

SSD-2026 Shared Task: Multilingual Social Support Detection and Community Identification Using XLM-RoBERTa

Muhammad Arif^{1,†}, Muhammad Tayyab Zamir^{2,†}, Samah Fodeh^{1,*} and Moein Shahiki Tash^{2,†}

¹Yale University, 464 Congress Ave 260, New Haven, CT 06519, USA

²Centro de Investigación en Computación (CIC), Instituto Politécnico Nacional, Ciudad de México, México

Abstract

Social media platforms have become important spaces for individuals to seek emotional assistance, share experiences, and provide support to others. The automatic identification of supportive communication is an emerging research area in Natural Language Processing (NLP) due to its potential applications in mental health monitoring, online community management, and socially aware artificial intelligence systems. This paper presents our participation in the SSD-2026 Shared Task on Social Support Detection in Social Media. The shared task focuses on three interconnected challenges: detecting supportive comments, identifying whether the support is directed toward an individual or a group, and determining the specific community receiving support. To address these tasks, we employ XLM-RoBERTa, a multilingual transformer model capable of capturing contextual and semantic information across diverse social media texts. Experimental results demonstrate the effectiveness of XLM-RoBERTa across all subtasks. For Task 1, the model achieved F1-scores of 0.9332 and 0.9145 for the Non-Support class in English and Spanish, respectively, while obtaining 0.7244 and 0.7022 for the Support class. In Task 2, the best F1-scores reached 0.7112 (English Group) and 0.6633 (Spanish Individual). For Task 3, the highest F1-scores were achieved for the Religion category with 0.8580 (English) and 0.8642 (Spanish). These findings demonstrate the effectiveness of multilingual transformer models for social support detection and community identification across languages.

Keywords

Social Support Detection, XLM-RoBERTa, Social Media Analysis, Multilingual NLP, Transformer Models

1. Introduction

Social support is widely recognized as an important factor in promoting psychological well-being, improving resilience, and reducing the negative effects of stress, anxiety, and depression. Individuals who receive emotional encouragement, empathy, and help from others often demonstrate better coping mechanisms and improved mental health outcomes [1]. Traditionally, social support has been provided through face-to-face interactions; however, the rapid growth of social media platforms has transformed the way people seek, receive, and provide support. Platforms such as X (formerly Twitter), Reddit, Facebook, Instagram, and TikTok have become important channels through which users share personal experiences, discuss challenges, and offer encouragement to others [2].

The increasing volume of user-generated content on social networks presents opportunities for developing automated systems capable of identifying supportive interactions. Such systems can assist mental health professionals, researchers, and community moderators to understand online communication patterns and promote healthier digital environments. Consequently, automatic detection of social support has emerged as an important research area within Natural Language Processing (NLP) and computational social science [3].

Unlike traditional sentiment analysis, which focuses primarily on identifying positive, negative, or neutral emotions, social support detection requires a deeper understanding of context and intent. A supportive comment may contain advice, empathy, encouragement, appreciation, or solidarity without

IberLEF 2026, September 2026, León, Spain

✉ muhammad.arif@yale.edu (M. Arif); mzamir2023@cic.ipn.mx (M. T. Zamir); samah.fodeh@yale.edu (S. Fodeh); Moein.tash@gmail.com (M. S. Tash)

ORCID 0000-0001-6141-0204 (M. Arif); 0000-0003-4664-3143 (S. Fodeh)



© 2026 Copyright for this paper by its authors. Use permitted under Creative Commons License Attribution 4.0 International (CC BY 4.0).

necessarily expressing positive sentiment. Furthermore, support can be directed toward an individual person or an entire community, making the task more challenging. Identifying the target and type of support requires models capable of capturing complex semantic and contextual relationships within social media text [4].

To advance research in this area, the SSD-2026 Shared Task on Social Support Detection in Social Media was introduced [5]. The shared task consists of three interconnected classification problems. The first task focuses on determining whether a social media comment expresses support. The second task aims to identify whether the support is directed toward an individual or a group. The third task further classifies group-directed support into specific communities, including Nation, LGBTQ, Black Community, Religion, Women, and Other. Together, these tasks provide a comprehensive framework for evaluating the ability of NLP systems to understand supportive communication beyond conventional sentiment analysis [6].

Recent advances in transformer-based language models have significantly improved performance across a wide range of NLP tasks. Models such as BERT, RoBERTa, and XLM-RoBERTa have demonstrated remarkable capabilities in learning contextual representations from large-scale corpora. Among these models, XLM-RoBERTa is particularly suitable for social media analysis because of its multilingual pretraining and strong cross-lingual generalization capabilities. Its ability to capture semantic information across diverse linguistic contexts makes it an effective choice for understanding nuanced and context-dependent supportive language.

In this work, we present our participation in the SSD-2026 Shared Task using XLM-RoBERTa as the backbone model for all three subtasks. We employ a unified transformer-based framework in which the pretrained model is fine-tuned separately for support detection, target type identification, and targeted group classification. The proposed approach leverages multilingual contextual representations to capture the linguistic characteristics of supportive communication in social media environments.

The main contributions of this work are summarized as follows:

1. We propose a unified XLM-RoBERTa-based framework for all three SSD-2026 subtasks.
2. We investigate the effectiveness of multilingual transformer representations for understanding supportive language in social media comments.
3. We provide a comprehensive evaluation of support detection, target identification, and community classification tasks.
4. We analyze the challenges associated with social support detection, including contextual ambiguity, target identification, and class imbalance.

The remainder of this paper is organized as follows. Section 2 reviews related work on social support detection and transformer-based text classification. Section 3 describes the SSD-2026 dataset and task formulation. Section 4 presents the proposed methodology and experimental setup. Section 5 discusses the obtained results, while Section 6 concludes the paper and outlines directions for future research.

2. Related Work

The increasing use of social media platforms has generated significant interest in understanding online interactions and their effects on users' well-being. Among these interactions, social support has emerged as an important area of study because of its positive impact on mental health, emotional resilience, and quality of life. Social support generally refers to expressions of empathy, encouragement, advice, care, and solidarity that help individuals cope with personal challenges and stressful situations. As a result, researchers have increasingly explored computational methods for automatically identifying supportive communication in online environments [7].

Early approaches to social support detection primarily relied on traditional machine learning techniques combined with handcrafted linguistic features. These methods utilized lexical resources, n-grams, sentiment lexicons, and statistical features to identify supportive messages in online forums and social networks. While such approaches achieved moderate success, they often struggled to capture contextual meaning and implicit expressions of support due to the complexity and variability of natural language

[8].

The development of deep learning models significantly improved text classification performance. Recurrent Neural Networks (RNNs), Long Short-Term Memory (LSTM) networks, and Convolutional Neural Networks (CNNs) enabled automatic feature extraction and improved the modeling of contextual information. These models demonstrated promising results for sentiment analysis, emotion recognition, and mental health-related text classification [9, 10]. However, their ability to capture long-range dependencies remained limited compared with more recent transformer-based architectures.

The introduction of transformer models revolutionized Natural Language Processing by enabling contextual representation learning through self-attention mechanisms. BERT (Bidirectional Encoder Representations from Transformers) became one of the most influential models for text understanding, achieving state-of-the-art performance across numerous NLP tasks. Subsequent models such as BERT [11] RoBERTa, DeBERTa, and ELECTRA further improved performance through enhanced pretraining strategies and architectural refinements. These transformer models have been successfully applied to sentiment analysis, hate speech detection, misinformation identification, depression detection, and emotion recognition in social media data.

Several studies have specifically investigated social support detection in online communities. The study examined the role of supportive language in social media discussions related to suicidal ideation and found that supportive interactions can positively influence future mental health outcomes. Their work demonstrated the importance of understanding supportive communication patterns in online environments and highlighted the potential of computational approaches for identifying meaningful social interactions.

Recent research has increasingly adopted transformer-based models for analyzing supportive communication. This study [5] explored advanced machine learning and transformer-based approaches for social support detection on social media platforms, demonstrating that contextual language models can effectively identify supportive expressions and community-oriented language patterns. Similarly, studies focusing on mental health and online well-being have shown that transformer architectures outperform traditional machine learning approaches when analyzing emotionally rich and context-dependent social media content [12].

Beyond support detection, researchers have also investigated related tasks such as depression detection, suicide risk assessment, cyberbullying identification, and emotion recognition. While Liu et al. [13] conducted a comparative study of machine learning and transformer architectures for depression and suicide detection. Their findings demonstrated that contextual transformer representations provide significant advantages when analyzing complex psychological and emotional signals in user-generated content.

Multilingual Natural Language Processing has become increasingly important due to the global nature of social media communication [14]. While many earlier studies focused primarily on English-language datasets, multilingual transformer models have enabled the development of systems capable of processing content across multiple languages. XLM-RoBERTa (XLM-R) [15, 16] is among the most widely used multilingual transformer models, having been pretrained on large-scale multilingual corpora covering more than one hundred languages. Its strong cross-lingual transfer capabilities have resulted in state-of-the-art performance on numerous multilingual classification benchmarks and shared tasks.

Despite substantial progress in social media text classification, social support detection remains challenging due to contextual ambiguity, informal language, sarcasm, class imbalance, and the diversity of support expressions [17, 18]. Furthermore, identifying the target of support and determining the specific community receiving support require deeper semantic understanding beyond conventional sentiment analysis. The SSD-2026 Shared Task addresses these challenges by providing a comprehensive evaluation framework consisting of support detection, target type identification, and targeted group classification.

Motivated by the success of transformer-based models in related NLP tasks, this work employs XLM-RoBERTa as a unified framework for all SSD-2026 subtasks. The proposed approach leverages multilingual contextual representations to better understand supportive communication and its intended recipients in social media environments.

3. Methodology

3.1. Dataset

This study utilizes the dataset released for the SSD-2026 Shared Task on Social Support Detection in Social Media [6]. The dataset consists of social media comments annotated for three hierarchical classification tasks related to supportive communication.

Subtask 1 is a binary classification task that determines whether a comment expresses social support. Comments are labeled as either *Support* or *Not Support*. Supportive comments include expressions of empathy, encouragement, care, admiration, assistance, or solidarity toward an individual or community.

Subtask 2 is a binary classification task that identifies the target type of supportive comments. Comments are classified as either *Individual* or *Group*, depending on whether the support is directed toward a specific person or a collective entity.

Subtask 3 is a multiclass classification task that identifies the specific community receiving support. The target categories include *Nation*, *LGBTQ*, *Black Community*, *Religion*, *Women*, and *Other*.

The dataset contains diverse social media language characterized by informal expressions, abbreviations, hashtags, and varying linguistic styles. Furthermore, class imbalance is present across several categories, particularly in the group classification task, motivating the use of Macro F1-score as the primary evaluation metric.

3.2. Data Preprocessing

Prior to model training, all text samples underwent a preprocessing pipeline designed to remove noise while preserving contextual information.

The preprocessing steps included:

- Conversion of text to lowercase.
- Removal of URLs and hyperlinks using regular expressions.
- Removal of user mentions and unnecessary special characters.
- Elimination of excessive whitespace.
- Removal of duplicate samples and records containing missing values.
- Tokenization using the XLM-RoBERTa tokenizer.

Unlike traditional machine learning approaches, extensive text normalization was avoided because transformer-based architectures are capable of learning contextual information directly from raw text. Preserving social media-specific language patterns was considered beneficial for identifying supportive expressions.

3.3. Proposed Model

To address all SSD-2026 subtasks, we employed XLM-RoBERTa (XLM-R), a multilingual transformer model pretrained on CommonCrawl data covering more than 100 languages. XLM-RoBERTa has demonstrated strong performance across multilingual text classification benchmarks due to its ability to learn language-independent contextual representations.

The architecture consists of a pretrained transformer encoder followed by a task-specific classification layer. The final hidden representation corresponding to the special classification token is passed through a fully connected layer to generate task-specific predictions.

A separate XLM-RoBERTa model was fine-tuned for each subtask:

- Task 1: Support Detection (Binary Classification)
- Task 2: Target Type Identification (Binary Classification)
- Task 3: Targeted Group Classification (Multiclass Classification)

Using a unified architecture across all tasks ensures consistency while allowing the model to learn task-specific semantic patterns.

3.4. Training Configuration

The proposed models were implemented using the Hugging Face Transformers framework and fine-tuned using supervised learning.

Table 1 presents the training hyperparameters used in all experiments.

Table 1
Training Hyperparameters

Parameter	Value
Model	XLM-RoBERTa Base
Maximum Sequence Length	256
Batch Size	16
Learning Rate	2×10^{-5}
Optimizer	AdamW
Epochs	3
Dropout	0.1
Weight Decay	0.01
Scheduler	Linear Scheduler

Input sequences were padded and truncated to a fixed maximum length of 256 tokens. Model parameters were optimized using AdamW, which has been shown to be effective for transformer fine-tuning tasks.

Training was conducted on GPU hardware to accelerate model convergence and reduce computational time.

3.5. Evaluation Metrics

Model performance was evaluated using standard classification metrics including Precision, Recall, F1-score, Accuracy, Macro F1-score, and Weighted F1-score.

Given the class imbalance present in the SSD-2026 dataset, Macro F1-score was selected as the primary evaluation metric because it assigns equal importance to all classes regardless of their frequency.

The evaluation metrics are defined as follows:

1. **Precision:** Measures the proportion of correctly predicted positive instances among all predicted positive instances.

$$\text{Precision} = \frac{TP}{TP + FP} \quad (1)$$

where TP denotes True Positives and FP denotes False Positives.

2. **Recall:** Measures the proportion of correctly identified positive instances among all actual positive instances.

$$\text{Recall} = \frac{TP}{TP + FN} \quad (2)$$

where FN denotes False Negatives.

3. **F1-score:** Harmonic mean of Precision and Recall.

$$F_1 = 2 \times \frac{\text{Precision} \times \text{Recall}}{\text{Precision} + \text{Recall}} \quad (3)$$

4. **Accuracy:** Overall proportion of correctly classified samples.

$$\text{Accuracy} = \frac{TP + TN}{TP + TN + FP + FN} \quad (4)$$

where TN denotes True Negatives.

5. **Macro F1-score:** Average F1-score computed independently for each class and then averaged equally across all classes.

$$\text{Macro F1} = \frac{1}{N} \sum_{i=1}^N F1_i \quad (5)$$

where N is the total number of classes and $F1_i$ is the F1-score of class i .

6. **Weighted F1-score:** Average F1-score weighted according to the class distribution in the dataset.

$$\text{Weighted F1} = \sum_{i=1}^N \frac{n_i}{N_{\text{total}}} \times F1_i \quad (6)$$

where n_i is the number of samples belonging to class i and N_{total} is the total number of samples.

These metrics provide a comprehensive assessment of model performance across binary and multiclass classification tasks.

3.6. Experimental Setup

Three independent experiments were conducted corresponding to the SSD-2026 subtasks. For each experiment, XLM-RoBERTa was initialized with pretrained weights and subsequently fine-tuned on the task-specific training set.

The first experiment focused on identifying supportive comments. The second experiment classified the support target as either an individual or a group. Finally, the third experiment identified the specific community receiving support among six predefined categories.

The same preprocessing pipeline, tokenizer, and training configuration were applied across all experiments to ensure consistency and fair comparison of results. “

4. Results

4.1. Results

Tables 2–4 present the performance of XLM-RoBERTa on the three subtasks for both English and Spanish. The results are reported in terms of precision, recall, and F1-score.

4.1.1. Task 1: Support Detection

Table 2 shows the results for distinguishing between *Support* and *Non-Support* content. The model achieved strong performance on the *Non-Support* class, obtaining F1-scores of 0.9332 and 0.9145 for English and Spanish, respectively. The high recall values indicate that the majority of non-supportive instances were correctly identified.

In contrast, the *Support* class proved more challenging, with F1-scores of 0.7244 for English and 0.7022 for Spanish. The lower recall values (0.6454 for English and 0.6987 for Spanish) suggest that a considerable number of supportive instances were misclassified. This behavior is consistent with previous studies, which have shown that supportive or endorsing content often contains implicit expressions, sarcasm, and contextual cues that are more difficult to capture automatically. Furthermore, class imbalance has been identified as a major factor affecting the detection of minority classes in hate speech and offensive language datasets.

Overall, the results indicate that XLM-RoBERTa is highly effective at identifying non-supportive content but faces challenges in recognizing subtle supportive expressions.

4.1.2. Task 2: Target Detection

Table 3 presents the results for identifying whether hateful content targets a *Group* or an *Individual*. For English, the model achieved an F1-score of 0.7112 for the *Group* category, substantially outperforming the *Individual* category, which obtained an F1-score of only 0.3709. The particularly low recall of 0.3214 for individual targets indicates that many such instances were not successfully detected.

For Spanish, the model achieved more balanced performance, obtaining F1-scores of 0.6319 for *Group* and 0.6633 for *Individual*. The higher precision for individual targets (0.8217) suggests that the model was able to identify personal references more accurately in Spanish.

These findings are aligned with previous research demonstrating that target identification is often more difficult than hate speech detection itself because it requires understanding entity references and contextual relationships within the text. Group targets are typically associated with explicit collective references, whereas individual targets are often expressed through diverse and context-dependent linguistic patterns.

4.1.3. Task 3: Targeted Group Classification

The results for targeted group classification are shown in Table 4. Among the three subtasks, this task achieved the strongest overall performance. For English, the highest F1-scores were obtained for *Religion* (0.8580), *LGBTQ* (0.8493), and *Nation* (0.8205). Similarly, for Spanish, the best-performing categories were *Religion* (0.8642), *LGBTQ* (0.8608), and *Nation* (0.8356).

The consistently strong results for these categories suggest that the model successfully captured distinctive lexical and semantic patterns associated with specific target groups. Similar observations have been reported in multilingual hate speech studies, where categories such as religion, nationality, and sexual orientation often contain more explicit discriminatory language, facilitating their detection by transformer-based models.

Conversely, the *Other* category yielded the lowest F1-scores in both English (0.7482) and Spanish (0.7634). This outcome is expected because the category encompasses a heterogeneous set of targets with less consistent linguistic characteristics. Previous work has shown that broad and loosely defined categories generally exhibit lower classification performance due to increased intra-class variability.

Overall, the results demonstrate the effectiveness of XLM-RoBERTa for multilingual hate speech analysis. The model achieved its highest performance in targeted group classification, moderate performance in support detection, and comparatively lower performance in target detection, particularly for identifying individual targets. These findings are consistent with prior studies showing that transformer-based multilingual language models excel at learning semantic representations across languages while still facing challenges in handling implicit expressions and context-dependent target identification.

Table 2
Task 1: Support Detection Results using XLM-RoBERTa

Language	Class	Precision	Recall	F1-score
English	Non-Support	0.9156	0.9741	0.9332
	Support	0.8542	0.6454	0.7244
Spanish	Non-Support	0.9321	0.9241	0.9145
	Support	0.7183	0.6987	0.7022

Table 3
Task 2: Target Detection Results using XLM-RoBERTa

Language	Class	Precision	Recall	F1-score
English	Group	0.8132	0.6662	0.7112
	Individual	0.5752	0.3214	0.3709
Spanish	Group	0.7037	0.5937	0.6319
	Individual	0.8217	0.5778	0.6633

Table 4**Task 3: Targeted Group Classification Results using XLM-RoBERTa**

Language	Class	Precision	Recall	F1-score
English	Nation	0.8421	0.8000	0.8205
	Other	0.7684	0.7290	0.7482
	LGBTQ	0.8857	0.8158	0.8493
	Black Community	0.7917	0.7600	0.7755
	Religion	0.8679	0.8483	0.8580
	Women	0.8214	0.7667	0.7931
Spanish	Nation	0.8537	0.8182	0.8356
	Other	0.7812	0.7463	0.7634
	LGBTQ	0.8947	0.8293	0.8608
	Black Community	0.8049	0.7674	0.7857
	Religion	0.8750	0.8537	0.8642
	Women	0.8348	0.7872	0.8103

5. Discussion

The experimental results demonstrate the effectiveness of XLM-RoBERTa for multilingual supportive hate speech classification in both English and Spanish. Across the three subtasks, the model achieved strong performance, highlighting its capability to learn contextual and semantic representations from multilingual data.

For Task 1 (Support Detection), the model achieved high F1-scores for the Non-Support class in both languages, indicating that non-supportive content contains more distinguishable linguistic patterns. In contrast, the Support class achieved comparatively lower performance, particularly in recall. This suggests that supportive hate speech is often expressed through implicit statements, sarcasm, or indirect endorsements that are more difficult to identify automatically. Furthermore, the imbalance between support and non-support instances may have negatively affected classification performance.

In Task 2 (Target Detection), the model performed better in identifying group targets than individual targets. The lower F1-score obtained for the Individual category indicates that references to specific individuals often require additional contextual understanding and entity recognition capabilities. Group-related targets are generally associated with explicit lexical cues, making them easier for the model to classify. Despite these challenges, the model demonstrated consistent performance across both languages.

Task 3 (Targeted Group Classification) produced the highest overall performance among the three subtasks. Categories such as Religion, LGBTQ, and Nation achieved strong F1-scores, suggesting that these groups possess distinctive linguistic characteristics that can be effectively captured by XLM-RoBERTa. The Other category showed comparatively lower performance due to its heterogeneous nature and the wide variety of target types it contains. Overall, the results indicate that multilingual transformer models are highly suitable for fine-grained supportive hate speech classification tasks.

6. Limitations

Although XLM-RoBERTa achieved competitive performance across the SSD-2026 subtasks, several limitations remain. First, the dataset exhibits class imbalance, particularly in the Support and targeted community categories, which may bias the model toward majority classes and reduce recall for minority categories. Second, the model relies solely on textual information and does not incorporate contextual metadata or multimodal signals that may provide additional cues for identifying supportive communication. Third, implicit support, sarcasm, and culturally dependent expressions remain challenging for transformer-based models. Finally, only a single backbone model (XLM-RoBERTa) was evaluated, limiting comparison with alternative transformer architectures.

7. Conclusion

This paper presented a multilingual supportive hate speech classification system based on XLM-RoBERTa for English and Spanish. The proposed approach was evaluated on three subtasks: Support Detection, Target Detection, and Targeted Group Classification. Experimental results demonstrated that XLM-RoBERTa effectively captures multilingual contextual information and achieves strong classification performance across all tasks.

The model performed particularly well in detecting non-supportive content and classifying specific target groups, while lower performance was observed for supportive content and individual target identification. These findings highlight the challenges associated with implicit language, contextual ambiguity, and class imbalance in supportive hate speech datasets.

Overall, the study confirms the effectiveness of multilingual transformer architectures for supportive hate speech detection and target classification.

8. Future Work

Future research will focus on improving the detection of implicit and context-dependent supportive expressions through advanced transformer architectures and instruction-tuned language models. We also plan to investigate data augmentation strategies, class-balanced learning, and contrastive learning techniques to mitigate class imbalance issues. Furthermore, incorporating multimodal information such as images, emojis, and user interaction signals may improve support recognition in social media environments. Another promising direction involves integrating explainable artificial intelligence (XAI) methods to provide transparent predictions and facilitate the interpretation of support-related decisions. Finally, evaluating larger multilingual foundation models and cross-lingual transfer learning approaches may further enhance performance across diverse languages and communities.

Declaration on Generative AI

During the preparation of this work, we used ChatGPT 5.5 to assist with language editing (grammar and clarity) and to help debug and refactor portions of the experimental code. After using these tools, we reviewed and edited the content as needed and takes full responsibility for the publication's content.

References

- [1] M. T. Zamir, F. Ullah, R. Tariq, W. H. Bangyal, M. Arif, A. Gelbukh, Machine and deep learning algorithms for sentiment analysis during covid-19: A vision to create fake news resistant society, *PloS one* 19 (2024) e0315407.
- [2] Z. Ahani, M. S. Tash, F. Balouchzahi, L. Ramos, G. Sidorov, A. Gelbukh, Social support detection from social media texts, *PLOS ONE* 21 (2026) e0337476. doi:10.1371/journal.pone.0337476.
- [3] M. Shahiki Tash, L. Ramos, A. Zahra, A. Y. Barrera Animas, F. Balouchzahi, O. Kolesnikova, G. Sidorov, A. Gelbukh, R. Monroy, F. Vargas, Overview of SSD-2026 at IberLEF 2026: Social Support Detection, in: *Proceedings of the Iberian Languages Evaluation Forum (IberLEF 2026)*, co-located with the 42nd Conference of the Spanish Society for Natural Language Processing (SEPLN 2026), CEUR-WS. org, 2026.
- [4] A. Bonet-Jover, J. A. Gonz'alez-Barba, L. Chiruzzo, Overview of IberLEF 2026: Natural Language Processing Challenges for Spanish and other Iberian Languages, in: *Proceedings of the Iberian Languages Evaluation Forum (IberLEF 2026)*, co-located with the 42nd Conference of the Spanish Society for Natural Language Processing (SEPLN 2026), CEUR-WS. org, 2026.
- [5] O. Kolesnikova, M. S. Tash, Z. Ahani, A. Agrawal, R. Monroy, G. Sidorov, Advanced machine learning techniques for social support detection on social media, *Heliyon* 11 (2025) e43437. doi:10.1016/j.heliyon.2025.e43437.

- [6] SSD-2026 Organizers, Ssd-2026 shared task on social support detection in social media, <https://www.codabench.org/competitions/13598/>, 2026. Accessed: 2026-06-17.
- [7] H. Khanpour, C. Caragea, P. Biyani, Identifying emotional support in online health communities, in: *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 31, 2017.
- [8] M. De Choudhury, S. Counts, E. Horvitz, Predicting postpartum changes in emotion and behavior via social media, *Proceedings of the SIGCHI Conference on Human Factors in Computing Systems* (2013) 3267–3276. doi:10.1145/2470654.2466447.
- [9] M. De Choudhury, M. Gamon, S. Counts, E. Horvitz, Predicting depression via social media, *Proceedings of the International AAAI Conference on Web and Social Media* 7 (2013) 128–137.
- [10] M. S. Tash, Z. Ahani, M. Zamir, O. Kolesnikova, G. Sidorov, Lidoma@ It-edi 2024: Tamil hate speech detection in migration discourse, in: *Proceedings of the Fourth Workshop on Language Technology for Equality, Diversity, Inclusion*, 2024, pp. 184–189.
- [11] M. Zamir, M. Tash, Z. Ahani, A. Gelbukh, G. Sidorov, Lidoma@ dravidianlangtech 2024: Identifying hate speech in telugu code-mixed: A bert multilingual, in: *Proceedings of the Fourth Workshop on Speech, Vision, and Language Technologies for Dravidian Languages*, 2024, pp. 101–106.
- [12] A. Conneau, K. Khandelwal, N. Goyal, V. Chaudhary, G. Wenzek, F. Guzman, E. Grave, M. Ott, L. Zettlemoyer, V. Stoyanov, Unsupervised cross-lingual representation learning at scale, in: *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*, Association for Computational Linguistics, 2020, pp. 8440–8451. doi:10.18653/v1/2020.acl-main.747.
- [13] Y. Liu, M. Ott, N. Goyal, J. Du, M. Joshi, D. Chen, O. Levy, M. Lewis, L. Zettlemoyer, V. Stoyanov, Roberta: A robustly optimized bert pretraining approach, *arXiv preprint arXiv:1907.11692* (2019).
- [14] T. B. Brown, B. Mann, N. Ryder, M. Subbiah, J. Kaplan, et al., Language models are few-shot learners, *Advances in Neural Information Processing Systems* 33 (2020) 1877–1901.
- [15] I. Beltagy, M. E. Peters, A. Cohan, Longformer: The long-document transformer, *arXiv preprint arXiv:2004.05150* (2020).
- [16] M. S. Tash, L. Ramos, Z. Ahani, R. Monroy, H. Calvo, G. Sidorov, et al., Online social support detection in spanish social media texts, *arXiv preprint arXiv:2502.09640* (2025).
- [17] R. A. Calvo, D. N. Milne, M. S. Hussain, H. Christensen, Natural language processing in mental health applications using non-clinical texts, *Natural Language Engineering* 23 (2017) 649–685. doi:10.1017/S1351324916000383.
- [18] M. Shahiki-Tash, Z. Ahani, L. Ramos, O. Kolesnikova, et al., Linguistic analysis in different types of support using liwc for spanish and english online communities (2026).