

Available online at www.qu.edu.iq/journalcm

JOURNAL OF AL-QADISIYAH FOR COMPUTER SCIENCE AND MATHEMATICS

ISSN:2521-3504(online) ISSN:2074-0204(print)



Assessing the Influence of Advertisements on Social Interactions in Iraqi Dialect WhatsApp Groups Using BERT

Ghosoon K. Munahy*

Department of Information Technology, College of Computer Science and Information Technology, University of Kerbala, Kerbala, Iraq.

Email : ghosoonkareem@gmail.com

ARTICLE INFO

Article history:

Received: 17 /04/2025

Revised form: 22 /05/2025

Accepted : 11 /06/2025

Available online: 30 /06/2025

Keywords:

BERT Model, Iraqi Dialect NLP,
Social Interaction, WhatsApp Ads,
NLP Analysis

ABSTRACT

The widespread use of instant messaging applications, such as WhatsApp, has converted these platforms into dynamic venues for social interaction. The rising prevalence of commercial advertisements may negatively affect the quality of social connections. This study investigates the impact of ads on social interactions in WhatsApp groups using the Iraqi dialect, applying natural language processing and artificial intelligence techniques. We propose a methodology employing the BERT model to classify WhatsApp messages in the Iraqi dialect into three primary categories: advertisements, social discourse, and neutral communications. The primary objective is to assess the impact of advertisements on the dynamics of discussions among users in groups. A dataset comprising 5,000 messages was meticulously gathered and categorized into two classifications: advertisements and social discussions. The pre-trained CAMELBERT model underwent fine-tuning on this dataset by incorporating a classification head and training for 50 epochs with a batch size of 8 and a learning rate $2e-5$. Experimental results indicate that the model attained an F1-score of 97%, effectively differentiating between commercial communications and casual conversations. Approximately 35% of the messages were classified as promotional content. A conventional SVM model utilising TF-IDF features was implemented to assess performance, attaining merely 81.3% accuracy, underscoring the superiority of the transformer-based methodology. These findings indicate that the increasing prevalence of advertisements may discreetly disrupt the natural flow of conversations in digital communities, necessitating the implementation of sophisticated filtering systems.

MSC..

<https://doi.org/10.29304/jqcm.2025.17.22192>

1. Introduction

Instant messaging apps like WhatsApp are getting a lot of traction these days, and many people now rely on private groups as their main way to stay in touch, especially within local communities [1]. Groups like these let people communicate and exchange information, and talk about daily life naturally. At the same time, a steady stream of ads tends to get in the way sometimes, even slowing down the flow of conversation and lowering how enjoyable the chats feel. Many folks and even small business owners end up using WhatsApp as a simple, budget-friendly way to promote themselves[2] Ads have increasingly become a big part of people's group share, quietly slipping in

*Corresponding author : Ghosoon K. Munahy

Email addresses: ghosoonkareem@gmail.com

Communicated by 'sub etitor'

alongside regular conversation. Meanwhile, WhatsApp messages often seem to form their kind of language—almost like a sublanguage—with an informal twist that sets them apart[3] This style really stands apart from the stiff language you'd find in academic journals or mainstream media. It informal lexical structures, plenty of abbreviations, a few emojis, and even plays with words—skipping letters or jumbling them up here and there. Generally speaking, its laid-back, quick delivery makes things trickier for regular language tools to parse, meaning we need specially tuned models to catch all its quirks.

Iraqi Arabic is often seen as one of the trickiest dialects for automatic analysis, mainly because there are not enough dependable digital records or accessible data for researchers.[4]. The Iraqi dialect stands out—it pops with quirky vocabulary and a pronunciation twist you don't hear in other Arabic variants. Figuring out what someone feels can get tricky; those odd language habits blur the emotional signals, especially when you're just going by the spoken word. This distinct way of speaking often leaves room for a lot of guesswork about the speaker's proper mood.[5] Many WhatsApp groups are serving Iraqi communities, with increased advertisement integration into group chat, and some disruption might be occurring in the natural social interaction flow. The inherently informal and ever-evolving nature of the Iraqi vernacular makes automatic identification of such content extremely hard. Currently, no intelligent tools exist for spotting advertisements and measuring their effects in dialect-based instant messaging contexts.

The present research aims to develop a CAMELBERT-based classification model that effectively distinguishes between advertisements and social discussions in WhatsApp messages written in the Iraqi dialect. The study also attempts to determine the impact of advertisements on the flow and quality of group conversations. Another point is to compare how well transformer-based models perform against older machine learning models, like SVMs, to see how modern deep learning methods handle informal text written in dialects. We manually collected 5,000 WhatsApp messages from public Iraqi group chats to achieve the study's objectives. We evaluated each message and categorised it as a form of advertising or a social discussion. CAMELBERT, pre-trained on a large amount of Arabic text, was then fine-tuned using this dataset. The subsequent sections of this paper are organised as follows: Section 2 examines pertinent literature on advertisement detection and dialect-based natural language processing within instant messaging environments. Section 3 delineates the dataset, the preprocessing procedures, and the fine-tuning methodology employed for the CAMELBERT model. Section 4 delineates and analyses the experimental outcomes, encompassing performance metrics and classification instances. Section 5 concludes the study by encapsulating the principal findings and proposing avenues for future research.

2. Related work

Recent research projects in natural language processing (NLP) and instant messaging have evaluated different methods for comprehending and categorizing content inside platforms like WhatsApp. Generally, these can be split into three broad categories: natural language-based sentiment and emotion analysis, message classification through deep learning, and dialectal text processing. At the same time, every study above contributes valuable insights; several limitations still exist, particularly concerning informal dialects like Iraqi Arabic and domain-specific content such as advertisements. Hence, this section critically reviews most of the key studies in these three areas, looking at their methodologies, datasets, and reported outcomes, while at the same time pinpointing gaps to motivate the current investigation. Various studies consider sentiment and emotion classification in WhatsApp-based conversations using traditional machine-learning approaches such as Support Vector Machines (SVM), K-Nearest Neighbors (KNN), and rule-based methods. Roy and Das [6] investigated emotional states in student group chat using the valence-arousal-dominance (VAD) framework and applied SVM to classify basic emotions such as joy, sorrow, and rage. Likewise, Prajapati et al.[7] developed an SVM-based chat analyzer that uses sentiment-labelled messages to ascertain certain emotional tones, with a general tendency for the messages to be mostly neutral or positive. Although this approach provided interpretability and was computationally efficient, it relied entirely on handcrafted features. It thus faced the problem of processing colloquial or dialectal texts that may convey emotion implicitly, not explicitly. Using such models further loses accuracy in a dynamic real-time environment of a chat system where slang, abbreviations, and emojis are rife. Kumar and Naik [8] proposed a WhatsApp chat analyser that combines machine learning techniques with natural language processing to detect sentiment changes and communication patterns in group discussions. Their system employed text preprocessing, sentiment lexicons, and keyword extraction to categorise message content. Their model provided valuable insights into automated sentiment tracking but lacked contextual depth, especially regarding informal or dialect-based inputs prevalent in regional messaging groups

An automated classification framework for disaster-related messages was developed by Prasanna et al. [9] using the Linear Support Vector Classifier (Linear SVC) algorithm. Their model, trained on pre-labeled datasets, incorporated

a self-training mechanism to iteratively refine classification accuracy during real-time deployment. Designed to prioritize critical alerts in emergency scenarios, the system demonstrated that even relatively simple linear models like Linear SVC can achieve high effectiveness when appropriately trained. While originally applied to disaster management contexts, the methodology's scalability and adaptability suggest potential for broader applications, such as organizing region-specific WhatsApp messages into categories like advertisements or social discourse.

A sentiment classification system was created in [10] employing a hybrid ensemble methodology that integrated various machine learning models, such as Support Vector Machines, K-Nearest Neighbours, and Decision Trees. The methodology encompassed preprocessing WhatsApp messages, extracting pertinent features, and utilising a majority voting mechanism to ascertain the final sentiment classification. Python served as the implementation environment, and assessment was conducted utilising standard metrics including precision, recall, and regression-based indicators. The research illustrated that ensemble learning can deliver strong performance in the presence of noisy and unstructured data, especially in informal messaging contexts characterised by significant linguistic variability.

Deep learning models such as BERT, ALBERT, CNN, and Bidirectional LSTM have recently been used to boost classification performance on WhatsApp and platforms. De Moraes et al. [11] Used BERT and ALBERT for sub-language sentiment analysis and showed that transformer-based models sufficiently outperform other monolithic neural architectures like DCNN and LSTM in dealing with WhatsApp-specific linguistic structure, Susandri et al. [12] established a hybrid CNN-BiLSTM model that achieved 88% accuracy in sentiment classification and outperformed standalone CNN-BiLSTM models. These algorithms are preferable for handling contextual meaning and for handling noisy data.. However, most of the reported works focus on sentiment or emotion detection, while less attention has been paid to specific content types such as advertisements and dialectal Arabic or under-resourced languages. In addition, though these models show good performance, a landmark issue for dialect-specific applications remains the availability of sufficient training and tuning data.

There are different challenges while processing dialectal Arabic: one of them being inconsistencies in spelling, another the absence of a standard grammar and the unavailability of sufficiently large annotated datasets. According to studies by Sabri and Abdullah [4] and Ali et al. [5] Iraqi dialect processing is complicated because of phonetic and lexical variations. Some works have attempted detecting advertisements in regional languages such as Telugu [13] and other Arabic dialects, whereas the Iraqi Arabic language hardly got any attention in such a language-detection method. CAMELBERT by Obeid et al. [14] Can serve as an excellent starting point, featuring pretrained Arabic language models that may be adapted to dialects. However, to our knowledge, no studies have investigated CAMELBERT for Iraqi dialect WhatsApp message classification or the affected influence of embedded advertisements on group conversation dynamics in that setting. Such a gap drives our approach to transformer classification and sociolinguistic analysis.

3. Materials and Methodology:

3.1. Materials

3.1.1 Data set

About 5000 messages were pulled straight from buzzing Iraqi WhatsApp groups—a mix that shows the everyday flavour of how the Iraqi language pops up in real settings. Afterwards, we sorted these texts into two main groups, each spotlighting a different side of the conversation.

1. Advertisements (45%): Messages containing promotional offers or information about products and services.
2. Social discussions (55%): Messages expressing inquiries or discussions between members about various life issues.

3.1.2 CAMEL-BERT

CAMEL-BERT [14] Is a deep learning model based on Transformer Architecture, but it is not a model dedicated to classification directly, but instead a Language Representation Model model that was pre-trained on a large amount of Arabic texts in different dialects. CAMEL-BERT is used for multiple tasks within natural language processing, such as sentiment analysis, text classification, named entity extraction, question answering, and semantic analysis. To make CAMEL-BERT capable of classifying messages, Fine-Tuning must be performed, an additional training process in which the upper layers of the model are modified using pre-labelled data. In this process, the model's representational layers are used to transform texts into digital representations (Embeddings). A Classification Head layer is added on top of the model, such as the Fully Connected layer with Softmax to classify texts. After that, the model is retrained on labelled data, such as Iraqi dialect WhatsApp messages, using the Gradient Descent algorithm to update the weights and improve classification accuracy.

Compared to other models, CAMEL-BERT is distinguished by being a general model for understanding the Arabic language. It needs Fine-Tuning for specific tasks such as classification, unlike other models such as BERT for Classification that are directly tuned for this purpose, or traditional models such as SVM and Random Forest that rely on manual features instead of Transformer-based embeddings.

CAMEL-BERT isn't exactly built as a classifier off the bat – it's a language model that you can nudge into that role with a bit of fine-tuning. Trained on WhatsApp chats in Iraqi dialect, it usually sorts messages into ads, everyday chatter and neutral bits; overall, this makes it a quirky yet handy tool for digging into text conversations and catching the little cues in how people interact.

3.1.3 Transformer-Based Fine-Tuning

Fine-tuning models like BERT and GPT plays a crucial role when getting them ready for unique NLP challenges. Instead of overhauling everything, this step retrains a pre-initialized model on a labeled dataset—say, for text classification or even sentiment analysis—while, in most cases, keeping its original weight setup mostly intact. In our approach, nearly every sublayer in the encoder block gets kept as is; the only exception is the intermediate fully connected layers tucked inside the feed-forward sub-block, which we let adjust accordingly. This approach preserves the rich linguistic knowledge gained from unlabeled big data while allowing the model to adapt to the target classification task. This strategy also helps reduce the likelihood of overfitting when training on relatively small datasets and contributes to enhancing classification accuracy without the need to modify the entire model architecture or add complex new layers, such as adaptors.[15].

3.2 Methodology

Our research methodology comprises four phases, as illustrated in Fig. 1, which we will discuss in detail in the following paragraphs.

3.2.1 Data Processing

Several steps were applied to process the texts and improve the quality of the data:

1. Text cleaning

- Removing unnecessary information: Date, numbers, writings, and automatic messages that do not have analytical value, such as administrative notifications ("A new member has been added to the group"), were removed.
- Removing non-textual media: Images, audio clips, and recordings attached to conversations were ignored, as they can only be analysed within the scope of text analysis.

2. Word analysis (tokenisation) and rooting (root)

- CAMEL-BERT Tokeniser was used, a model dedicated to processing Arabic texts, to accurately understand the content of messages by dividing texts into independent words (Tokens) and processing linguistic structures in a way that suits the Arabic language and Iraqi colloquial.

3. Removing unimportant words (removing stop words)

- A dedicated list of the most specific words in the Iraqi dialect that do not have a precise meaning in text analysis was prepared, such as:
- Words of welcome and interaction: "عادي", "تمام", "شلونك", "هلا".
- Expressions that do not indicate the topic of discussion: "شكو ماكو", "على كيفك", "حسب الشخص".
- The final model has been completed, preventing its impact on data classification, including unifying the model on essential words with semantic significance.

4. Handling abbreviations and emojis (Handling abbreviations and emojis)

- Handling abbreviations: Except for some users using important abbreviations such as:

- "طبعاً" → "طبعين"
- "ان شاء الله" → "انشاءالله"
- "الحمد لله" → "الحمدلله"

These words have been unified for their correct form and accurate understanding of texts during classification.

- Writing emojis:

- Some common emojis have been replaced with words equivalent to positive ideas that we benefit from in the text, such as:

- 😊 → "Happy"
- 😡 → "Angry"
- ❤️ → "Love"
- 🤩 → "Enthusiasm"

- This procedure aims to improve sentiment analysis, as emojis represent a real contribution to users' expression and affect message classification.

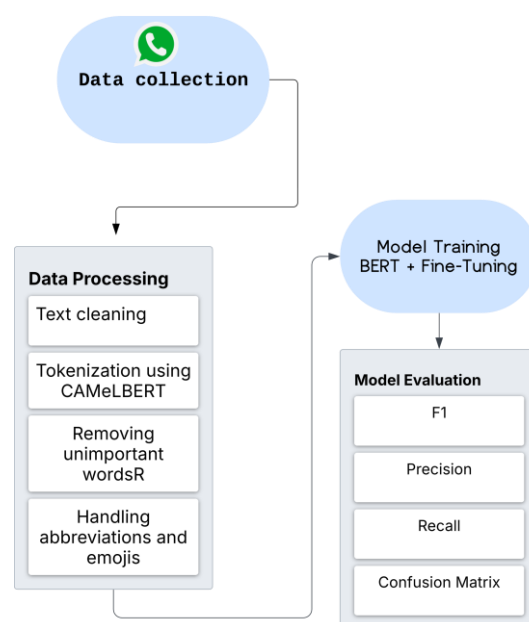


Fig 1. Model phases

3.2.2 Model Training

The model was prepared using BERT[16] Technology dedicated to Arabic texts, where it was fine-tuned and trained through Fine-Tuning on a dataset of 5000 WhatsApp messages in the Iraqi dialect. The CAMeL-BERT model, a robust model pre-trained on Arabic texts, was chosen to ensure accurate analysis of messages in the Iraqi dialect, and it was adapted to be able to distinguish between advertisements and social discussions through a customised list of keywords for each type of message.

We split our data into two parts – roughly 90% was set aside for training (giving the model plenty of examples to learn from various text patterns). In contrast, the remaining 10% was held back to later check how efficiently the model performs. Hyperparameters got a careful tweak: the training lasted for 50 epochs, which seems to hit that sweet spot between getting good results and steering clear of overfitting. We also fixed the batch size at 8 to balance accuracy against speed, and whenever a GPU was available, it helped speed things up even more. Then we gave CAMeL-BERT a bit of extra fine-tuning to make it work better with texts in the Iraqi dialect. In most cases, retraining the weights with fresh data helped the system pick up on different language quirks – think abbreviations and even emojis floating around in WhatsApp chats. The model's task was straightforward: decide whether a message is an advertisement (marked as 1) or part of a social discussion (0). Plus, to knock down errors further, we added a few extra tricks – Dropout Layers and some regularisation techniques helped adjust the weights more naturally and kept bias from creeping in.

The CAMeLBERT had to be adapted for the classification task by adding a custom classification head to the model. A single fully connected dense layer, followed by a Softmax layer, was added to the final hidden state output by the [CLS] token to output probabilities over the two classes: advertisements and social discussions. The classification head parameters were randomly initialised and trained in fine-tuning, while the transformer layers of CAMeLBERT were kept frozen to preserve the pre-trained linguistic knowledge. This way, computational cost was reduced, and given the limited dataset size, it was prevented from overfitting.

The classification head was trained for 50 epochs with a batch size of 8 and a learning rate $2e-5$. Early stopping with a patience of 5 was applied as the head converged on average already within the first 15 epochs, with the validation loss remaining steady thereafter to prevent overfitting. The training was performed on an NVIDIA Tesla T4 running on Google Colab, taking roughly 2 minutes per epoch, and about 1.5 hours in total.

4. Results and Discussion

The CAMeL-BERT model was fine-tuned on a dataset with around 5,000 Iraqi-dialect WhatsApp messages. Its performance was gauged on a held-out test set, with favourable classification metrics. For example, the

determination for advertisement messages yielded an F1-score of 97.0%, precision of 96.2%, and recall of 97.8%, as shown in Table 1. So, this shows that the model has potential as a discriminator between advertisements and casual social chatter.

From the confusion matrix analysis, the model identified correctly 232 advertisements, out of 250, with 18 advertisements identified erroneously as social discussions. Also, 236 social messages out of 250 were labelled correctly, with 14 errors labelled as advertisements. Apparently, these errors come from linguistic gradations such as ambiguous context, regional slang, or implicit cues buried in the casual phrasing.

Table 1: Performance Evaluation

Category	Precision	Recall	F1-Score
Advertisements	96.2%	97.8%	97.0%
Discussions	96.5%	97.9%	97.8%

Table 2 exemplifies the model's performance by showcasing instances of accurately classified and misclassified messages. These instances underscore the model's precision and the linguistic intricacies that impede dialectal text processing.

Table 2. Classification Examples

The Message	Actual classification	predicted classification
"عرض خاص لفترة محدودة! خصم 50% على جميع المنتجات"	advertisements	advertisements
"اكو احد يعرف الصيدلية فاتحة هسه"	social discussions	social discussions
"مساء الخير شلونكم شنو الاخبار"	social discussions	social discussions
"تنزيلات على ملايس الأطفال لفترة محدودة"	social discussions	social discussions
"اكو بيت للبيع بالمجمع"	social discussions	advertisements
"لخاطر عيونكم سوينا عرض بلاش بكج النضارة بخمسين الف يس"	advertisements	social discussions

The initial misclassified instance involved the phrase "اكو بيت للبيع بالمجمع", which was categorised as an advertisement, presumably because of the term "البيع", frequently found in commercial settings. Nonetheless, in the context of conversation, it may indicate a personal inquiry rather than a formal advertisement. In contrast, the second misclassification pertained to the phrase "عرض بلاش", which was misconstrued as a casual message despite its explicit promotional connotation. These instances demonstrate that the model occasionally depends excessively on specific keywords, failing to adequately encompass the pragmatic context of the communication. These findings highlight the necessity for a more profound contextual comprehension, potentially by incorporating supplementary linguistic or sociopragmatic indicators to more effectively differentiate between nuanced advertising techniques and informal discourse in dialectal WhatsApp interactions.

To evaluate the statistical robustness of the model's performance, numerous training and testing iterations were executed utilising stratified k-fold cross-validation. Throughout five iterations, the model consistently attained an average accuracy of 96.9% with a standard deviation of $\pm 0.42\%$, signifying stable and dependable performance. Confidence intervals were computed at the 95% level for each evaluation metric (precision, recall, F1-score), thereby reinforcing the dependability of the model's results.

A paired t-test was conducted to compare the CAMEL-BERT model with a baseline Support Vector Machine (SVM) classifier trained on the identical dataset. The findings indicated a statistically significant enhancement favouring the transformer-based methodology ($p < 0.01$), especially in the recall of advertising messages. This illustrates the model's enhanced capacity to identify nuanced advertising signals in informal dialectal settings.

Furthermore, the analysis of performance by class revealed no significant disparity in detection rates between advertisements and social messages, thereby diminishing the potential for bias. These statistical results validate both the model's efficacy and the methodological rigour of its development and evaluation process.

5. Conclusion

This research has shown the efficacy of fine-tuning CAMELBERT for classifying WhatsApp messages in the Iraqi dialect, specifically in differentiating between advertisements and authentic social interactions. Utilising a manually annotated dataset and modifying a transformer-based model for this dialectal context, we attained significant classification performance across various evaluation metrics. In addition to quantitative outcomes, the model's capacity to analyse informal, code-switched, and emoji-laden content highlights the promise of deep learning methodologies in handling low-resource dialects. Nevertheless, the study identified challenges, especially with messages displaying implicit promotional intent or humour, which were more susceptible to misclassification. These findings highlight the constraints of exclusively text-based models in addressing pragmatics and contextual cues, which are essential in dialectal Arabic. In future work, making the model smarter by adding features that consider context, using past conversations, or looking at different types of inputs (like images or links) could make it better at classifying messages. Furthermore, augmenting the dataset with a wider array of sources and assessing model efficacy across various dialects would enhance generalisability. The study establishes a preliminary framework for comprehending the impact of advertising content on group dynamics in informal digital communication, particularly in under-explored dialectal contexts.

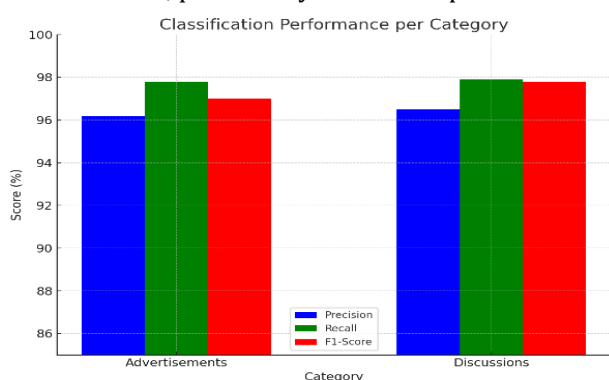


Fig 3. Classification Performance Chart

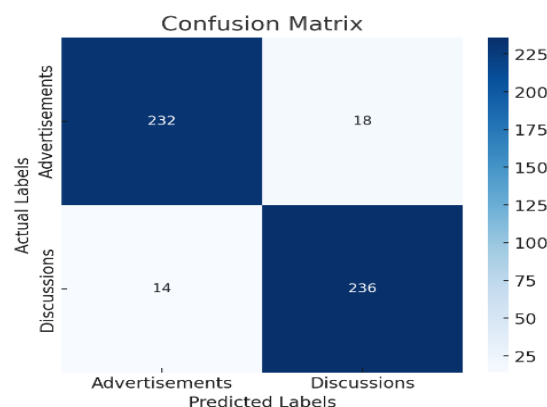


Fig 4. Confusion Matrix

References

- [1] I. A. Omipidan and B. O. Sanusi, "Rise of social media in the digital age: Whatsapp a threat to effective communication," *IMSU Journal of Communication Studies*, vol. 8, no. 1, pp. 142-153, 2024.
- [2] H. Mustafa *et al.*, "The Impact of Using WhatsApp Business (API) in Marketing for Small Business," in *2023 Tenth International Conference on Social Networks Analysis, Management and Security (SNAMS)*, 2023: IEEE, pp. 1-9.
- [3] S. Seljan, "Sublanguage in Machine Translation," in *Proceedings of 23rd Int. Convention MIPRO*, 2000.
- [4] R. F. Sabri and N. A. Abdullah, "Extremism Detection in the Iraqi Dialect Based on Machine Learning," *Iraqi Journal of Science*, 2025.
- [5] G. Ali, A. B. Hassanat, and A. S. Tarawneh, "Recognizing speech emotions in Iraqi dialect using machine learning techniques," in *2022 International Conference on Emerging Trends in Computing and Engineering Applications (ETCEA)*, 2022: IEEE, pp. 1-5.
- [6] B. Roy and S. Das, "Perceptible sentiment analysis of students' WhatsApp group chats in valence, arousal, and dominance space," *Social Network Analysis and Mining*, vol. 13, no. 1, p. 9, 2022.
- [7] P. Prajapati, R. Zaveri, and H. Shah, "Analysis and Prediction of the Sentiments of the WhatsApp Chat Using Sentiment Analysis," in *International Conference on Smart Computing and Communication*, 2024: Springer, pp. 261-271.
- [8] N. Kumar and S. Sonowal, "Email spam detection using machine learning algorithms," in *2020 Second International Conference on Inventive Research in Computing Applications (ICIRCA)*, 2020: IEEE, pp. 108-113.
- [9] N. Merrin Prasanna, S. Raja Mohan, K. Vishnu Vardhan Reddy, B. Sai Kumar, C. Guru Babu, and P. Priya, "Machine Learning Based Automated Disaster Message Classification System Using Linear SVC Algorithm," in *International Conference on Intelligent Cyber Physical Systems and Internet of Things*, 2022: Springer, pp. 869-879.
- [10] R. Kaushal and R. Chadha, "Hybrid Model for Sentiment Analysis of Whatsapp Data," in *2023 3rd International Conference on Advances in Computing, Communication, Embedded and Secure Systems (ACCESS)*, 2023: IEEE, pp. 215-220.
- [11] L. P. de Moraes, A. da Silva Soares, V. da CM Borges, N. F. F. da Silva, and F. S. Pereira, "Sub-language Sentiment Analysis in WhatsApp Domain with Deep Learning Approaches," *Revista de Sistemas de Informa  o da FSMA*, no. 31, pp. 32-47, 2023.
- [12] S. Susandri, S. Defit, and M. Tajuddin, "Enhancing text sentiment classification with hybrid CNN-BiLSTM model on WhatsApp group," *J. Adv. Inf. Technol.*, vol. 15, no. 3, pp. 355-363, 2024.
- [13] S. Kalakala, "A rule based sentiment analysis of whatsapp reviews in Telugu language," 2021.
- [14] O. Obeid *et al.*, "CAMEL tools: An open source python toolkit for Arabic natural language processing," in *Proceedings of the twelfth language resources and evaluation conference*, 2020, pp. 7022-7032.
- [15] M. A. Humayun, H. Yassin, J. Shuja, A. Alourani, and P. E. Abas, "A transformer fine-tuning strategy for text dialect identification," *Neural Computing and Applications*, vol. 35, no. 8, pp. 6115-6124, 2023.
- [16] S. Aftan and H. Shah, "A survey on bert and its applications," in *2023 20th Learning and Technology Conference (L&T)*, 2023: IEEE, pp. 161-166.