer-generated feedback, creating a democratized information ecosystem that significantly influences purchasing behaviors [1],[2]. This paradigm shift has rendered online reviews indispensable assets for businesses, directly impacting brand perception, consumer acquisition, and revenue generation.

The exponential proliferation of internet-mediated commerce has engendered unprecedented reliance on digital testimonials. Empirical evidence indicates that consumers extensively consult online reviews prior to transactional decisions, with favorable ratings substantially augmenting conversion metrics [3]. Conversely, adverse evaluations can substantially diminish consumer interest and compromise business viability. This commercial imperative has unfortunately incentivized deceptive practices wherein malicious actors disseminate inauthentic reviews either to amplify their own offerings or to undermine competitors' market positions [4].

1.2 The Challenges of Fake Reviews

The proliferation of counterfeit reviews has evolved into a systemic challenge confronting digital commerce ecosystems. Malicious actors employ diverse manipulation tactics including:

- **Positive Spamming:** Spreading excessively praising reviews to artificially enhance product visibility.
- **Negative Spamming:** Publishing harmful evaluations to undermine competitive offerings.
- **Content Duplication:** Repurposing existing reviews to amplify deceptive narratives.
- **Coordinated Campaigns:** Composing collective review manipulation through organized networks [5],[6].

Research indicates that between 10-20% of online reviews may be fabricated, varying across platforms and product categories [7],[8]. This pervasive deception erodes consumer confidence, distorts market competition, and imposes economic burdens on both consumers and legitimate businesses.

1.3 Machine Learning in Fake Review Detection

The evolution of data science, particularly the maturation of machine learning and deep learning paradigms, has catalyzed substantial progress in automated fake review identification. Contemporary computational approaches can systematically analyze linguistic patterns, behavioral signatures, and textual characteristics to differentiate authentic from deceptive content [9],[10]. Despite these advancements, persistent challenges including concept drift, class imbalance, domain adaptation limitations, and adversarial vulnerabilities continue to hamper optimal performance.

1.4 Research Objectives and Questions

This systematic literature review pursues three primary objectives: to systematically identify and analyze feature extraction methodologies commonly employed in fake review detection, to critically examine machine learning algorithms applied to this domain, and to investigate deep learning architectures and ensemble techniques.

To achieve these objectives, the review addresses three research questions:

1. What feature extraction methods and machine learning algorithms are predominantly utilized for fake review classification?
2. What are the comparative performance characteristics of various machine learning and deep learning approaches?
3. What are the principal challenges and limitations confronting current fake review detection systems?

1.5 Significance of the Review

This systematic literature review makes several key contributions to the field of fake review detection. It provides a comprehensive synthesis of research published between 2019 and 2024, offering a comparative evaluation of algorithms, features, and performance metrics. The review identifies critical research gaps and methodological limitations in current approaches, proposes evidence-based directions for advancing the field, and delivers practical insights for researchers and practitioners developing detection systems.

2. Methodology

2.1 Search Strategy

A systematic literature review protocol was implemented following established guidelines. The search strategy encompassed:

Search Keywords: "fake review detection," "spam review detection," "deceptive review detection," "opinion spam,"

Academic Databases: Google Scholar, ResearchGate, ScienceDirect, IEEE Xplore, ACM Digital Library.

Temporal Scope: Publications from 2019 to 2024 to capture recent methodological developments.

Inclusion Criteria: Papers were included if they applied machine learning or deep learning techniques for fake review detection, were published in peer-reviewed journals or conference proceedings, provided clear methodological description and performance evaluation, and demonstrated relevance to the research questions.

2.2 Selection Process

The paper selection process comprised four stages:

1. **Initial Identification:** 50 papers identified from reputable journals and conferences.
2. **Screening:** Papers assessed based on titles, abstracts, and keywords.
3. **Eligibility Assessment:** Full-text review for relevance and methodological rigor.
4. **Final Selection:** The final selection of 31 papers was based on the diversity of machine learning and deep learning techniques employed, the variety of feature extraction approaches utilized, relevance to the research questions, and publication in reputable venues.

2.3 Data Extraction

Data extraction employed a structured framework capturing:

- **Bibliographic Information:** Authors, year, publication venue.
- **Research Objectives:** Problem addressed and study goals.
- **Methodology:** Feature extraction, algorithms, computational approaches.
- **Dataset Characteristics:** Source, size, domain, labeling methodology.
- **Validation Protocols:** Cross-validation, train-test splits, evaluation metrics.
- **Performance Metrics:** Accuracy, F1-score, precision, recall, AUC.
- **Key Findings:** Principal contributions and insights.
- **Research Gaps:** Limitations and future work directions.

5. Literature Review and Analysis

5.1 Feature Extraction Methodologies

Feature extraction constitutes a foundational element in fake review detection, categorizable into three principal paradigms: review-centric (textual) features, reviewer-centric (behavioral) features, and hybrid approaches integrating both modalities.

5.1.1 Review-Centric Features

Textual feature extraction focuses on linguistic and stylistic characteristics distinguishing authentic from deceptive content.

Linguistic and Stylometric Features: Abri et al. (2020) investigated 15 linguistic features including redundancy, pausality, and lexical diversity, identifying redundancy, pausality, and adjective frequency as most discriminative. Multilayer Perceptron (MLP) achieved 79.09% accuracy using only four key features. Zeng et al. (2019) proposed a review structure-based ensemble model (RSBE) leveraging

emotional intensity concentrations at review boundaries. Their bidirectional LSTM with self-attention achieved F1-scores of 85.7% for hotel reviews, 85.8% for restaurant reviews, and 85.0% for doctor reviews.

N-gram Features: Siagian and Aritsugi (2020) investigated word and character n-gram combinations, finding that joint utilization consistently outperformed individual feature types. Truthful reviews exhibited greater spelling errors than deceptive ones. Syntactic n-grams demonstrated robustness in cross-domain and cross-type evaluations. Martinez-Torres and Toral (2019) employed polarity-oriented unique attributes for hotel review classification. SVM achieved F1-scores of 0.881 using both comprehensive word sets and reduced feature subsets (134 attributes).

Sentiment Features: Shan et al. (2021) examined review inconsistency, proposing 22 features spanning rating-sentiment, content, and language inconsistency. Random Forest achieved 92.9% accuracy, 93.6% precision, 92.8% recall, and 93.2% F1-score. Moon et al. (2021) conducted survey-based text categorization, developing an "All Terms" model achieving 74.6% accuracy. Linguistic indicators of fabricated reviews included emotional exaggeration, exclamation mark usage, first-person pronoun prevalence, and present/future tense orientation

5.1.2 Reviewer-Centric Features

Behavioral features examine reviewer activity patterns and characteristics indicative of deceptive behavior.

User Activity Features: Barbado et al. (2019) developed the Fake Feature Framework (F3), categorizing features into user-centric (Personal, Social, Reviewing Activity, Trusting) and review-centric categories. AdaBoost achieved F1-score of 82% with user-centric features, while review-centric features alone performed poorly ($F1 < 60\%$).

Temporal and Rating Patterns: Hussain et al. (2020) compared behavioral (SRD-BM) and linguistic (SRD-LM) methods, with SRD-BM achieving 93.1% accuracy using 13 behavioral attributes including content similarity, review burstiness, and activity window. Chopra and Dixit (2023) examined biased ratings (Push Ratings and Nuke Ratings), finding significant impact on CFRS performance, particularly in larger datasets. Push ratings exhibited greater prevalence in fake reviews.

Network-Based Features: Wu et al. (2021) implemented a semi-supervised probabilistic ensemble capturing individual reviewer behaviors and reviewer network structures, achieving F-measures between 0.837 and 0.866. They identified strong homophily effects, indicating spammers' heightened connectivity with other spammers. Noekhah et al. (2020) proposed ten novel features including Sentiment-Rate Difference and Review Group Agreement, enhancing model accuracy through multi-iterative graph-based spamicity propagation. Xue et al. (2019) demonstrated that behavioral features alone are insufficient for detecting sophisticated spam, with aspect-specific opinion deviations constituting stronger spam indicators.

5.1.3 Hybrid Features

Research consistently demonstrates that unimodal approaches whether exclusively textual or behavioral are insufficient for accurate fake review detection, necessitating integrated feature combinations [13]. Alshehri (2024) integrated content-based, behavior-based, and product-based features with PU-learning. Behavioral and product-based features proved more discriminative than purely textual attributes. Budhi et al. (2021) proposed 133 hybrid features combining content-based, metadata, and behavioral elements. CNN and Gradient Boosting performed optimally, with user-perspective behavior features contributing most significantly to prediction accuracy. Alsubari et al. (2021) demonstrated CNN-LSTM hybrid models

achieving 89% accuracy in cross-domain settings. Sun et al. (2024) integrated BERT for text analysis with transformers for reviewer behavior modeling, achieving 89.39% accuracy with 89.48% F1-score.

5.2 Machine Learning Approaches

5.2.1 Supervised Learning

Supervised learning constitutes the predominant paradigm, utilizing fully labeled datasets. SVM, KNN, and LR are among the most frequently employed algorithms [7],[16],[18],[28],[30].

Traditional Machine Learning: Tufail et al. (2022) proposed the SKL approach (SVM, KNN, LR). SVM with 80-20 split achieved 95% accuracy on Yelp and 89.03% on TripAdvisor, consistently outperforming KNN and LR. Alsubari et al. (2022) applied SVM, Naive Bayes, Random Forest, and AdaBoost, with Random Forest achieving 95% accuracy. Truthful reviews contained more nouns and adjectives, while fake reviews exhibited greater verb and adverb usage. Ahmed and Muhammad (2019) compared boosting approaches (XGBoost, AdaBoost, GBM) with traditional classifiers. XGBoost with Chi2 features achieved 95.8% accuracy, significantly outperforming conventional methods. Gutierrez-Espinoza et al. (2020) investigated ensemble learning, finding AdaBoost with MLP achieving 77.3% test accuracy, outperforming stand-alone classifiers.

Semi-Supervised Learning: Semi-supervised approaches address the costly labeling challenge in real-world datasets. Alshehri (2024) employed PU-learning (Positive-Unlabeled learning), achieving F-measure of 0.84 for deceptive reviews with limited labeled data. Ligthart et al. (2021) analyzed semi-supervised learning approaches, with Self-training plus Naive Bayes achieving 93% accuracy with only 20% labeled data. Co-training, TSVM, and Label Propagation did not consistently outperform supervised baselines. Bian et al. (2021) utilized PU-learning categorizing high-confidence reviews as positive samples, achieving best F1-score of 80.0%. Jing-Yu and Ya-Jun (2022) employed AspamGAN with attention mechanism, achieving 86.56% accuracy and 88.18% F1-score with 50% labeled data. Wu et al. (2021) implemented semi-supervised probabilistic ensemble with ten behavioral features, achieving F-measures between 0.837 and 0.866. Kumaran and Chowdary (2019) applied Naive Bayes, Logistic Regression, and SVM with dictionary-based features, achieving 94.968% accuracy on IMDB reviews.

5.2.2 Unsupervised Learning

Unsupervised approaches operate without labeled data, treating fake reviews as anomalies. Noekhah et al. (2020) proposed MGSD, an unsupervised graph-based model achieving 93% accuracy on crowdsourced data and 95.3% on Ott's dataset. Gao et al. (2021) developed unsupervised spam detection with attention mechanism for movie reviews, achieving 87.03% accuracy and 71.13% F1-score for deceptive review detection. Tang et al. (2020) utilized GANs to improve spam detection without historical behavior data, achieving F1-scores of 83.4% (hotel) and 75.1% (restaurant).

5.3 Deep Learning Architectures

Deep learning techniques leverage neural networks to automatically learn features and structural patterns from review data. Shahariar et al. (2019) compared MLP, CNN, and LSTM, with LSTM achieving 94.57% accuracy, outperforming traditional ML classifiers. Alsubari et al. (2021) employed hybrid CNN-LSTM for multi-domain detection, achieving 89% accuracy in cross-domain settings. Sun et al. (2024) integrated BERT with transformers for reviewer behavior modeling, achieving 89.39% accuracy and 0.9531 AUC. Budhi et al. (2021) evaluated MLP, LR, DT, CNN, and SVM with 10-fold cross-validation, with CNN and Gradient Boosting performing optimally post-sampling. Sadiq et al. (2021) trained deep learning models

for star rating prediction, with CNN achieving 89% accuracy; approximately 25% of user ratings exhibited bias. Alawadh et al. (2023) employed Universal Sentence Encoder for semantic embedding, demonstrating robustness of semantically-aware features.

5.4 Ensemble and Boosting Methods

Ensemble and boosting approaches combine multiple models to enhance performance and robustness. Gutierrez-Espinoza et al. (2020) found AdaBoost with MLP outperformed stand-alone classifiers. Ahmed and Muhammad (2019) demonstrated XGBoost achieving 95.8% accuracy, significantly outperforming traditional classifiers. Barbado et al. (2019) reported AdaBoost achieving F1-score of 82% with user-centric features.

5.5 Persistent Challenges

Concept Drift: Mohawesh et al. (2021) conducted comprehensive concept drift analysis, finding classification accuracy declined over time (53.38-68.54% across datasets). Strong negative correlation exists between concept drift and classification performance.

Class Imbalance: Budhi et al. (2021) addressed class imbalance, finding without sampling, fake class accuracy was extremely low (0.08-17.58%). Sampling improved fake detection to 89% but reduced genuine accuracy to 65-69%.

Domain Adaptation: Zeng et al. (2019) found poor cross-domain performance (F1 = 72.9% for Restaurant). Siagian and Aritsugi (2020) observed that features optimized for one domain may not generalize to others.

Adversarial Vulnerabilities: Tang et al. (2020) highlighted the need for adversarial robustness. Gutierrez-Espinoza et al. (2020) emphasized addressing adversarial attacks in detection systems.

5.6 Summary of Findings

Tabel 1: Summary of Findings

Aspect	Key Findings
Features	Behavioral features outperform textual alone; hybrid approaches yield optimal results
Algorithms	SVM and ensemble methods (AdaBoost, XGBoost) consistently demonstrate strong performance
Deep Learning	LSTM and BERT establish new benchmarks but demand substantial computational resources
Semi-Supervised	Effective for limited labeled data scenarios
Challenges	Concept drift, class imbalance, domain adaptation, and adversarial evasion persist

6. Discussion and Analysis

4.1 Algorithmic Performance Comparison

Tabel 2: Algorithmic Performance Comparison

Algorithm	Accuracy Range	Strengths	Limitations
SVM	72-95% (Tufail et al., 2022; Alshehri, 2024)	High-dimensional data handling, robust	Computational cost, limited interpretability

KNN	70-91% (Shahariar et al., 2019)	Simple, non-parametric	Slow inference, distance metric sensitivity
Logistic Regression	80-89% (Tufail et al., 2022)	Interpretable, efficient	Linear decision boundary assumption
Random Forest	71-95% (Alsubari et al., 2022)	Non-linear handling, feature importance	Less interpretable, slower inference
XGBoost	78-96% (Ahmed and Muhammad, 2019)	Highly accurate, scalable	Tuning complexity
LSTM	85-96% (Shahariar et al., 2019)	Context capture, sequence learning	Resource intensive, data hungry
BERT/ Transformers	85-95% (Sun et al., 2024)	Superior semantic understanding	Extremely resource intensive

4.2 Research Gaps and Limitations

Despite significant advancements in fake review detection, several critical research gaps and methodological limitations persist. Most existing models ignore the temporal evolution of fake review strategies, failing to account for concept drift that diminishes model performance over time (Mohawesh et al., 2021). Domain adaptation remains a substantial challenge, as models trained on one domain perform poorly on others, yet limited cross-domain comparative studies exist (Siagian and Aritsugi, 2020). Scalability concerns persist due to limited evaluation on real-time, large-scale systems, with most approaches tested only in supervised contexts (Budhi et al., 2021). Deep learning models lack transparency, highlighting the need for more interpretable deception features (Zeng et al., 2019). Models remain vulnerable to evasion attacks, necessitating improved adversarial robustness (Tang et al., 2020). Dataset limitations are evident, as most studies use small datasets that limit generalizability, highlighting the need for diverse datasets and real-world validation (Martinez-Torres and Toralb, 2019; Gutierrez-Espinoza et al., 2020). Feature diversity is restricted, with research largely limited to English reviews and specific feature types, indicating a need for multilingual detection and behavioral feature integration (Shahariar et al., 2019). Developing scalable solutions for real-time review moderation is essential (Alshehri, 2024), while the emerging threat of AI-generated fake reviews remains inadequately addressed. Furthermore, unsupervised and semi-supervised methods have seen limited exploration compared to supervised approaches, representing an important area for future investigation.

6. Conclusion

This systematic literature review synthesizes contemporary research on fake review detection utilizing machine learning and deep learning techniques. The review reveals that behavioral features demonstrate superior discriminative power compared to textual features alone, while hybrid approaches integrating behavioral and linguistic characteristics consistently achieve optimal accuracy, with user-centric features, particularly reviewing activity patterns, exhibiting the strongest discriminative capabilities. SVM and ensemble methods (AdaBoost, XGBoost) demonstrate consistent performance across heterogeneous datasets, with ensemble learning and deep learning models outperforming traditional machine learning approaches. Deep learning architectures such as LSTM and BERT establish performance benchmarks while demanding substantial computational resources and large datasets, and hybrid CNN-LSTM architectures demonstrate robust cross-domain performance. In terms of contributions, this review provides a comprehensive systematic analysis of literature from 2019 to 2024, synthesizing findings from

31 peer-reviewed studies, and offers a comparative evaluation of feature extraction methodologies and algorithmic approaches. Based on these findings, researchers are recommended to prioritize adaptive models for concept drift, develop robust domain adaptation techniques, explore unsupervised and semi-supervised methods, and address adversarial robustness. Practitioners should implement hybrid feature approaches and ensemble methods while monitoring concept drift regularly. Platform providers are advised to integrate real-time detection systems and ensure balanced datasets. Future research priorities should emphasize generalizability, real-time detection, adversarial robustness, and ethical considerations to develop more trustworthy and effective systems.

References

1. N. Hussain, H. Turab Mirza, I. Hussain, F. Iqbal, and I. Memon, "Spam Review Detection Using the Linguistic and Spammer Behavioral Methods," *IEEE Access*, vol. 8, pp. 53801–53816, 2020, doi: 10.1109/ACCESS.2020.2979226.
2. G. Shan, L. Zhou, and D. Zhang, "From conflicts and confusion to doubts: Examining review inconsistency for fake review detection," *Decis. Support Syst.*, vol. 144, pp. 0–34, 2021, doi: 10.1016/j.dss.2021.113513.
3. S. Moon, M. Y. Kim, and D. Iacobucci, "Content analysis of fake consumer reviews by survey-based text categorization," *Int. J. Res. Mark.*, vol. 38, no. 2, pp. 343–364, 2021, doi: 10.1016/j.ijresmar.2020.08.001.
4. 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, 2020, doi: 10.1109/ICMLA51294.2020.00063.
5. Z. Y. Zeng, J. J. Lin, M. S. Chen, M. H. Chen, Y. Q. Lan, and J. L. Liu, "A review structure based ensemble model for deceptive review spam," *Inf.*, vol. 10, no. 7, pp. 1–15, 2019, doi: 10.3390/info10070243.
6. A. H. A. M. Siagian and M. Aritsugi, "Robustness of Word and Character N-gram Combinations in Detecting Deceptive and Truthful Opinions," *J. Data Inf. Qual.*, vol. 12, no. 1, pp. 1–24, 2020, doi: 10.1145/3349536.
7. H. Tufail, M. U. Ashraf, K. Alsubhi, and H. M. Aljahdali, "The Effect of Fake Reviews on e-Commerce during and after Covid-19 Pandemic: SKL-Based Fake Reviews Detection," *IEEE Access*, vol. 10, pp. 25555–25564, 2022, doi: 10.1109/ACCESS.2022.3152806.
8. S. N. Alsubari, S. N. Deshmukh, M. H. Al-Adhaileh, F. W. Alsaade, and T. H. H. Aldhyani, "Development of Integrated Neural Network Model for Identification of Fake Reviews in E-Commerce Using Multidomain Datasets," *Appl. Bionics Biomech.*, vol. 2021, 2021, doi: 10.1155/2021/5522574.
9. R. Barbado, O. Araque, and C. A. Iglesias, "A framework for fake review detection in online consumer electronics retailers," *Inf. Process. Manag.*, vol. 56, no. 4, pp. 1234–1244, 2019, doi: 10.1016/j.ipm.2019.03.002.
10. H. Xue, Q. Wang, B. Luo, H. Seo, and F. Li, "Content-aware trust propagation toward online review spam detection," *J. Data Inf. Qual.*, vol. 11, no. 3, 2019, doi: 10.1145/3305258.
11. A. B. Chopra and V. S. Dixit, "Detecting biased user-product ratings for online products using opinion mining," *J. Intell. Syst.*, vol. 32, no. 1, 2023, doi: 10.1515/jisys-2022-9030.

12. S. Noekhah, N. binti Salim, and N. H. Zakaria, "Opinion spam detection: Using multi-iterative graph-based model," *Inf. Process. Manag.*, vol. 57, no. 1, p. 102140, 2020, doi: 10.1016/j.ipm.2019.102140.
13. M. Abdulqader, A. Namoun, and Y. Alsaawy, "Fake Online Reviews: A Unified Detection Model Using Deception Theories," *IEEE Access*, vol. 10, no. November, pp. 128622–128655, 2022, doi: 10.1109/ACCESS.2022.3227631.
14. P. Sun et al., "Fake Review Detection Model Based on Comment Content and Review Behavior," *Electron.*, vol. 13, no. 21, 2024, doi: 10.3390/electronics13214322.
15. P. Bian, L. Liu, and P. Sweetser, "Detecting Spam Game Reviews on Steam with a Semi-Supervised Approach," *ACM Int. Conf. Proceeding Ser.*, vol. 1, no. 1, pp. 1–14, 2021, doi: 10.1145/3472538.3472547.
16. M. R. Martinez-Torres and S. L. Toral, "A Machine Learning approach for the Identification of the Deceptive Reviews in the Hospitality Sector using Unique Attributes and Sentiment Orientation," *Tourism Management*, vol. 75, 2019.
17. R. Mohawesh, S. Tran, R. Ollington, and S. Xu, "Analysis of concept drift in fake reviews detection," *Expert Syst. Appl.*, vol. 169, p. 114318, 2021, doi: 10.1016/j.eswa.2020.114318.
18. A. Ligthart, C. Catal, and B. Tekinerdogan, "Analyzing the effectiveness of semi-supervised learning approaches for opinion spam classification," *Appl. Soft Comput.*, vol. 101, p. 107023, 2021, doi: 10.1016/j.asoc.2020.107023.
19. H. M. Alawadh, A. Alabrah, T. Meraj, and H. T. Rauf, "Semantic Features-Based Discourse Analysis Using Deceptive and Real Text Reviews," *Inf.*, vol. 14, no. 1, pp. 1–16, 2023, doi: 10.3390/info14010034.
20. S. N. Alsubari et al., "Data analytics for the identification of fake reviews using supervised learning," *Comput. Mater. Contin.*, vol. 70, no. 2, pp. 3189–3204, 2022, doi: 10.32604/cmc.2022.019625.
21. D. S. N. Kumaran and C. H. Chowdary, "Detection of fake online reviews using semi-supervised and supervised learning," in 2nd Int. Conf. Electr. Comput. Commun. Eng. ECCE 2019, 2019, doi: 10.1109/ECACE.2019.8679186.
22. A. H. Alshehri, "An Online Fake Review Detection Approach Using Famous Machine Learning Algorithms," *Comput. Mater. Contin.*, vol. 78, no. 2, pp. 2767–2786, 2024, doi: 10.32604/cmc.2023.046838.
23. C. Jing-Yu and W. Ya-Jun, "Semi-Supervised Fake Reviews Detection based on AspamGAN," *J. Artif. Intell. Capsul. Networks*, vol. 4, no. 1, pp. 17–36, 2022, doi: 10.36548/jaicn.2022.1.002.
24. Z. Wu, G. Liu, J. Wu, and J. Tan, "Are Neighbors Alike? A Semi-supervised Probabilistic Ensemble for Online Review Spammers Detection," *SSRN*, 2021.
25. X. Tang, T. Qian, and Z. You, "Generating behavior features for cold-start spam review detection with adversarial learning," *Inf. Sci. (Ny)*, vol. 526, pp. 274–288, 2020, doi: 10.1016/j.ins.2020.03.063.
26. Y. Gao, M. Gong, Y. Xie, and A. K. Qin, "An Attention-Based Unsupervised Adversarial Model for Movie Review Spam Detection," *IEEE Trans. Multimed.*, vol. 23, pp. 784–796, 2021, doi: 10.1109/TMM.2020.2990085.
27. G. Satia Budhi, R. Chiong, Z. Wang, and S. Dhakal, "Using a hybrid content-based and behaviour-based featuring approach in a parallel environment to detect fake reviews," *Electron. Commer. Res. Appl.*, vol. 47, 2021, doi: 10.1016/j.elerap.2021.101048.

28. G. M. Shahariar, S. Biswas, F. Omar, F. M. Shah, and S. Binte Hassan, "Spam Review Detection Using Deep Learning," 2019 IEEE 10th Annu. Inf. Technol. Electron. Mob. Commun. Conf. IEMCON 2019, pp. 27–33, 2019, doi: 10.1109/IEMCON.2019.8936148.
29. S. Sadiq, M. Umer, S. Ullah, S. Mirjalili, V. Rupapara, and M. Nappi, "Discrepancy detection between actual user reviews and numeric ratings of Google App store using deep learning," *Expert Syst. Appl.*, vol. 181, no. June, 2021, doi: 10.1016/j.eswa.2021.115111.
30. L. Gutierrez-Espinoza, F. Abri, A. S. Namin, K. S. Jones, and D. R. W. Sears, "Fake Reviews Detection through Ensemble Learning," pp. 1–8, 2020, [Online]. Available: <http://arxiv.org/abs/2006.07912>
31. S. Ahmed and F. Muhammad, "Using Boosting Approaches to Detect Spam Reviews," 1st Int. Conf. Adv. Sci. Eng. Robot. Technol. 2019, ICASERT 2019, doi: 10.1109/ICASERT.2019.8934467.