

Overview of Shared Task on Caste and Migration Hate Speech Detection

Saranya Rajiakodi¹, Bharathi Raja Chakravarthi², Rahul Ponnusamy³,
Prasanna Kumar Kumaresan³, Sathiyaraj Thangasamy⁴, Bhuvaneswari Sivagnanam¹,
Charmathi Rajkumar⁵

¹Central University of Tamil Nadu, India

²School of Computer Science, University of Galway, Ireland

³Data Science Institute, University of Galway, Ireland

⁴Department of Tamil, Sri Krishna Adithya College of Arts and Science, Tamil Nadu, India

⁵The American College, Madurai, Tamil Nadu, India

Abstract

We present an overview of the first shared task on "Caste and Migration Hate Speech Detection." The shared task is organized as part of LT-EDI@EACL 2024. The system must delineate between binary outcomes, ascertaining whether the text is categorized as a caste/migration hate speech or not. The dataset presented in this shared task is in Tamil, which is one of the under-resource languages. There are a total of 51 teams participated in this task. Among them, 15 teams submitted their research results for the task. To the best of our knowledge, this is the first time the shared task has been conducted on textual hate speech detection concerning caste and migration. In this study, we have conducted a systematic analysis and detailed presentation of all the contributions of the participants as well as the statistics of the dataset, which is the social media comments in Tamil language to detect hate speech. It also further goes into the details of a comprehensive analysis of the participants' methodology and their findings.

1 Introduction

Social media platforms have become an integral part of the daily life of today's society. It gives new meaning to how we communicate, connect, and share information among people (Kumaresan et al., 2023). This digital landscape allows users to share their opinions, how they live, and their work worldwide. It has a huge impact on society, which makes it possible to influence everything around us. It encompasses the way of communication, how we perceive the world, disaster response, education, and how it makes information accessible to everyone, as well as giving our support to the things or people that are overlooked by society (Ponnusamy et al., 2023; Chakravarthi, 2023).

However, social media platforms also present significant challenges to people, like discrimination and bias against certain kinds of people, including those based on caste/migration. It replicates the societal fractures openly or anonymously, enabling harassment, bullying, and exclusion based on caste/migration. The caste system also significantly influences instances of homicide and violence targeting inter-caste marriages (Sathi, 2023). Caste systems are in a hierarchical order from high to low. Here, low-level caste people have been subjected to many kinds of discrimination in many firms (Goraya, 2023). Racial and extractive capitalism has exploited the people into hard labor, their dehumanizing treatment, and their psychological trauma experiences (Dulhunty, 2023).

With the rapid rise of social media (Lande et al., 2023; Priyadharshini et al., 2022), the shadows of age-old social prejudice, such as caste-based discrimination, bias, and hate speech, stepped into the digital realm, which especially targeted lower-caste people, thus leading to their psychological trauma and trampling on the dignity of humans. There are some cases where caste/migration-based discrimination leads to violence. This leads to a way of effectively addressing and detecting hate speech regarding caste/migration. With the help of machine learning, deep learning, and natural language processing technology, one can detect caste/migration-based hate speech. However, the quality of the detection depends on the quality of the training data without any biases and critical concerns. The task involves the use of a newly created dataset for detecting caste or migration-based hate speech in the Tamil language.

In this overview paper, we have discussed the shared task of caste and migration hate speech detection in Tamil at LT-EDI@EACL 2024. In the subsequent sections, we have discussed the task description, dataset statistics, methodologies used by the participants to detect caste/migration hate speech in Tamil, and their results and ranking.

2 Related work

In recent years, social media has developed rapidly, leading to advantages and disadvantages in its use. The common one is hate speech towards a certain kind of race, religion, sexual orientation, gender, caste, and beyond. This issue leads many researchers to research to research a methodology that can detect hate speech related to specific targets (Ibrohim and Budi, 2023). In particular, (Ryzhova et al., 2022) presents XLM RoBERTa as a base model to detect religion-based hate speech on English, Russian, and Hindi datasets. Afterward, the base model was finetuned with the different datasets, and to improve the model, a text attack algorithm was applied. (Parvareh, 2023) focuses on hate speech towards Afghan immigrants in Iran and also shows the subtle ways of spreading hate speech without directly using any hateful words. (Breazu, 2023) focuses on the hate speech comments on YouTube regarding the Roma community, and it also points out how the obvious and subtle way of hate comments and discrimination spread against the Roma community. It also highlighted the term “entitlement racism,” which is when individuals believe it’s justifiable to propagate racial animosity. (Nave and Lane, 2023) highlights how online platforms lack the detailed guidelines to successfully combat online hate speech when incorporating HRDD (Human Rights Due Diligence) into their terms of service (TOS). In order to protect marginalized people, it advocates online platforms, including the TOS, with European human rights norms.

3 Task Description

The primary goal of this task is to develop an automated classification system that can accurately determine the presence of hate speech on caste or migration within textual content on social media platforms. This shared task is held through the Codalab competition¹. The participants will pro-

vided with training, development, and test datasets in the Tamil language. To access and contribute to the data, navigate to Codalab and select the "Participate" tab. To the best of our knowledge, this represents the first shared task dedicated to identifying hate speech concerning caste/migration. The annotation has been done in two categories.

- **Hate speech on caste/migration:** Comments that contain a variety of harmful, derogatory, discriminatory content as well as mocking and ridicule, Delegitimization content aimed at certain caste or migrated people.
- **Not Hate speech on caste/migration:** Comments that do not contain text aimed at any caste or migrated people

Sample for Hate speech on caste:

நி தமிழனே கிடையாது.
வடுக பள்ளன் னு செப்பேடு
சொல்லுது. நி எப்படி டா
வேளாளர் ஆக முடியும்

Figure 1: Tamil Language- Hate speech on caste

English Translation for Figure 1: You are not Tamizhian. Seppedu says Vaduka Pallan. How can you become a vaellalar?

- **Tamil-English:** Nee yaen mukkulathor aah pirika ninaikura
- **English Translation for Above text:** Why do you want to separate the MUkkulathor?

Sample for Hate speech on migration:

டேய் வடக்கன அடிச்சு
விரட்டுங்கட நா பறையன்
தாண்ட இருந்தாலும் தமிழ்
இனத்தை சேர்தவண்டா
சொன்னா கேளுங்க டா அவ
நமக்குள்ள வேண்டா தவவு
செஞ்சு

Figure 2: Tamil Language- Hate speech on migration

English Translation for Figure 2: Banish the North Indians. Na Parayan Thanda. However, I

¹<https://codalab.lisn.upsaclay.fr/competitions/16089>

am from the Tamil race. Listen, sir. Don't involve them with us.

- **Tamil-English:** polaikavantha marvadikku yavvalavu thimuru. Nam thamar aachithan enime.
- **English Translation for Above text:** How arrogant is the Marvadi who came to work. Nam thamar aachithan enime.

4 Dataset Description

This dataset was methodically collected from a social media platform with a concentrated focus on the YouTube platform by utilizing the YouTube-comment-scraper tool to collect the YouTube video comments in support of the shared task on 'Caste/Migration Hate Speech Detection' at LT-EDI@EACL 2024. It represents the first dataset focusing on caste and migration hate speech in low-resource languages, with a particular emphasis on Tamil. A total of 7,875 comments were collected, and each of the comments was meticulously annotated for the presence of caste/migration hate speech (labeled as '1') and the absence of such hate speech (labeled as '0').

A few samples of the dataset have been discussed in the previous section. The samples contain some words that point at some caste or migrated people such as *vadakan*, *marvadi*, *parayan*, *vaduka pallan*, *vaellalar*, *mukkulathor*. The offensive texts that may come along with the above words or beyond may have targeted a certain community. Since there are many castes, some tend to discriminate against other castes which leads to these kinds of offensive comments.

The dataset was segmented into training, development, and test data subsets to help with the thorough analysis and to facilitate the model training. The accompanying Table 1 provides the detailed distribution of comments across these subsets, providing the dataset's structure and composition.

4.1 Training Phase:

Initially, participants were provided with both training and validation data for caste/migration hate speech detection model development. They could run preliminary evaluations and fine-tune the model settings. There are 51 teams participated and accessed the data.

4.2 Evaluation Phase:

The second phase involved releasing test sets in Tamil for system evaluation. Participating teams submitted their predicted results for assessment through Google Forms. The submission will be evaluated with macro average F1-score. The results should be submitted on the google form in the form of zip.

5 Participant's Methodology

- **BITS_GraphAI:** This team used two pre-trained transformers, TamilSBERT-STS and Indic-SBERT(Deode et al., 2023)(Mirashi et al., 2024), and fine-tuned them on the given data with triplet and cosine similarity loss, respectively. In triplet loss, triples are of the form (anchor, positive, negative), where the anchor and positive sentences are of the same class, whereas negatives are from other classes. Triplet loss helps the model to discern nuanced relationships. To further enhance classification, we used TextGCN architecture(Yao et al., 2019) by feeding the fine-tuned sBERT embeddings as feature input and text graph as an adjacency matrix. All three models performed well, showing slightly better results from sBERT-enhanced TextGCN graph-based learning. This team achieved rank 4 with a macro F1 score of 0.77.
- **SSN-Nova (Reddy et al., 2024):** This team delves into an array of boosting techniques, encompassing Adaboost, XGBoost,(Demir and Sahin, 2023) and a comparative analysis with a voting classifier that aggregates multiple traditional models. This methodology provides a comprehensive exploration of ensemble methods, leveraging the strengths of boosting algorithms and traditional models to enhance the overall predictive capabilities of their system. They secured rank 12 with the macro F1 score of 0.59
- **Kubapok (Pokrywka and Jassem, 2024):** The Team employed a systematic approach in their endeavor, utilizing solely the data provided by the organizers. Their methodology centered around employing various models, including 'l3cube/pune-kannada'(Deode et al., 2023)(Mirashi et al., 2024), 'microsoft-mdeberta-v3'(He et al., 2020) (utilized twice), and 'xlm-roberta.' These models were trained

Sets	Caste/Migration Hate Speech	Not	Total
Train	2,052	3,303	5,355
Development	351	594	945
Test	602	973	1,575
Total			7,875

Table 1: Dataset Description

using standard Hugging Face scripts for text classification, with adjustments such as a warm-up ratio of 0.1 and 30 epochs. Notably, they aggregated both training and development data, creating fresh random splits for each model. The selection of the optimal epoch checkpoint was based on the development F1 score. A key aspect of their strategy involved averaging the probabilities generated by all four models, with a threshold of 0.5 for class selection. This team achieved rank 2 with a macro F1 score of 0.81.

- **KEC_AI_DSNLP:** (Shanmugavadivel et al., 2024) This team created a machine learning model such as KNN, Decision trees and Naive Baiyes to classify the hate speech text and got 0.65 macro F1 and secured the rank 9.
- **CUET_NLP_GoodFellows:** This team used two BERT models to accomplish their task. They are mBERT and XLM-R, and both of these models are finetuned. Along with these, the team also used fine-tuned random classifier.
- **selam:** This team employed a Support Vector Machine (SVM) approach within the realm of Natural Language Processing (NLP). The primary goal was to develop a robust text classification system capable of predicting whether a given text contains caste/migration hate speech and scored macro F1 of 0.62 and secured rank 10.
- **KEC_DL_KSK:** Team employed the sampling methods such as SMOTE and random oversampler to balance the datasets. They have used machine learning algorithms, namely, Random Forest, Support Vector Machine, Naive Bayes, Logistic Regression, and Decision Tree, along with word embedding techniques like TF-IDF, Word2Vec, Doc2Vec, and FastText. the got the 0.49 macro f1 score.

- **bytesizedllm:** This team has utilized the embeddings generated from a subset of AI4Bharat’s data, encompassing 100,000 randomized lines. These embeddings were created using their custom-built subword tokenizers for Telugu (with a size of 7.6 MB) and Tamil (with a size of 1.3 MB) languages. They employed a Bidirectional Long Short-Term Memory (BiLSTM) classifier to perform classification tasks. The model was trained on labeled datasets and scored 0.61 macro F1 score.

- **Word Wizards:** This team utilized Labse(Pei et al., 2022), a pre-trained language representation model specifically designed for understanding the text in multiple languages. Labse employs a Siamese encoder architecture, capable of generating high-quality sentence embeddings by encoding text into a fixed-size vector representation. Following the encoding process, a K-Nearest Neighbors (KNN) model was implemented to perform various tasks, likely involving similarity searches or classification based on the encoded representations. KNN is a simple yet effective algorithm used for classification and regression tasks, particularly in scenarios where data points are mapped in a high-dimensional space. This combination likely facilitated tasks involving semantic similarity, clustering, or classification of text data based on the learned representations from Labse embeddings. They achieved rank 13 with the macro F1 score of 0.54.

- **Lidoma:** (Tash et al., 2024) This team employed deep learning and machine learning models like convolutional neural networks and support vector machine algorithms to classify hate speech detection on caste/migration in the Tamil language and secured that 6th rank with the macro F1 score of 0.76.

- **Transformers:** (Singhal and Bedi, 2024) This

Teams	Macro F1	Rank
Transformers_run3 (Singhal and Bedi, 2024)	0.82	1
kubapok_run1 (Pokrywka and Jassem, 2024)	0.81	2
CUET_NLP_Manning_run3 (Alam et al., 2024)	0.80	3
BITS_Graph4NLP_run1	0.77	4
Algorithmalliance_run1 (Sangeetham et al., 2024)	0.76	5
lidoma_run2 (Tash et al., 2024)	0.76	6
CUET_NLP_GoodFellows_run2	0.75	7
quartet_run1 (H et al., 2024)	0.73	8
KEC_AI_DSNLP_run2 (Shanmugavadivel et al., 2024)	0.65	9
selam_run1	0.62	10
byteSizedllm_run1	0.61	11
SSN-nova_run3 (Reddy et al., 2024)	0.59	12
WordWizards_tamil_run1	0.54	13
KEC_DL_KSK_run2	0.49	14
Habesha_run1	0.38	15

Table 2: Rank List Based on Average macro F1 Score

team utilized an ensemble model that comprises XLMroberta, a multilingual bert base model, and muril cased model(Subramanian et al., 2022). The accuracy of these models was highest compared to all the other models tested. A combination of these models improved the overall accuracy. All the models were trained on the text without cleaning since the performance of all the models suffered after cleaning.

- **CUET_NLP_Manning** (Alam et al., 2024): This team employed six machine learning(LR, SVM, SGD, XGB, ENSEMBLE, RF) models, 3 deep learning models(BiLSTM, Attention, and BiLSTM-CNN) and three transformer-based models (M-BERT, XLM-R, and Tamil-BERT) with TF-IDF and fasttext embeddings. Among the models, the transformer-based model yields better results with the highest evaluation score. This team achieved rank 3 with a score of 0.80.
- **Habesha**:This Team utilized a model built upon BERT transformers for creating embeddings. Additionally, They integrated LSTM (Long Short-Term Memory) deep learning techniques to facilitate classification tasks. This combined architecture allows for the effective representation of input data through transformer-based embeddings while leveraging the sequential learning capabilities of LSTM for accurate classification and scoring

0.38 macro F1.

- **ALGORITHM ALLIANCE** (Sangeetham et al., 2024): This team has applied several supervised machine learning algorithms such as Logistic Regression, Support Vector Machine(SVM), Multi-Layer Perceptron(MLP), Random Forest Classifier(RFC), Decision Tree, KNN as their classification models to the highest accuracy. Among these, SVM yielded the highest scores which is 0.76 macro F1 score, and secured rank 5.
- **Quartet** (H et al., 2024): This team used machine learning models such as Logistic Regression, Support Vector Machine (SVM), Random Forest Classifier, Decision Tree Classifier, and Naive Bayes for classification and TF-IDF for feature representation. Support Vector Machine yielded a better accuracy with a 0.73 macro F1 score and ranked 8th.

6 Results

The hate speech detection shared task for caste/migrants was conducted for one of the low-resource Languages that is Tamil. As mentioned in the former part, many participants have contributed to the shared task. A total of 51 teams participated in the shared task of detecting hate speech on caste/migration in the Tamil Language. Among them, 15 teams submitted their results. The ranking and Evaluation of the shared tasks was based on the

average macro F1 score. Table 2 shows the rankings of the teams that participated in the task. Here we have accentuated the top three teams that participated in the shared task and got the top rankings. The team "Transformers"(Singhal and Bedi, 2024) ranked first in the shared task with the macro F1 score of 0.82 using the ensemble methods combining various transformers-based model. "kubapok" Team ranked second among the participants with the macro F1 score of 0.81 which has been resulted from utilizing the microsoft-mdeberta-v3, and xlm-roberta. CUET_NLP_Manning(Alam et al., 2024) ranked third with the macro F1 score of 0.80 by using machine learning, deep learning, and transformer-based models like mBERT, Tamil BERT, xlm- R models.

7 Conclusion

We presented the overview of the shared task on caste and migration hate speech in social media comments using a dataset in Tamil. It represents an important step toward the creation of healthy online communities. We can mitigate hate speech on certain individuals, castes, and migration by exploiting advanced technologies, algorithms, and computational tools.

8 Ethical Considerations

In conducting this study, which involves the utilization of YouTube comments, we have taken into account many ethical implications while collecting the YouTube comments we made sure that the privacy of commentators was well protected and that the comments were not used in a way that could have caused any harm. In addition, we ensured that data distribution is prohibited for anything but academic and non-commercial research uses Thus we have done our research ethically and responsibly to reduce harm, protect privacy, and make meaningful contributions to the field of study.

Acknowledgements

This work was conducted with the financial support of the Science Foundation Ireland Centre for Research Training in Artificial Intelligence under Grant No. 18/CRT/6223, supported in part by a research grant from Science Foundation Ireland (SFI) under Grant Number SFI/12/RC/2289_P2(Insight_2).

References

- Md Ashraful Alam, Hasan Mesbaul Ali Taher, Jawad Hossain, Shawly Ahsan, and Mohammed Moshuiul Hoque. 2024. CUET_NLP_Manning@LT-EDI 2024: Transformer-based Approach on Caste and Migration Hate Speech Detection. In *Proceedings of the Fourth Workshop on Language Technology for Equality, Diversity and Inclusion*, Malta. European Chapter of the Association for Computational Linguistics.
- Petre Breazu. 2023. Entitlement racism on youtube: White injury—the licence to humiliate roma migrants in the uk. *Discourse, Context & Media*, 55:100718.
- Bharathi Raja Chakravarthi. 2023. Detection of homophobia and transphobia in youtube comments. *International Journal of Data Science and Analytics*, pages 1–20.
- Selçuk Demir and Emrehan Kutlug Sahin. 2023. An investigation of feature selection methods for soil liquefaction prediction based on tree-based ensemble algorithms using adaboost, gradient boosting, and xgboost. *Neural Computing and Applications*, 35(4):3173–3190.
- Samruddhi Deode, Janhavi Gadre, Aditi Kajale, Ananya Joshi, and Raviraj Joshi. 2023. L3cube-indicsbert: A simple approach for learning cross-lingual sentence representations using multilingual bert. *arXiv preprint arXiv:2304.11434*.
- Annabel Dulhunty. 2023. When extractive and racial capitalism combine—indigenous and caste based struggles with land, labour and law in india. *Geoforum*, 147:103887.
- Sampreet Singh Goraya. 2023. How does caste affect entrepreneurship? birth versus worth. *Journal of Monetary Economics*, 135:116–133.
- Shaun Allan H, Samyuktaa Sivakumar, Rohan R, Nikilesh Jayaguptha, and Durairaj Thenmozhi. 2024. Quartet@LT-EDI 2024: A Support Vector Machine Approach For Caste and Migration Hate Speech Detection. In *Proceedings of the Fourth Workshop on Language Technology for Equality, Diversity and Inclusion*, Malta. European Chapter of the Association for Computational Linguistics.
- Pengcheng He, Xiaodong Liu, Jianfeng Gao, and Weizhu Chen. 2020. Deberta: Decoding-enhanced bert with disentangled attention. *arXiv preprint arXiv:2006.03654*.
- Muhammad Okky Ibrohim and Indra Budi. 2023. Hate speech and abusive language detection in indonesian social media: Progress and challenges. *Heliyon*.
- Prasanna Kumar Kumaresan, Kishore Kumar Pon-nusamy, Kogilavani S V, Subalalitha Cn, Ruba Priyadharshini, and Bharathi Raja Chakravarthi. 2023. VEL@LT-EDI: Detecting homophobia and transphobia in code-mixed Spanish social media comments. In *Proceedings of the Third Workshop on*

- Language Technology for Equality, Diversity and Inclusion*, pages 233–238, Varna, Bulgaria. INCOMA Ltd., Shoumen, Bulgaria.
- Kaustubh Lande, Rahul Ponnusamy, Prasanna Kumar Kumaresan, and Bharathi Raja Chakravarthi. 2023. [KaustubhSharedTask@LT-EDI 2023: Homophobia-transphobia detection in social media comments with NLPAUG-driven data augmentation](#). In *Proceedings of the Third Workshop on Language Technology for Equality, Diversity and Inclusion*, pages 71–77, Varna, Bulgaria. INCOMA Ltd., Shoumen, Bulgaria.
- Aishwarya Mirashi, Srushti Sonavane, Purva Lingayat, Tejas Padhiyar, and Raviraj Joshi. 2024. L3cube-indicnews: News-based short text and long document classification datasets in indic languages. *arXiv preprint arXiv:2401.02254*.
- Eva Nave and Lottie Lane. 2023. Countering online hate speech: How does human rights due diligence impact terms of service? *Computer Law & Security Review*, 51:105884.
- Vahid Parvaresh. 2023. Covertly communicated hate speech: A corpus-assisted pragmatic study. *Journal of Pragmatics*, 205:63–77.
- Yijie Pei, Siqi Chen, Zunwang Ke, Wushour Silamu, and Qinglang Guo. 2022. Ab-labse: Uyghur sentiment analysis via the pre-training model with bilstm. *Applied Sciences*, 12(3):1182.
- Jakub Pokrywka and Krzysztof Jassem. 2024. kubapok@LT-EDI 2024: Evaluating Transformer Models for Hate Speech Detection in Tamil. In *Proceedings of the Fourth Workshop on Language Technology for Equality, Diversity and Inclusion*, Malta. European Chapter of the Association for Computational Linguistics.
- Rahul Ponnusamy, Malliga S, Sajeetha Thavareesan, Ruba Priyadharshini, and Bharathi Raja Chakravarthi. 2023. [VEL@LT-EDI-2023: Automatic detection of hope speech in Bulgarian language using embedding techniques](#). In *Proceedings of the Third Workshop on Language Technology for Equality, Diversity and Inclusion*, pages 179–184, Varna, Bulgaria. INCOMA Ltd., Shoumen, Bulgaria.
- Ruba Priyadharshini, Bharathi Raja Chakravarthi, Subalalitha Cn, Thenmozhi Durairaj, Malliga Subramanian, Kogilavani Shanmugavadivel, Siddhanth U Hegde, and Prasanna Kumaresan. 2022. Overview of abusive comment detection in tamil-acl 2022. In *Proceedings of the Second Workshop on Speech and Language Technologies for Dravidian Languages*, pages 292–298.
- A Ankitha Reddy, Ann Maria Thomas, Pranav Moorthi, and B Bharathi. 2024. SSN-Nova@LT-EDI 2024: POS Tagging, Boosting Techniques and Voting Classifiers for Caste And Migration Hate Speech Detection. In *Proceedings of the Fourth Workshop on Language Technology for Equality, Diversity and Inclusion*, Malta. European Chapter of the Association for Computational Linguistics.
- Anastasia Ryzhova, Dmitry Devyatkin, Sergey Volkov, and Vladimir Budzko. 2022. Training multilingual and adversarial attack-robust models for hate detection on social media. *Procedia Computer Science*, 213:196–202.
- Saisandeep Sangeetham, Shreyamanisha C Vinay, Kavin Rajan G, Abishna A, and B Bharathi. 2024. Algorithm Alliance@LT-EDI-2024: Caste and Migration Hate Speech Detection. In *Proceedings of the Fourth Workshop on Language Technology for Equality, Diversity and Inclusion*, Malta. European Chapter of the Association for Computational Linguistics.
- Sreerekha Sathi. 2023. Marriage murders and anti-caste feminist politics in india. In *Women's Studies International Forum*, volume 100, page 102816. Elsevier.
- Kogilavani Shanmugavadivel, Malliga Subramanian, Aiswarya M, Aruna T, and Jeevaananth S. 2024. KEC AI DSNLP@LT-EDI-2024:Caste and Migration Hate Speech Detection using Machine Learning Techniques. In *Proceedings of the Fourth Workshop on Language Technology for Equality, Diversity and Inclusion*, Malta. European Chapter of the Association for Computational Linguistics.
- Kriti Singhal and Jatin Bedi. 2024. Transformers@LT-EDI-EACL2024: Caste and Migration Hate Speech Detection in Tamil Using Ensembling on Transformers. In *Proceedings of the Fourth Workshop on Language Technology for Equality, Diversity and Inclusion*, Malta. European Chapter of the Association for Computational Linguistics.
- Malliga Subramanian, Rahul Ponnusamy, Sean Benhur, Kogilavani Shanmugavadivel, Adhithiya Ganesan, Deepti Ravi, Gowtham Krishnan Shanmugasundaram, Ruba Priyadharshini, and Bharathi Raja Chakravarthi. 2022. Offensive language detection in tamil youtube comments by adapters and cross-domain knowledge transfer. *Computer Speech & Language*, 76:101404.
- M Shahiki Tash, Z Ahani, M T Zamir, O Kolesnikova, and G Sidorov. 2024. Lidoma@LT-EDI 2024:Tamil Hate Speech Detection in Migration Discourse. In *Proceedings of the Fourth Workshop on Language Technology for Equality, Diversity and Inclusion*, Malta. European Chapter of the Association for Computational Linguistics.
- Liang Yao, Chengsheng Mao, and Yuan Luo. 2019. Graph convolutional networks for text classification. In *Proceedings of the AAAI conference on artificial intelligence*, volume 33, pages 7370–7377.