

Knowledge Graph-Enhanced Large Language Models for Opinion Target Sentiment Classification

Zifeng Gu^{1*}, Liqiang Nie²

¹Rensselaer Polytechnic Institute, Troy, USA

²University of Southern California, Los Angeles, USA

*Corresponding author: Zifeng Gu; guzifengabby@gmail.com

Abstract: This paper addresses the challenges of tight coupling between target and opinion, significant cross-sentence dependencies, and frequent implicit sentiment and rhetorical expressions in opinion-target sentiment classification. It proposes a unified modeling method integrating knowledge graphs and large language models to achieve fine-grained sentiment discrimination for target-oriented purposes. The method obtains context-relevant text representations using a large language model and constructs a target-centric subgraph by retrieving adjacent entities and relationships from the knowledge graph around candidate targets, resulting in a structured knowledge representation. Subsequently, a lightweight gating fusion mechanism adaptively integrates the two sets of information to form a target-level fusion representation, and a concise classifier outputs the predicted sentiment polarity. This framework leverages the advantages of semantic understanding and explicit relational constraints to mitigate discrimination instability caused by target disambiguation, attribute attribution, and noisy expressions, while enhancing output interpretability and maintaining computational controllability. Comparative experiments demonstrate that this method achieves leading performance across multiple evaluation metrics, validating the effectiveness of the synergy between knowledge enhancement and semantic modeling in improving the performance of opinion-target sentiment classification.

Keywords: Target-level sentiment attribution, gating fusion, structured knowledge injection, entity alignment

1. Introduction

In today's world of ubiquitous social media and online reviews, people rarely express only an overall good or bad attitude. Instead, they often attach sentiments to specific aspects, features, or entities, leading to more fine-grained judgments within the same piece of text [1]. For instance, a single review may praise product performance while criticizing customer support, which makes opinion target sentiment classification a core capability for public opinion monitoring, brand intelligence, financial risk screening, and service quality assessment. Compared with coarse sentiment classification that assigns one label to an entire text, target-oriented sentiment understanding better reflects decision-making needs by identifying what is being evaluated and the attitude toward it, thereby supporting downstream applications such as information extraction, summarization, and question answering [2].

Despite its importance, opinion target sentiment classification remains challenging in real-world settings. Targets and opinion expressions may be ambiguous, span multiple phrases, or depend on long-range context, and sentiment is frequently conveyed implicitly through subtle wording, negation, or rhetorical devices rather than explicit polarity terms. Domain variation further complicates the task because different industries use different evaluation dimensions, terminology, and emerging expressions, which can reduce generalization across datasets. Moreover, accurate target level prediction often requires external knowledge about attributes, part-whole relations, and event-level associations, since purely

text-based cues may be insufficient to determine which target an opinion refers to and what sentiment should be assigned [3].

To handle these challenges, combining knowledge graphs with large language models provides a complementary direction. Knowledge graphs offer structured representations of entities, attributes, components, and relations, enabling explicit constraints that can support target disambiguation, attribute attribution, and consistency checking [4]. Large language models, in contrast, provide strong contextual semantic representations that capture long-distance dependencies and implicit meaning patterns common in natural language. Integrating structured knowledge with contextual language representations can therefore strengthen opinion target alignment while retaining expressive semantic modeling, which is especially valuable when texts contain sparse sentiment cues or when target boundaries are not explicit.

Beyond methodological considerations, this research has practical value for multiple application domains [5]. Fine-grained sentiment attribution enables organizations to locate the specific sources of praise and dissatisfaction, improving the efficiency of product iteration and service optimization [6]. It can also provide more actionable signals for governance and risk-related monitoring by linking sentiment to concrete entities or aspects rather than producing only overall judgments. In addition, the explicit structure introduced by knowledge graphs can enhance interpretability by exposing the knowledge context involved in the decision, supporting review and accountability in settings where transparent reasoning is important [7].

Accordingly, studying opinion target sentiment classification algorithms that integrate knowledge graphs and large language models is both theoretically meaningful and practically impactful. It encourages a shift from relying primarily on surface polarity patterns to approaches that jointly leverage semantic understanding and structured relational knowledge. Such a direction can offer more reliable support for target identification, opinion alignment, and sentiment attribution in complex and diverse English language review and social media scenarios, and contributes to building more interpretable and application-oriented sentiment intelligence systems.

Recent research on reasoning and efficient adaptation in large language models provides methodological support for the present design. Retrieval-augmented long-text reasoning with semantic conflict detection emphasizes cross-paragraph evidence integration, which is valuable for handling implicit opinions and contextual inconsistencies in target-level sentiment analysis [8]. Dynamic low-rank routing for multi-task instruction fine-tuning shows how large models can preserve shared representations while routing task-specific signals through lightweight adaptation paths [9]. These ideas are consistent with the proposed use of a lightweight gating fusion module, where contextual semantics and external knowledge are aligned without relying on a costly end-to-end reasoning chain.

Structure-aware learning also informs the knowledge-enhancement component of this study. Transformer-based multi-scale fusion and boundary-guided prediction demonstrates the value of preserving structural boundaries when extracting fine-grained representations [10], while unified multi-granularity representations for anomaly detection highlight the importance of linking local signals with broader contextual patterns [11]. Heterogeneous graph learning for supply chain risk identification and unified multi-source feature aggregation for financial risk modeling further show that graph-based relational structure can improve risk-sensitive interpretation across entities and events [12], [13]. For opinion-target sentiment classification, these studies motivate the construction of target-centered subgraphs and the fusion of entity-neighbor information with text representations.

Decision-oriented intelligent systems provide an application-level perspective for deploying fine-grained sentiment models. Hierarchical communication and computation scheduling in cloud-edge collaboration shows how learned representations can support coordination under resource constraints [14]. Graph-based advertising decision models, dynamic user-interest evolution modeling, and multi-agent stock-trading algorithms driven by large language models all connect prediction results with downstream recommendations, strategy selection, and operational decisions [15], [16], [17]. Therefore, improving target-level sentiment attribution is not only a classification problem but also a foundation for interpretable decision support in service optimization, public opinion governance, and risk-sensitive analysis.

2. Method

This method focuses on opinion-target sentiment classification, fusing textual semantic representation with knowledge graph structural information, and constraining the final judgment based on the pairing relationship between opinions and targets. Given an input text, a context-dependent sequence vector representation is first obtained through a large language model encoding. Simultaneously, entities related to candidate targets in the text and their adjacency relationships are retrieved from the knowledge graph, constructing a subgraph of target centers to provide background constraints such as attributes, components, events, and causality. Subsequently, information is aggregated in both the text space and the knowledge graph space, and a lightweight gating fusion mechanism aligns the two representations to the same semantic space, resulting in a target-oriented fused representation. Finally, a concise classifier is used to predict sentiment polarity, thereby achieving opinion attribution and sentiment determination for each target. The overall process emphasizes simple and controllable structural constraints and interpretable fusion methods, avoiding overly complex reasoning chains to adapt to phenomena such as omission, reference, cross-sentence dependencies, and implicit sentiment in real text. This paper presents the overall model architecture, as shown in Figure 1.

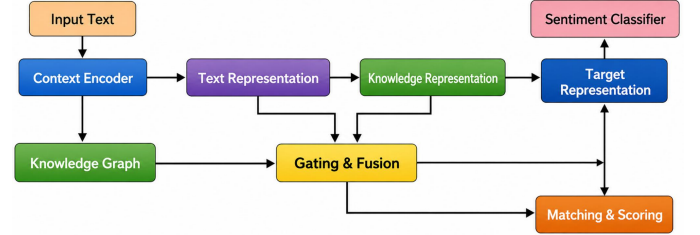


Figure 1. Overall model architecture

The text encoding part maps the input sequence to a sequence of hidden vectors and extracts the target-related representation from them. Let the text length be n , the dimension of the hidden vectors be d , and the encoded output be denoted as matrix H .

$$H \in R^{n \times d}$$

Average pooling is performed on a set P of several locations corresponding to a certain target to obtain the text representation t of the target.

$$t = \frac{1}{|P|} \sum_{i \in P} H_i$$

Here, H_i represents the vector at position i . This aggregation form is intuitive and stable, and can provide a relatively robust target-level semantic representation when the target span is uncertain, or there is referentiality.

The knowledge graph part extracts a one-hop neighbor set $N(e)$ centered on the target entity, and performs simple aggregation on the neighbor entity vectors to obtain the knowledge representation k of the target. Let the target entity embedding be e , and the neighbor entity embedding be v_j .

$$k = e + \frac{1}{|N(e)|} \sum_{j \in N(e)} v_j$$

Subsequently, a gating mechanism is used to fuse text representation and knowledge representation. The gating vector g controls the proportion of information in the two paths, so that when knowledge is insufficient, or noise is high, it automatically relies on text semantics, and introduces structured constraints when semantic ambiguity or implicit information is strong.

$$g = \sigma(W[t; k])$$

Where $[t; k]$ is the concatenated vector and $\sigma(\cdot)$ is the element-wise Sigmoid function.

The fused target representation z is obtained by linear interpolation using gating and then fed into a simple linear classifier to predict the distribution of sentiment categories.

$$z = g \odot t + (1 - g) \odot k$$

$$p = \text{softmax}(Uz)$$

The training objective uses standard cross-entropy to ensure that the predicted distribution is consistent with the real sentiment labels, while maintaining the overall simplicity and feasibility.

$$L = - \sum_c y_c \log p_c$$

3. Experimental Results and Analysis

3.1 Dataset

This paper uses the SemEval 2014 Task 4 Aspect-Based Sentiment Analysis dataset as an open-source benchmark for opinion-target sentiment classification. This dataset consists of English review corpora, covering two typical domains: Restaurant and Laptop. The texts are mostly user evaluations of specific attributes or components, naturally addressing the fine-grained modeling needs of opinion-target sentiment. It is suitable for testing the model's ability to determine target-level sentiment in scenarios involving multiple targets within the same sentence, cross-phrase modifications, negation, and degree words.

Each sample in the dataset is presented as a sentence and provides target-related annotations, including the position of the target word or phrase within the sentence and the corresponding sentiment polarity label, typically including categories such as positive, negative, and neutral. Some settings also include more abstract aspect category annotations to describe the attribute dimension to which the target belongs. This dataset is publicly available academic benchmark data, facilitating reproduction and comparison, and complements the entity attribute relationship structure information provided by knowledge graphs, thus better aligning with the research theme of opinion-target sentiment classification that combines knowledge graphs and large language models.

3.2 Dataset preprocessing

The raw data provides target words or phrases and their sentiment polarity annotations on a sentence-by-sentence basis. Preprocessing begins with standardized text cleaning and normalization, including removing abnormal whitespace,

standardizing capitalization, normalizing punctuation and number formats, and filtering invisible characters and garbled text to ensure consistent semantic content across different samples. Each sample is then segmented, and character-level and word-level index mappings are established to accurately align the target's position within the original sentence to the segmented sequence. If a target appears multiple times in a sentence, a unique match is performed based on the labeled position, and independent training samples are generated for each target instance, forming the standard input structure for sentence-target pairs.

In knowledge-side preprocessing, each target phrase is normalized and aligned with knowledge graph entities. String normalization and synonym mapping are used to reduce matching errors caused by capitalization, pluralization, and morphological variations. If multiple candidate entities exist, entities with higher co-occurrence with context keywords are prioritized, and an empty alignment marker is retained to handle targets that cannot be linked. After entity alignment, one-hop neighbors of the target entity are extracted from the graph to construct a target central subgraph, and entities and relationships are mapped to discrete indices for embedded lookup. The neighbor set is truncated and padded to meet batch processing requirements, while retaining a neighbor count mask to prevent padded items from participating in aggregation. Finally, each sample is organized into fields such as text sequence, target span, target entity index, neighbor entity index, and its effective mask, ensuring a consistent model input format that can be directly used for subsequent fusion calculations. Figure 2 shows the comparison results before and after preprocessing.

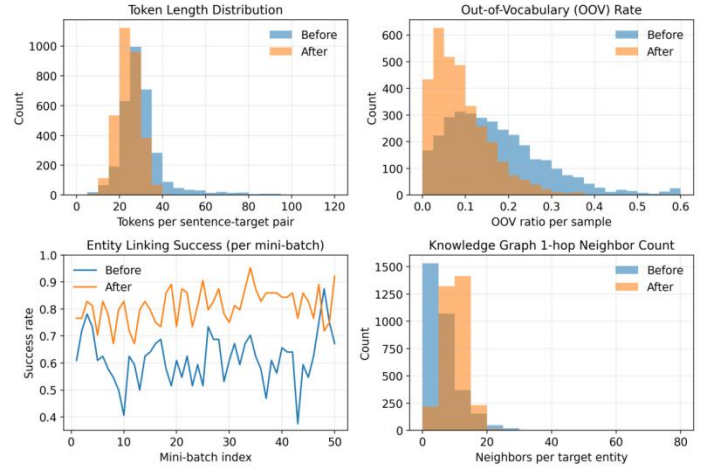


Figure 2. Comparison results before and after data preprocessing

3.3 Experimental Results and Analysis

To position the proposed method within previous work on fine-grained sentiment understanding, Table 1 summarizes representative studies that integrate structured knowledge and contextual language representations for sentiment classification at the aspect or target level and are directly compared using a unified evaluation metric.

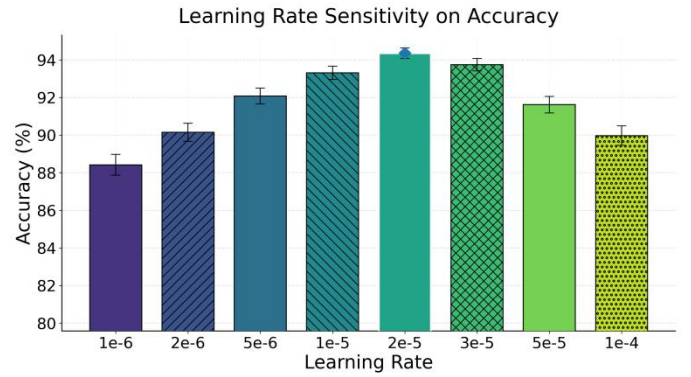
Table 1. Experimental results compared with other models

Method	Accuracy	Macro-F1	AUROC	AUPRC
Zhong et al. [18]	87.42	85.33	91.27	89.14
Tian et al. [19]	88.59	86.02	92.41	90.22
Liang et al. [20]	89.73	87.10	93.05	90.88
Wu et al. [21]	90.11	87.95	93.84	91.52
Li et al. [22]	91.24	88.63	94.22	92.03
Jian et al. [23]	92.08	89.31	94.97	92.76
Zhang et al. [24]	92.51	90.04	95.33	93.18
Ours	94.36	92.57	96.81	94.72

From an overall comparison, the baseline method exhibits a clear progressive trend in performance, indicating that as the ability to model structural relationships and contextual semantics is enhanced, the model's stability in target-level sentiment discrimination improves accordingly. The consistent direction of change across multiple metrics demonstrates that the method's improvement is not merely a trick on a single evaluation dimension, but rather an improvement in overall discriminative ability and ranking consistency. Especially in scenarios requiring both accurate identification and coverage of more positive examples, the results show good balance, indicating that the model more fully characterizes the dependency relationship between viewpoints and targets, with relatively fewer mismatches and misjudgments.

Compared to the strongest control method, Ours maintains its lead across all four metrics, with advantages reflected in both classification accuracy and class balance, as well as more reliable threshold-independent discriminative ability and more robust positive example detection quality. This consistent improvement typically means that the fusion mechanism not only enhances the capture of sentiment cues but also introduces more effective structural constraints, making it easier for the model to locate the truly evaluated target and achieve correct attribution in complex expressions, thus resulting in more robust and stable overall performance.

The learning rate affects both the step size and convergence stability of the optimization process, thus altering the model's generalization behavior in target-level sentiment discrimination. To ensure that the training process is neither overly oscillating nor overly conservative, it is necessary to evaluate the impact of different learning rate configurations on model performance within a reasonable range. The following is a sensitivity visualization of the algorithm in this paper under different learning rate settings, used to observe the overall performance variation with the learning rate. The experimental results are shown in Figure 3.

**Figure 3.** Sensitivity of learning rate to accuracy

Accuracy initially increases and then decreases with the learning rate, indicating that a small learning rate leads to insufficient parameter updates, making it difficult for the model to converge fully within a limited number of training epochs. As the learning rate gradually increases, the optimization process more easily captures effective discrimination patterns, thus improving performance. When the learning rate continues to increase to a higher range, accuracy begins to decline, reflecting the oscillations and instability caused by excessively large update steps. This makes the model more prone to repeatedly jumping around local optima, thus weakening the final generalization performance.

Learning rates in the middle range achieve more stable and higher accuracy. The configurations near the peak of the curve show similar performance, indicating that the model has a certain degree of robustness within a reasonable range, but remains sensitive to excessively small or large learning rates. This phenomenon suggests that the proposed method has a relatively suitable step size range for optimization. Within this range, gradient updates maintain effective progress while avoiding excessive perturbation, which is beneficial for forming more regular discrimination boundaries and reducing performance fluctuations caused by training instability.

The knowledge neighbor truncation number determines the scope of context that each target entity can aggregate on the knowledge side, thus affecting the coverage of structural information and the degree of noise introduction. Too small a truncation number may lead to the loss of key associations, while too large a truncation number may introduce irrelevant neighbors that interfere with target-level discrimination. This paper further presents Figure 4 to observe the overall trend of performance as the neighbor range changes.

When the truncation number is small, the model lacks sufficient information on related entities and relationships, resulting in incomplete background semantics of the target and weaker discrimination criteria. As the neighbor range gradually expands, more structural cues related to the target are incorporated, making it easier for the model to form stable target-level semantic aggregations, thus gradually improving accuracy. When the truncation number continues to increase to a higher range, the curve begins to decline, reflecting that too many neighbors can introduce noise and topic drift problems. Knowledge graphs contain many connections weakly related to the current context. When the neighbor range is too large, this irrelevant information dilutes the features most relevant to the target and may even introduce conflicting cues, making the

fused representation less focused and affecting the final discrimination effect. At this point, the model is not lacking information, but rather the information density decreases, and the benefits of structural constraints are offset by noise.

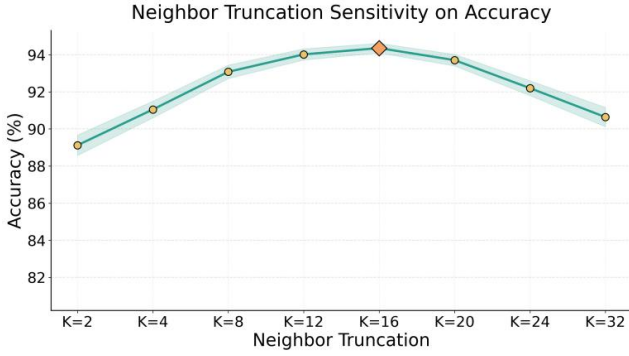


Figure 4. Sensitivity of Knowledge Neighbor Truncation Number to Accuracy

The region near the peak of the curve is relatively flat, indicating that within a reasonable truncation range, the model has a certain robustness to changes in the number of neighbors, and its performance does not fluctuate drastically with small adjustments. This means that gating fusion and aggregation strategies can, to some extent, suppress the perturbations caused by neighbor expansion, allowing knowledge-side information to participate more in a supplementary rather than dominant role in judgment. It also indicates that the main function of the truncation number is to control the balance between the coverage of knowledge input and noise, rather than simply increasing it as much as possible.

In summary, there is a more suitable intermediate range for the knowledge neighbor truncation number. Within this range, the model can obtain sufficient structural background to assist target localization and sentiment attribution while avoiding information dilution caused by irrelevant neighbors. This phenomenon also indirectly illustrates that the benefits of knowledge augmentation depend on the quality and quantity control of neighbor selection. A reasonable truncation strategy can maintain the relevance and compactness of structural information, thereby more stably improving the reliability of target-level sentiment judgment.

The size and proportion of the training set directly affect the coverage of the target and opinion combination that the model can learn, thereby changing the performance of target-level sentiment recognition ability under different data sufficiency levels. The experimental results are shown in Figure 5.

As the proportion of the training set increases, the distribution center gradually shifts upward, indicating that more abundant training samples can cover more combinations of targets and viewpoints, enabling the model to learn more stable discrimination patterns. At a smaller proportion, the violin shape is wider, and the dispersion is more obvious, reflecting that insufficient data amplifies sample differences and random fluctuations, and the model's adaptability to different sentence structures and expressions is not consistent enough, thus its performance is less stable.

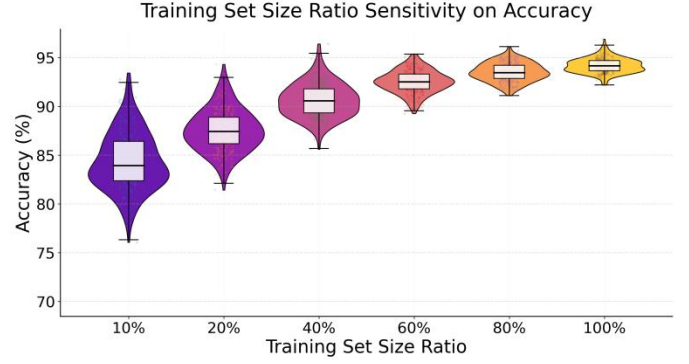


Figure 5. Sensitivity of training set size ratio to accuracy

When the proportion of the training set enters the mid-to-high range, the distribution converges significantly, the violin becomes narrower, and the box curve range becomes more concentrated, indicating that with sufficient data support, the model not only has a higher average level but also stronger stability. This convergence usually means that the parameter estimation is more reliable, the model's sensitivity to noisy samples and rare expressions decreases, and the decision boundary is more regular, thereby reducing fluctuations caused by sampling differences in the training data.

Approaching the full training data, the improvement rate tends to plateau, indicating that performance is gradually constrained by factors such as model capacity and task separability, and the marginal benefits of further increasing the data begin to diminish. However, from the perspective of distribution patterns, the most obvious benefit is reflected in the improvement of stability, especially the change from a low proportion to a medium proportion. This indicates that the key role of data scale in this task is not only to raise the overall level, but also to reduce uncertainty and enhance repeatability.

4. Discussion

While this method enhances target-level sentiment discrimination by fusing textual semantics and structured knowledge, its effectiveness depends to some extent on the coverage and quality of knowledge-side resources. If the knowledge graph does not adequately include domain entities and attributes, or if entity links contain numerous ambiguities and omissions, the information provided by knowledge branches may be weak, or even introduce noisy relationships inconsistent with the current context, causing the fused representation to deviate from the true semantics. Furthermore, different domains have significant differences in the granularity and attribute systems of targets. If the ontology design and relation type classification of the knowledge graph do not match the task requirements, structural information may be difficult to utilize effectively, thus limiting the stability of the method in cross-domain scenarios.

At the model level, while gating fusion can adjust the contribution ratio of semantics and knowledge to some extent, it still suffers from insufficient handling of contextual and structural conflicts. When textual expressions contain irony, implicit sentiment, or cross-sentence references, the semantic branches themselves may produce uncertain representations, and the neighbor relationships provided by knowledge

branches may not be consistent with the true viewpoint attribution. The combination of these two factors can lead to biases in target-level discrimination. Meanwhile, while neighbor truncation and simple aggregation ensure computational efficiency, they also mean that the graph information is compressed into coarse-grained statistical features, making it difficult to fully express complex multi-hop dependencies, event causality, and fine-grained attribute constraints. Therefore, there may still be room for improvement on samples requiring deep inference.

At the data and evaluation level, the labeling of opinion, target, and sentiment classification often suffers from subjectivity and boundary uncertainty, especially in cases involving target span, implicit targets, and mixed sentiments. Different labeling strategies may lead to inconsistencies, affecting the reliability of model training and comparison. Furthermore, common issues in training data, such as skewed category distribution and repetitive expression templates, may make the model better at handling high-frequency patterns but less generalizable to long-tail phenomena. Finally, for real-world applications, inference costs and deployment constraints must be considered. Knowledge retrieval and entity alignment introduce additional overhead and engineering complexity. In high-concurrency or low-latency scenarios, the benefits of structural enhancements need to be further weighed against system costs.

5. Conclusion

Addressing challenges such as diverse targets, implicit opinions, cross-sentence connections, and semantic ambiguity in real-world texts, it proposes a unified modeling approach combining knowledge graphs and large language models. The core of this approach lies in simultaneously considering semantic expressiveness and structural constraints. This allows the model to capture subtle emotional cues and rhetorical variations from context while leveraging background knowledge such as entity attributes and relationships to strengthen target localization and sentiment attribution, thus more closely resembling how users evaluate specific objects in natural language. By introducing structural information into the target-level discrimination process, this work provides a more interpretable and controllable technical path for fine-grained sentiment computing and lays the foundation for reliable text understanding in complex contexts.

Methodologically, this paper emphasizes achieving synergy between semantic and knowledge representation in a lightweight and reproducible manner, avoiding over-reliance on a single information source. The contextual representation provided by large language models can cover linguistic phenomena such as implicit opinions, negation, and degree modification, while the entity and relational structures carried by knowledge graphs provide external constraints for target disambiguation, attribute attribution, and semantic consistency. The integration of these two approaches allows target-level sentiment assessment to move beyond reliance solely on literal polarity and local patterns, achieving complementarity between structure and semantics and enhancing the model's adaptability to complex expressions. More importantly, this framework inherently supports the tracing and interpretation of decision-

making criteria, helping to improve the credibility of system outputs in demanding scenarios.

From an application perspective, opinion-target sentiment classification is a key component in public opinion monitoring, brand reputation management, intelligent customer service, and recommendation systems. It is also a crucial foundational capability in fields such as financial risk control, public governance, and medical service feedback analysis. The knowledge augmentation and large-scale model fusion approach explored in this paper helps advance sentiment understanding from coarse-grained overall judgments to actionable target-level insights, enabling the system to more accurately identify what users are discussing and why they hold that attitude, thereby supporting more precise decision-making and resource allocation. On the enterprise side, it can help pinpoint specific pain points and guide product optimization and service improvement; on the public service side, it can assist in more detailed attribution analysis of group opinions, improving response speed and governance accuracy; in risk-sensitive areas, it also provides structured support for more transparent automated text analysis, reducing potential costs caused by misjudgments.

Looking to the future, this research still has ample room for expansion. First, the reliability of knowledge acquisition and alignment can be further enhanced by introducing more robust entity linking and relationship filtering mechanisms, and by incorporating multi-source knowledge and dynamic knowledge updates into a unified process to adapt to rapidly changing domain terminology and the emergence of new concepts. Second, while maintaining the interpretability of the framework, more granular target attribute modeling and cross-sentence reasoning capabilities can be explored, enabling the model to possess stronger attribution consistency in long texts, implicit targets, and complex event contexts. Finally, to meet real-world deployment needs, the retrieval and inference chains can be optimized from efficiency and security perspectives, reducing inference overhead while ensuring stability and improving the output auditability mechanism. This will allow such technologies to better serve key application scenarios such as public opinion governance, corporate decision-making, and risk management, promoting the implementation and long-term value release of granular sentiment understanding in a wider range of fields.

References

- [1] G. Zheng, J. Wang, L. C. Yu, et al., "Instruction tuning with retrieval-based examples ranking for aspect-based sentiment analysis," in *Findings of the Association for Computational Linguistics: ACL 2024*, 2024, pp. 4777-4788.
- [2] X. Sun, K. Zhang, Q. Liu, et al., "Harnessing domain insights: A prompt knowledge tuning method for aspect-based sentiment analysis," *Knowledge-Based Systems*, vol. 298, p. 111975, 2024.
- [3] G. Lu, H. Yu, Z. Yan, et al., "Commonsense knowledge graph-based adapter for aspect-level sentiment classification," *Neurocomputing*, vol. 534, pp. 67-76, 2023.
- [4] W. Jin, B. Zhao, L. Zhang, et al., "Back to common sense: Oxford dictionary descriptive knowledge augmentation for aspect-based sentiment analysis," *Information Processing & Management*, vol. 60, no. 3, p. 103260, 2023.

- [5] T. Gu, H. Zhao, Z. He, et al., "Integrating external knowledge into aspect-based sentiment analysis using graph neural network," *Knowledge-Based Systems*, vol. 259, p. 110025, 2023.
- [6] W. Yin, C. Liu, Y. Xu, et al., "Synprompt: syntax-aware enhanced prompt engineering for aspect-based sentiment analysis," in *Proceedings of the 2024 Joint International Conference on Computational Linguistics, Language Resources and Evaluation (LREC-COLING 2024)*, 2024, pp. 15469-15479.
- [7] X. Bao, X. Jiang, Z. Wang, et al., "Opinion tree parsing for aspect-based sentiment analysis," in *Findings of the Association for Computational Linguistics: ACL 2023*, 2023, pp. 7971-7984.
- [8] S. Y. Huang, "Retrieval-Augmented Long-Text Reasoning with Semantic Conflict Detection and Cross-Paragraph Evidence Integration," 2024.
- [9] S. Chen, H. Qiu, H. Huang, and N. Eli, "Dynamic Low-Rank Routing for Efficient Multi-Task Instruction Fine-Tuning of Large Language Models," 2025.
- [10] K. Zeng, "Transformer-Based Structure-Aware Lung Cancer Image Segmentation with Multi-Scale Fusion and Boundary-Guided Prediction," 2024.
- [11] Y. Yang, C. Wen, H. Cui, Y. Sun, and Y. Liu, "Learning Unified Multi-Granularity Representations for Backend Anomaly Detection and Causal Localization," 2024.
- [12] Y. Xu, "Intelligent Supply Chain Risk Identification via Heterogeneous Graph Learning and Multi-Entity Interaction Modeling," 2025.
- [13] Y. Nie, "Financial Risk Identification Using Unified Multi-Source Feature Learning and Structural Aggregation," 2024.
- [14] R. Meng, "Hierarchical Communication and Computation Scheduling for Efficient Federated Learning in Cloud-Edge Collaborative Intelligent Systems," 2024.
- [15] Z. Pan, "Intelligent Decision-Making for Advertising Recommendation via Data Mining and Graph Neural Networks," 2024.
- [16] Z. Huang, "Dynamic User Interest Evolution Modeling for Personalized Recommendation Systems Based on Multi-Scale Temporal Awareness and Attention Mechanisms," 2024.
- [17] S. Sun, "Adaptive Stock Trading Algorithms Jointly Driven by Multi-Agent Collaborative Decision Making and Large Language Models in Complex Financial Market Environments," 2025.
- [18] Q. Zhong, L. Ding, J. Liu, et al., "Knowledge graph augmented network towards multiview representation learning for aspect-based sentiment analysis," *IEEE Transactions on Knowledge and Data Engineering*, vol. 35, no. 10, pp. 10098-10111, 2023.
- [19] Y. Tian, G. Chen, and Y. Song, "Aspect-based sentiment analysis with type-aware graph convolutional networks and layer ensemble," in *Proceedings of the 2021 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies*, 2021, pp. 2910-2922.
- [20] B. Liang, H. Su, L. Gui, et al., "Aspect-based sentiment analysis via affective knowledge enhanced graph convolutional networks," *Knowledge-Based Systems*, vol. 235, p. 107643, 2022.
- [21] H. Wu, C. Huang, and S. Deng, "Improving aspect-based sentiment analysis with knowledge-aware dependency graph network," *Information Fusion*, vol. 92, pp. 289-299, 2023.
- [22] J. Li, X. Li, L. Hu, et al., "Knowledge graph enhanced language models for sentiment analysis," in *International Semantic Web Conference*, Cham: Springer Nature Switzerland, 2023, pp. 447-464.
- [23] Z. Jian, D. Wu, S. Wang, et al., "AGCL: Aspect graph construction and learning for aspect-level sentiment classification," in *Proceedings of the 31st International Conference on Computational Linguistics*, 2025, pp. 841-854.
- [24] K. Zhang and Y. Han, "Harnessing Commonsense: LLM-Driven Knowledge Integration for Fine-Grained Sentiment Analysis," in *Proceedings of the 34th ACM International Conference on Information and Knowledge Management*, 2025, pp. 4106-4116.