

Sarcasm Detection on Twitter Using Deep Learning and Natural Language Processing Techniques

Arti Verma

M.Tech Scholar

Department of Computer science & Engineering
Radharaman Engineering College, Bhopal, India.
Email: avartiverma3001@gmail.com

Rakesh Shivhare

Assistant Professor

Department of Computer science & Engineering
Radharaman Engineering College, Bhopal, India.
Email: rirtcollege@gmail.com

Abstract: Sarcasm detection has emerged as one of the most challenging tasks in Natural Language Processing (NLP) due to the complex interplay between literal expressions and intended meanings. Sarcasm is frequently observed across e-commerce platforms, social media, and online reviews, where its misinterpretation can negatively impact sentiment analysis and opinion mining applications. Traditional models primarily relied on manual feature engineering and focused only on the textual content of expressions, often ignoring the contextual and semantic cues necessary for accurate detection. This limitation results in poor generalization, particularly in handling implicit sarcasm. With the advent of deep learning (DL) and transformer-based architectures, significant progress has been made in modeling contextual information and capturing complex linguistic patterns. In this paper, we provide a comprehensive review of existing sarcasm detection approaches on online social media platforms such as Twitter, Facebook, blogs, and product reviews. The reviewed studies cover various machine learning, deep learning, and hybrid models, with a comparative analysis based on accuracy, precision, recall, and F1-score. Furthermore, we highlight existing challenges, including context dependency, multimodal sarcasm cues, and low-resource language limitations, while identifying potential directions for future research to develop more robust and explainable sarcasm detection systems.

Keywords: Sarcasm Detection, Sentiment Analysis, Natural Language Processing, Deep Learning, Transformer Models, Twitter, Social Media Analytics.

1. Introduction

Twitter has emerged as one of the largest web-based platforms for people to express opinions, share thoughts, and report real-time events. Over the past decade, the volume of Twitter content has grown exponentially, making it a significant example of “big data.” According to official statistics, Twitter hosts more than 288 million active users and processes over 500 million tweets daily. Organizations and researchers widely use these vast data streams to study public opinions on topics such as political events [1], product reviews [2], and movies [3]. Despite its value, analyzing Twitter data presents significant challenges. Tweets are short (limited to 140/280 characters) and often written in informal language, making automatic sentiment analysis complex. The presence of sarcasm further complicates the process, as sarcasm occurs when the literal meaning of a statement differs from its intended meaning. Understanding and accurately detecting sarcasm in tweets is, therefore, critical for applications such as sentiment analysis, opinion mining, and business intelligence. Natural Language Processing (NLP), a rapidly growing field in Artificial Intelligence (AI), plays a vital role in this task. Sentiment analysis, also known as opinion mining, identifies subjective opinions and emotions in user-generated content [4]. Social networks like Twitter and Facebook are the primary sources for such analysis, providing massive, constantly updated datasets that hold valuable insights into public sentiment.

2. Background and Theoretical Foundations

Text-based sarcasm detection involves understanding the gap between the explicit meaning of a statement and the speaker’s intended meaning. Sarcasm is commonly defined as expressing the opposite of what is meant, creating ambiguous and non-straightforward data. For example, the sentence “I love going to the dentist!” typically expresses dissatisfaction rather than pleasure. Detecting sarcasm is inherently challenging because textual data lacks tone, facial expressions, and non-verbal cues, which humans often rely on to interpret sarcasm [5]. Furthermore, sarcasm can vary significantly across cultures, languages, and personal communication styles, adding another layer of complexity. Traditional sarcasm detection approaches employed handcrafted features such as n-grams, lexical patterns, and sentiment flips, but these methods have shown limited accuracy due to their inability to generalize across diverse

contexts. The basic methodology of sarcasm detection is presented in Fig 1. Recent advances in deep learning have addressed some of these limitations. Architectures such as Convolutional Neural Networks (CNNs), Recurrent Neural Networks (RNNs), Long Short-Term Memory (LSTM) models, and transformer-based models like BERT, RoBERTa, and GPT have demonstrated substantial improvements in sarcasm detection accuracy [6]-[12].

- However, several research gaps remain unresolved:
- Contextual sarcasm often depends on user-specific history or previous conversations.
- Most datasets are domain-specific, limiting the transferability of models.
- Multimodal sarcasm detection (e.g., text + emojis + images) is still underexplored.

Thus, this review aims to analyze state-of-the-art sarcasm detection techniques, identify research gaps, and highlight future opportunities for designing robust, context-aware models.

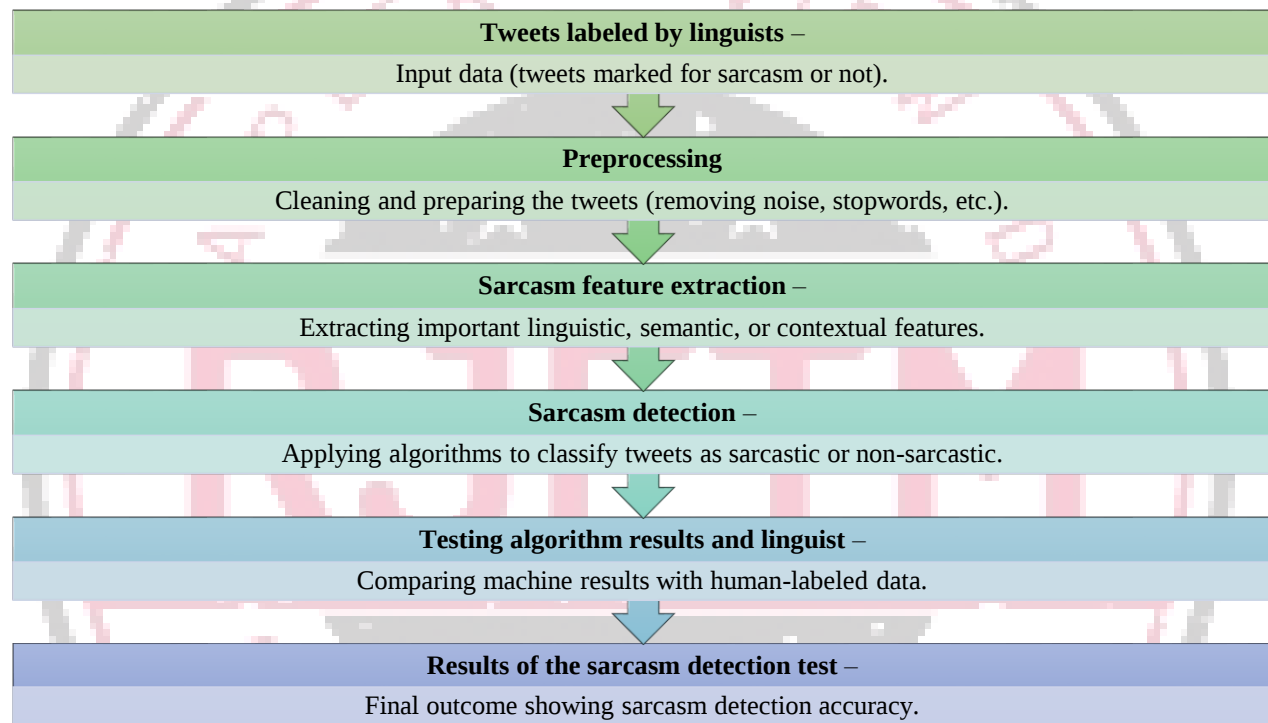


Figure 1. The basic Methodology of Sarcasm Detection

3. Literature Review

Recent studies have explored a range of techniques for sarcasm detection, from classical machine learning approaches to advanced deep learning and transformer-based frameworks [13]-[37]. Table I summarizes prominent research efforts, methodologies, and reported performance.

3.1 Traditional Machine Learning Approaches

Early research on sarcasm detection relied heavily on machine learning algorithms such as Naïve Bayes, Support Vector Machines (SVM), Random Forests, and Logistic Regression. These models typically depended on manually engineered features like n-grams, sentiment polarity flips, and lexical cues. While they showed moderate success, their generalization ability across datasets remained limited. Senthilkumar et al. [13] proposed an automated sarcasm-context detection method using the Weighted Random Forest approach. Trained on the MUSTARD dataset and tested on SARC, the model achieved an F1-score of 60.1%. The study highlights the potential for conversational AI systems to recognize and respond to sarcasm effectively. Pradhan et al. [22] developed a sarcasm detection model using 13 handcrafted linguistic features capturing meaning, lexical diversity, and readability. Tested across multiple machine

learning algorithms, the ensemble model outperformed others, achieving 93% accuracy and improving the F1-score by 5% on the News Headline Dataset. Thaokar et al. [26] investigated sarcasm detection using word-level features such as N-grams, negation words, and PoS tags. Compared machine learning, hybrid, and deep learning models across three benchmark datasets. Results showed up to 92% accuracy with Random Forest classifiers and 87% accuracy using hybrid deep learning models. Sharma et al. [33] explored sarcasm detection in Twitter data using machine learning classifiers, including logistic regression, naive Bayes, SVM, decision trees, and ensemble models. The ensemble approach achieved a peak accuracy of 75%, demonstrating effectiveness for real-time sarcasm detection in social media applications. Vinoth et al. [37] introduced an Intelligent ML-based Sarcasm Detection and Classification (IMLB-SDC) framework integrating TF-IDF-based feature engineering, chi-square and information gain feature selection, and SVM classification optimized via Particle Swarm Optimization (PSO). The model achieved outstanding performance with precision = 94.7%, recall = 95.2%, and F1-score = 94.9%, outperforming recent state-of-the-art techniques.

3.2 Deep Learning-Based Approaches

Deep learning techniques such as Convolutional Neural Networks (CNNs), Recurrent Neural Networks (RNNs), and Long Short-Term Memory (LSTM) networks have demonstrated superior performance compared to traditional methods. These models automatically learn semantic and contextual representations from raw data, reducing the need for extensive feature engineering. Saleem et al. [27] proposed a deep learning-based sarcasm detection model using BoW, TF-IDF, and word embeddings combined with an LSTM network. Tested on a Twitter dataset of sarcasm, irony, and regular tweets, the model achieved an accuracy of 99.01% using precision, recall, F1, and accuracy as evaluation metrics. Kumar et al. [28] investigated context-aware sarcasm detection using three models: TF-IDF with ensemble voting classifiers (Naive Bayes, Gradient Boosting, Random Forest). Semantic + pragmatic features with top-200 TF-IDF features and five baseline classifiers. Bi-directional LSTM with GloVe embeddings. The Bi-LSTM model achieved the highest accuracy of 86.32% on Twitter and 82.91% on Reddit. Pal and Prasad [34] developed an LSTM-based sarcasm detection model using GloVe embeddings achieving 91.94% accuracy. Sarcastic sentences were also integrated into sentiment analysis corpora and evaluated using various classifiers, where Linear SVM achieved the highest precision (92.5%) on unigram, bigram, and trigram models. Sinha and Yadav [35] proposed a hybrid CNN-LSTM model for sarcasm detection in news headlines, combining the strengths of both architectures. The model significantly outperformed traditional models, achieving 97% accuracy. Goel et al. [36] Designed an ensemble deep learning model combining LSTM, GRU, and CNN with multiple pretrained embeddings (fastText, Word2Vec, GloVe). Achieved 96% accuracy on the News Headlines dataset and 73% on Reddit. A Weighted Average Ensemble performed best, attaining 99% accuracy on News Headlines and 82% on Reddit, improving precision and stability. Srinu et al. [17] introduced a novel SD-GOALD technique, integrating a deep belief network (DBN) with the Grasshopper Optimization Algorithm (GOA) for hyperparameter tuning. Using GloVe embeddings, the model demonstrated superior performance over recent approaches, improving sarcasm detection accuracy significantly on benchmark datasets. Pandey et al. [23] proposed a hybrid CNN-H model that combines character and word embeddings to improve performance on multilingual datasets. Evaluated on five benchmark datasets (English, Hindi, and Hinglish), the model achieved an F1-score of 93%, outperforming conventional CNNs and other state-of-the-art models, particularly in handling noisy social media data.

3.3 Transformer-Based and Pretrained Models

Recent advancements leverage transformer architectures like BERT, RoBERTa, GPT, and XLNet for sarcasm detection. These models exploit self-attention mechanisms to capture long-range dependencies and contextual nuances, achieving state-of-the-art performance on several benchmarks. Prakash [14] proposed a hybrid deep learning model combining LSTM, BiLSTM, CNN, and BERT with transformer augmentation for detecting sarcasm in social media text. Evaluated on benchmark datasets, the model achieved 96% accuracy, 97% precision, 93% recall, and 95% F1-score. The approach demonstrates significant improvement over baseline models and benefits applications like social media monitoring, opinion mining, and customer feedback analysis. Hassan et al. [19] developed a novel Urdu sarcasm detection model using a dataset of 12,910 manually annotated tweets. Proposed a multilingual BERT (mBERT) + BiLSTM + Multi-Head Attention (MHA) hybrid architecture, outperforming traditional deep learning models (CNN, LSTM, GRU, BiLSTM, CNN-LSTM) using fastText embeddings. Achieved an accuracy of 79.51%

and an F1-score of 80.04%, making it one of the first effective approaches for Urdu sarcasm detection. Sukhavasi et al. [21] Designed an attention-based transformer model integrating text and emoji analysis. Key contributions include: Preprocessing techniques: stop word removal, lemmatization, stemming, tokenization. Feature extraction: ATF-IDF for text and Emoji-to-Vector Modeling (E-VM). Modeling: GT-BiCNet for text, ALABerT (enhanced transformer-based model) for classification. Optimization: Enhanced Pelican Optimization Algorithm (EpoA). Achieved state-of-the-art results with 99.1% accuracy on the English Twitter dataset and 98.1% accuracy on the Hindi Twitter dataset, outperforming existing methods. Javed et al. [25] conducted a comparative analysis of BERT and LSTM models for sarcasm detection on social media, focusing on improving sentiment classification accuracy, particularly for Twitter data. Findings show that BERT outperforms LSTM due to better contextual understanding, highlighting its suitability for sarcasm detection in short, informal texts.

3.4 Hybrid or Ensemble Models

Ambreen et al. [15] proposed a CCNN-ELLSTM hybrid model leveraging both linguistic and visual cues (text + emoji embeddings) for sarcasm detection in English tweets. Achieved 97.57% accuracy and an F1-score of 0.9718, demonstrating that combining emoji-based contextual information significantly enhances performance. Sabera et al. [16] conducted a comparative analysis of feature extraction techniques (Word2Vec, GloVe, BERT, RoBERTa) combined with classifiers (SVM, XGBoost, Random Forest) using the KDD framework. The RoBERTa + SVM combination performed best, achieving 88% accuracy, highlighting the advantage of transformer-based contextual embeddings. However, the model struggles with implicit sarcasm due to limited contextual cues. Gedela et al. [20] introduced a voting-based ensemble model integrating BERT embeddings with four deep learning architectures (CNN, BiLSTM, CNN-LSTM-parallel, CNN-LSTM-sequential). The outputs were classified using multiple ML classifiers and finalized through majority voting. Achieved 94.89% accuracy on the news headlines dataset and an F1-score of 80.49% on the Reddit corpus, improving state-of-the-art results. Deora et al. [24] proposed the SNTDD model combining ML and DL techniques for detecting sarcasm in news headlines and tweets. Achieved 97.4% accuracy and an F1-score of 94.4%, showing strong performance. The study also highlights the psychological impact of sarcasm on mental health and its link to depression detection. Kumar et al. [29] proposed a multi-task deep learning framework to jointly model sentiment classification and sarcasm detection, leveraging the correlation between the two tasks. Achieved an F1-score of 94%, outperforming standalone sarcasm classifiers by 3%, indicating that shared representations improve sentiment understanding. Murthy et al. [30] introduced Tensor-DNN-50, a real-time sarcasm detection framework capable of multi-class classification. Achieved 96.4% accuracy for binary classification, 94.32% for four-class, and 99.1% for seven-class sarcasm detection, outperforming existing techniques like SVM, att-RNN, and EANN. Rajani et al. [31] proposed a deep neural network with multi-task learning to improve both sentiment analysis and sarcasm detection simultaneously. Demonstrated a 3% improvement over existing models with an F1-score of 97%, reinforcing the positive interdependence between these tasks. Kumar et al. [32] developed first-of-its-kind Telugu and Tamil sarcasm datasets using Twitter data and experimented with various models, including ML classifiers and deep neural networks (CNN, LSTM). Achieved 95.68% accuracy for Telugu and 95.37% for Tamil, addressing the low-resource language gap in sarcasm detection.

Table II. Prominent Research Efforts

Ref. No.	Author (et al.)	Methodology Used	DL Technique Used	Result
[13]	Senthilkumar et al.	Weighted Random Forest on MUsTARD & SARC datasets	Attention-based LSTM + RF	F1 = 60.1%
[14]	Prakash et al.	Transformer augmentation + hybrid optimized DL (LSTM, BiLSTM, CNN, BERT)	Hybrid Transformer + BERT	Acc = 96%, F1 = 95%
[15]	Ambreen et al.	CCNN + ELLSTM with emoji embeddings	CCNN + ELLSTM	Acc = 97.57%, F1 = 0.9718
[16]	Sabera et al.	Feature extraction (Word2Vec, GloVe, BERT, RoBERTa) + SVM/XGBoost/RF	RoBERTa + SVM	Acc = 88%
[17]	Srinu et al.	Grasshopper Optimization with DBN	DBN + GOA optimization	Better than baselines

[18]	Hassan et al.	Urdu sarcasm dataset + mBERT-BiLSTM-MHA	mBERT + BiLSTM + MHA	Acc = 79.51%, F1 = 80.04%
[19]	Gedela et al.	Voting-based ensemble with CNN, BiLSTM, CNN-LSTM	BERT embeddings + CNN/BiLSTM	Acc = 94.89%, F1 = 80.49%
[20]	Sukhavasi et al.	Transformer with text + emoji features	GT-BiCNet + ALABerT	Acc = 99.1% (Eng), 98.1% (Hindi)
[21]	Pradhan et al.	Linguistic handcrafted features + ML ensemble	Ensemble + DNN	Acc = 93%
[11]	Chakraborty et al.	Feature-rich classifiers (lexical, pragmatic, semantic)	Logistic Regression, SVM	Acc = 85%
[1]	Amir et al.	User embeddings + neural networks	CNN + Embeddings	Acc = 91%, F1 = 0.90
[24]	Mishra et al.	Cognitive features + attention RNN	Attention RNN	F1 = 0.84

4. Datasets for Sarcasm Detection

Several benchmark datasets have been introduced for sarcasm detection research, ranging from Twitter-based collections to Reddit and Amazon reviews. Table II summarizes the most commonly used datasets, their sources, languages, and annotation methodologies.

Table II. Commonly Used Datasets

Dataset	Source/Domain	Approx. Size	Annotation/Notes	Ref
MUSARD	Multimodal Sarcasm in Dialogues	~3–5k clips	Annotated sarcasm in dialogue (often multimodal)	[13],[18]
SARC (Reddit)	Reddit comments	>1M	Self-annotated/heuristic sarcasm markers	[13],[18],[20],[28]
News Headlines	News headline corpora	~10–30k	Headline sarcasm labels (gold standard)	[16],[20],[22],[35],[36]
SemEval 2015 Task 11	Twitter	~10k	Sentiment/context task; used for sarcasm experiments	[28]
Urdu Twitter Sarcasm	Twitter (Urdu)	12,910 (re-annotated)	Manual re-annotation into sarcasm/non-sarcasm	[19]
Tamil/Telugu Sarcasm	Twitter (Tamil, Telugu)	~10–20k each (est.)	First datasets for these languages	[32]
Ghosh Tweet	Twitter	—	Sarcasm tweet dataset (benchmark)	[23]
Riloff Tweet	Twitter	—	Benchmark sarcasm dataset	[23]
Hinglish (Hi-En)	Social media	—	Code-mixed sarcasm dataset	[23]
Generic Twitter (Eng/Hindi)	Twitter	—	Used for emoji + text approaches	[21]

5. Evaluation Metrics

Evaluation of sarcasm detection models typically relies on standard classification metrics such as accuracy, precision, recall, and F1-score, as presented in table III. However, given sarcasm's subjective nature, F1-score and Cohen's Kappa are often preferred to measure model robustness and agreement.

Table III.

Metric	Formula	What it Captures	Notes / Limitations
Accuracy	$(TP + TN)/(TP + TN + FP + FN)$	Overall correctness	Can be misleading on imbalanced datasets

Precision	$TP/(TP + FP)$	How many predicted sarcastic are truly sarcastic	High precision → few false alarms
Recall	$TP/(TP + FN)$	How many sarcastic items are correctly found	High recall → few misses
F1-Score	$2 \cdot (P \cdot R)/(P + R)$	Balance between precision and recall	Preferred when classes are imbalanced
Macro-F1	Avg F1 over classes	Treats classes equally	Useful with class imbalance
Cohen's Kappa	$(po - pe)/(1 - pe)$	Agreement beyond chance	Sometimes reported for labeling consistency

6. Challenges and Future Research Directions

Despite significant progress in sarcasm detection, several challenges remain that hinder the development of highly accurate and generalizable models.

- Sarcasm often relies heavily on implicit context, cultural norms, and user-specific background knowledge. Without additional contextual signals—such as user history, conversational threads, or multimedia cues—models may misinterpret sarcastic content.
- Platforms like Twitter include non-standard spellings, abbreviations, emojis, and hashtags, which make sarcasm detection harder. Additionally, character limitations lead to incomplete linguistic cues, forcing models to infer meaning from sparse data.
- Most sarcasm datasets are small, domain-specific, and monolingual. This restricts the performance of models when applied to cross-domain or multilingual scenarios, especially where sarcasm manifests differently across cultures and languages.
- Modern social media posts often combine text, images, GIFs, and videos. Detecting sarcasm purely from textual content fails to capture additional visual and auditory cues.
- Sarcastic statements represent only a small fraction of social media posts, leading to class imbalance. Models tend to overfit to the majority (non-sarcastic) class, degrading overall performance.

To address these challenges, future research on sarcasm detection should focus on the following areas:

- Integrating user-level, conversation-level, and temporal context can improve sarcasm interpretation. Leveraging knowledge graphs and user interaction histories can help models better understand implicit references.
- Future work should explore combining text, images, videos, and audio cues using multimodal deep learning architectures. Transformer-based models like CLIP and BLIP could be extended for sarcasm understanding in richer contexts.
- Developing multilingual sarcasm detection frameworks and domain-adaptive transfer learning techniques can improve model generalization, enabling robustness across different languages, cultures, and social media platforms.
- To mitigate the scarcity of labeled data, few-shot, zero-shot, and self-supervised learning approaches can be employed. Leveraging large language models (LLMs) such as GPT, BERT, and RoBERTa in semi-supervised frameworks holds significant promise.
- Integrating explainability frameworks will enhance user trust and allow practitioners to understand model decisions. Techniques like SHAP, LIME, and attention-based visualization could make sarcasm detection models more interpretable.
- Deploying sarcasm detection in real-time streaming environments (e.g., Twitter feeds) will require optimizing models for latency, scalability, and energy efficiency. Lightweight transformer variants and edge AI solutions should be explored.

7. Conclusion

Sarcasm remains a complex linguistic phenomenon that poses significant challenges for Natural Language Processing (NLP) systems. Despite the growing interest in sarcasm detection within the research community, accurately identifying sarcastic expressions remains challenging due to their context-dependent nature and the subtle interplay between literal and intended meanings. Existing studies typically rely on either large-scale, tag-based datasets, which often suffer from noisy annotations, or small, manually curated datasets, which provide high-quality labels but lack sufficient volume to train deep learning models effectively. This paper explored a variety of NLP and deep learning (DL)-based approaches to address these limitations and provided insights into state-of-the-art techniques for sarcasm detection in social media platforms such as Twitter. Through the review of recent models—ranging from transformer-based contextual embeddings to multi-task learning frameworks and ensemble architectures—we highlighted the strengths and shortcomings of existing methods. The findings emphasize that effective sarcasm detection requires context-aware modeling, multimodal feature integration (e.g., text, emojis, and sentiment), and balanced datasets to reduce label noise while maintaining scalability. The insights from this study can serve as a reference for future researchers and practitioners in designing robust, explainable, and generalizable sarcasm detection systems capable of handling diverse linguistic patterns across domains and languages.

REFERENCES

- [1] S. Amir *et al.*, “Modelling context with user embeddings for sarcasm detection in social media,” *arXiv preprint*, arXiv:1607.00976, 2016.
- [2] D. Bamman and N. Smith, “Contextualized sarcasm detection on Twitter,” in *Proc. Int. AAAI Conf. Web and Social Media*, vol. 9, no. 1, pp. 574–582, 2015.
- [3] A. Joshi, P. Bhattacharyya, and M. J. Carman, “Automatic sarcasm detection: A survey,” *ACM Comput. Surv.*, vol. 50, no. 5, pp. 1–22, Sep. 2017.
- [4] P. Liu, W. Chen, G. Ou, T. Wang, D. Yang, and K. Lei, “Sarcasm detection in social media based on imbalanced classification,” in *Proc. Int. Conf. Web-Age Inf. Manage. (WAIM)*, Cham, Switzerland: Springer, 2014, pp. 459–471.
- [5] D. K. Sharma, B. Singh, S. Agarwal, N. Pachauri, A. A. Alhussan, and H. A. Abdallah, “Sarcasm detection over social media platforms using hybrid ensemble model with fuzzy logic,” *Electronics*, vol. 12, no. 4, p. 937, Feb. 2023.
- [6] Y. A. Kolchinski and C. Potts, “Representing social media users for sarcasm detection,” *arXiv preprint*, arXiv:1808.08470, 2018.
- [7] P. Liu, W. Chen, G. Ou, T. Wang, D. Yang, and K. Lei, “Sarcasm detection in social media based on imbalanced classification,” in *Lecture Notes in Computer Science*, vol. 8485, F. Li, G. Li, S. Hwang, B. Yao, and Z. Zhang, Eds. Cham, Switzerland: Springer, 2014, pp. 459–471.
- [8] D. Vinoth and P. Prabhavathy, “An intelligent machine learning-based sarcasm detection and classification model on social networks,” *J. Supercomput.*, vol. 78, pp. 10575–10594, 2022, doi: 10.1007/s11227-022-04312-x.
- [9] Y. Zhang *et al.*, “Learning multitask commonness and uniqueness for multimodal sarcasm detection and sentiment analysis in conversation,” *IEEE Trans. Artif. Intell.*, vol. 5, no. 3, pp. 1349–1361, Mar. 2024, doi: 10.1109/TAI.2023.3298328.
- [10] Y. Kumar and N. Goel, “AI-based learning techniques for sarcasm detection of social media tweets: State-of-the-art survey,” *SN Comput. Sci.*, vol. 1, no. 318, 2020, doi: 10.1007/s42979-020-00336-3.
- [11] M. Bedi, S. Kumar, M. S. Akhtar, and T. Chakraborty, “Multi-modal sarcasm detection and humor classification in code-mixed conversations,” *IEEE Trans. Affect. Comput.*, vol. 14, no. 2, pp. 1363–1375, Apr. 2023, doi: 10.1109/TAFFC.2023.1234567.
- [12] E. Lunando and A. Purwarianti, “Indonesian social media sentiment analysis with sarcasm detection,” in *Proc. Int. Conf. Adv. Comput. Sci. Inf. Syst. (ICACSIS)*, Sanur Bali, Indonesia, 2013, pp. 195–198, doi: 10.1109/ICACSIS.2013.6761575.
- [13] K. K. Senthilkumar *et al.*, “Twitter sarcasm detection using natural language processing and deep learning techniques,” in *Proc. Global Conf. Commun. Inf. Technol. (GCCIT)*, Bangalore, India, 2024, pp. 1–8, doi: 10.1109/GCCIT63234.2024.10862514.
- [14] O. Prakash, “Sarcasm detection for sentiment analysis using deep learning enhancement,” *Procedia Comput. Sci.*, vol. 258, pp. 203–212, 2025, doi: 10.1016/j.procs.2025.01.023.

- [15] S. A. Mohd Ibrahim, M. M. Deshpande, and V. N. Pawar, "Automated sarcasm detection in English tweets using CCNN and ELLSTM with text and emoji embeddings," in *Proc. Int. Conf. Intell. Innov. Technol. Comput., Electr. Electron. (IITCEE)*, Bangalore, India, 2025, pp. 1–8, doi: [10.1109/IITCEE64140.2025.10915496](https://doi.org/10.1109/IITCEE64140.2025.10915496).
- [16] D. N. Sabera and D. A. Kristiyanti, "Comparative analysis of large language models as feature extraction methods in sarcasm detection using classification algorithms," in *Proc. Int. Conf. Electron. Represent. Algorithm (ICERA)*, Yogyakarta, Indonesia, 2025, pp. 352–357, doi: [10.1109/ICERA66156.2025.11087290](https://doi.org/10.1109/ICERA66156.2025.11087290).
- [17] N. Srinu, K. Sivaraman, and M. Sriram, "Enhancing sarcasm detection through grasshopper optimization with deep learning-based sentiment analysis on social media," *Int. J. Inf. Technol.*, vol. 17, pp. 1785–1791, 2025, doi: [10.1007/s41870-024-02057-9](https://doi.org/10.1007/s41870-024-02057-9).
- [18] K. K. Senthilkumar *et al.*, "Twitter sarcasm detection using natural language processing and deep learning techniques," in *Proc. Global Conf. Commun. Inf. Technol. (GCCIT)*, Bangalore, India, 2024, pp. 1–8, doi: [10.1109/GCCIT63234.2024.10862514](https://doi.org/10.1109/GCCIT63234.2024.10862514).
- [19] M. E. Hassan *et al.*, "Detection of sarcasm in Urdu tweets using deep learning and transformer-based hybrid approaches," *IEEE Access*, vol. 12, pp. 61542–61555, 2024, doi: [10.1109/ACCESS.2024.3393856](https://doi.org/10.1109/ACCESS.2024.3393856).
- [20] R. T. Gedela, U. Baruah, and B. Soni, "Deep contextualised text representation and learning for sarcasm detection," *Arab. J. Sci. Eng.*, vol. 49, pp. 3719–3734, 2024, doi: [10.1007/s13369-023-08170-4](https://doi.org/10.1007/s13369-023-08170-4).
- [21] V. Sukhavasi and V. Dondeti, "Effective Automated Transformer Model Based Sarcasm Detection Using Multilingual Data," *Multimedia Tools and Applications*, vol. 83, pp. 47531–47562, 2024, doi: [10.1007/s11042-023-17302-9](https://doi.org/10.1007/s11042-023-17302-9).
- [22] J. Pradhan, R. Verma, S. Kumar, and V. Sharma, "An Efficient Sarcasm Detection Using Linguistic Features and Ensemble Machine Learning," *Procedia Computer Science*, vol. 235, pp. 1058–1067, 2024, doi: [10.1016/j.procs.2024.04.100](https://doi.org/10.1016/j.procs.2024.04.100).
- [23] R. Pandey, A. Kumar, J. P. Singh, *et al.*, "A Hybrid Convolutional Neural Network for Sarcasm Detection from Multilingual Social Media Posts," *Multimedia Tools and Applications*, vol. 84, pp. 15867–15895, 2025, doi: [10.1007/s11042-024-19672-0](https://doi.org/10.1007/s11042-024-19672-0).
- [24] B. S. Deora and S. S. Sarangdevot, "Deep Learning Approaches for Early Detection of Depression Using Sarcasm and Twitter Dataset," in *Intelligent Sustainable Systems (Worlds4 2024)*, Lecture Notes in Networks and Systems, vol. 1180, Springer, Singapore, 2025, doi: [10.1007/978-981-97-9324-2_44](https://doi.org/10.1007/978-981-97-9324-2_44).
- [25] T. Javed, M. A. Nouman, and R. Zahid, "BERT Model Adoption for Sarcasm Detection on Twitter Data," *VFAST Transactions on Software Engineering*, vol. 12, no. 3, pp. 177–198, 2024, doi: [10.21015/vtse.v12i3.1908](https://doi.org/10.21015/vtse.v12i3.1908).
- [26] C. Thaokar, J. K. Rout, M. Rout, *et al.*, "N-Gram Based Sarcasm Detection for News and Social Media Text Using Hybrid Deep Learning Models," *SN Computer Science*, vol. 5, p. 163, 2024, doi: [10.1007/s42979-023-02506-5](https://doi.org/10.1007/s42979-023-02506-5).
- [27] H. Saleem, A. Naeem, K. Abid, and N. Aslam, "Sarcasm Detection on Twitter Using Deep Handcrafted Features," *Journal of Computing & Biomedical Informatics*, vol. 4, no. 2, pp. 117–127, 2023. [Online]. Available: <https://jcbi.org/index.php/Main/article/view/128>.
- [28] A. Kumar and G. Garg, "Empirical Study of Shallow and Deep Learning Models for Sarcasm Detection Using Context in Benchmark Datasets," *Journal of Ambient Intelligence and Humanized Computing*, vol. 14, pp. 5327–5342, 2023, doi: [10.1007/s12652-019-01419-7](https://doi.org/10.1007/s12652-019-01419-7).
- [29] A. Kumar, S. R. Sangwan, A. K. Singh, and G. Wadhwa, "Hybrid Deep Learning Model for Sarcasm Detection in Indian Indigenous Language Using Word-Emoji Embeddings," *ACM Transactions on Asian and Low-Resource Language Information Processing*, vol. 22, no. 5, pp. 1–20, May 2023, doi: [10.1145/3519299](https://doi.org/10.1145/3519299).
- [30] J. S. Murthy and G. M. Siddesh, "A Real-Time Framework for Automatic Sarcasm Detection Using Proposed Tensor-DNN-50 Algorithm," in *Proceedings of the 4th International Conference on Frontiers in Computing and Systems (COMSYS 2023)*, Lecture Notes in Networks and Systems, vol. 974, Springer, Singapore, 2024, doi: [10.1007/978-981-97-2611-0_8](https://doi.org/10.1007/978-981-97-2611-0_8).
- [31] B. Rajani, S. Saxena, B. S. Kumar, and G. Narang, "Sarcasm Detection and Classification Using Deep Learning Model," in *Soft Computing: Theories and Applications (SoCTA 2023)*, Lecture Notes in Networks and Systems, vol. 971, Springer, Singapore, 2024, doi: [10.1007/978-981-97-2089-7_34](https://doi.org/10.1007/978-981-97-2089-7_34).
- [32] R. P. Kumar, *et al.*, "Sarcasm Detection in Telugu and Tamil: An Exploration of Machine Learning and Deep Neural Networks," in *Proc. 14th Int. Conf. Computing, Communication and Networking Technologies (ICCCNT)*, Delhi, India, 2023, pp. 1–7, doi: [10.1109/ICCCNT56998.2023.10306775](https://doi.org/10.1109/ICCCNT56998.2023.10306775).
- [33] S. Sharma and N. Joshi, "An Optimized Approach for Sarcasm Detection Using Machine Learning Classifier," in *Data Science and Applications (ICDSA 2023)*, Lecture Notes in Networks and Systems, vol. 821, Springer, Singapore, 2024, doi: [10.1007/978-981-99-7814-4_7](https://doi.org/10.1007/978-981-99-7814-4_7).

- [34] M. Pal and R. Prasad, "Sarcasm Detection Followed by Sentiment Analysis for Bengali Language: Neural Network & Supervised Approach," in *Proc. Int. Conf. Advances in Intelligent Computing and Applications (AICAPS)*, Kochi, India, 2023, pp. 1–7, doi: [10.1109/AICAPS57044.2023.10074510](https://doi.org/10.1109/AICAPS57044.2023.10074510).
- [35] S. Sinha and V. K. Yadav, "Sarcasm Detection in News Headlines Using Deep Learning," in *Proc. Int. Conf. Recent Advances in Science and Engineering Technology (ICRASET)*, B G Nagara, India, 2023, pp. 1–6, doi: [10.1109/ICRASET59632.2023.10420344](https://doi.org/10.1109/ICRASET59632.2023.10420344).
- [36] P. Goel, R. Jain, A. Nayyar, *et al.*, "Sarcasm Detection Using Deep Learning and Ensemble Learning," *Multimedia Tools and Applications*, vol. 81, pp. 43229–43252, 2022, doi: [10.1007/s11042-022-12930-z](https://doi.org/10.1007/s11042-022-12930-z).
- [37] D. Vinoth and P. Prabhavathy, "An Intelligent Machine Learning-Based Sarcasm Detection and Classification Model on Social Networks," *Journal of Supercomputing*, vol. 78, pp. 10575–10594, 2022, doi: [10.1007/s11227-022-04312-x](https://doi.org/10.1007/s11227-022-04312-x).

