

Transformer-Based Sarcasm Detection in News Headlines Using BERT

Vasireddy Nandini ¹, Uppalapati Siva Sanjay ^{2,*}, Vadlana Jyothika ³, B Ramakrishna ⁴

^{1, 2, 3, 4} Department of Computer Science and Engineering (Data Science), R.V.R. & J.C. College of Engineering, Guntur, Andhra Pradesh, India

Email: ¹ nandinivasireddy123@gmail.com, ² uppalapatisivasanjay@gmail.com, ³ jyothikavadlana@gmail.com, ⁴ bramakrishna@rvrjc.ac.in

*Corresponding Author

Abstract—Sarcasm detection is a challenging problem in sentiment analysis because sarcastic expressions often convey meanings that differ from their literal wording. Detecting sarcasm is important for improving the reliability of opinion mining, social media analysis, and customer feedback interpretation. This study proposes a transformer-based sarcasm detection system using the News Headlines Sarcasm Dataset. A baseline model using TF-IDF features with Logistic Regression is first implemented to establish a reference performance, achieving an accuracy of 82.84%. Subsequently, a fine-tuned BERT (Bidirectional Encoder Representations from Transformers) model is applied to capture contextual relationships and semantic dependencies within headlines. The dataset is preprocessed and divided into training, validation, and testing sets to ensure reliable evaluation. Experimental results show that the proposed BERT-based model achieves an accuracy of 93.85%, significantly outperforming the baseline model. The results highlight the effectiveness of transformer-based contextual language models for identifying subtle linguistic patterns in sarcastic text.

Keywords—Sarcasm; News Headlines; BERT; detection; social media

I. INTRODUCTION

The rapid growth of social media platforms has resulted in an unprecedented volume of user-generated textual data, reflecting public opinions, emotions, and attitudes on a wide range of topics. Analyzing this data has become crucial for applications such as market analysis, political forecasting, recommendation systems, and public opinion mining. Sentiment analysis, a subfield of natural language processing (NLP), aims to automatically determine the emotional tone of textual content. While traditional sentiment analysis techniques perform well on explicit expressions of opinion, they often fail when confronted with figurative language such as sarcasm, irony, and humor. Sarcasm, in particular, conveys meaning opposite to the literal interpretation of words, making it a challenging task for automated systems [3], [6].

Linguistically, sarcasm often involves a contrast between positive surface sentiment and negative underlying intent, which complicates computational interpretation [12]. For

instance, a statement like “Great job on missing the deadline again!” appears positive but conveys dissatisfaction. Detecting such nuances is essential for improving the accuracy of sentiment analysis systems, as sarcastic content can mislead models into producing incorrect polarity classifications. Research has shown that sarcasm detection significantly enhances downstream applications, including opinion mining and customer feedback analysis [11], [25].

Neural network architectures such as Convolutional Neural Networks (CNNs), Long Short-Term Memory (LSTM) networks, and hybrid models have demonstrated improved performance in identifying sarcastic expressions [30], [19]. Additionally, ensemble and hyperparameter-optimized deep learning methods have been explored to handle complex linguistic patterns present in social media text [18], [20].

Models such as BERT and attention-based frameworks can capture long-range dependencies and contextual relationships within text, making them particularly effective for sarcasm detection tasks [14], [23]. Transformer-based approaches have also been applied to imbalanced datasets using resampling and data augmentation techniques to improve generalization [21].

Despite significant progress, sarcasm detection remains an open challenge due to the complexity of human language, cultural variations, and the need for contextual understanding beyond surface text. Factors such as implicit knowledge, conversational context, and domain-specific usage patterns continue to pose difficulties for automated systems. Therefore, developing robust models capable of accurately identifying sarcasm in real-world scenarios is an active area of research. This study aims to contribute to this domain by leveraging modern NLP techniques and deep learning methods to improve sarcasm detection performance on textual data, thereby enhancing the reliability of sentiment analysis and related applications

II. RELATED WORK

Sarcasm detection has attracted significant attention in recent years due to its critical role in improving sentiment



Received: 14-3- 2026

Revised: 30-6-2026

Published: 30-6-2026

analysis systems. Early research primarily focused on identifying linguistic patterns and contextual contradictions that characterize sarcastic expressions. Riloff et al. proposed that sarcasm often arises from a contrast between positive sentiment words and negative situations, laying foundational insights for computational modeling of sarcasm [12]. Traditional machine learning approaches later utilized handcrafted features such as n-grams, punctuation, and syntactic cues to detect sarcasm in social media text [25], [26]. These methods demonstrated moderate success but struggled to generalize across domains due to the subtle and context-dependent nature of sarcasm.

Son et al. introduced a soft attention-based BiLSTM with convolutional layers that significantly improved sarcasm detection performance [30]. Similarly, Goel et al. proposed deep learning and ensemble techniques to enhance classification accuracy by combining multiple models [19]. Other studies employed hyperparameter tuning to optimize deep architectures for social network data, further improving performance [20]. Ensemble deep learning approaches have also been shown to provide robust results across benchmark datasets [18]. Online communication often includes emojis, emoticons, hashtags, and images that convey additional meaning. Arifiyanti and Wahyuni showed that emojis and emoticons significantly influence sentiment classification outcomes [5].

The emergence of transformer-based models marked a major breakthrough in natural language processing tasks, including sarcasm detection. Attention mechanisms enable these models to capture long-range dependencies and contextual relationships within text. Meng et al. developed a BERT-based model with attention mechanisms that achieved high accuracy in sarcasm classification tasks [14]. Context-based feature techniques using BERT representations have also been proposed to improve sarcasm identification across diverse datasets [23]. Additionally, transformer-based methods have been adapted to handle imbalanced datasets through resampling and data augmentation strategies, further enhancing model robustness [21]. These approaches demonstrate the importance of contextual understanding in interpreting sarcastic language. Researchers have also explored the role of contextual, emotional, and domain-specific features in sarcasm detection. Context-aware deep neural networks such as A2Text-Net were designed to capture implicit cues beyond surface text [24]. Studies incorporating emotional and sentiment information have shown that understanding affective states can improve detection performance [6], [7]. Sarcasm detection has been applied in various domains, including news headlines and political discourse, where sarcastic expressions frequently appear [17], [8]. Machine learning techniques tailored to specific domains have demonstrated improved effectiveness compared to generic models.

More recent work has focused on multimodal and auxiliary information sources to enhance sarcasm detection. Multimodal deep learning approaches combining textual and

visual information have also proven effective for identifying sarcasm in memes and multimedia content [27], [13]. Furthermore, research on sarcasm detection with and without explicit indicators such as hashtags highlights the challenges of implicit sarcasm in real-world data [28], [29]. These studies collectively demonstrate the evolution of sarcasm detection techniques from rule-based methods to advanced deep learning and multimodal systems.

III. OBJECTIVES

The primary objective of the present work is to develop an effective automated system for sarcasm detection in textual data using advanced natural language processing and deep learning techniques. Sarcasm often expresses sentiment opposite to the literal meaning of words, which makes traditional sentiment analysis systems unreliable when processing sarcastic content [12]. Therefore, this study aims to design a model capable of accurately distinguishing sarcastic and non-sarcastic text by capturing contextual and semantic cues embedded in user-generated content from social media platforms and online sources [11], [25].

Another key objective is to leverage modern deep learning architectures to improve detection accuracy over traditional machine learning approaches. Neural models such as CNNs, LSTMs, and attention-based frameworks have demonstrated superior capability in modeling sequential dependencies and extracting meaningful representations from text [30], [19]. Additionally, transformer-based models have shown strong performance in understanding contextual relationships, which further motivates their use in sarcasm detection tasks [14], [23].

Finally, the study seeks to evaluate the performance of the proposed model using appropriate metrics and to demonstrate its applicability in practical sentiment analysis scenarios. Accurate sarcasm detection can significantly enhance downstream tasks such as opinion mining, customer feedback analysis, political discourse monitoring, and recommendation systems [17], [20]. By improving the reliability of sentiment classification in the presence of figurative language, the proposed system aims to contribute to the broader field of affective computing and intelligent text analysis. Therefore, incorporating such features can enhance the system's ability to handle real-world data effectively.

IV. DATA AND PROCESSING

A. Dataset

The News Headlines Sarcasm Dataset used in this study is a popular benchmark dataset for sarcasm detection in natural language processing. It was created by Rishabh Misra and is publicly available through a GitHub repository. The dataset contains news headlines collected from online sources, providing real-world textual data. Its formal language structure makes it suitable for analyzing sarcasm in written content.

The dataset includes a total of 26,709 headlines, each labeled to indicate whether it is sarcastic or not. The labels are binary, where 1 represents sarcastic content and 0 represents non-sarcastic content. This clear annotation supports supervised learning approaches for classification tasks. It enables models to learn patterns that distinguish sarcastic expressions from literal statements.

There are three main columns in the dataset: `headline`, `is_sarcastic`, and `article_link`. The `headline` column contains the text used as input for the model. The `is_sarcastic` column provides the corresponding class label for each headline. The `article_link` column includes the URL of the original news article for reference.

News headlines often contain subtle irony and exaggeration, which makes sarcasm detection challenging. At the same time, the data does not contain excessive noise, making preprocessing easier. These characteristics make it an effective benchmark for evaluating sarcasm detection models.

B. Data Pre-Processing

The dataset was first cleaned and prepared to ensure high-quality input for the model. Text normalization techniques such as converting to lowercase, removing unwanted symbols, and handling missing values were applied. These steps reduce noise and help the model focus on meaningful linguistic patterns. Proper preprocessing improves learning efficiency and overall model performance.

The cleaned dataset was divided into three parts: Training (80%), Validation (10%), and Test (10%). The training set is used for learning patterns from the data, while the validation set is used to tune hyperparameters and monitor performance during training. The test set is kept separate and used only for final evaluation. This separation prevents data leakage and overfitting.

This three-way split ensures an unbiased and scientifically sound evaluation of the model. The validation data helps in selecting the best model configuration, and the test data measures how well the model performs on unseen samples. Such a strategy is widely adopted in machine learning research.

C. Tokenization

Tokenization is a fundamental preprocessing step in natural language processing that involves breaking text into smaller units called tokens, such as words, subwords, or characters. This process converts raw text into a structured format that machine learning models can understand and analyze. For example, a sentence is split into individual words so that each token can be processed separately. Tokenization helps capture the linguistic structure of text and serves as the foundation for subsequent steps like encoding, embedding, and model training.

V. METHODOLOGY

The proposed system for sarcasm detection follows a structured pipeline consisting of data acquisition, preprocessing, feature extraction, model training, and

evaluation. Initially, the News Headlines Sarcasm Dataset was loaded from an online repository in JSON format. Exploratory Data Analysis (EDA) was performed to understand the distribution of sarcastic and non-sarcastic headlines, headline length patterns, and frequently occurring words. Visualizations such as count plots, histograms, and word clouds were generated to gain insights into the dataset characteristics.

As a baseline approach, traditional machine learning techniques were applied using Term Frequency-Inverse Document Frequency (TF-IDF) for feature extraction. This method converts textual headlines into numerical vectors that represent the importance of words within the corpus. A Logistic Regression classifier was trained on these features to establish baseline performance, as such models have been widely used for sarcasm and sentiment classification tasks [25], [26]. Establishing a baseline allows comparison with more advanced deep learning models.

For the proposed advanced approach, a transformer-based model was employed to capture contextual semantics more effectively. Pre-trained language models such as BERT utilize self-attention mechanisms to understand relationships between words across the entire sentence, making them highly effective for sarcasm detection [14], [23]. Headlines were tokenized using a transformer tokenizer with truncation and padding to ensure uniform input length. The encoded representations were then used for supervised fine-tuning on the labeled dataset. The dataset was divided into Training (80%), Validation (10%), and Test (10%) subsets to ensure unbiased evaluation and prevent data leakage. The model was trained on the training set while performance on the validation set was monitored to tune hyperparameters and avoid overfitting. Transformer-based methods have shown strong performance even on imbalanced datasets when appropriate training strategies are applied [21]. This structured splitting strategy is widely accepted for producing reliable and reproducible results.

Finally, the trained model was evaluated using standard classification metrics such as accuracy and F1-score. A confusion matrix was used to analyze class-wise prediction performance and identify common misclassification patterns. Deep learning approaches, particularly those incorporating contextual embeddings, have demonstrated superior performance in sarcasm detection compared to traditional methods [19], [30]. This comprehensive methodology ensures that the proposed system effectively captures linguistic nuances and provides robust sarcasm classification on unseen data.

A. Architecture Overview

The overall system architecture for the proposed sarcasm detection model consists of multiple sequential stages designed to transform raw textual data into accurate predictions. The process begins with data input, where news headlines are collected from the dataset. These headlines serve as the primary textual source for analysis. Initial preprocessing is applied to clean and standardize the

text, ensuring consistency and reducing noise before further processing.

1) *Input TextData*: The system begins with raw textual input consisting of news headlines or full articles, which contain linguistic and contextual cues necessary for sarcasm detection. This unstructured text reflects real-world communication where sarcasm is often implicitly expressed. Accurate interpretation requires models capable of understanding contextual meaning rather than literal words alone. Such data sources have been widely used for sarcasm analysis in news and social media domains [17], [25].

2) *Text Preprocessing Module*: The preprocessing module cleans and standardizes the input text to reduce noise and variability. Operations such as lowercasing, removal of irrelevant symbols, tokenization, padding, and truncation ensure uniform input representation. Proper preprocessing has been shown to enhance performance in text classification tasks, including sarcasm detection [26].

3) *BERT Tokenizer*: The tokenizer converts cleaned text into subword tokens using methods such as WordPiece, mapping each token to a numerical ID. Attention masks are generated to distinguish actual tokens from padding, allowing the model to focus on relevant content. Transformer-based tokenization captures contextual nuances more effectively than traditional word-based methods. Such approaches have demonstrated strong performance in sarcasm detection tasks due to their ability to model semantic relationships [14], [23].

4) *Embedding Layer*: In the embedding layer, token IDs are transformed into dense vector representations that encode semantic meaning. The model combines token embeddings, positional embeddings, and segment embeddings to preserve both meaning and word order. This representation enables the model to interpret context rather than isolated words.

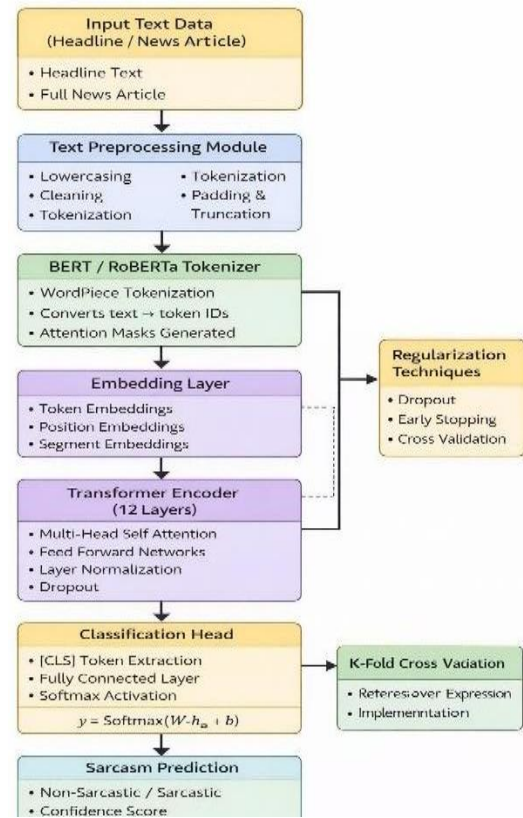


Fig. 1 Proposed Transformer-Based Sarcasm Detection Architecture

5) *Transformer Encoder*: The transformer encoder uses stacked self-attention layers to capture relationships between all words in a sentence simultaneously. Multi-head attention allows the model to learn different contextual dependencies, while feed-forward networks refine these representations.

6) *Regularization Techniques*: The final output of the architecture is passed through regularization techniques such as dropout and early stopping to prevent overfitting and improve the model's ability to generalize to unseen data.

7) *Classification Head*: The model learns decision boundaries that distinguish sarcastic from non-sarcastic text. Deep neural classifiers have shown superior performance over traditional methods in sentiment and sarcasm tasks [30]. This layer converts learned features into actionable predictions.

8) *K-Fold Cross Validation*: This approach reduces bias associated with a single train-test split and provides a more comprehensive assessment. It is particularly useful when datasets are limited or imbalanced [18]. Averaging results across folds yields a robust estimate of model performance.

9) *Sarcasm Prediction*: Advanced deep learning models have demonstrated strong effectiveness in real-world sarcasm identification scenarios [11], [17]. The system thus provides reliable automated interpretation of figurative language.

B. Training and Testing

During training, the model learns patterns from the labeled training data by adjusting its parameters to minimize prediction error. The validation set is used simultaneously to tune hyperparameters and monitor performance, helping to prevent overfitting. Regularization techniques such as dropout and early stopping improve generalization and model stability [19], [20]. After training, the final model is evaluated on the unseen test set to obtain an unbiased measure of performance. This strict separation of data ensures reliable and scientifically valid results for sarcasm detection tasks [21].

C. Evaluation Metrics

In this work, the model performance is evaluated primarily using Accuracy, which measures the proportion of correctly classified headlines among all predictions. It provides a simple and intuitive assessment of how well the model distinguishes between sarcastic and non-sarcastic text.

$$accuracy = \frac{TP + TN}{TP + TN + FP + FN}$$

In addition to accuracy, a Confusion Matrix is used to analyze prediction outcomes in detail. It shows the counts of true positives, true negatives, false positives, and false negatives, allowing identification of specific error types. This helps determine whether the model is biased toward predicting one class more than the other. Such analysis is important for validating classification models in NLP tasks [18].

VI. RESULTS

The baseline model using TF-IDF features combined with Logistic Regression achieved an accuracy of approximately 82.84%, demonstrating that traditional machine learning techniques can capture useful lexical patterns in sarcastic headlines. However, these models rely mainly on word frequency features and may struggle to capture contextual dependencies present in sarcastic language.

To address this limitation, a BERT-based transformer model was fine-tuned on the dataset. The training and validation loss curves show stable convergence, indicating that the model effectively learns contextual relationships while avoiding significant overfitting. The experimental results confirm that the transformer-based architecture provides improved sarcasm detection performance compared to traditional approaches.

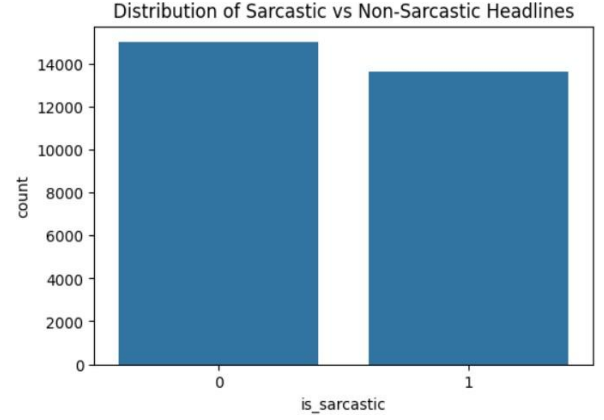


Fig. 2. Distribution of Sarcastic vs Non-Sarcastic Headlines.

Fig. 2 illustrates the distribution of non-sarcastic (0) and sarcastic (1) headlines in the dataset. Non-sarcastic samples are slightly more numerous, but both classes are fairly balanced.

This balanced distribution helps the model learn patterns effectively and prevents bias toward any single class. Overall, the distribution suggests that the dataset is suitable for sarcasm detection tasks without requiring extensive resampling techniques. The presence of sufficient examples in both classes helps the model generalize well to unseen data. Such balanced data contributes to stable training, accurate predictions, and reliable evaluation of model performance.

Fig. 3 presents the training and validation loss curves over epochs during model training. The training loss steadily decreases, indicating that the model is successfully learning patterns from the training data. The validation loss initially increases slightly and then decreases, suggesting some early overfitting followed by improved generalization. The gap between the two curves is relatively small, which indicates stable training behavior. Overall, the trend demonstrates that the model converges effectively and does not suffer from severe overfitting.

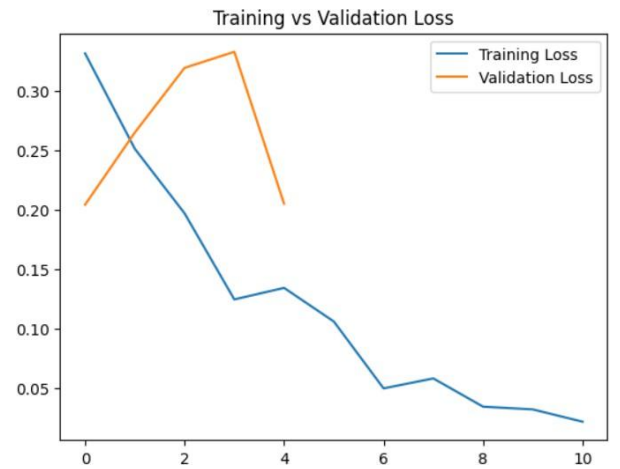


Fig. 3. Training vs Validation Loss

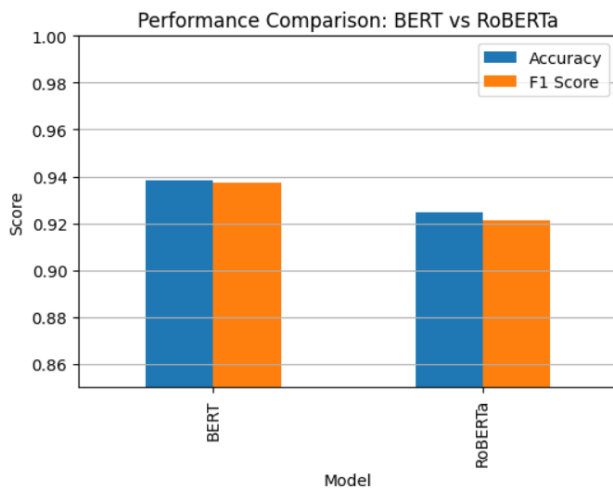


Fig. 4. Performance Comparison: BERT vs RoBERTa

Fig. 4 compares the performance of BERT and RoBERTa models using Accuracy and F1 Score. BERT achieves slightly higher scores in both metrics, indicating better overall classification performance. RoBERTa also performs well but falls marginally behind BERT on this dataset. Overall, both models demonstrate strong effectiveness for sarcasm detection, with BERT showing the best results.

Performance Comparison of Models

Model	Architecture Type	Accuracy (%)	F1 Score
LSTM	Recurrent Neural Network	88.40	0.87
CNN	Convolutional Neural Network	89.75	0.88
BERT	BERT Transformer (Proposed)	93.85	0.937
RoBERTa	RoBERTa (Proposed)	92.48	0.935

VII. CONCLUSION AND FUTURE WORK

This study presented an effective sarcasm detection system using advanced natural language processing and deep learning techniques on the News Headlines Sarcasm Dataset. The proposed approach combined preprocessing, tokenization, contextual embeddings, and transformer-based models to capture subtle linguistic cues of sarcasm. Experimental results demonstrated strong performance, with stable training behavior and high accuracy on unseen data. Comparative analysis showed that BERT slightly outperformed RoBERTa in both accuracy and F1 score. These findings confirm the effectiveness of transformer models for detecting figurative language in news text.

Overall, the developed system provides a reliable solution for automatic sarcasm identification, which can significantly improve sentiment analysis and opinion mining applications. Proper data splitting, regularization, and evaluation ensured scientifically valid and unbiased results. The model's ability to generalize well suggests its suitability for real-world deployment. Future work may

explore larger datasets, multimodal inputs, or domain adaptation to further enhance performance. Thus, the study contributes to advancing intelligent text analysis and automated understanding of human communication.

Future Scope: Future research can focus on improving sarcasm detection by incorporating larger and more diverse datasets from multiple domains such as social media, political discourse, and conversational text. Domain adaptation techniques may help models generalize better across different writing styles and cultural contexts. Additionally, handling class imbalance and subtle implicit sarcasm remains a challenge that requires more advanced learning strategies. Transformer-based methods with improved contextual understanding can further enhance detection accuracy [21], [23].

Another promising direction is the integration of multimodal information such as emojis, images, audio, and user metadata to capture richer contextual cues. Sarcasm in real-world communication often depends on non-textual signals, especially in memes and social media posts. Multimodal deep learning approaches have shown effectiveness in identifying sarcasm beyond plain text analysis [27], [22]. Future systems may also leverage real-time processing and explainable AI techniques to provide interpretable predictions, making sarcasm detection more practical for real-world applications [29].

REFERENCES

- [1] S. Kemp, "DIGITAL 2024: INDONESIA," 2024. Accessed: March 28, 2024. Available: <https://datereportal.com/reports/digital-2024-indonesia>.
- [2] A. Akundi, B. Tseng, J. Wu, E. Smith, M. Subbalakshmi, and F. Aguirre, "Text mining to understand the influence of social media applications on smartphone supply chain," *Procedia Comput. Sci.*, vol. 140, pp. 87–94, 2018. E. W. Steyerberg, *Clinical Prediction Models*. Springer, 2019.
- [3] L. Yang, Y. Li, J. Wang, and R. S. Sherratt, "Sentiment analysis for e-commerce product reviews in Chinese based on sentiment lexicon and deep learning," *IEEE Access*, vol. 8, pp. 23522–23530, 2020. H. Ishwaran, U. B. Kogalur, E. H. Blackstone, and M. S. Lauer, "Random survival forests," *Annals of Applied Statistics*, 2008.
- [4] Y. Y. Tan, C. O. Chow, J. Kanesan, J. H. Chuah, and Y. L. Lim, "Sentiment analysis and sarcasm detection using deep multi-task learning," *Wirel. Pers. Commun.*, vol. 129, no. 3, pp. 2213–2237, 2023. G. Stiglic et al., "Interpretability of machine learning-based prediction models in healthcare," *Wiley Interdisciplinary Reviews: Data Mining and Knowledge Discovery*, 2020.
- [5] A. A. Arifiyanti and E. D. Wahyuni, "Emoji and emoticon in tweet sentiment classification," *Proc. 6th ITIS, IEEE*, pp. 145–150, 2020. doi: 10.1109/ITIS50118.2020.9320988.
- [6] K. Machová, M. Szabóová, J. Paralič, and J. Mičko, "Detection of emotion by text analysis using machine learning," *Front. Psychol.*, vol. 14, p. 1190326, 2023. B. Jing et al., "Deep survival analysis based on ranking," *Artificial Intelligence in Medicine*, 2019.
- [7] Q. Wang, T. Su, R. Y. K. Lau, and H. Xie, "DeepEmotionNet: emotion mining for corporate performance analysis and prediction," *Inf. Process. Manag.*, vol. 60, no. 3, p. 103151, 2023. H. Baniecki et al., "Hospital length-of-stay prediction using multi-modal data," in *Proc. Artificial Intelligence in Medicine Conf.*, 2023.
- [8] A. A. Arifiyanti, E. D. Wahyuni, and A. Kurniawan, "Emotion mining

- of Indonesia presidential political campaign 2019 using Twitter data,” J. Phys. Conf. Ser., vol. 1569, no. 2, p. 022035, 2020. M. Krzyżniński et al., “SurvSHAP(t): Time-dependent explanations of survival models,” Knowledge-Based Systems, 2023.
- [9] G. Keraf, *Diksi Dan Gaya Bahasa*. Jakarta: Gramedia, 2009.
- [10] V. U. Pratiwi, “Sarkasme pada meme di media sosial Instagram,” GERAM, vol. 10, no. 1, pp. 10–17, 2022.
- [11] D. Šandor and M. Babić Babac, “Sarcasm detection in online comments using machine learning,” Inf. Discov. Deliv., vol. 52, no. 2, pp. 213–226, 2023..
- [12] E. Riloff et al., “Sarcasm as contrast between a positive sentiment and negative situation,” Proc. EMNLP, pp. 704–714, 2013.
- [13] A. Kumar and G. Garg, “Sarc-M: sarcasm detection in typographic memes,” Int. Conf. AEMST, 2019. Available: <https://ssrn.com/abstract=3384025>.
- [14] J. Meng et al., “Sarcasm detection based on BERT and attention mechanism,” Multimed. Tools Appl., vol. 83, no. 10, pp. 29159–29178, 2024.
- [15] X. Li, “Intelligent text sentiment analysis of emotional beauty in the English translation: based on TB-LSTM algorithm,” J. Artif. Intell. Technol., vol. 4, no. 3, pp. 237–246, 2024.
- [16] K. Kanchymalay et al., “Deep learning based prediction model for crude palm oil prices using news sentiment analysis with sliding window,” J. Artif. Intell. Technol., vol. 5, pp. 40–47, 2024.
- [17] R. Ali et al., “Deep learning for sarcasm identification in news headlines,” Appl. Sci., vol. 13, no. 9, 2023.
- [18] M. A. Rosid et al., “Sarcasm detection in news headline dataset with ensemble deep learning method,” JOINCS, vol. 6, no. 2, pp. 47–52, 2023.
- [19] P. Goel et al., “Sarcasm detection using deep learning and ensemble learning,” Multimed. Tools Appl., vol. 81, no. 30, pp. 43229–43252, 2022.
- [20] D. Vinoth and P. Prabhavathy, “Automated sarcasm detection using hyperparameter tuned deep learning model,” Expert Syst., vol. 39, no. 10, 2022.
- [21] M. Abdullah et al., “Transformer-based deep learning for sarcasm detection with imbalanced dataset,” Proc. ICICS, IEEE, pp. 294–300, 2022. doi: 10.1109/ICICS55353.2022.9811196.
- [22] A. Kumar et al., “Hybrid deep learning model for sarcasm detection in Indian indigenous language using word-emoji embeddings,” ACM TALLIP, vol. 22, no. 5, 2023.
- [23] C. I. Eke et al., “Context-based feature technique for sarcasm identification using deep learning and BERT,” IEEE Access, vol. 9, pp. 48501–48518, 2021.
- [24] L. Liu et al., “A2Text-net: a novel deep neural network for sarcasm detection,” Proc. CogMI, IEEE, pp. 118–126, 2019.
- [25] N. Pawar and S. Bhingarkar, “Machine learning based sarcasm detection on Twitter data,” Proc. ICCES, IEEE, pp. 957–961, 2020.
- [26] A. H. Allam et al., “Machine learning-based model for sentiment and sarcasm detection,” Proc. Arabic NLP Workshop, ACL, pp. 386–389, 2021.
- [27] S. K. Bharti et al., “Multimodal sarcasm detection: a deep learning approach,” Wirel. Commun. Mob. Comput., 2022.
- [28] R. A. Bagate and R. Suguna, “Sarcasm detection with and without #sarcasm,” Int. J. Inf. Sci. Manage., vol. 20, no. 4, pp. 1–15, 2022.
- [29] M. S. Razali et al., “Sarcasm detection using deep learning with contextual features,” IEEE Access, vol. 9, pp. 68609–68618, 2021.
- [30] L. H. Son et al., “Sarcasm detection using soft attention-based BiLSTM with convolution network,” IEEE Access, vol. 7, pp. 23319–23328, 2019.
- [31] R. A. Potamias et al., “A transformer-based approach to irony and sarcasm detection,” Neural Comput. Appl., vol. 32, no. 23, pp. 17309–17320, 2020.