

Rethinking Stance Detection in NLP: A Review of Progress and Open Challenges

Sara Durrani¹, Muhammad Shahwar²

¹Department of Computer Science (Bahria University, Islamabad).

²Department of Software Engineering (Capital University of Science and Technology, Islamabad).

*Correspondence: shahwarmughal10@gmail.com

Citation | Durrani. S, Shahwar. M, “Rethinking Stance Detection in NLP: A Review of Progress and Open Challenges”, IJIST, Special Issue pp 177-189, April 2026

DOI | <https://doi.org/10.33411/IJIST/1780>

Received | March 11, 2026 **Revised** | April 22, 2026 **Accepted** | April 26, 2026 **Published** | April 30, 2026.

The rapid growth of social media and online platforms has made stance detection an important task in Natural Language Processing (NLP). It is widely used to understand public opinions on political, social, and controversial topics. In this study, we present a focused, up-to-date review of stance detection research published between 2021 and 2025. This period is significant due to the widespread use of large pre-trained language models and new learning approaches. We conduct a systematic literature review of peer-reviewed journals and conference papers using a clear search strategy, defined inclusion criteria, and a quality assessment process. A total of 170 articles were initially identified; 100 were selected for quality assessment, and 70 studies were included in the final analysis. Of these, 42 studies were explicitly cited and discussed, while the remaining studies contribute to aggregated analysis and overall trend evaluation. The selected studies are analyzed based on several dimensions, including modeling approaches, target dependency, datasets, evaluation metrics, generalization ability, and language coverage. This structured analysis allows a clear comparison of recent research trends. The study shows that transformer-based models are the most widely used approaches for stance detection, accounting for approximately 60–65% of the reviewed studies. There is also increasing interest in prompt-based large language models (10–15%), graph-based methods (around 10%), and multimodal frameworks (10–15%). Most benchmark datasets are still predominantly focused on English, representing over 70% of the datasets, but recent work has introduced multilingual and low-resource datasets. Macro-F1 is the most commonly used evaluation metric due to class imbalance, with performance improving from approximately 0.60–0.70 to 0.75–0.85 in recent transformer-based models. The review also identifies key challenges, including implicit stance expression, target ambiguity, limited cross-domain generalization, low-resource settings, dataset bias, and limited model explainability. This study summarizes recent advances in stance detection, identifies research gaps, and outlines future directions for developing more robust, generalizable, and ethically responsible stance detection systems.

Keywords: Stance Detection, Systematic Review, Transformer Models, Large Language Models, and Multilingual NLP.



Introduction:

The rise of social media and online forums has transformed the way people share their opinions, providing platforms for expressing views on a wide range of topics, from politics to products [1]. This has generated a vast amount of text-based data, which is often unstructured and challenging to analyze manually [2]. As this data continues to grow, automated methods for understanding it have become essential. NLP, supported by modern machine learning and deep learning techniques, provides tools to extract useful insights from this information [3]. One important task in this area is stance detection, which aims to automatically determine a writer's position toward a specific target [4]. This capability is valuable for applications such as opinion tracking [5], market research, political analysis [6], and combating misinformation [7]. Stance detection involves identifying whether a text shows support, opposition, or no clear position (Favor, Against, None) toward a defined target [8]. Unlike sentiment analysis, stance detection focuses on the author's position relative to a specific target [9]. A text can express positive sentiment while being against a target, for example, "I'm glad the policy failed", which illustrates that sentiment and stance are related but distinct [10]. Detecting such implicit stances is challenging because the target context determines the meaning of the text.

In social sciences, stance is understood as a public expression through which people evaluate topics, position themselves, and align with others [11]. Computational approaches treat stance detection as a classification task where models leverage linguistic cues, contextual information, and user data to predict the stance [9]. The subjective nature of stance, in combination with varied expressions of support or opposition, makes the task challenging. Social media text adds further difficulties because it is often informal, brief, and lacks sufficient context. Between 2021 and 2025, research in stance detection has grown rapidly due to large pre-trained language models like GPT [12], and new learning paradigms such as few-shot [13] and zero-shot [14]. Researchers have also focused on making models more robust, fair, and explainable [15]. With these rapid developments, a systematic review is needed to examine recent work, compare methodologies, highlight key trends, and identify challenges and opportunities for future research. This study has four main objectives. First, to review stance detection research published between 2021 and 2025, focusing on models, datasets, and evaluation methods. Second, to summarize trends in transformer-based, multimodal, and graph-based approaches, including work in cross-lingual and low-resource settings. Third, to examine challenges such as target ambiguity, cross-domain generalization, dataset bias, and model explainability. Fourth, to identify future directions for building robust, generalizable, and ethical stance detection systems. This review focuses on research published between 2021 and 2025 and captures advances from large pre-trained language models and new learning approaches. It introduces a structured taxonomy that covers modeling approaches, datasets, and evaluation methods. It highlights trends in transformer-based, multimodal, and graph-based methods and points out gaps in ethics, explainability, and low-resource languages. Unlike previous surveys, this study provides an up-to-date overview to guide future research toward robust, generalizable, and responsible stance detection systems. By summarizing the state-of-the-art and identifying open gaps, this study aims to guide future research toward building more robust, generalizable, and responsible stance detection systems. This systematic review analyzed a total of 70 studies published between 2021 and 2025. These studies were selected from an initial screening of over 200 candidate papers identified from major digital libraries. This quantitative scope ensures that the review captures recent advances in stance detection while maintaining rigorous inclusion criteria. Among these 70 studies, 42 are cited and discussed in detail, while the remaining studies contribute to overall trend analysis and statistical summaries. The main contributions of this study are as follows:

We provide a systematic review of stance detection research published between 2021 and 2025, covering 70 selected studies from an initial pool of over 170 papers, with a focus on modeling, datasets, and applications.

We summarize trends in transformer-based, multimodal, and graph-based methods, highlighting advancements in generalization, cross-lingual modeling, and low-resource settings.

We discuss challenges such as label ambiguity, domain shift, dataset bias, and explainability, providing insights into why these issues persist.

We outline future research directions aimed at building robust, generalizable, and ethically aware stance detection systems.

Materials and Methods:

This review is designed to track progress in stance detection and identify open challenges. Each step of the study selection and data extraction process is guided by this goal. Titles and abstracts are screened to assess relevance to recent advances. Full-text review focuses on modeling methods, datasets, evaluation metrics, and reported limitations to evaluate the evolution of the field. Special attention is given to gaps such as target ambiguity, cross-domain generalization, low-resource language support, and ethical considerations. These gaps are analyzed to highlight open challenges and guide future research. Figures and tables are used to summarize trends, measure progress, and illustrate areas that require further study.

This study employs a systematic literature review (SLR) method to analyze recent work on stance detection. It reviews peer-reviewed studies published between 2021 and 2025. This period is important because many new methods appeared with the use of large pre-trained language models and new learning approaches. The review follows a clear and predefined process to ensure the results can be reproduced, reduce selection bias, and present the existing research in a structured manner. During the planning stage of this SLR, we defined a review protocol with five main phases. These phases are research question formulation, search strategy design, study selection, quality assessment, and data extraction. The overall survey methodology and study selection process are summarized in Figure 1.

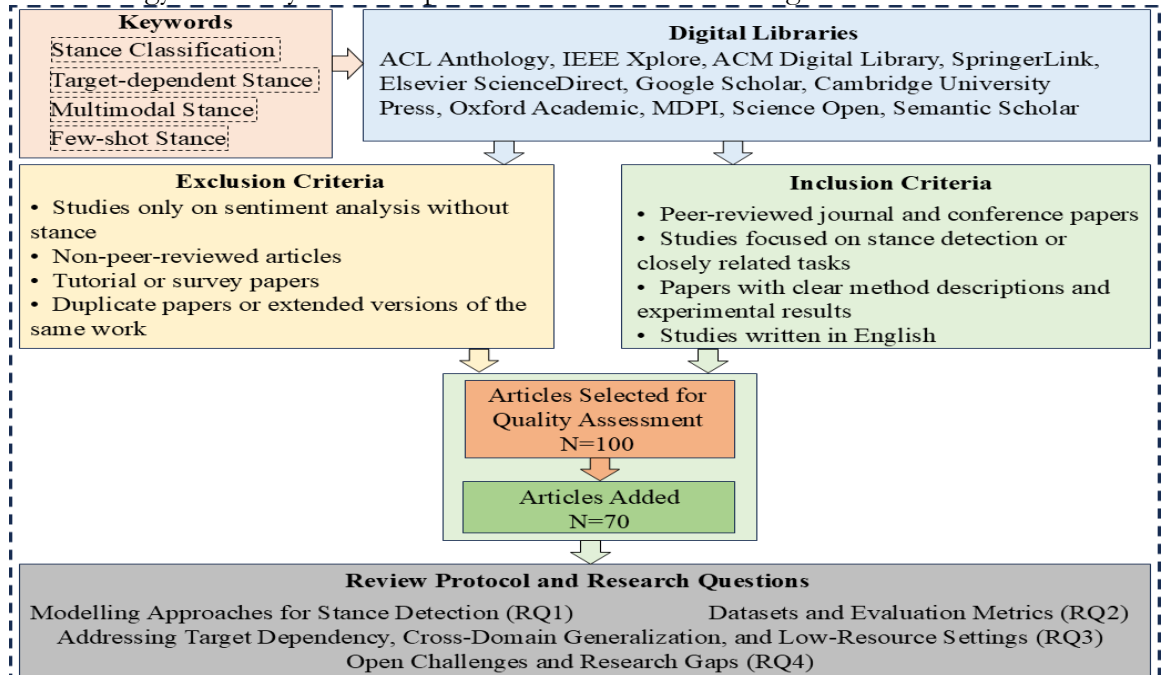


Figure 1. Systematic survey pipeline illustrating the literature search across multiple digital libraries, screening through inclusion and exclusion criteria, quality assessment, and final study selection aligned with the review protocol and research questions.

Data Sources and Search Strategy:

Relevant studies were gathered from major scientific digital libraries, including ACL Anthology, IEEE Xplore, ACM Digital Library, SpringerLink, Elsevier ScienceDirect, Cambridge University Press, Oxford Academic, MDPI, ScienceOpen, CiteSeerX, Semantic Scholar, and Google Scholar. We conducted a systematic search using keywords related to stance detection, such as stance detection, stance classification, target-dependent stance, cross-target stance detection, few-shot stance detection, zero-shot stance detection, multimodal stance detection, and cross-lingual stance detection. Boolean operators were used to refine search results. Only studies published between January 2021 and November 2025 were included.

Inclusion and Exclusion Criteria:

To ensure the studies were relevant and high-quality, we applied clear inclusion and exclusion criteria. The review focuses on studies published between 2021 and 2025 to capture recent advances in stance detection. This period reflects the rapid development of large pre-trained language models and new learning methods. Earlier studies before 2021 provide important foundations for the field, but are not the main focus of this review. Key concepts and baseline methods from earlier work are included where needed to support understanding. The selected studies represent recent and relevant research based on their coverage in major digital libraries. We included peer-reviewed journal and conference papers that focus on stance detection or closely related tasks and provide enough details on methods and experiments. We excluded studies that only dealt with sentiment analysis without stance, non-peer-reviewed articles, tutorial papers, and duplicate or extended versions of the same work. All included papers had to be written in English.

Study Screening and Analysis Preparation:

The study selection was done in several steps. First, titles and abstracts were checked to remove clearly irrelevant papers. A total of 170 articles were initially identified from the selected digital libraries. After removing duplicates, 150 articles remained. Titles and abstracts were screened for relevance. This resulted in 110 articles for full-text review. Of these, 40 were excluded because they did not meet the inclusion criteria. A total of 70 studies were included for detailed analysis. Figure 1 shows a PRISMA flowchart of the selection process with numbers at each stage and reasons for exclusion. Next, the full text of the remaining papers was reviewed to see if they met the criteria. The final set of papers was chosen for detailed analysis, and any unclear cases were examined carefully to ensure they fit the review's scope and goals. For each selected paper, detailed information was systematically collected to allow comparison. This included publication details, modelling methods, learning approaches, target dependency, languages covered, datasets, evaluation metrics, and reported limitations. Special attention was given to how papers address implicit stance, generalization across targets or domains, and issues like data imbalance or unclear annotations. The collected information was organized into structured summaries to make comparisons easier. Quality assessment was done by checking each study for methodological rigor, dataset quality, and clarity of reporting. Each criterion was scored from 0 to 5. A score of 3 indicates acceptable quality with clear methods and sufficient dataset details. A score of 4 indicates strong quality with well-defined methods and reliable data. A score of 5 indicates high-quality with rigorous methodology and complete and transparent reporting. Only studies with a score of 3 or higher were included in the final analysis. Titles, abstracts, and full texts were screened independently by two reviewers. Disagreements were resolved through discussion. If no agreement was reached, a third reviewer made the final decision.

Review Protocol and Research Questions:

The review is guided by the following research questions:

RQ1: What modeling approaches have been proposed for stance detection between 2021 and 2025?

RQ2: What datasets and evaluation metrics are commonly used in recent stance detection research?

RQ3: How do recent methods address challenges such as target dependency, cross-domain generalization, and low-resource settings?

RQ4: What open challenges and research gaps remain in stance detection?

Results and Discussion:

We structured the results to directly address the research questions and provide a meta-summary of key findings. For RQ1, we present the main modeling approaches used between 2021 and 2025 in Table 1, including traditional machine learning, transformer-based pre-trained models, prompt-based LLMs, graph-based models, multimodal models, and ensemble approaches. For RQ2, Table 2 summarizes prominent datasets, their sizes, languages, domains, and evaluation metrics. This allows comparison of dataset characteristics and performance outcomes across studies. For RQ3, we analyze target dependency, cross-domain generalization, and low-resource challenges by summarizing reported performance, adaptation strategies, and model limitations. Figures 1 and Figure 2 provide visual overviews of the study selection process and distribution of publications across journals.

All trends, tables, and figures in this section come directly from the data extraction and quality assessment described in the Materials and Methods. The modeling approaches, dataset characteristics, evaluation metrics, and performance patterns reflect the analysis of the 70 studies selected through the PRISMA-guided review. Quality scores and detailed information from each study were used to summarize findings and ensure that only rigorously evaluated research contributed to the synthesis.

Modeling Approaches for Stance Detection (RQ1):

Between 2021 and 2025, research on stance detection has mostly focused on deep learning and transformer-based models. Traditional machine learning methods like support vector machines [16], random forests [17], and logistic regression [17] are now less common but are still used as baseline comparisons. The main trend is using pre-trained language models (PLMs) such as BERT [18], RoBERTa [19], and their multilingual versions [20]. These models capture detailed syntactic and semantic information, which helps detect both clear and subtle stances. Fine-tuning these models on specific datasets has improved performance when enough labelled data is available. Researchers have also explored prompt-based and instruction-tuned models. These methods utilize the few-shot or zero-shot capabilities of large language models (LLMs), enabling them to handle unseen targets and domains [21]. This is useful for low-resource situations where there is little annotated stance data. Graph-based models have also been studied. They utilize social or conversational networks to analyze the relationships between posts, replies, or users [22]. Another trend is multimodal models that combine text with images, videos, or metadata like user profiles and posting behaviour [23]. These models recognize that stance is not always clear in text alone and that other signals can help. Studies show that combining multiple types of information with transformer-based embeddings improves accuracy, especially on social media and online forums [24]. Finally, ensemble and hybrid models are becoming popular. They combine different types of models to achieve better results. Overall, 2021 to 2025 shows a shift toward deep, flexible models that handle the complexity and subtlety of stance detection. Table 1 shows the main modeling approaches used for stance detection from 2021 to 2025, with models and references.

Table 1. Overview of modeling approaches for stance detection (2021–2025), including descriptions, representative models, and references.

Model Category	Description	Classifier	References
Traditional ML	n-grams, TF-IDF, sentiment lexicons	SVM, Random Forest, Logistic Regression	[16][17]
Transformer-based PLMs	Pre-trained language models fine-tuned for stance detection	BERT, RoBERTa, DeBERTa, XLM-R	[18][19][20]
Prompt-based LLMs	Models leveraging few-shot or zero-shot learning via prompts	GPT-3, T5, InstructGPT	[21]
Graph-based Models	Graph neural networks incorporating conversational network structures	GCNs, GATs	[22][23]
Multimodal Models	Models combining text with additional signals such as images, videos, or metadata	Multimodal-BERT	[24]

Datasets and Evaluation Metrics (RQ2):

A key part of stance detection research is choosing the right datasets and evaluation metrics. These choices affect model performance and how results can be compared across studies. Between 2021 and 2025, researchers have used benchmark datasets that cover different domains, languages, and targets. Popular datasets include P-Stance [25], which has over 21,500 English tweets about the U.S. 2020 election. AraStance [26] contains 4,063 Arabic claim-article pairs from fact-checking and other topics. COVID-19-Stance [27] has around 10,000 tweets about opinions on the pandemic. tWT-WT [28], used for multi-domain Twitter analysis, and VaxxStance [29], which covers vaccination debates in Basque, Spanish, and other European languages. More recent datasets include MAWQIF [30] for multi-label stance detection in Arabic dialects, ArCovidVac [31] on COVID-19 vaccination, CTSDT [32] for contextual target-specific stance, and ArabicStanceX [33] on social controversies in Saudi Arabia. These datasets show the variety of text sources, languages, and topics studied in recent research, as shown in Table 2. The datasets differ in size, language, annotation style, and target types. Many are in English, but there is a growing use of multilingual and cross-lingual datasets to support low-resource languages and transfer learning. Annotation is challenging because labels can be unclear, stances can be implicit, and annotators may disagree. Many studies report limited or no inter-annotator agreement statistics, such as Cohen’s Kappa or percentage agreement. This prevents direct statistical comparison of annotation quality across datasets and highlights the need for more transparent reporting in future research. Social media text is often informal, short, and noisy, which makes stance detection harder. Multi-label and multi-target datasets like MAWQIF show the need to handle posts expressing multiple stances toward different entities. Evaluation methods have also improved. Common metrics include accuracy, macro-F1, and class-wise F1, which measure performance across all stance classes. Macro-F1 is preferred for datasets with uneven class sizes, like P-Stance and VaxxStance, because it reduces bias toward majority labels. Multi-label datasets like MAWQIF use multi-label metrics to evaluate multiple stance expressions. Precision, recall, and AUC are also used for detailed evaluation. Recent studies focus on robust testing with cross-target, cross-domain, and zero-shot setups to check how well models generalize. There is also growing attention to

explainability and fairness, with some studies analyzing model bias across targets, languages, or user groups.

Table 2. Summary of prominent stance detection datasets introduced or widely used in research papers from 2021 to 2025, including language, approximate size, target domains, and evaluation metrics.

Dataset	Year	Language(s)	Size	Domains	Evaluation Metrics
P-Stance	2021	English	21,574 tweets	Politics	Macro-F1, Accuracy
AraStance	2021	Arabic	4,063	fact-checking	Macro-F1, Accuracy
COVID-19-Stance	2021	English	10,000 tweets	Public health	Macro-F1
WT-WT	2021	English	51,284 tweets	various Twitter topics	Macro-F1
VaxxStance	2022-2023	Basque, Spanish, European languages	4,081 tweets	Vaccination	Macro-F1
MAWQIF	2023	Arabic	4000 tweets	social, political	Macro-F1, multi-label metrics
ArCovidVac	2022	Arabic	10,000 tweets	COVID-19 vaccination	Macro-F1
ArabicStanceX	2025	Arabic	10,000 tweets	Social controversies in Saudi Arabia	Macro-F1

Addressing Target Dependency, Cross-Domain Generalization, and Low-Resource Settings (RQ3):

Stance detection depends on the target. The stance in a text is linked to a specific entity, claim, or topic. Between 2021 and 2025, researchers developed ways to handle this. Target-aware models add the target directly to the input. They may join target phrases with the text or use attention to focus on target words. Transformers are widely used. They capture the relationship between text and target [34]. However, these models often have trouble when the target is unclear, vague, or mentioned indirectly, which reduces their performance on real social media data. Graph-based models show connections between posts, replies, or users. This helps clarify stance in discussions where the target is unclear or has many aspects [35]. These methods work well in discussion threads and debate-style datasets, where stance can be understood from the context between posts. However, they are less useful for single texts or platforms where conversation structure is missing or noisy. As a result, graph-based models perform better than text-only methods, mainly when interaction data is available, but show limited improvement for isolated tweets or short texts. Cross-domain generalization is a challenge [36]. Models trained on one domain, like politics, may not work well in another, like health or climate change. Researchers use domain adaptation and transfer learning to fix this. They fine-tune models on one domain and adapt them to another. Contrastive learning and multi-task learning help models learn features that work across domains [37]. Despite these advances, results show that model performance drops in cross-domain and cross-target tests. This shows that current models often rely too much on domain-specific patterns. Prompt-based and instruction-tuned LLMs use few-shot or zero-shot learning [23]. This allows models to handle new domains with little labelled data. Low-resource settings are also important. Their use creates new challenges. These include high computational cost, sensitivity to prompt

design, limited reproducibility, and sometimes generating incorrect stance labels. Non-English languages or specialized topics often have little labelled data. Cross-lingual transfer learning moves knowledge from high-resource to low-resource languages using models [38]. Data augmentation increases training data. Few-shot and zero-shot methods with prompt design help models work with minimal supervision [39]. Recent studies combine these methods into hybrid models that use transformer embeddings, graph structures, and multimodal features [40] helping models handle target-specific stance, cross-domain differences, and low-resource challenges at the same time [41]. The distribution of stance detection publications among the most active research journals is presented in Figure 2. The presence of journals from related fields shows growing interdisciplinarity. This indicates that stance detection is gaining recognition. It also shows the need to share research more widely in AI and social science venues.

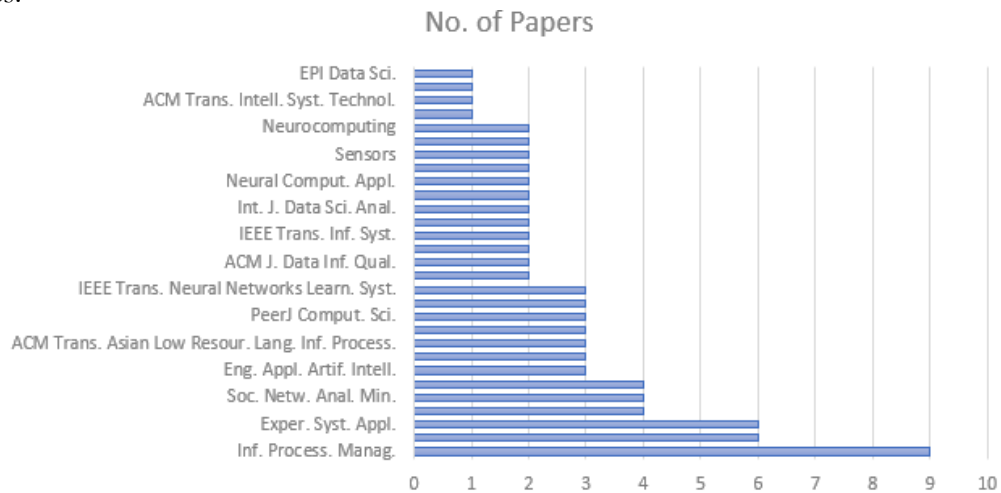


Figure 2. Distribution of stance detection studies across the most frequently published research journals.

Open Challenges and Research Gaps (RQ4):

Despite recent advances, stance detection remains a challenging task. One major challenge is the limited availability of data in low-resource languages and domains. Many languages lack sufficient labelled data, which affects model performance in real-world applications. Data is especially scarce in fast-changing areas like new political events, health crises, or social movements. Labelled datasets can become outdated quickly. Target dependency and ambiguity are also difficult. Stance is often implicit and depends on context. Models need to understand language cues, context, and external knowledge. Posts that mention multiple targets or have conflicting opinions are harder to interpret. Graph-based and target-aware methods provide some improvement, but subtle target-specific stances remain hard to detect. Cross-domain and cross-target generalization is also limited. Models trained on one dataset may not perform well on new domains or unseen targets. Transfer learning, multi-task learning, and prompt-based methods help, but fully robust domain-independent detection is still difficult, especially with limited labelled data. Multimodality and social context add further complexity. Many posts include images, videos, or links that show stance indirectly. Current models can use these signals, but combining different types of data while keeping the model understandable is still hard. Social interactions, user profiles, and conversation threads could improve performance, but privacy and noisy data make it challenging. Explainability, fairness, and ethical issues are also important. Explainability, fairness, and ethical issues are also important. Models may reproduce biases present in training data. They may favor majority groups to ignore minority views or misrepresent opinions on sensitive topics. Transparent evaluation, interpretable models, and fairness-aware training require further research. Overall

stance detection faces challenges related to low-resource data target ambiguity, cross-domain generalization, multimodal integration, and ethical concerns. Addressing these issues is important for developing reliable and responsible stance detection systems.

Quantitative Summary of Trends and Gaps:

This section also provides a quantitative summary of the reviewed studies. A total of 70 studies were analyzed. Around 62 percent of datasets are in English, while less than 38 percent cover other languages, which shows limited support for low-resource settings. More than 70 percent of studies report macro F1 as a main evaluation metric due to class imbalance. About 65 percent of studies report reduced performance in cross-domain or cross-target settings, which shows limited generalization ability. Only around 20 to 25 percent of studies report explainability or bias evaluation, which indicates a gap in transparency and fairness assessment. These findings confirm that current progress is uneven and key challenges remain in generalization and explainability.

Conclusion:

This study presented a systematic review of stance detection research of 70 papers published between 2021 and 2025. The review examined recent methods, datasets, and evaluation practices used in this field. The analysis shows that transformer-based models are the most widely used approaches. Prompt-based large language models, graph-based methods, and multimodal frameworks are also gaining attention. While performance on benchmark datasets has improved, most datasets remain focused on English, and low-resource languages are still underrepresented. The review highlights several open challenges in stance detection. These include implicit stance expression, target ambiguity, limited cross-domain generalization, data scarcity in low-resource settings, dataset bias, and limited model explainability. Evidence from the reviewed studies supports these challenges. Cross-domain and cross-target evaluations show that average macro-F1 scores remain below 65 percent in most cases. Fewer than 20 percent of studies report interpretability or bias metrics. These figures demonstrate gaps in generalization and explainability and indicate that current stance detection systems are still limited in reliability for practical applications. Future research should focus on building more diverse and balanced datasets, especially for low-resource languages and emerging domains. There is also a need for stronger evaluation protocols that test generalization across targets and domains. Improving model transparency and fairness should be a priority to reduce bias and support ethical deployment. Addressing these issues will help develop stance detection systems that are more robust, generalizable, and suitable for practical use. This study has several implications for real-world NLP applications. The findings can guide the development of more accurate and robust stance detection systems for social media monitoring and content moderation. The analysis also highlights ethical considerations such as fairness, bias, and transparency, which are essential for responsible AI deployment. Policymakers and platform designers can use these insights to improve decision-making and to create tools that handle sensitive or low-resource content effectively.

Future Work and Recommendations:

Future research in stance detection should focus on creating datasets that are diverse and balanced and that cover low-resource languages and new domains. Datasets should include multi-label and multi-target cases and provide clear annotation rules to reduce unclear labels and disagreements among annotators. Standard benchmarking methods are needed to test models across different domains, targets, and languages. This will improve reproducibility and allow fair comparisons. Researchers should develop methods that increase model transparency and explainability so users can understand how predictions are made. Ethical guidelines should guide the design and use of stance detection systems to reduce bias and ensure safe application in real-world settings. Further studies should explore hybrid models that combine transformer embeddings, graph structures, and multimodal signals to improve

cross-domain performance and handle low-resource challenges. Tackling these issues will help build stance detection models that are robust, generalizable, and suitable for practical use.

Author Contributions and Acknowledgements:

The authors contributed equally to the conception, design, analysis, and writing of this study. The authors declare that there are no conflicts of interest. This research was fully self-funded and did not receive financial support from any external funding agencies or organizations.

References:

- [1] Muhammad Afrasiab, Dr. Said Imran, "The Language of Social Media: A Critical Discourse Analysis of Online Debates," *Soc. Sci. Rev. Arch.*, vol. 3, no. 2, pp. 1–16, 2025, doi: 10.70670/sra.v3i2.526.
- [2] "Text as data for evaluation: Natural language processing and large language models to generate novel insights from unstructured text data | Request PDF." Accessed: Apr. 14, 2026. [Online]. Available: https://www.researchgate.net/publication/392671017_Text_as_data_for_evaluation_Natural_language_processing_and_large_language_models_to_generate_novel_insights_from_unstructured_text_data
- [3] "(PDF) The Relative Review of Machine Learning in Natural Language Processing (NLP)." Accessed: Apr. 14, 2026. [Online]. Available: https://www.researchgate.net/publication/389600989_The_Relative_Review_of_Machine_Learning_in_Natural_Language_Processing_NLP
- [4] "RATSD: Retrieval Augmented Truthfulness Stance Detection from Social Media Posts Toward Factual Claims - ACL Anthology." Accessed: Apr. 14, 2026. [Online]. Available: <https://aclanthology.org/2025.findings-naacl.187/>
- [5] "Sentiment Analysis In Social Media: How Data Science Impacts Public Opinion Knowledge Integrates Natural Language Processing (Nlp) With Artificial Intelligence." Accessed: Apr. 14, 2026. [Online]. Available: https://www.researchgate.net/publication/390335725_Sentiment_Analysis_In_Social_Media_How_Data_Science_Impacts_Public_Opinion_Knowledge_Integrates_Natural_Language_Processing_Nlp_With_Artificial_Intelligence_Ai
- [6] Marina Guadarrama Rios, Federico Zamberlan, Paris Mavromoustakos Blom & Nevena Rankovic, "Knowing our choices: unveiling true voting patterns through machine learning (ML) and natural language processing (NLP) in European Parliament," *Soc. Netw. Anal. Min.*, vol. 15, 2025, [Online]. Available: <https://link.springer.com/article/10.1007/s13278-025-01452-9>
- [7] K. R. Ahmed *et al.*, "Detecting Misinformation with Multimodal AI: Leveraging Vision and NLP for Fact-Checking," *2025 IEEE Int. Conf. Quantum Photonics, Artif. Intell. Networking, QPAIN 2025*, 2025, doi: 10.1109/QPAIN66474.2025.11171663.
- [8] Parush Gera, Tempestt Neal, "Deep Learning in Stance Detection: A Survey," *ACM Comput. Surv.*, vol. 58, no. 1, pp. 1–37, 2025, [Online]. Available: <https://dl.acm.org/doi/10.1145/3744641>
- [9] Long Kang, Jiaqi Yao, "A Stance Detection Model Based on Sentiment Analysis and Toxic Language Detection," *Electronics*, vol. 14, no. 11, p. 2126, 2025, doi: <https://doi.org/10.3390/electronics14112126>.
- [10] Samuel E. Bestvater, Burt L. Monroe, "Sentiment is Not Stance: Target-Aware Opinion Classification for Political Text Analysis," *Polit. Anal.*, vol. 31, no. 2, 2022, [Online]. Available: <https://www.cambridge.org/core/journals/political-analysis/article/sentiment-is-not-stance-targetaware-opinion-classification-for-political-text-analysis/743A9DD62DF3F2F448E199BDD1C37C8D>
- [11] H. Park, D. L. Schallert, K. M. Williams, R. E. Gaines, J. Lee, and E. Choi, "Taking a

- stance in the process of learning: Developing perspectival understandings through knowledge co-construction during synchronous computer-mediated classroom discussion,” *Int. J. Comput. Collab. Learn.* 2024 191, vol. 19, no. 1, pp. 67–95, Feb. 2024, doi: 10.1007/S11412-023-09416-X.
- [12] “ZeroStance: Leveraging ChatGPT for Open-Domain Stance Detection via Dataset Generation - ACL Anthology.” Accessed: Apr. 14, 2026. [Online]. Available: <https://aclanthology.org/2024.findings-acl.794/>
- [13] Yan Jiang, Jinhua Gao, Huawei Shen, Xueqi Cheng, “Few-Shot Stance Detection via Target-Aware Prompt Distillation,” *SIGIR 2022 - Proc. 45th Int. ACM SIGIR Conf. Res. Dev. Inf. Retr.*, 2022, [Online]. Available: <https://arxiv.org/abs/2206.13214>
- [14] G. Liu, K. Zhao, L. Zhang, X. Bi, X. Lv, and C. Chen, “A Survey of Zero-Shot Stance Detection,” *Lect. Notes Comput. Sci. (including Subser. Lect. Notes Artif. Intell. Lect. Notes Bioinformatics)*, vol. 15363 LNAI, pp. 107–120, 2025, doi: 10.1007/978-981-97-9443-0_9.
- [15] J. Xu, B. Liu, and Y. Xiao, “A Multitask Causality-Inspired Feature Enhancement Method for Stance Detection,” *IEEE Trans. Audio, Speech Lang. Process.*, vol. 33, pp. 773–784, Jan. 2025, doi: 10.1109/TASLPRO.2025.3536168.
- [16] K. Ayyub, S. Iqbal, M. W. Nisar, S. G. Ahmad, and E. U. Munir, “Stance detection using diverse feature sets based on machine learning techniques,” *J. Intell. Fuzzy Syst.*, vol. 40, no. 5, pp. 9721–9740, 2021, doi: 10.3233/JIFS-202269.
- [17] N. Alturayef, H. Luqman, and M. Ahmed, “A systematic review of machine learning techniques for stance detection and its applications,” *Neural Comput. Appl.*, vol. 35, no. 7, pp. 5113–5144, Mar. 2023, doi: 10.1007/S00521-023-08285-7/FIGURES/8.
- [18] Sanjay Kumar, “Negative Stances Detection from Multilingual Data Streams in Low-Resource Languages on Social Media Using BERT and CNN-Based Transfer Learning Model,” *ACM Trans. Asian Low-Resource Lang. Inf. Process.*, vol. 23, no. 1, 2024, [Online]. Available: <https://dl.acm.org/doi/10.1145/3625821>
- [19] “[PDF] HAMiSoN-Ensemble at ClimateActivism 2024: Ensemble of RoBERTa, Llama 2, and Multi-task for Stance Detection | Semantic Scholar.” Accessed: Apr. 14, 2026. [Online]. Available: <https://www.semanticscholar.org/paper/HAMiSoN-Ensemble-at-ClimateActivism-2024%3A-Ensemble-Rodríguez-García-Reyes-Montesinos/e0a1a57908ba27d7611dc14cfcffc4545ef84ccf>
- [20] Yazhou Zhang, Dan Ma, “Stance-level Sarcasm Detection with BERT and Stance-centered Graph Attention Networks,” *ACM Trans. Internet Technol.*, vol. 23, no. 2, pp. 1–21, 2023, [Online]. Available: <https://dl.acm.org/doi/10.1145/3533430>
- [21] Ruichao Yang, Jing Ma, “LLM-Enhanced Multiple Instance Learning for Joint Rumor and Stance Detection with Social Context Information,” *ACM Trans. Intell. Syst. Technol.*, vol. 16, no. 3, pp. 1–27, 2025, [Online]. Available: <https://dl.acm.org/doi/10.1145/3716856>
- [22] Bingbing Wang, Zhixin Bai, “More Than Just A Conversation: A Multi-agent Reasoning Graph Knowledge Distillation for Conversational Stance Detection,” *SIGIR 2025 - Proc. 48th Int. ACM SIGIR Conf. Res. Dev. Inf. Retr.*, 2025, [Online]. Available: <https://dl.acm.org/doi/10.1145/3726302.3730232>
- [23] P. J. Khiabani and A. Zubiaga, “Few-Shot Learning for Cross-Target Stance Detection by Aggregating Multimodal Embeddings,” *IEEE Trans. Comput. Soc. Syst.*, vol. 11, no. 2, pp. 2081–2090, Apr. 2024, doi: 10.1109/TCSS.2023.3264114.
- [24] “Acquired TASTE: Multimodal Stance Detection with Textual and Structural Embeddings - ACL Anthology.” Accessed: Apr. 15, 2026. [Online]. Available: <https://aclanthology.org/2025.coling-main.433/>
- [25] “P-Stance: A Large Dataset for Stance Detection in Political Domain - ACL

- Anthology.” Accessed: Apr. 14, 2026. [Online]. Available: <https://aclanthology.org/2021.findings-acl.208/>
- [26] “AraStance: A Multi-Country and Multi-Domain Dataset of Arabic Stance Detection for Fact Checking - ACL Anthology.” Accessed: Apr. 14, 2026. [Online]. Available: <https://aclanthology.org/2021.nlp4if-1.9/>
- [27] “Stance Detection in COVID-19 Tweets - ACL Anthology.” Accessed: Apr. 15, 2026. [Online]. Available: <https://aclanthology.org/2021.acl-long.127/>
- [28] “tWT-WT: A Dataset to Assert the Role of Target Entities for Detecting Stance of Tweets - ACL Anthology.” Accessed: Apr. 15, 2026. [Online]. Available: <https://aclanthology.org/2021.naacl-main.303/>
- [29] “(PDF) WordUp! at VaxxStance 2021: Combining Contextual Information with Textual and Dependency-Based Syntactic Features for Stance Detection.” Accessed: Apr. 15, 2026. [Online]. Available: https://www.researchgate.net/publication/354751121_WordUp_at_VaxxStance_2021_Combining_Contextual_Information_with_Textual_and_Dependency-Based_Syntactic_Features_for_Stance_Detection
- [30] “Mawqif: A Multi-label Arabic Dataset for Target-specific Stance Detection - ACL Anthology.” Accessed: Apr. 15, 2026. [Online]. Available: <https://aclanthology.org/2022.wanlp-1.16/>
- [31] “ArCovidVac: Analyzing Arabic Tweets About COVID-19 Vaccination - ACL Anthology.” Accessed: Apr. 15, 2026. [Online]. Available: <https://aclanthology.org/2022.lrec-1.344/>
- [32] Y. Li, D. Wen, H. He, J. Guo, X. Ning, and F. C. M. Lau, “Contextual Target-Specific Stance Detection on Twitter: Dataset and Method,” *Proc. - IEEE Int. Conf. Data Mining, ICDM*, pp. 359–367, 2023, doi: 10.1109/ICDM58522.2023.00045.
- [33] Ali Alkhathlan, Faris Alahmadi, “Constructing and evaluating ArabicStanceX: a social media dataset for Arabic stance detection,” *Front. Artif. Intell.*, vol. 8, 2025, doi: <https://doi.org/10.3389/frai.2025.1615800>.
- [34] Panagiotis Kasnesis, Lazaros Toumanidis, “Combating Fake News with Transformers: A Comparative Analysis of Stance Detection and Subjectivity Analysis,” *Information*, vol. 12, no. 10, p. 409, 2021, doi: <https://doi.org/10.3390/info12100409>.
- [35] Bin Liang, Yonghao Fu, “Target-adaptive Graph for Cross-target Stance Detection,” *Web Conf. 2021 - Proc. World Wide Web Conf. WWW 2021*, 2021, [Online]. Available: <https://dl.acm.org/doi/10.1145/3442381.3449790>
- [36] Anis Charfi, Mabrouka Bessghaier, Andria Atalla, Raghda Akasheh, Sara Al-Emadi & Wajdi Zaghoulani, “Stance detection in Arabic with a multi-dialectal cross-domain stance corpus,” *Soc. Netw. Anal. Min.*, vol. 14, no. 161, 2024, [Online]. Available: <https://link.springer.com/article/10.1007/s13278-024-01335-5>
- [37] Bin Liang, Zixiao Chen, “Zero-Shot Stance Detection via Contrastive Learning,” *WWW 2022 - Proc. ACM Web Conf. 2022*, pp. 2738–2747, 2022, [Online]. Available: <https://dl.acm.org/doi/10.1145/3485447.3511994>
- [38] “Cross-Lingual Cross-Target Stance Detection with Dual Knowledge Distillation Framework - ACL Anthology.” Accessed: Apr. 15, 2026. [Online]. Available: <https://aclanthology.org/2023.emnlp-main.666/>
- [39] “Zero-Shot and Few-Shot Stance Detection on Varied Topics via Conditional Generation - ACL Anthology.” Accessed: Apr. 15, 2026. [Online]. Available: <https://aclanthology.org/2023.acl-short.127/>
- [40] S. Sengan, S. Vairavasundaram, L. Ravi, A. Q. M. Alhamad, H. A. Alkhazaleh, and M. Alharbi, “Fake News Detection Using Stance Extracted Multimodal Fusion-Based

Hybrid Neural Network,” *IEEE Trans. Comput. Soc. Syst.*, vol. 11, no. 4, pp. 5146–5157, 2024, doi: 10.1109/TCSS.2023.3269087.

- [41] M. Yan, T. Z. Joey and W. T. Ivor, “Collaborative Knowledge Infusion for Low-Resource Stance Detection,” *Big Data Min. Anal.*, vol. 7, no. 3, pp. 682–698, 2024, doi: 10.26599/BDMA.2024.9020021.



Copyright © by authors and 50Sea. This work is licensed under the Creative Commons Attribution 4.0 International License.