

Weakly Supervised Aspect Category Detection Method Based on Label Hierarchical Filtering Mechanism

Yan Xiang, Anlan Zhang, Jinlei Shi, and Yuxin Huang*

Abstract

Aspect category detection is crucial for aspect-level sentiment analysis. Traditionally, supervised methods require a large number of labeled samples to achieve effectiveness, whereas unsupervised methods rely heavily on human judgment, which can compromise identification accuracy. To address these limitations, we propose a weakly supervised aspect category detection method that uses a hierarchical label filtering mechanism. Our approach begins by assigning preliminary pseudo-labels to comments based on topic similarity. Subsequently, unreliable pseudo-labeled samples are filtered out using semantic similarity. Finally, high-confidence training samples are selected using both high and low thresholds. Using the hierarchical label filtering mechanism employed in these three stages, we constructed a high-confidence training dataset, which was used to train a topic representation-enhanced classifier for aspect category detection. The proposed method, which was evaluated on three publicly available datasets, outperformed existing baseline models, thereby reducing the need for human intervention.

Keywords

Aspect Category Detection, Aspect-Level Sentiment Analysis, Label Filtering, Topic Model, Weakly Supervised Method

1. Introduction

With the rapid development of the Internet, an enormous number of online comments have accumulated. To analyze the opinions in these comments, aspect-based sentiment analysis (ABSA) [1] was proposed. Aspect category detection (ACD) [2], a foundational task of ABSA, aims to classify comments into predefined categories based on their key parts. For example, the aspect categories of the comment “The service is good although rooms are pretty expensive” are “service” and “price,” respectively. This categorization is essential for understanding user opinions.

Conventional supervised ACD methods depend on a substantial volume of labeled data, which can be time consuming and costly to obtain. Consequently, researchers have attempted to use unsupervised or weakly supervised methods to accomplish ACD tasks. One unsupervised method is topic modeling [3], which leverages word co-occurrence statistics to generate aspect keywords.

Another approach involves clustering-based models [4,5], which cluster relevant comments and match them to predefined aspect categories. However, even with the advancements in these models, a manual mapping process for aspect categories is still required, and the performance of category detection remains

※ This is an Open Access article distributed under the terms of the Creative Commons Attribution Non-Commercial License (<http://creativecommons.org/licenses/by-nc/3.0/>) which permits unrestricted non-commercial use, distribution, and reproduction in any medium, provided the original work is properly cited.

Manuscript received October 12, 2023; first revision January 26, 2024; second revision April 9, 2024; accepted April 22, 2024.

*Corresponding Author: Yuxin Huang (huangyuxin2004@163.com)

School of Information Engineering and Automation, Kunming University of Science and Technology, Kunming, China (50691012@qq.com, 2646196719@qq.com, 1185356966@qq.com, huangyuxin2004@163.com)

limited.

Given the limitations of these methods, weakly supervised methods have garnered extensive attention in recent years. A prevalent approach involves iterative training for ACD guided by seed words [3]. This approach focuses on learning the embedding space for sentences and seed words to establish similarities between sentences and aspects. However, aspect representations are often constrained by a fixed initial number of seed words. To expand the initial seed word set, some researchers have leveraged the vocabulary inherent to pretrained models [2]. However, relying solely on weakly supervised data for training can lead to suboptimal model performance depending on the limitations of the seed-word representations.

To overcome this challenge, an approach suggested implementing a pseudo-labeling strategy [6]. In this method, labeled data are used to train an initial model, which is then applied to unlabeled data to generate pseudo-labels, thereby expanding the training dataset [7]. However, filtering out noisy labels during this process is crucial. To address this issue, we propose a framework for selecting pseudo-labels in a stepwise manner. The primary contributions are summarized as follows:

- A weakly supervised ACD method that utilizes a hierarchical label-filtering mechanism is introduced. This method iteratively identifies high-confidence pseudo-labeled comments from the unlabeled data by combining topics and semantic similarities. Consequently, these high-quality pseudo-labels enhance the training of the ACD classifier.
- Following the acquisition of a high-confidence pseudo-labeled dataset, topic words are integrated into the bidirectional encoder representations from transformers (BERT) classifier to characterize the connections between comments and topic words. This process is crucial for predicting the aspect categories.
- In experiments on three benchmark datasets, the proposed method outperformed several competitive methods.

2. Related Works

2.1 Supervised Aspect Category Detection

ACD can be categorized into supervised, unsupervised, and weakly supervised methods, based on the availability of labels [5,8]. In the supervised approach, the task typically takes the form of sequence labeling in which methods such as conditional random fields and recurrent neural networks are employed. However, these models depend heavily on a substantial amount of domain-specific training [9]. To address the limited training data in the target domain, researchers have introduced cross-domain models. However, designing effective cross-domain models is particularly challenging, given the inherent differences between the source and target domains [10].

2.2 Unsupervised Aspect Category Detection

The introduction of unsupervised methods provides a solution to this issue. Earlier unsupervised ACD approaches relied heavily on the latent Dirichlet allocation (LDA) topic model [11-13] to generate word distributions for each aspect category using a Dirichlet prior. Wang et al. [14] introduced an augmented restricted Boltzmann machine (RBM) that integrated prior knowledge to jointly capture the general and sentiment aspects of comments. Subsequently, the aspect-based auto-encoder (ABAE) model [15] and its variants [16] were introduced, demonstrating significant performance improvements. These models capture word co-occurrence patterns to identify the coherent aspects in the text. Shi et al. [16] recently

applied self-supervised representation learning using contrastive algorithms to enhance aspect and comment segment representations, thereby achieving improved outcomes. Yang et al. [17] devised a set of aspect category experts by assigning each expert the responsibility of coding for a specific aspect category. The primary limitation of these methods is the need for manual intervention, such that the subjectivity of human judgment significantly affects the results.

2.3 Weakly Supervised Aspect Category Detection

In recent years, weakly supervised methods for ACD have attracted widespread attention owing to their ability to model meaningful aspects using only a small amount of domain knowledge. Handcrafted mapping rules [2] and seed-driven methods [5] have emerged as two prominent approaches to weakly supervised ACD. Huang et al. [18] leveraged seed words to generate aspect representations and used a convolutional neural network (CNN) model to align reviews with corresponding aspects. Nguyen et al. [19] introduced an aspect detection encoder that transformed comments and aspects into a low-dimensional embedding space. Tulkens and Van Cranenburgh [20] identified aspects by measuring the cosine similarity between pretrained aspect representations and label names. However, the effectiveness of seed words is often constrained, making it challenging to enhance the performance of these weakly supervised methods.

3. Methodology

3.1 Task Definition

Given a dataset with multiple aspect categories C , each category contains a few labeled comments and numerous unlabeled comments. Letting x denote one of the unlabeled comments and s_j^i denote the j -th labeled comment of the i -th category, with the entire dataset as $\{xx^i\}_{i=1}^C$, where $i = 1, 2, \dots, C$; $j = 1, 2, \dots, |s^i|$, $|s^i|$ is the number of labeled comments of category i . The task is to predict the aspect category of the unlabeled comments using the provided data.

3.2 Overview of the Model

The model consists of two main components. 1) Construction of the training set, which employs a hierarchical filtering mechanism. This process involves combining the topic and semantic similarity of comments to iteratively filter and identify high-confidence pseudo-labeled comments from unlabeled comments. 2) Training a topic-enhanced classifier based on the constructed training set, with comments and their corresponding topic terms concatenated and input into a BERT classifier to detect aspect categories. The overall workflow of the model is shown in Fig. 1.

3.3 Training Set Construction based on Hierarchical Filtering Mechanism

The training dataset was constructed in three steps:

- 1) Pseudo-labels were initially assigned to unlabeled comments based on their topic similarity values with labeled comments.
- 2) The probabilities of pseudo-labeled comments belonging to different categories were calculated based on their semantic similarity values with labeled comments.

3) High-confidence pseudo-labeled comments were selected for the training set using hierarchical filtering rules.

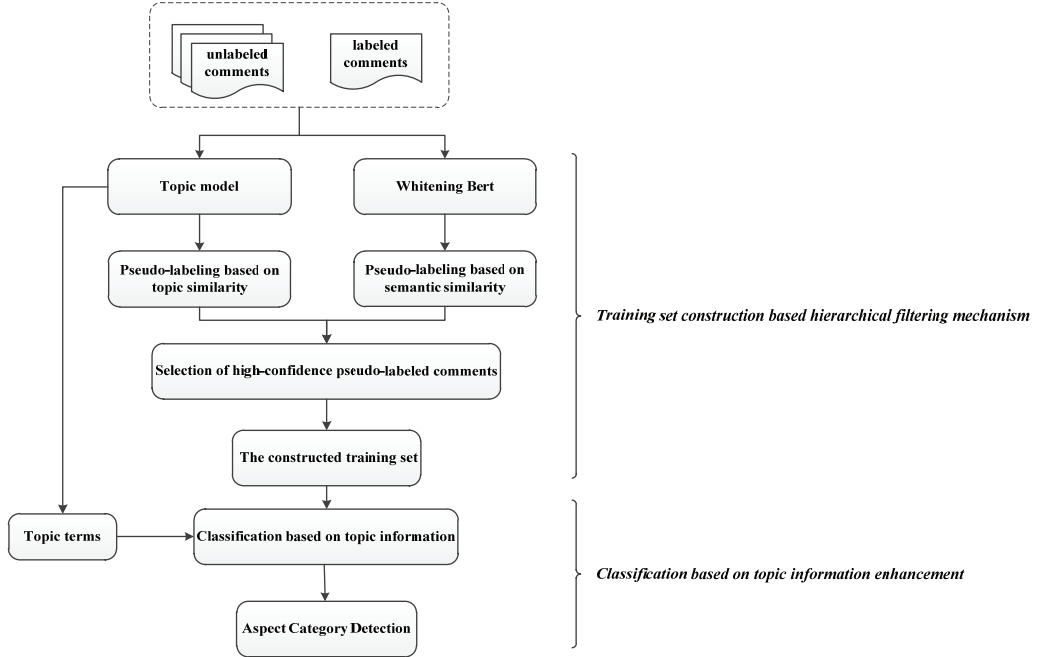


Fig. 1. Flow chart of the aspect category detection model.

3.3.1 Initial pseudo-labeling based on topic similarity

To assign pseudo-labels to unlabeled comments, we first apply the topic model ABAE [13] with the topic number K on the dataset $\{xx^i\}_{i=1}^C$, and obtain the topic distribution vector $p_i \in \mathbb{R}^K$ of the i -th comment, as well as its top m topic terms $\{w_j\}_{j=1}^m$. Subsequently the cosine similarity [21] of the topic distribution vector is used to measure the distance between unlabeled and labeled comments. Specifically, the topic similarity $sim_t(\cdot)$ of the two comments is calculated using Eq. (1):

$$sim_t(p_j^i, q) = \frac{p_j^i}{\|p_j^i\|} \cdot \frac{q}{\|q\|}, \quad (1)$$

where P_j^i is the topic probability vector of the j -th labeled comment belonging to category i , and q is the topic distribution vector of the unlabeled comments x . The similarity values between the unlabeled comment x and all the labeled comments in the i -th category are calculated and averaged to obtain the topic similarity $SimT^i$ of x belonging to the i -th category, as follows:

$$SimT^i = \frac{\sum_{j=1}^{|s^i|} sim_t(p_j^i, q)}{|s^i|}. \quad (2)$$

Finally, a pseudo-label is assigned to the unlabeled comment x according to the highest value of $SimT^i$. Based on the above process, the unlabeled comments are classified into different aspect categories.

3.3.2 Further pseudo-labeling based on semantic similarity

To eliminate the potential errors in the pseudo-labeled samples thus obtained, relatively reliable pseudo-labeled samples are filtered out by combining semantic similarities. Whitenning BERT [22] is first performed to calculate the semantic similarity values between the unlabeled comment x and all the labeled comments in the i -th category, which are averaged to obtain the similarity $SimB^i$ of x belonging to the i -th category, as follows:

$$SimB^i = \frac{\sum_{j=1}^{|s^i|} BERT - Whitenning(x, s_j^i)}{|s^i|}, \quad (3)$$

where $BERT - Whitenning$ represents the whitenning BERT operator, which involves a linear transformation of the output features from the BERT model. This operator reduces the redundancy among features and enhances the effectiveness of computing semantic similarity [22]. $SimB^i$ is then normalized to obtain the probability $Score^i$ that x belongs to category i , as follows:

$$Score^i = \frac{e^{SimB^i}}{1 + e^{SimB^i}}. \quad (4)$$

After assigning another pseudo-label to the unlabeled comment based on the highest value of $Score^i$, this new pseudo-label is compared with the initial label obtained through topic similarity in the previous stage. If the new pseudo-label is inconsistent with the initial label, it is considered unreliable and the corresponding pseudo-labeled comment is excluded from the training set. However, if the new pseudo-label is consistent with the initial label, we proceed to the next step for further evaluation.

3.3.3 Selection of pseudo-labeled comments with high confidence

Assuming that both the pseudo-labels obtained by the topic and semantic similarity are of category c , the confidence level of the pseudo-label is determined as follows, based on the probability $Score^c$ that x belongs to category c . Fig. 2 illustrates the entire filtering process.

- 1) If $Score^c$ is greater than the high threshold τ_h , it indicates that the label is reliable, and the comment x and its corresponding label c can be directly added to the training set.
- 2) If $Score^c$ is lower than the high threshold τ_h , we further consider whether the probability that the comment belongs to the other categories is lower than the threshold τ_l . If all the probabilities are indeed lower than τ_l , the comment can be added to the training set.

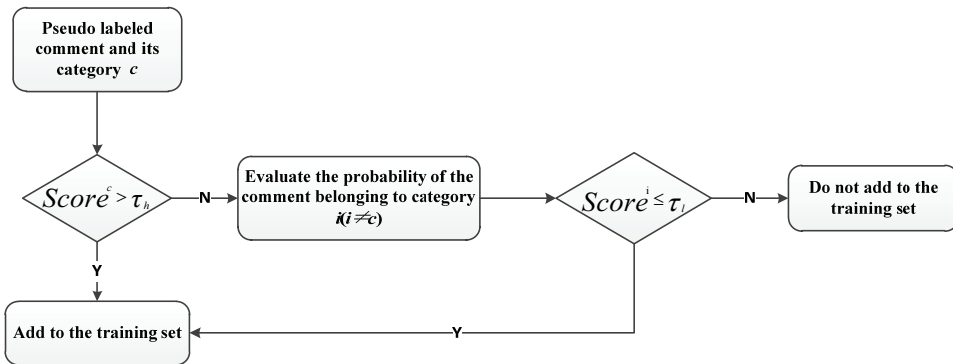


Fig. 2. Flowchart of selecting pseudo-labeled comments with high confidence.

Specifically, we used PN^c to indicate that comment x belongs to category c .

$$PN^c = \begin{cases} 1, & Score^c \geq \tau_h \\ \prod_{i=1}^c I[Score^i \leq \tau_l]_{i \neq c}, & Score^c < \tau_h \end{cases} \quad (5)$$

where $I(\cdot)$ is an indicator function that returns a value of one when the condition is satisfied. If $PN^c = 1$, then comments and their corresponding labels can be added to the training set. In contrast to previous methods, this approach considers both high and low probabilities, ensuring that the selected pseudo-labeled samples in the training set are both abundant and reliable. Through the aforementioned steps subsections 3.3.1 to 3.3.3, we construct a training set.

3.4 Classification based on Topic Information Enhancement

After obtaining a high-confidence training set, we trained the classifier using a fine-tuning method with a pretrained language model. To enhance the capacity of the model to represent aspect categories, we constructed a sequence by merging comments with their corresponding topic terms. The sequence was encoded using the BERT model [23]. The input was as follows: "[CLS] + Comment + [SEP] + Topic Terms + [SEP]". Here, "[CLS]" serves as a special identifier, and "[SEP]" is used as a separator to distinguish comments from topic terms. By adopting this approach, interactions are established between individual words in comments and the corresponding topic terms. Consequently, hidden vector h_j of the [CLS] output from the last layer serves as a characteristic representation of the j -th comment. Finally, the characteristic representation h_j of the comment is sent to a fully connected layer and a softmax layer to obtain the classification result.

$$\hat{y}_j = \text{softmax}(\mathbf{W}h_j + b), \quad (6)$$

where $\mathbf{W}$ and b denote the learnable weight and bias, respectively; $\hat{y}_j \in \mathbb{R}^C$ is the predicted probability that the j -th comment belongs to different aspect categories; C is the number of classes. The training loss can be expressed as:

$$\mathcal{L} = -\frac{1}{N} \sum_{j=1}^N \sum_{i=1}^C y_{ji} \log \hat{y}_{ji}, \quad (7)$$

where N represents the number of training samples; C denotes the number of aspect categories; y_{ji} is the true label probability that the j -th sample belongs to the i -th category; and $\hat{y}_{ji}$ is the probability predicted by the model that the j -th sample belongs to the i -th category.

4. Experiments

4.1 Dataset

Three datasets were used in the experiments: Restaurant, Bags, and Keyboards. The Restaurant dataset (<https://huggingface.co/datasets/Charitharth/SemEval2014-Task4>) is from Citysearch New York restaurant comments and the latter two datasets are from Amazon product comments (https://cseweb.ucsd.edu/~jmcauley/datasets/amazon_v2/). Table 1 presents detailed statistical results on these datasets.

Table 1. Statistics on three datasets

Aspect categories	Vocabulary number	Number of unlabeled comments	Number of labeled comments
Food / Staff / Ambience	45023	52574	1490
Looks / Price / Quality	6438	584332	641
Layout / Looks / Noise / Price	6904	603379	681

4.2 Experimental Settings

Five labeled samples were randomly selected for each category; punctuation, stop words, and words that occurred fewer than 10 times in the dataset were removed. We used ABAE [15] as the topic model and set the number of topics to 14. In the pseudo-label screening, the high threshold τ_h was set to 0.5 and the low threshold τ_l was set to 0.15. The classifier and BERT-whitening were configured with the “BERT-base” setting [23], with a hidden vector dimension of 768. During training, the batch size was 250 and optimization was conducted using *adaptive moment estimation* (Adam) as the optimizer with a learning rate of 0.01. The model underwent 15 update iterations and a dropout layer was incorporated to prevent overfitting. The experimental environment was Windows 10, equipped with an Intel Core i5-6300HQ @2.3 GHz, 32 GB RAM, and NVIDIA GeForce GTX 960M GPU. Python 3.7 was the programming language, and the programming framework was PyTorch 1.11.

4.3 Baseline Models

SERBM: Sentiment-aspect extraction based on RBM [14] jointly extracts comment aspects and sentiment polarity in an unsupervised manner based on RBM. In this model, the hidden vector dimension was set to 10.

W2VLDA: Word2vec-based LDA [11] is a topic-based modeling approach that combines word embeddings with LDA. It automatically aligns discovered topics with predefined aspect names by utilizing seed words provided by the user for different aspects.

ABAE [15] is an unsupervised neural topic model that employs neural word embedding to acquire coherent aspects. The model parameters were optimized using the Adam optimization technique with a learning rate of 0.001. The experimental process encompassed 15 iterations, and a batch size of 50 was employed to expedite training.

ABAE_{init} [24] modifies the aspect embedding vectors in ABAE by substituting them with the centroid of the relevant seed word embeddings and fixing the aspect embedding vector during training. The experimental parameters were set in line with those of ABAE.

LDA-Anchors [25] uses the latent Dirichlet to obtain topic word distributions via seed words.

MATE: The multi-seed aspect extractor [24] is a weakly supervised neural model that combines a seed aspect extractor trained under a multi-task objective with a multi-instance learning sentiment predictor to identify and extract useful comments.

MATE+MT: The multi-task (MT) counterpart of MATE [24] is a weakly supervised auto-encoder extension of the ABAE model that initializes the aspect-embedding matrix through aspect-specific seed words and subsequently refines them during training.

TS-*: Teacher-student (TS-*) [26] is a weakly supervised training framework in which the teacher network is a seed-word-based bag-of-words classifier, and the student network uses word2vec embedding and the BERT model to encode text fragments.

AE-CSA: Aspect extraction via context-enhanced Sememe attentions [27] is an unsupervised neural

framework that enhances lexical semantics using semantic elements. The entire framework is similar to that of an auto-encoder reconstructing sentence representations and learning aspects using latent variables.

CAT: Continual adapter tuning [2] was introduced as a basic heuristic model that encompasses comparative attention and an automated aspect assignment method.

4.4 Experimental Results and Analysis

4.4.1 Comparison with the baseline models

The experimental results of the different models on the Bags, Keyboards, and Restaurant datasets are listed in Table 2, where we observe that: 1) compared with the suboptimal models on the three datasets, the Macro-F1 score of the proposed model increased by 0.7%, 0.6%, and 2.1%, respectively, proving that our model can compete with other baseline models. 2) Unsupervised models, such as ABAE and SERBM, which do not utilize labeled data, require manual judgment of the correspondence between aspect words and aspect categories. This significantly affects performance. 3) Weakly supervised models such as MATE and MATE-MT surpass unsupervised ABAE by leveraging aspect words for more accurate aspect category predictions, suggesting that the topic information provided by aspect words is important for ACD. Based on these observations, it can be inferred that our model achieves substantial enhancements by leveraging semantic and thematic similarities to obtain more relevant samples, filtering out the low-confidence samples based on the confidence threshold.

Table 2. Macro-F1 score (%) of different models

Model	Macro-F1 score (%)		
	Bags dataset	Keyboards dataset	Restaurant dataset
ABAE	38.1	38.6	77.5
SERBM	-	-	74.5
W2VLDA	-	-	72
AE-CSA	-	-	81.9
CAt	-	-	82.5
LDA-Anchors	33.5	34.7	-
ABAE _{init}	41.6	41.3	-
MATE	46.2	43.5	-
MATE-MT	48.6	45.3	-
TS-Teacher	55.1	52.0	-
TS-Stu-W2V	59.3	57.0	-
TS-Stu-BERT	61.4	57.5	-
Proposed	62.1	58.1	84.6

4.4.2 Effectiveness of the label hierarchical filtering mechanism

To evaluate the impact of the hierarchical label filtering mechanism on the model, we performed ablation experiments on the Restaurant dataset. We compared the results of two approaches: the model utilizing initial pseudo-labeling based on topic similarity and the model employing the complete label filtering mechanism. The F1 scores for both methods are shown in Fig. 3.

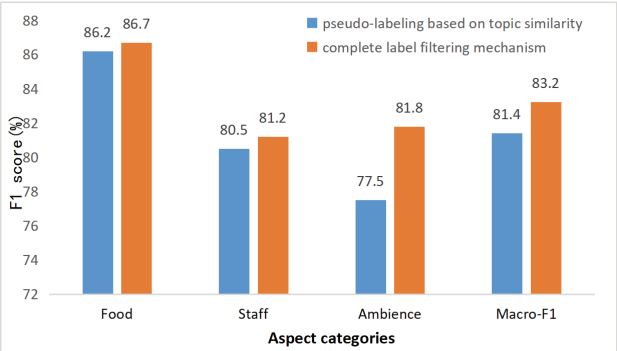


Fig. 3. Experimental results of the hierarchical filtering mechanism.

The hierarchical filtering approach resulted in noticeable improvements in specific categories, specifically, a 0.5% improvement for the “Food” category, a 0.7% improvement for the “Staff” category, and a substantial 4.2% improvement for the “Ambience” category. Overall, the stepwise filtering mechanism led to a 1.8% improvement in the Macro-F1 score compared with the method that relied exclusively on topic similarity for pseudo-label filtering. This observation suggests that certain incorrect pseudo-labeled comments annotated by the topic modeling process exert a substantial influence on model performance. Applying hierarchical filtering and eliminating some of these inaccurate comments, increases the reliability of the training samples, thereby improving model performance.

4.4.3 Effect of the number of labeled samples on the model

In this section, we describe the ablation experiments performed on the Restaurant dataset to assess the impact of the amount of labeled data on the model. The experimental setup involved randomly selecting 5, 10, and 15 labeled samples for each category. After determining the training set using a label hierarchical filtering mechanism, we trained the BERT classifier on topic enhancement. The results of these experiments are presented in Fig. 4.

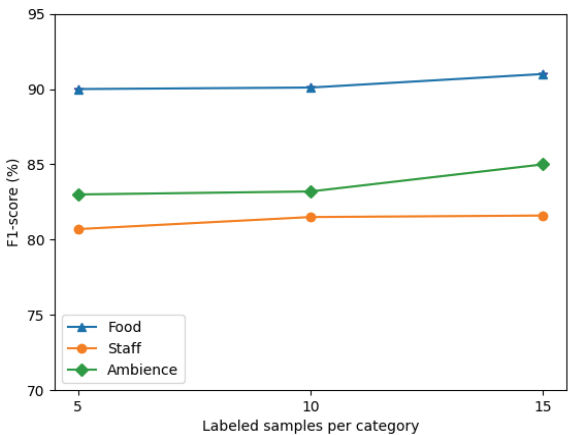


Fig. 4. Experimental results on the number of labeled samples.

In Fig. 4, the performance increases to varying degrees for all three aspect categories by adding more labeled data. Notably, the most substantial improvement is observed for the “Ambience” category. This

indicates that introducing more labeled samples with classification information enables the model to identify additional pseudo-labeled comments related to specific aspect categories, thereby enhancing overall performance.

4.4.4 Validity of topic information

Ablation experiments were performed on the Restaurant dataset to assess the influence of topic information on the classifier. We evaluated the model performance both with and without topic information augmentation by BERT, using the same constructed training set.

In Fig. 5, a performance improvement of 1.8% is observed when topic information is incorporated into each comment sentence. This topic information directs the model's attention toward aspects related to specific themes. The inclusion of specific thematic information helps the model better understand and differentiate the meanings of aspects, thereby enhancing its ability to identify aspect categories.

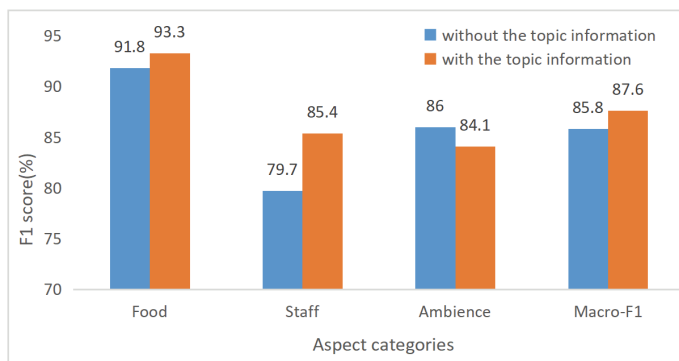


Fig. 5. Experimental results on the validity of topic information.

5. Conclusion

In this paper, we propose a weakly supervised ACD model that includes the following steps. First, limited labeled data are utilized to assign pseudo-labels to extensive unlabeled data by employing topic-based similarity. Second, filtering rules are applied to select comments labeled with high confidence. Finally, the ACD classifier enhances its ability to identify aspect categories by integrating a high-quality pseudo-labeled dataset with topic terms. The experimental results demonstrated that the proposed model achieves more accurate ACD than existing models and reduces the subjectivity introduced by human intervention. Ablation experiments demonstrated the effectiveness of the proposed hierarchical label-filtering mechanism. Moreover, incorporating thematic information is useful because it directs the attention of the classifier to words relevant to the aspect categories. In this study, BERT-whitening was employed for semantic similarity computations. In future studies, we plan to synergize the understanding capabilities of large language models with hierarchical filtering mechanisms to achieve enhanced performance.

Conflict of Interest

The authors declare that they have no competing interests.

Funding

This work was supported by the National Natural Science Foundation of China (Grant Nos. 62162037, 62266027, U21B2027, 62266028, and 62241604), the Yunnan Provincial Major Science and Technology Special Plan Projects (Grant No. 202302AD080003), and the General Projects of Basic Research in Yunnan Province (Grant No. 202301AT070444).

References

- [1] M. Pontiki, D. Galanis, H. Papageorgiou, I. Androutsopoulos, S. Manandhar, "SemEval-2014 Task 4: aspect based sentiment analysis," in *Proceedings of the 8th International Workshop on Semantic Evaluation (SemEval2014)*, Dublin, Ireland, 2014, pp. 27-35. <https://doi.org/10.3115/v1/S14-2004>
- [2] Y. Li, P. Wang, Y. Wang, Y. Dai, Y. Wang, L. Pan, and Z. Xu, "Leveraging only the category name for aspect detection through prompt-based constrained clustering," in *Findings of the Association for Computational Linguistics: EMNLP 2022*. Stroudsburg, PA: Association for Computational Linguistics, 2022, pp. 1352-1364. <https://doi.org/10.18653/v1/2022.findings-emnlp.97>
- [3] B. Ozyurt and M. A. Akcayol, "A new topic modeling based approach for aspect extraction in aspect based sentiment analysis: SS-LDA," *Expert Systems with Applications*, vol. 168, article no. 114231, 2021. <https://doi.org/10.1016/j.eswa.2020.114231>
- [4] S. Xiong and D. Ji, "Exploiting flexible-constrained K-means clustering with word embedding for aspect-phras grouping," *Information Sciences*, vol. 367, pp. 689-699, 2016. <https://doi.org/10.1016/j.ins.2016.07.002>
- [5] S. U. S. Chebolu, P. Rosso, S. Kar, and T. Solorio, "Survey on aspect category detection," *ACM Computing Surveys*, vol. 55, no. 7, article no. 132, 2022. <https://doi.org/10.1145/3544557>
- [6] Z. Liu, Z. Ma, J. Wang, J. Li, J. Sun, and C. Bao, "SPCPFS: a pseudo-label filtering strategy with fusion of perplexity and confidence," in *Proceedings of the 2022 5th International Conference on Machine Learning and Natural Language Processing*, Sanya, China, 2022, pp. 383-389. <https://doi.org/10.1145/3578741.3578833>
- [7] D. Mekala and J. Shang, "Contextualized weak supervision for text classification," in *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*, Virtual Event, 2020, pp. 323-333. <https://doi.org/10.18653/v1/2020.acl-main.30>
- [8] W. Zhang, X. Li, Y. Deng, L. Bing, and W. Lam, "A survey on aspect-based sentiment analysis: tasks, methods, and challenges," *IEEE Transactions on Knowledge and Data Engineering*, vol. 35, no. 11, pp. 11019-11038, 2023. <https://doi.org/10.1109/TKDE.2022.3230975>
- [9] R. Seoh, I. Birlle, M. Tak, H. S. Chang, B. Pinette, and A. Hough, "Open aspect target sentiment classification with natural language prompts," in *Proceedings of the 2021 Conference on Empirical Methods in Natural Language Processing*, Virtual Event, Punta Cana, Dominican Republic, 2021, pp. 6311-6322. <https://doi.org/10.18653/v1/2021.emnlp-main.509>
- [10] F. Zhao, Y. Shen, Z. Wu, and X. Dai, "Label-driven denoising framework for multi-label few-shot aspect category detection," in *Findings of the Association for Computational Linguistics: EMNLP 2022*. Stroudsburg, PA: Association for Computational Linguistics, 2022, pp. 2390-2402. <https://doi.org/10.18653/v1/2022.findings-emnlp.177>
- [11] A. Garcia-Pablos, M. Cuadros, and G. Rigau, "W2VLDA: almost unsupervised system for aspect based sentiment analysis," *Expert Systems with Applications*, vol. 91, pp. 127-137, 2018. <https://doi.org/10.1016/j.eswa.2017.08.049>

- [12] L. Zheng, Z. He, and S. He, "A novel probabilistic graphic model to detect product defects from social media data," *Decision Support Systems*, vol. 137, article no. 113369, 2020. <https://doi.org/10.1016/j.dss.2020.113369>
- [13] J. Cheevaprawatdomrong, A. Schofield, and A. T. Rutherford, "More than words: collocation retokenization for latent Dirichlet allocation models," in *Findings of the Association for Computational Linguistics: ACL 2022*. Stroudsburg, PA: Association for Computational Linguistics, 2022, pp. 2696-2704. <https://doi.org/10.18653/v1/2022.findings-acl.212>
- [14] L. Wang, K. Liu, Z. Cao, J. Zhao, and G. De Melo, "Sentiment-aspect extraction based on restricted Boltzmann machines," in *Proceedings of the 53rd Annual Meeting of the Association for Computational Linguistics and the 7th International Joint Conference on Natural Language Processing (Volume 1: Long Papers)*, Beijing, China, 2015, pp. 616-625. <https://doi.org/10.3115/v1/P15-1060>
- [15] R. He, W. S. Lee, H. T. Ng, and D. Dahlmeier, "An unsupervised neural attention model for aspect extraction," in *Proceedings of the 55th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, Vancouver, Canada, 2017, pp. 388-397. <https://doi.org/10.18653/v1/P17-1036>
- [16] T. Shi, L. Li, P. Wang, and C. K. Reddy, "A simple and effective self-supervised contrastive learning framework for aspect detection," *Proceedings of the AAAI Conference on Artificial Intelligence*, vol. 35, no. 15, pp. 13815-13824, 2021. <https://doi.org/10.1609/aaai.v35i15.17628>
- [17] L. Yang, H. Li, L. Li, C. Yuan, and S. T. Xia, "ACEs: unsupervised multi-label aspect detection with aspect-category experts," in *Proceedings of 2022 International Joint Conference on Neural Networks (IJCNN)*, Padua, Italy, 2022, pp. 1-8. <https://doi.org/10.1109/IJCNN55064.2022.9892431>
- [18] J. Huang, Y. Meng, F. Guo, H. Ji, and J. Han, "Weakly-supervised aspect-based sentiment analysis via joint aspect-sentiment topic embedding," *Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing (EMNLP)*, Virtual Event, 2020, pp. 6989-6999. <https://doi.org/10.18653/v1/2020.emnlp-main.568>
- [19] T. N. Nguyen, K. H. Nguyen, Y. I. Song, and T. D. Cao, "An uncertainty-aware encoder for aspect detection," in *Findings of the Association for Computational Linguistics: EMNLP 2021*. Stroudsburg, PA: Association for Computational Linguistics, 2021, pp. 797-806. <https://doi.org/10.18653/v1/2021.findings-emnlp.69>
- [20] S. Tulkens and A. Van Cranenburgh, "Embarrassingly simple unsupervised aspect extraction," in *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*, Virtual Event, 2020, pp. 3182-3187. <https://doi.org/10.18653/v1/2020.acl-main.290>
- [21] E. Yu, L. Du, Y. Jin, Z. Wei, and Y. Chang, "Learning semantic textual similarity via topic-informed discrete latent variables," in *Proceedings of the 2022 Conference on Empirical Methods in Natural Language Processing*, Abu Dhabi, United Arab Emirates, 2022, pp. 4937-4948. <https://doi.org/10.18653/v1/2022.emnlp-main.328>
- [22] W. Yin and L. Shang, "Efficient nearest neighbor emotion classification with BERT-whitening," in *Proceedings of the 2022 Conference on Empirical Methods in Natural Language Processing*, Abu Dhabi, United Arab Emirates, 2022, pp. 4738-4745. <https://doi.org/10.18653/v1/2022.emnlp-main.312>
- [23] J. Devlin, M. W. Chang, K. Lee, and K. Toutanova, "BERT: pre-training of deep bidirectional transformers for language understanding," in *Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 1 (Long and Short Papers)*, Minneapolis, MN, USA, 2019, pp. 4171-4186. <https://doi.org/10.18653/v1/N19-1423>
- [24] S. Angelidis and M. Lapata, "Summarizing opinions: aspect extraction meets sentiment prediction and they are both weakly supervised," in *Proceedings of the 2018 Conference on Empirical Methods in Natural Language Processing*, Brussels, Belgium, 2018, pp. 3675-3686. <https://doi.org/10.18653/v1/D18-1403>
- [25] J. Lund, C. Cook, K. Seppi, and J. Boyd-Graber, "Tandem anchoring: a multiword anchor approach for interactive topic modeling," in *Proceedings of the 55th Annual Meeting of the Association for Computational*

Linguistics (Volume 1: Long Papers), Vancouver, Canada, 2017, pp. 896-905. <https://doi.org/10.18653/v1/P17-1083>

- [26] G. Karamanolakis, D. Hsu, and L. Gravano, “Leveraging just a few keywords for fine-grained aspect detection through weakly supervised co-training,” in *Proceedings of the 2019 Conference on Empirical Methods in Natural Language Processing and the 9th International Joint Conference on Natural Language Processing (EMNLP-IJCNLP)*, Hong Kong, China, 2019, pp. 4611-4621. <https://doi.org/10.18653/v1/D19-1468>
- [27] L. Luo, X. Ao, Y. Song, J. Li, X. Yang, Q. He, and D. Yu, “Unsupervised neural aspect extraction with Sememes,” in *Proceedings of the 28th International Joint Conference on Artificial Intelligence (IJCAI-19)*, Macau, China, 2019, pp. 5123-5129. <https://doi.org/10.24963/ijcai.2019/712>



Yan Xiang <https://orcid.org/0000-0002-6475-638X>

She is an associate professor at the Faculty of Information Engineering and Automation at the Kunming University of Science and Technology. She graduated with a B.E. degree in engineering and an M.S. degree in science from Wuhan University, China. She presided over one general basic research project in the Yunnan Province, one Yunnan Provincial Department of Education project, and one of the special plan sub-projects of the Yunnan Provincial Major Science and Technology. She has published over 20 papers as first or corresponding author. Her primary research interests include text mining and sentiment analysis.



Anlan Zhang <https://orcid.org/0009-0005-2247-3153>

She is an M.S. candidate at the Faculty of Information Engineering and Automation, Kunming University of Science and Technology. Her research interests include natural language processing and sentiment analysis.



Jinlei Shi <https://orcid.org/0000-0002-6712-998X>

He is an M.S. candidate at the Faculty of Information Engineering and Automation, Kunming University of Science and Technology. His research interests include text mining and sentiment analysis.



Yuxin Huang <https://orcid.org/0000-0003-1277-6212>

He is currently an associate professor at Kunming University of Science and Technology, China. He received the Ph.D. degree from Kunming University of Science and Technology, China, in 2021. His research interests include information retrieval and text generation.