

Unmasking Fake Opinions through Behavioral Analysis and Machine Learning: Identifying Genuine Users vs. Fraudulent Actors

Ravishankar S. Ulle^{1,*}, S. Yoganathan²

^{1,2}Department of Management Studies, CMS Business School, Jain University, Bangalore, Karnataka, India.
dr.ravishanakrulle@cms.ac.in¹, dr.s.yoganathan@cms.ac.in²

Abstract: The proliferation of fake online opinions undermines consumer trust and distorts decision-making processes. Traditional detection methods relying on content analysis face limitations, such as difficulty in identifying sophisticated fraudulent behavior and adapting to new patterns. This paper investigates the potential of combining behavioral analysis with machine learning (ML) to improve the detection of fake reviews. Through a comprehensive review of existing literature, we explore behavioral analysis techniques for identifying suspicious activities and the application of ML algorithms for automated detection. We propose a conceptual framework focusing on reviewer behavior, review content, and review authenticity as primary variables while considering platform characteristics and product categories as moderating factors and reviewer motivation as a mediating factor. The integration of these dimensions aims to capture the nuances of fraudulent activities and enhance detection accuracy. By identifying key research gaps, such as the lack of real-time detection methods and insufficient focus on behavioral indicators, this review formulates targeted research questions to guide future studies. Our findings suggest that the synergy between behavioral analysis and ML holds promise for developing robust systems to unmask fake online opinions. This research contributes to advancing detection methods and restoring consumer trust in online platforms.

Keywords: Fake Online Opinions; Behavioral Analysis; Machine Learning (ML); Fraudulent Actors; Content Analysis; Sentiment Analysis; Reviewer Motivation; Support Vector Machine (SVM); Adaptive Boosting (AB).

Received on: 21/11/2023, **Revised on:** 25/01/2024, **Accepted on:** 11/03/2024, **Published on:** 03/06/2024

Journal Homepage: <https://www.fmdbpublish.com/user/journals/details/FTSML>

DOI: <https://doi.org/10.69888/FTSML.2024.000235>

Cite as: R. S. Ulle, and S. Yoganathan, “Unmasking Fake Opinions through Behavioral Analysis and Machine Learning: Identifying Genuine Users vs. Fraudulent Actors,” *FMDB Transactions on Sustainable Management Letters*, vol. 2, no. 2, pp. 51–64, 2024.

Copyright © 2024 R. S. Ulle, and S. Yoganathan, licensed to Fernando Martins De Bulhão (FMDB) Publishing Company. This is an open access article distributed under [CC BY-NC-SA 4.0](https://creativecommons.org/licenses/by-nc-sa/4.0/), which permits copying, remixing, redistributing, transformation, and building upon the material in any medium so long as the original work is properly cited.

1. Introduction

The internet has become a central hub for consumer reviews, influencing purchasing decisions in a major way. However, this trust in online opinions can be easily manipulated by the presence of fake reviews. These fabricated opinions, often created by bots or inauthentic accounts, can paint an unrealistic picture of a product or service [19]. The prevalence of fake online opinions is a growing concern. Studies suggest that a significant portion of online reviews – estimates range from 10% to 30% – may be misleading or completely fabricated [20]. This infiltration of fake reviews erodes consumer trust in online information. When consumers encounter a sea of inauthentic praise or negativity, they become skeptical of all reviews, making it difficult to distinguish genuine opinions from manipulative ones [21]. This loss of trust has a ripple effect, impacting consumer decision-

*Corresponding author.

making. Consumers rely on reviews to compare products, assess quality, and identify potential problems. Fake reviews distort this process, leading consumers to make uninformed choices. In the worst-case scenario, they may be swayed towards a subpar product or service or miss out on a genuinely good one [22]. Ultimately, the prevalence of fake online opinions undermines the entire review system, hindering its ability to serve as a valuable resource for consumers [23]. While content analysis has been the initial approach to tackling fake reviews, it has limitations that struggle to keep pace with increasingly sophisticated tactics. Here's why solely relying on content analysis falls short:

- **Keyword Focus:** Traditional methods often rely on identifying specific keywords or phrases associated with fake reviews. However, fraudulent actors can easily adapt their language to bypass these filters. They can use synonyms, rephrase sentences, or even resort to sentiment manipulation where the overall tone is positive or negative, but the content lacks specifics [24].
- **Limited Context:** Content analysis often struggles to understand the context of a review. For instance, a single negative review might be legitimate, but a flurry of negative reviews with similar language and timing could be a red flag. Traditional methods may miss these connections, overlooking patterns that indicate coordinated inauthentic activity [25].
- **Evolving Tactics:** Fraudulent actors constantly refine their techniques. As detection methods based on specific content become ineffective, they resort to more nuanced tactics like incorporating genuine-sounding details or mimicking real user behavior. Content analysis often fails to adapt to these ever-changing strategies [26].
- **Labor Intensive:** Manually reviewing content for red flags can be a time-consuming and resource-intensive process. This makes it difficult to scale content analysis for large platforms with vast amounts of user-generated content [27].

As the limitations of content analysis become more apparent, researchers are exploring alternative methods for identifying fake online opinions. One promising approach is behavioral analysis [28]. This method goes beyond the content of the review itself and focuses on the behaviour of the user who posted it. Behavioral analysis examines a user's activity patterns on the platform. This can include factors like:

- **Review Frequency:** Does the user leave an unusually high number of reviews in a short period?
- **Rating Consistency:** Do their ratings consistently skew positive or negative, regardless of the product or service?
- **Time Between Reviews:** Are there unusually short gaps between reviews, suggesting automated activity?
- **Purchase History:** Has the user actually purchased the product they are reviewing?

By analyzing these behavioural patterns, researchers can identify inconsistencies that might point toward a fraudulent actor. For instance, a user leaving a glowing review for a product they haven't purchased or posting a series of negative reviews within minutes of each other could be red flags [29]. The potential of behavioral analysis lies in its ability to uncover hidden patterns and inconsistencies that content analysis might miss. By examining user behavior, we can gain a more holistic understanding of the reviewer's authenticity and intentions [30]. This deeper level of analysis holds promise for creating more robust detection methods that can stay ahead of evolving tactics employed by those leaving fake online opinions [31].

While behavioural analysis offers valuable insights, manually sifting through vast amounts of user data to identify patterns can be overwhelming [32]. This is where machine learning (ML) comes into play [33]. ML algorithms can be trained on large datasets of user behavior and review characteristics associated with both genuine and fake opinions [34]. Once trained, these algorithms can automatically analyze user behavior and review content, identifying patterns and inconsistencies that might be indicative of fraudulent activity [35]. This automation significantly reduces the manual effort required for detection and allows for scaling up the process to handle the massive volume of online reviews generated daily [36].

Furthermore, ML algorithms can continuously learn and improve over time. As they are exposed to new data and encounter new tactics employed by fraudulent actors, they can adapt and refine their detection methods [37]. This continual learning process ensures that the system remains effective in the face of ever-evolving threats. In essence, machine learning acts as a powerful tool that automates and streamlines the process of analyzing user behavior and reviewing content for signs of inauthenticity [38]. This allows researchers and platforms to develop more robust and scalable detection methods that can effectively combat the growing problem of fake online opinions.

2. Literature Review

2.1. Moving Beyond Content: A Look at Behavioral Analysis for Detecting Fake Reviews

While content analysis has served as the initial bulwark against deceptive online reviews, its limitations become increasingly apparent in the face of ever-evolving tactics employed by those leaving fake reviews. To address this challenge, researchers

are turning to behavioral analysis, a method that delves deeper than the review content itself, scrutinizing the reviewer's activity for signs of inauthenticity. This multifaceted approach leverages several key techniques. One method involves analyzing reviewer activity patterns [39]. A sudden influx of reviews, particularly skewed towards positive or negative sentiment, from a new user, could be an attempt to manipulate product perception. Similarly, unusually short intervals between reviews might indicate automated activity, a hallmark of fake reviews. Purchase history can also offer valuable insights [40]. Inconsistencies between a reviewer's purchase history and their reviews warrant investigation. For instance, a user leaving a detailed review for a product they haven't purchased raises red flags. Platforms can leverage purchase data to verify reviewers' claims and identify potential fraudulent actors [41].

Social network analysis becomes particularly revealing on platforms that incorporate social networking features. If a user with a limited social network suddenly leaves a flurry of reviews for similar products, it could suggest a coordinated inauthentic activity, where multiple fake accounts work together to manipulate reviews [42]. Engagement metrics, which extend beyond the review content itself, can also be helpful. Does the reviewer respond to comments or questions? Do they interact with other users on the platform? A lack of engagement, particularly for reviews expressing strong opinions, could indicate a fake account. While content analysis has its place, its effectiveness is amplified when combined with behavioral analysis [43]. For example, a reviewer consistently leaving positive reviews that contain negative sentiment words (e.g., "disappointed, but overall good") could be a sign of inauthentic activity. This combined approach helps identify reviews that might appear positive on the surface but harbor hidden negativity [44]. By employing these techniques in tandem, behavioral analysis offers a promising path forward in the fight against fake reviews, ensuring the online review system remains a reliable source of information for consumers [45]. These techniques, used individually or in combination, can paint a more comprehensive picture of a reviewer's authenticity. By analyzing behavior patterns and inconsistencies, platforms can develop more robust detection methods to identify and combat fake online opinions [46].

Molina et al., [1] identify seven types of content under "fake news": false news, polarized content, satire, misreporting, commentary, persuasive information, and citizen journalism. These types are contrasted with "real news" using a taxonomy of operational indicators across four domains: message, source, structure, and network. This taxonomy aims to clarify the nature of online news content and enhance the accuracy of detection algorithms. Zannettou et al., [2] highlight the urgency of addressing this problem. It proposes a way to categorize the different types of false information, the actors spreading it, and their motives. The article also reviews existing research on how people perceive false information, how it spreads, and how to detect and contain it. It emphasizes the particular dangers of political misinformation, which can spread faster and have more severe consequences. Finally, the article proposes future research directions to help us combat the spread of false information online.

Luceri et al., [3] examine the detection of key players in state-sponsored information operations (IOs) on social media platforms, specifically Twitter. By utilizing similarity graphs based on behavioural pattern data, the study identifies coordinated IO drivers using network properties that have been underutilized until now. Analyzing a comprehensive dataset of 49 million tweets from six countries, which includes multiple verified IOs, the research highlights the limitations of traditional network filtering techniques in consistently identifying IO drivers across campaigns. To address these limitations, the paper proposes a framework based on node pruning, which proves more effective when combining multiple behavioral indicators across different networks. Additionally, the study introduces a supervised machine-learning model that utilizes a vector representation of the fused similarity network. This model achieves a precision exceeding 0.95, effectively classifying IO drivers on a global scale and reliably predicting their temporal activities. The findings of this study are significant in the fight against deceptive influence campaigns on social media, providing better tools and methods to understand and detect such operations.

Iacobucci et al., [10] investigate whether priming users with information about deepfake (DF) media enhances their ability to recognize such content. Deepfake videos, created using advanced AI, are highly realistic and possess significant deceptive potential. In response, practitioners and institutions are developing debunking strategies to mitigate the spread of misleading DF videos. The research focuses on two main questions: does priming users with definitions and potential harms of DFs improve their recognition of these videos? Additionally, does an individual's susceptibility to epistemically suspect beliefs (bullshit receptivity) affect the effectiveness of this priming? The results show that educational and cultural strategies can effectively counter DF deception, but primarily for individuals less prone to believing misleading claims. A serial mediation analysis further reveals that better DF recognition reduces users' intentions to share such content, thus addressing the issue of DF virality at its root. The study concludes that a simple, well-reasoned digital literacy intervention could significantly enhance society's defense against the threats posed by DFs. This finding emphasizes the importance of digital literacy in mitigating the impact of deepfake media.

Al-Adhaileh and Alsaade [4], the primary goal of this paper is to distinguish between fake and truthful product reviews through a seven-phase methodology. This approach involves reviewing online products, analyzing linguistic features using Linguistic Inquiry and Word Count (LIWC), preprocessing data to clean and normalize it, embedding words using Word2Vec, and

employing artificial deep-learning algorithms for classification. The study evaluates two deep-learning neural network models, bidirectional long-short-term memory (BiLSTM) and convolutional neural network (CNN), using standard Yelp product reviews. Results indicate that the BiLSTM model outperforms the CNN model in detecting fake reviews, achieving higher accuracy. This research underscores the importance of leveraging advanced techniques such as deep learning to combat the proliferation of fake reviews online. By employing sophisticated methodologies to analyze linguistic features and embedding words, the study aims to provide a reliable framework for distinguishing between genuine and deceptive product reviews. The findings suggest that the BiLSTM model holds promise for effectively identifying fake opinions, offering valuable insights for e-commerce platforms and social media companies seeking to maintain trust and integrity in their review systems.

Alsubari et al., [5], the prevalence of fake reviews, also known as deceptive opinions, poses a significant challenge in online marketing transactions, where e-commerce platforms provide customers with the opportunity to post reviews and comments about products or services. This abundance of reviews complicates the task for new customers seeking to distinguish between truthful and fake opinions, potentially leading to deception, financial losses, and reputational damage for companies. In response, this paper aims to develop an intelligent system for detecting fake reviews on e-commerce platforms by leveraging n-grams of review text and sentiment scores provided by reviewers. The proposed methodology involves utilizing a standard fake hotel review dataset for experimentation and employing data preprocessing techniques alongside a term frequency-inverse document frequency (TF-IDF) approach for feature extraction and representation [47]. Four different supervised machine-learning techniques, including naïve Bayes (NB), support vector machine (SVM), adaptive boosting (AB), and random forest (RF), are employed for detection and classification. These techniques are trained and tested on a dataset collected from the Trip Advisor website, with classification results showing testing accuracy and F1-score of 88% (NB), 93% (SVM), 94% (AB), and 95% (RF). Comparisons with existing works using the same dataset indicate that the proposed methods outperform comparable approaches in terms of accuracy, demonstrating the effectiveness of the developed system in identifying fake reviews on e-commerce platforms [48].

Oh and Park [6], detecting deception in online comments, particularly regarding social issues, presents a significant challenge, as humans often struggle to discern fake opinions accurately. While efforts have been largely focused on identifying fake consumer reviews, techniques for detecting deceptive opinions on social issues remain underexplored. Addressing this gap, this study aims to develop an automated machine-learning technique for determining the trustworthiness of comments in online discussions. The research introduces a large-scale ground truth dataset comprising 866 truthful and 869 deceptive comments on social issues, representing one of the first attempts to detect comment deception in Asian languages, specifically Korean. The proposed machine-learning technique achieves an impressive accuracy of nearly 81% in identifying untruthful opinions about social issues, surpassing the performance of human judges. This breakthrough offers promising potential for improving the reliability of online discourse and combating deceptive practices in social media environments.

Elmoghy et al., [7] extract the behavioral attributes of reviewers alongside review characteristics. The study conducts several experiments on a real Yelp dataset of restaurant reviews, comparing the performance of machine learning classifiers, including KNN, Naive Bayes (NB), and Logistic Regression. Results indicate that Logistic Regression achieves the highest accuracy among the classifiers tested, demonstrating superior capability in distinguishing between fake and genuine reviews. This research contributes to ongoing efforts in the dynamic field of fake review detection, offering insights into effective machine-learning techniques for enhancing the authenticity and reliability of online reviews on E-commerce platforms. Bhatt et al., [8] present a comprehensive review of research in cognitive behavior analysis, examining various parameters, including physical characteristics, emotional behaviors, data collection methods, unimodal and multimodal datasets, and AI/ML modeling techniques. It discusses challenges and future research directions in utilizing AI/ML for inferring human behavior, contributing to advancements in behavioral science and forensic investigations.

Alhazbi [9] argues that the behaviors of sponsored troll accounts differ from those of ordinary users due to their extrinsic motivation, making them distinguishable through machine learning techniques based on their activities on social media platforms. The study proposes a set of behavioral features to detect political troll accounts on Twitter. It develops four classification models based on decision trees, random forest, Adaboost, and gradient boost algorithms. Using a dataset of Saudi trolls disclosed by Twitter in 2019, the models achieve an overall classification accuracy of up to 94.4%. Moreover, the models demonstrate the ability to identify Russian trolls with an accuracy of up to 72.6%, even without specific training on this dataset. These findings suggest that although the strategies of coordinated trolls may vary across organizations, they exhibit common behaviors that can be effectively identified through machine learning approaches, highlighting the potential for automated detection of sponsored troll accounts on social media platforms [49].

Patel and Patel [11] explore various data mining techniques, including supervised, unsupervised, and semi-supervised approaches, for fake review detection based on different features. By leveraging these techniques, researchers aim to develop effective methods to identify and filter out fraudulent reviews, thereby enhancing the reliability of online review systems and protecting consumers from misinformation. Sultana and Palaniappan [12] delve into deception detection in customer reviews,

employing various supervised machine learning methods. A machine learning model utilizing the stochastic gradient descent algorithm is proposed for spam review detection, integrating bagging and boosting techniques to reduce bias and variance. Additionally, feature selection using regular expressions is employed to select the most relevant features. Experiments conducted on a hotel review dataset demonstrate the effectiveness of the proposed approach in detecting fake reviews.

Bhargava and Choudhary [18] propose a novel approach to address this group-level manipulation. The methodology leverages a technique called "deep-walk" to create a behavioral representation for each suspected fake reviewer account. By analyzing online activity patterns, deep-walk helps identify characteristic behaviors associated with deceptive practices. Subsequently, the research employs a combination of supervised and unsupervised machine learning algorithms to analyze the behavioral data and uncover groups of fake reviewers working in concert. This research offers a significant contribution to the field of online review trustworthiness. By effectively detecting and eliminating group-level manipulation, the proposed approach paves the way for a more reliable and secure social media environment. This, in turn, fosters trust and transparency in online product reviews, empowering users to make informed decisions based on authentic customer experiences. Ahmed et al., [13] review explores the use of machine learning classifiers to detect fake news automatically. These algorithms can analyze vast amounts of data, searching for patterns and linguistic cues that differentiate factual news from fabricated stories. By implementing such automated detection systems, we can strive towards a more reliable online environment.

Manaskasemsak et al., [14] propose two novel graph partitioning approaches, BeGP and BeGPX, to distinguish fake reviewers from genuine ones. BeGP constructs a behavioral graph where reviewers are connected based on shared characteristic features indicative of similar behavior. The algorithm initiates with a small subgraph of known fake reviewers and expands it by including connected suspicious reviewers, hypothesizing that their reviews are untruthful. BeGPX enhances fake review detection by incorporating semantic content and emotions expressed in reviews. It utilizes deep neural networks to learn word embeddings and lexicon-based emotion indicators for graph construction. Both approaches are evaluated on real-world review datasets from Yelp.com, outperforming state-of-the-art methods in accurately identifying fake reviewers within the k-first order of rankings. BeGPX exhibits significant improvement, even with limited labeled data, demonstrating its effectiveness in enhancing fake review detection.

Alsaad and Joshi, [15], user-generated reviews wield significant influence in e-commerce, impacting organizations' revenue and reputation. The trustworthiness of these reviews is paramount, as customers often rely on them to make purchasing decisions. However, fraudulent reviews created by individuals hired by online companies aim to deceive customers and manipulate their decisions. Despite extensive research in the past two decades, there remains a lack of comprehensive literature surveys addressing the methods and challenges of detecting fake reviews. To address this gap, this survey consolidates publicly available datasets and their acquisition methods for detecting fraudulent reviews. It scrutinizes current approaches for feature engineering in fake review analysis, as well as the application of deep learning and classical machine learning for fake review classification. By identifying inconsistencies and limitations in existing methods, this survey aims to contribute to the advancement of fraud detection in online reviews.

Mohseni et al., [16], in the current era marked by the prevalence of fake news and misinformation, combatting their propagation poses significant challenges. Algorithms used in news feeds and search platforms may inadvertently contribute to the widespread dissemination of false information. To address this issue, our research explores the effectiveness of integrating an Explainable AI assistant into news review platforms to counter the spread of fake news. We developed a new reviewing and sharing interface, curated a dataset of news stories, and trained four interpretable fake news detection algorithms. These algorithms were designed to study the impact of algorithmic transparency on end users. Through multiple controlled crowd-sourced studies, we evaluated the effectiveness of these Explainable AI systems. We analyzed the interactions between user engagement, mental models, trust, and performance measures during the explanation process. The results of our study indicate that explanations provided by the AI assistant helped participants develop appropriate mental models of the system and adjust their trust levels based on their perceptions of the model's limitations.

Wang et al., [17], in recent times, online reviews have become instrumental in guiding purchase decisions by providing customers with valuable insights into products or services. However, the proliferation of fake reviews created by spammers to artificially promote or degrade the quality of goods or services presents a significant challenge. Such fraudulent behavior can mislead customers and lead to erroneous decisions. In this paper, we propose the utilization of two novel feature sets, namely readability features and topic features, along with supervised machine learning algorithms to address this issue using real-life data from Yelp. Our findings reveal that these new features outperform traditional n-gram features in detecting spam reviews. Additionally, incorporating behavioral features about reviewers and their reviews significantly enhances the classification accuracy of genuine Yelp review data. Furthermore, balancing the number of reviewers improves the overall classification performance, highlighting the efficacy of our approach in combating opinion spam.

2.2. Leveraging Machine Learning to Combat Fake Reviews: A Behavioral Analysis Approach

While content analysis has served as the initial line of defense against deceptive online reviews, its limitations become increasingly apparent. To address this challenge, researchers are turning to machine learning (ML) to bolster behavioural analysis [50]. This approach delves deeper, scrutinizing reviewer activity beyond the content of the review itself to identify signs of inauthenticity. ML algorithms play a crucial role in automating and scaling up the detection of fake reviews. Here's a look at some common approaches:

- **Anomaly Detection:** These algorithms excel at identifying data points that deviate significantly from established patterns. In the context of fake reviews, anomaly detection algorithms can analyze reviewer behavior patterns such as review frequency, rating consistency, and time between reviews. Significant deviations from these patterns can flag potential fraudulent activity [51].
- **Supervised Learning:** This supervised learning approach involves training ML models on labeled datasets consisting of genuine and fake reviews. These models learn to identify features associated with each category, such as reviewer activity patterns, sentiment inconsistencies within reviews (e.g., positive reviews with negative words), and the use of specific language patterns often found in fake reviews. Once trained, the model can analyze new reviews and predict their authenticity with a high degree of accuracy [52].
- **Sentiment Analysis as a Supporting Tool:** Sentiment analysis, while not without limitations, becomes a valuable tool when combined with other techniques. ML algorithms can analyze the sentiment expressed in a review and compare it to the reviewer's overall rating. For instance, a review overflowing with positive sentiment but containing a high number of negative words could be indicative of a fake review.

2.3. Effectiveness and Challenges of ML-powered Behavioral Analysis

ML algorithms offer significant advantages for automated detection. They can efficiently process vast amounts of data, identify complex patterns in reviewer behavior, and continuously learn and improve over time. However, their effectiveness depends on the quality and size of the training data. Additionally, fraudulent actors constantly adapt their tactics, requiring the models to be updated regularly to maintain accuracy. It's an ongoing arms race where researchers constantly refine the models to stay ahead of evolving tactics employed by those leaving fake reviews. Overall, machine learning provides a powerful set of tools for detecting fake opinions. By combining different algorithms like anomaly detection, supervised learning, and sentiment analysis, platforms can develop robust and scalable detection systems. However, staying ahead of evolving tactics and ensuring the quality of training data remain ongoing challenges.

3. Conceptual Design

3.1. Independent Variables: Unveiling Fake Reviews

Independent variables in your research design act as the potential causes or influences you're examining. In this case, you're investigating what factors might indicate a fake online opinion. Here's a breakdown of the two key independent variables you've identified:

- **Reviewer Behavior:** This category delves into the actions and patterns exhibited by the reviewers themselves. It goes beyond the content of the review and focuses on user activity.
- **Review Frequency:** Analyzing the number of reviews a user leaves within a specific timeframe can be revealing. A sudden surge in reviews, especially all positive or negative, could suggest an attempt to manipulate product perception. Unusually short intervals between reviews might also indicate automated activity.
- **Rating Consistency:** Does the user consistently leave extreme ratings (all positive or negative) regardless of the product or service? This lack of variation could be a red flag.
- **Time Between Reviews:** Examining the time gaps between a user's reviews can provide insights. Abnormally short intervals could suggest automated posting, while very long gaps might not be indicative of genuine user behavior in all situations.
- **Purchase History:** Inconsistencies between a reviewer's purchase history and their reviews raise suspicion. For instance, a user leaving a detailed review for a product they haven't purchased is unlikely to be a genuine customer.
- **Review Content:** This variable focuses on the actual content of the review itself. While limitations exist in relying solely on content analysis, it can still provide valuable clues when combined with reviewer behavior. Here's what to consider:
- **Sentiment Analysis Scores:** Machine learning algorithms can analyze the overall sentiment expressed in the review, identifying positive, negative, or neutral language.

- **Word Usage Patterns:** Are there specific words or phrases commonly found in fake reviews? Research can identify red flags like excessive use of promotional language, generic positive or negative terms, or an abundance of exclamation points.
- **Presence of Red Flags:** Studies have identified specific content patterns associated with fake reviews. These might include irrelevant details, excessive self-promotion, or inconsistencies in the review itself (e.g., mentioning features not present in the product).

3.2. Dependent Variable: Unveiling the Truth - Review authenticity

The dependent variable in your research is the outcome you're trying to predict or explain – in this case, the authenticity of an online review. Here's a closer look:

- **Review Authenticity:** This is a binary variable with two possible values: genuine (real user, authentic opinion) or fake (fraudulent actor, manipulated opinion).
- **Determining Authenticity:** There are two main approaches to establish the dependent variable:
- **Ground Truth Data:** This is the ideal scenario where you have access to verified data on review authenticity. This could involve human verification by experts or confirmation of fraudulent campaigns by platforms. However, obtaining such data can be challenging and resource-intensive.
- **Advanced Detection Algorithms:** Machine learning algorithms can be trained on large datasets of labeled reviews (genuine and fake) to predict the authenticity of new reviews. These algorithms rely on identifying patterns in reviewer behavior and reviewing content associated with each category.

3.3. Moderating Variables: Nuances in the Online Landscape

While reviewer behavior and review content are crucial factors in identifying fake reviews, their effectiveness can be influenced by other aspects of the online environment. These are known as moderating variables, and they can affect the relationship between the independent and dependent variables. Here's a look at the two moderating variables you've identified:

- **Platform Characteristics:** The platform where the review is posted can play a significant role. Different platforms have unique user demographics and review moderation practices and functionalities. These factors can influence the effectiveness of detection methods:
- **User Behavior Patterns:** User behavior on an e-commerce platform might differ from that on a social media platform. For instance, users on e-commerce sites might be more likely to leave reviews after completing a purchase, while those on social media might be swayed by social influence. Detection methods need to be adapted to account for these behavioral variations across platforms.
- **Review Moderation Practices:** Platforms have varying degrees of review moderation. Some platforms have stricter policies and employ automated filters, making it more difficult for fraudulent actors to post fake reviews. Conversely, platforms with lax moderation might require more sophisticated detection methods.
- **Product Category:** The type of product being reviewed can also influence the landscape of fake reviews:
- **Fraudulent Actor Strategies:** The tactics employed by fraudulent actors might differ depending on the product category. For instance, fake reviews for electronics might focus on technical specifications, while those for clothing might emphasize fit and style. Detection methods need to be tailored to identify red flags specific to each product category.
- **Effectiveness of Detection Approaches:** The effectiveness of detection methods can vary depending on the product. For example, analyzing purchase history might be more relevant for physical products on e-commerce platforms, while analyzing social network connections might be more insightful for reviews on social media platforms related to services.

3.4. Mediating Variable (Optional): Unveiling the Why - Reviewer Motivation

While not essential for every research design, including a mediating variable can provide valuable insights into the "why" behind fake reviews. Here's how reviewer motivation can act as a mediating variable in your research:

3.4.1. Reviewer Motivation

This variable delves into the underlying reasons why someone might post a fake review. Understanding these motivations can help explain the relationship between reviewer behavior and review authenticity (the independent and dependent variables). Here are some potential motivations:

- **Financial Gain:** Fraudulent actors might be incentivized by financial rewards, either for posting positive reviews for specific products or leaving negative reviews to damage a competitor.
- **Promoting a Competitor:** Businesses might resort to posting fake positive reviews for their products or negative reviews for competing products to manipulate market perception.
- **Damaging a Brand:** Motivations can be malicious, with individuals aiming to damage a brand's reputation through fake negative reviews.
- **Social Influence:** In some cases, reviewers might be swayed by social influence or a desire to "fit in" with a particular online community, leading them to post fake reviews that align with the dominant sentiment.

3.4.2. Mediating the Relationship

Reviewer motivation can act as a mediator by explaining the link between reviewer behavior and review authenticity. For instance, a user leaving a suspiciously high number of positive reviews within a short period (independent variable) might be motivated by financial gain (mediating variable) and, therefore more likely to be posting fake reviews (dependent variable).

3.4.3. Tailoring Detection Strategies

Identifying reviewer motivation can help tailor detection methods. For example, focusing on analyzing purchase patterns might be more effective in identifying reviews motivated by financial gain. At the same time, social network analysis could be more insightful for reviews influenced by social pressure.

4. Research Gap and Proposed Research Questions

4.1. Data Availability and Quality

- Limited access to ground truth data: Obtaining verified data on review authenticity (genuine vs. fake) can be challenging. This limits the ability to train and validate machine learning models effectively.
- Biases in training data: Existing datasets might be biased toward certain types of fake reviews or user behavior patterns. This can lead to models that are less effective in detecting new or evolving tactics employed by fraudulent actors.

4.2. Model Development and Refinement

Addressing evolving tactics: Fraudulent actors constantly adapt their strategies. Existing models might struggle to keep pace with these changes, requiring continuous development and refinement. Explainability and interpretability of ML models: While machine learning can be powerful, understanding "why" a model classifies a review as fake can be challenging. This lack of interpretability can hinder efforts to improve detection methods and identify new red flags.

4.3. Integration and Practical Applications

- Combining behavioral analysis with other techniques: More research is needed on how to effectively combine behavioral analysis with other detection methods like content analysis and social network analysis for a more holistic approach.
- Tailoring detection methods for specific platforms and products: Existing research might not adequately address the nuances of different online platforms and product categories. Developing detection methods that can adapt to these variations is crucial for real-world applications.
- Scalability and cost-effectiveness: Implementing large-scale detection systems can be resource-intensive. Research is needed to develop cost-effective and scalable solutions for platforms to handle the massive volume of online reviews.

4.4. Ethical Considerations

User privacy concerns: Balancing the need for effective detection with user privacy is a critical consideration. Research is needed to develop methods that can identify fake reviews without compromising user data. Potential for bias: Detection methods can inadvertently perpetuate biases against certain user groups or product categories. Research should address these concerns and ensure fair and unbiased detection practices.

4.5. Specific Research Questions: Unveiling Fake Reviews

Building upon the identified gaps in existing research, here are some specific research questions you can explore:

Effectiveness of Behavioral Analysis Techniques

- **RQ 1.1:** How does the effectiveness of analyzing reviewer activity patterns (frequency, rating consistency, time between reviews) compare to analyzing purchase history in identifying fake reviews on e-commerce platforms?
- **RQ 1.2:** Can combining social network analysis techniques with traditional behavioral analysis improve the accuracy of detecting fake reviews on social media platforms promoting services?

Performance of Machine Learning Algorithms:

- **RQ 2.1:** How does the performance of a supervised learning model trained on reviewer behavior and sentiment analysis data compare to an anomaly detection model for identifying fake reviews on a review platform with limited access to ground truth data?
- **RQ 2.2:** Can incorporating reviewer location data into machine learning models improve the detection of fake reviews for geographically restricted products or services?

Moderating Effects of Platform and Product:

- **RQ 3.1:** Does the effectiveness of machine learning models trained on e-commerce review data translate to accurately detecting fake reviews on social media platforms for the same product category (e.g., clothing)?
- **RQ 3.2:** How do platform-specific review moderation practices moderate the prevalence and nature of fake reviews for different product categories (e.g., electronics vs. travel experiences)?

Reviewer Motivation as a Mediating Variable

- **RQ 4.1:** Can analyzing reviewer profiles and past behavior patterns alongside traditional detection methods help identify the underlying motivation (financial gain, brand damage) behind fake reviews, leading to more targeted detection strategies?
- **RQ 4.2:** Does understanding reviewer motivation through surveys mediate the relationship between reviewer behavior and review authenticity, allowing for a more nuanced understanding of fake review patterns?

These research questions delve deeper into specific aspects of behavioral analysis, machine learning, and the moderating influences of platform and product characteristics. Additionally, exploring reviewer motivation as a mediating variable can provide valuable insights into the "why" behind fake reviews. By investigating these questions, you can contribute to the development of more robust and adaptable detection methods that can stay ahead of evolving tactics employed by fraudulent actors.

5. Methodology

5.1. Identifying Relevant Sources

This research review delves into the potential of behavioral analysis and machine learning (ML) as tools to combat fake online opinions. To achieve this, a well-defined methodology has been laid out. The search strategy focuses on scouring renowned academic databases like Scopus, Web of Science, and ScienceDirect for relevant studies. These databases are known for their scholarly articles on information technology, e-commerce, and consumer behavior. To identify the most pertinent research, a combination of keywords will be used. These keywords target studies related to fake reviews, behavioral analysis, and machine learning techniques like anomaly detection and sentiment analysis. The selection criteria ensure the review captures the latest advancements in the field. Only peer-reviewed articles published within the last 5-10 years will be included, guaranteeing research quality and relevance.

Most importantly, the studies must explicitly explore the use of behavioral analysis and/or machine learning to detect fake online opinions. Once the studies are selected, key information will be extracted from each one. This information will include details about the research methodology used (data collection methods and analysis techniques employed). More importantly, the findings on the effectiveness of these techniques in detecting fake reviews will be closely examined. Additionally, any limitations and gaps identified in the existing research will be noted. To analyze the extracted data, a thematic approach will be

used. This will involve identifying trends, recurring themes, and any potential contradictions across the studies. The analysis will then evaluate the effectiveness of different approaches, considering factors like the platform where the reviews are posted, the product category being reviewed, and the motivations behind the fake reviews. A narrative synthesis approach will be used to synthesize the data. This approach involves summarizing and critically evaluating the findings from the selected studies to provide a comprehensive overview of the current research landscape on detecting fake reviews using behavioral analysis and machine learning.

The review will not only explore the effectiveness of these techniques but also discuss the ethical considerations involved in detecting fake reviews. This includes concerns about user privacy and potential biases that might exist within the detection methods themselves. By following this comprehensive methodology, the research aims to conduct a systematic and critical review of the current research on leveraging behavioral analysis and machine learning to combat fake online opinions. This review can contribute valuable insights to the ongoing efforts to maintain trust and credibility in the online review system.

6. Implications and Discussion

This research review on leveraging behavioral analysis and machine learning for detecting fake online reviews has the potential to yield significant outcomes that can empower both researchers and developers in the fight against this pervasive issue. We anticipate gleaning valuable insights across several key areas. Firstly, the review is poised to identify the most effective behavioural analysis techniques for detecting fraudulent reviewers. This might involve pinpointing unusual patterns in review frequency, rating consistency, time intervals between reviews, and purchase history inconsistencies.

By analyzing the effectiveness of various techniques across different studies, we can illuminate a data-driven path for identifying suspicious user behavior. Secondly, the review aims to provide a comprehensive evaluation of the effectiveness of various machine learning algorithms in automated fake review detection. This will shed light on the strengths and weaknesses of different approaches, such as supervised learning, anomaly detection, and sentiment analysis algorithms. Additionally, we can expect to gain insights into the impact of training data quality on the accuracy of these models, highlighting the importance of robust and up-to-date datasets for optimal performance. Furthermore, the review will emphasize the need for continuous adaptation of these algorithms, ensuring they remain effective against evolving tactics employed by fraudulent actors.

Moving beyond the core methods, the review will delve into the moderating effects of platform characteristics and product categories. This analysis promises valuable knowledge on how platform functionalities, user demographics, and moderation practices influence the prevalence and nature of fake reviews within different online environments. Additionally, we can expect to learn how the effectiveness of detection methods might vary depending on the type of product being reviewed. For instance, the strategies used to identify fake reviews for electronics might differ from those employed for travel experiences. Finally, the review ventures beyond simply identifying fake reviews to understand the "why" behind them.

By exploring reviewer motivation as a mediating variable, we hope to gain insights into the underlying reasons for fake reviews, such as financial gain, brand damage, or social influence. This knowledge can then be harnessed to develop more targeted detection methods that address specific types of fraudulent activity. The anticipated outcomes of this research review hold significant value for researchers and developers working to combat fake online reviews. By illuminating promising techniques, evaluating machine learning algorithms, understanding moderating effects, and exploring reviewer motivation, we can pave the way for a future where online reviews remain a reliable source of information for consumers. This comprehensive approach offers a promising path forward in the fight against online deception.

7. Ethical Implications

This research review investigates the potential of behavioural analysis and machine learning (ML) for detecting fraudulent online reviews. While these methods offer a compelling avenue for creating a more trustworthy online review landscape, it is crucial to acknowledge and address the ethical implications associated with their implementation.

7.1. Privacy Concerns in the Digital Age

A primary area of concern lies in the realm of user privacy. Behavioral analysis and ML models rely heavily on user data to function effectively. This review will delve into the data collection practices employed to ensure user privacy is protected. Key questions to be addressed include: How are user data anonymization techniques implemented? What are the established limitations on how this data can be utilized? Additionally, the review will explore the need for transparency in data collection practices and user control over their data usage. Consumers have a fundamental right to comprehend how their online behavior is leveraged in detecting fake reviews, and they should have the ability to control how their data is collected and employed.

7.2. Mitigating Algorithmic Bias for Fair and Accurate Detection

Another significant ethical consideration concerns algorithmic bias. ML algorithms are inherently limited by the quality of data they are trained on. If the training data harbors inherent biases, the resulting model can perpetuate those biases in its detection methods. This review will explore strategies to ensure fairness and accuracy in detection methods. This includes investigating techniques to mitigate bias against specific user groups or product categories. Furthermore, the review will emphasize the importance of developing interpretable models. Understanding the rationale behind an ML model's decision to classify a review as fake is critical for ensuring responsible and unbiased detection. Without interpretability, there's a risk of unfairly penalizing legitimate reviews.

7.3. Striking a Balance: Protecting Free Speech and Preventing Malicious Use

The fight against fake reviews should not come at the expense of silencing genuine opinions. Overly aggressive detection methods could inadvertently flag legitimate user reviews, hindering free expression. The review will explore the need for safeguards to prevent such scenarios. For instance, establishing clear criteria for flagging reviews and allowing users to appeal flagged reviews can help ensure legitimate voices are heard. Another potential pitfall is the misuse of detection methods for market manipulation. Competitors could potentially exploit these methods to suppress negative reviews about their rivals, gaining an unfair market advantage. The review will discuss the importance of ethical implementation to prevent such misuse. Clear guidelines and regulations can help ensure that these powerful tools are used responsibly and ethically.

8. Conclusion

This systematic review has shed light on the potential of behavioral analysis and machine learning (ML) as powerful tools for combating fake online reviews. The identified techniques for analyzing user activity patterns and the promising applications of ML algorithms offer a path forward for creating a more trustworthy online review landscape. However, it is crucial to acknowledge the importance of ethical considerations like user privacy and algorithmic bias in the development and implementation of these methods. By addressing these concerns and pursuing the proposed research directions, such as multimodal detection, scalable solutions, privacy-preserving techniques, user perception studies, and continuous evaluation, we can work towards a future where online reviews remain a reliable source of information for consumers. This will ultimately foster a more transparent and ethical online marketplace that benefits everyone.

8.1. Key Findings: Illuminating Deceptive Behavior

The review identified promising techniques for detecting fraudulent reviewers through behavioural analysis. These techniques focus on scrutinizing user activity patterns, including review frequency, rating consistency, time intervals between reviews, and purchase history inconsistencies. By examining these patterns, researchers can gain valuable insights into potentially deceptive behavior. The review also found that machine learning algorithms hold significant promise for automated detection of fake reviews. However, their effectiveness hinges on factors like the quality of training data and the need for continuous adaptation. As fraudulent tactics evolve, ML models must be continuously updated to maintain accuracy.

8.2. Tailoring Detection Methods and Understanding Motivations

The review highlighted the importance of considering platform characteristics and product categories when developing detection methods. The prevalence and nature of fake reviews can vary significantly across different online environments and product types. Tailoring detection strategies to specific platforms (e.g., e-commerce vs. social media) and products (e.g., electronics vs. travel experiences) can enhance their effectiveness. Furthermore, exploring reviewer motivation as a mediating variable emerged as a promising avenue for understanding the "why" behind fake reviews. By identifying motivations like financial gain, brand damage, or social influence, researchers can inform the development of more targeted detection methods that address specific types of fraudulent activity.

8.3. Building a More Trustworthy Online Landscape

This review offers a valuable framework for researchers and developers working to combat fake reviews. The identified gaps in the existing research provide a roadmap for future investigations aimed at significantly improving the effectiveness of detection methods. Additionally, by emphasizing the importance of addressing ethical considerations like user privacy and algorithmic bias, the review promotes the responsible development and implementation of these techniques.

8.4. Future Research Directions: A Roadmap for Continuous Improvement

Several promising avenues exist for future research in this field. Here are some key areas for exploration:

- **Multimodal Detection:** Investigate the effectiveness of combining behavioral analysis with other detection methods, such as content analysis and social network analysis, to create a more comprehensive approach.
- **Scalable Solutions:** Develop cost-effective and scalable ML models suitable for real-world implementation on large online platforms, enabling widespread adoption of these detection techniques.
- **Privacy-Preserving Techniques:** Explore user privacy-preserving techniques for data collection and analysis in the context of fake review detection. Balancing the need for effective detection with the protection of user privacy is crucial.
- **User Perception:** Conduct user studies to understand how consumers perceive and respond to the issue of fake reviews and the detection methods employed by online platforms. This can inform the development of user-centric solutions.
- **Continuous Evaluation:** Continuously evaluate and refine detection methods to stay ahead of evolving tactics employed by fraudulent actors. The fight against fake reviews requires ongoing vigilance and adaptation.

By pursuing these research directions, we can contribute to a future where online reviews remain a trusted source of information for consumers. This will ultimately foster a more transparent and reliable online marketplace for all stakeholders.

Acknowledgment: We are deeply grateful to CMS Business School, Jain University for their expertise and dedication, which greatly enriches this work.

Data Availability Statement: The data for this study can be made available upon request to the corresponding author.

Funding Statement: This manuscript and research paper were prepared without any financial support or funding

Conflicts of Interest Statement: The authors have no conflicts of interest to declare. This work represents a new contribution by the authors, and all citations and references are appropriately included based on the information utilized.

Ethics and Consent Statement: This research adheres to ethical guidelines, obtaining informed consent from all participants. Confidentiality measures were implemented to safeguard participant privacy.

References

1. M. D. Molina, S. S. Sundar, T. Le, and D. Lee, "'Fake news' is not simply false information: A concept explication and taxonomy of online content," *Am. Behav. Sci.*, vol. 65, no. 2, pp. 180–212, 2021.
2. S. Zannettou, M. Sirivianos, J. Blackburn, and N. Kourtellis, "The Web of false information: Rumors, fake news, hoaxes, clickbait, and various other shenanigans," *ACM J. Data Inf. Qual.*, vol. 11, no. 3, pp. 1–37, 2019.
3. L. Luceri, V. Pantè, K. Burghardt, and E. Ferrara, "Unmasking the web of deceit: Uncovering coordinated activity to expose information operations on Twitter," *arXiv [cs.SI]*, 2023, Press.
4. M. H. Al-Adhaileh and F. W. Alsaade, "Detecting and Analysing Fake Opinions Using Artificial Intelligence Algorithms," *Intelligent Automation & Soft Computing*, vol. 32, no. 1, pp. 643–655, 2022.
5. S. N. Alsubari et al., "Data analytics for the identification of fake reviews using supervised learning," *Computers, Materials & Continua*, vol. 70, no. 2, pp. 3189–3204, 2022.
6. Y. W. Oh and C. H. Park, "Machine cleaning of online opinion spam: Developing a machine-learning algorithm for detecting deceptive comments," *Am. Behav. Sci.*, vol. 65, no. 2, pp. 389–403, 2021.
7. A. M. Elmogy, U. Tariq, A. Mohammed, and A. Ibrahim, "Fake reviews detection using supervised machine learning," *Int. J. Adv. Comput. Sci. Appl.*, vol. 12, no. 1, pp. 601–606, 2021.
8. P. Bhatt et al., "Machine learning for cognitive behavioral analysis: datasets, methods, paradigms, and research directions," *Brain Inform.*, vol. 10, no. 1, p. 18, 2023.
9. S. Alhazbi. "Behavior-based machine learning approaches to identify state-sponsored trolls on Twitter". *IEEE access*, vol. 8, no.10, pp.195132-195141, 2020.
10. S. Iacobucci, R. De Cicco, F. Michetti, R. Palumbo, and S. Pagliaro, "Deepfakes unmasked: The effects of information priming and bullshit receptivity on deepfake recognition and sharing intention," *Cyberpsychol. Behav. Soc. Netw.*, vol. 24, no. 3, pp. 194–202, 2021.
11. N. A. Patel and R. Patel, "A survey on fake review detection using machine learning techniques," in *2018 4th International Conference on Computing Communication and Automation (ICCCA)*, Greater Noida, India, 2018.

12. N. Sultana and S. Palaniappan, "Deceptive opinion detection using machine learning techniques," *Int. J. Inf. Eng. Electron. Bus.*, vol. 12, no. 1, pp. 1–7, 2020.
13. A. A. A. Ahmed, A. Aljabouh, P. K. Donepudi, and M. S. Choi, "Detecting fake news using machine learning: A systematic literature review," *vol. 58, no. 1, pp. 1932–1939*, 2021.
14. B. Manaskasemsak, J. Tantisuwankul, and A. Rungsawang, "Fake review and reviewer detection through behavioral graph partitioning integrating deep neural network," *Neural Comput. Appl.*, vol. 35, no. 2, pp. 1169–1182, 2023.
15. M. M. B. Alsaad and H. Joshi, "Trends of Consumer Behavior Analysis in E-Commerce for Fake/Spam Opinion Detection," *Mathematical Statistician and Engineering Applications*, vol. 71, no. 3s, pp. 612–623, 2022.
16. S. Mohseni et al., "Machine learning explanations to prevent overtrust in fake news detection," *Proceedings of the International AAAI Conference on Web and social media*, vol. 15, no. 5, pp. 421–431, 2021.
17. X. Wang, X. Zhang, C. Jiang, and H. Liu, "Identification of fake reviews using semantic and behavioral features," in *2018 4th International Conference on Information Management (ICIM)*, Oxfordshire, England, 2018.
18. S. Bhargava and S. Choudhary, "Behavioral analysis of depressed sentimentality over twitter: Based on supervised machine learning approach," *SSRN Electron. J.*, vol. 3, no. 7, pp. 1011–1017, 2018.
19. A. Qazi, N. Hasan, R. Mao, M. E. M. Abo, S. K. Dey, and G. Hardaker, "Machine learning-based opinion spam detection: A systematic literature review," *IEEE Access*, vol. 1, no. 99, pp. 1–2, 2024.
20. C. R. Balaji, R. Annavarapu, and A. Bablani, "Machine learning algorithms for social media analysis: A survey," *Comput. Sci. Rev.*, vol. 40, no. 10, p. 100395, 2021.
21. D. Mahat and S. Mathema, "Gender Perspective on Compensation of Health Institution in Ramechhap District of Nepal," *Nepal Journal of Multidisciplinary Research*, vol. 1, no. 1, pp. 30–40, 2018.
22. D. Mahat, "Unveiling Faculty Performance through Student Eyes: A Comparative Study of Perspectives," *Formosa Journal of Multidisciplinary Research*, vol. 2, no. 9, pp. 1575–1596, 2023.
23. C. Das, R. Sengupta, A. Patil, and R. Yadav, "Effect of automation in banks as regulatory intervention in the Indian bourses: New evidence from India," *European Economic Letters*, vol. 14, no. 2, pp. 1167–1174, 2024.
24. F. Wahab, M. J. Khan, and M. Y. Khan, "The impact of climate induced agricultural loan recovery on financial stability; evidence from the emerging economy Pakistan," *Journal of Contemporary Macroeconomic Issues*, vol. 3, no. 2, pp. 1–12, 2022.
25. F. Wahab, M. J. Khan, M. Y. Khan, and R. Mushtaq, "The impact of climate change on agricultural productivity and agricultural loan recovery; evidence from a developing economy," *Environ. Dev. Sustain.*, vol. 26, no. 8, pp. 24777–24790, 2024.
26. J. Cao, G. Bhuvaneswari, T. Arumugam, and A. B. R., "The digital edge: Examining the relationship between digital competency and language learning outcomes," *Frontiers in Psychology*, vol. 14, no. 6, p. 11, 2023.
27. J. Rehman, M. Kashif, and T. Arumugam, "From the land of Gama: Event attachment scale (EAS) development exploring fans' attachment and their intentions to spectate at traditional gaming events," *International Journal of Event and Festival Management*, vol. 14, no. 3, pp. 363–379, 2023.
28. R. Jadhavrao, R. Sengupta, A. Patil, and R. S. Yadav, "Stock market trend prediction using artificial intelligence algorithm (RNN-LSTM) - Comparison with current techniques and research on its effectiveness in forecasting in Indian market and US market," *Empirical Economics Letters*, vol. 22, no. S2, pp. 124–133, 2023.
29. K. M. Ashifa and P. Ramya, "Health afflictions and quality of work life among women working in fireworks industry," *International Journal of Engineering and Advanced Technology*, vol. 8, no. 6, pp. 1723–1725, 2019.
30. K. M. Ashifa, "Developmental initiatives for persons with disabilities: Appraisal on village-based rehabilitation of Amar Seva Sangam," *Indian Journal of Public Health Research and Development*, vol. 10, no. 12, pp. 1257–1261, 2019.
31. K. U. Kiran and T. Arumugam, "Role of programmatic advertising on effective digital promotion strategy: A conceptual framework," *Journal of Physics: Conference Series*, vol. 1716, no. 12, p. 012032, 2020.
32. A. Kolte, A. Patil, and H. Mal, "Monetary policy in tough times-case study approach," *International Journal of Intelligent Enterprise*, vol. 8, no. 1, pp. 105–122, 2021.
33. A. Kolte, M. Rossi, O. Torriero, A. Patil, and A. Pawar, "Balance of payment crisis: Lessons from Indian payment crisis for developing economies," *International Journal of Behavioural Accounting and Finance*, vol. 6, no. 3, pp. 262–279, 2021.
34. M. A. Sanjeev, A. Thangaraja, and P. K. S. Kumar, "Multidimensional scale of perceived social support: Validity and reliability in the Indian context," *International Journal of Management Practice*, vol. 14, no. 4, p. 472, 2021.
35. M. A. Sanjeev, S. Khademizadeh, T. Arumugam, and D. K. Tripathi, "Generation Z and intention to use the digital library: Does personality matter?," *The Electronic Library*, vol. 40, no. 1/2, pp. 18–37, 2021.
36. M. Kanwal, N. A. Khan, and A. A. Khan, "A machine learning approach to user profiling for data annotation of online behavior," *Comput. Mater. Contin.*, vol. 78, no. 2, pp. 2419–2440, 2024.
37. M. Lee, Y. H. Song, L. Li, K. Y. Lee, and S.-B. Yang, "Detecting fake reviews with supervised machine learning algorithms," *Serv. Ind. J.*, vol. 42, no. 13–14, pp. 1101–1121, 2022.

38. R. Kenny, B. Fischhoff, A. Davis, K. M. Carley, and C. Canfield, "Duped by bots: Why some are better than others at detecting fake social media personas," *Hum. Factors*, vol. 66, no. 1, pp. 88–102, 2024.
39. R. Rafique, R. Gantassi, R. Amin, J. Frnda, A. Mustapha, and A. H. Alshehri, "Deep fake detection and classification using error-level analysis and deep learning," *Sci. Rep.*, vol. 13, no. 1, p. 7422, 2023.
40. R. Shrivastava, "Effect of substantial competent team on training after ERP implementation, continual system improvement, department related to decision support and business performance – A study," *J. Contemp. Issues Bus. Gov.*, vol. 27, no. 2, pp. 690–698, 2021.
41. R. Shrivastava, "Role of artificial intelligence in future of education," *Int. J. Prof. Bus. Rev.*, vol. 8, no. 1, pp. 1–15, 2023.
42. R. Singh et al., "Antisocial Behavior Identification from Twitter Feeds Using Traditional Machine Learning Algorithms and Deep Learning," *EAI Endorsed Transactions on Scalable Information Systems*, vol. 10, no. 4, p.17, 2023.
43. S. Gupta, N. Pande, T. Arumugam, and M. A. Sanjeev, "Reputational impact of COVID-19 pandemic management on World Health Organization among Indian public health professionals," *Journal of Public Affairs*, Oct. 2022, Press.
44. S. Hameed, S. Madhavan, and T. Arumugam, "Is consumer behaviour varying towards low and high involvement products even sports celebrity endorsed?," *International Journal of Scientific & Technology Research*, vol. 9, no. 3, p.12, 2020.
45. S. Verma, N. Garg, and T. Arumugam, "Being ethically resilient during COVID-19: A cross-sectional study of Indian supply chain companies," *The International Journal of Logistics Management*, Aug. 2022, Press.
46. R. Sengupta and A. Patil, "An analysis of the impact of the merger on the performance of Indian Public Sector banks merged in 2020 using the CAMEL Model Analysis: An emerging nation's perspective," *Kanpur Philosophers*, vol. 9, no. 2, pp. 1493–1512, 2022.
47. R. Sengupta and A. A. Patil, "Green Finance Products and Investments in the Changing Business World," in *Handbook of Research on Sustainable Consumption and Production for Greener Economies*, IGI Global, USA, pp. 344–357, 2023.
48. T. Arumugam, B. L. Lavanya, V. Karthik, K. Velusamy, U. K. Kommuri, and D. Panneerselvam, "Portraying women in advertisements: An analogy between past and present," *The American Journal of Economics and Sociology*, vol. 81, no. 1, pp. 207–223, 2022.
49. V. M. Aragani, "Leveraging AI and Machine Learning to Innovate Payment Solutions: Insights into SWIFT-MX Services," *International Journal of Innovations in Scientific Engineering*, vol. 17, no. 1, pp. 56-69, 2023.
50. V. M. Aragani, "Securing the Future of Banking: Addressing Cybersecurity Threats, Consumer Protection, and Emerging Technologies," *International Journal of Innovations in Applied Sciences and Engineering*, vol. 8, no.1, pp. 178-196, Nov. 11, 2022.
51. V. M. Aragani, "Unveiling the Magic of AI and Data Analytics: Revolutionizing Risk Assessment and Underwriting in the Insurance Industry," *International Journal of Advances in Engineering Research*, vol. 24, no. 6, pp. 1-13, 2022.
52. R. Yadav, A. Patil, and R. Sengupta, "An analysis of Satyam case using bankruptcy and fraud detection models," *SocioEconomic Challenges*, vol. 7, no. 4, pp. 24–35, 2023.