

Dynamic Order Template Prediction for Generative Aspect-Based Sentiment Analysis

Anonymous ACL submission

Abstract

Aspect-based sentiment analysis (ABSA) assesses sentiments towards specific aspects within texts, resulting in detailed sentiment tuples. Previous ABSA models often use static templates to predict all of the elements in the tuples, and these models often fail to accurately capture dependencies between elements. Multi-view prompting method improves the performance of ABSA by predicting tuples with various templates and then ensembling the results. However, this method suffers from inefficiencies and out-of-distribution errors. In this paper, we propose a Dynamic Order Template (DOT) method for ABSA, which dynamically generates necessary views for each instance based on instance-level entropy. Ensuring the diverse and relevant view generation, our proposed method improves F1-scores on ASQP and ACOS datasets while significantly reducing inference time.

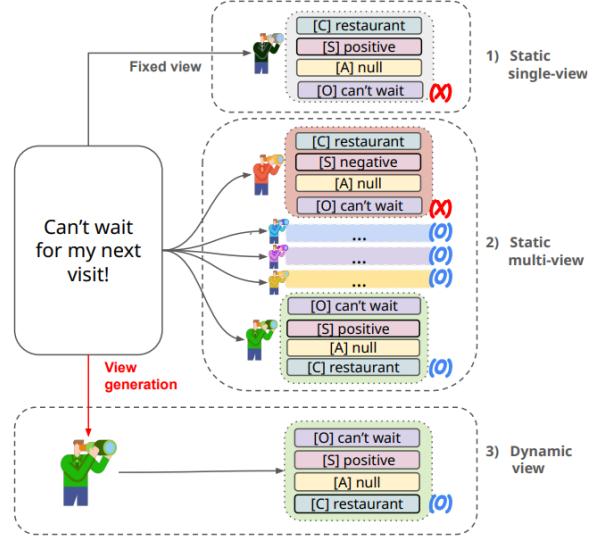


Figure 1: Comparison of three different generative ABSA methods. 1) static single-view (T5-paraphrase), 2) static multi-view (MvP), and 3) dynamic-view prediction (ours).

1 Introduction

Aspect-based sentiment analysis (ABSA) aims to identify sentiment for the aspects in a given text rather than simply classifying the overall sentiment of the entire text. ABSA research evolves to generate quadruples consisting of four elements: 1) Aspect (*A*), 2) Category (*C*) for the type of *A*, 3) Opinion (*O*) for *A*, and 4) Sentiment (*S*) for *A*. Many recent studies such as T5-paraphrase (Zhang et al., 2021c) tackle this problem using generative models. These approaches usually get review sentences as input and output the span of quadruples as fixed order paraphrased form, such as "*C* is *S* because *A* is *O*" (Zhang et al., 2021a). However, this static single-order template cannot express the dependence between elements as in Figure 1 due to the autoregressive nature of transformer (Vaswani et al., 2017). Moreover, the model's output can heavily depend on the order of generating each element (Hu et al., 2022a).

Multi-view prompting (Gou et al., 2023) (MvP) deals with this issue by constructing order templates as a channel for "viewing" different perspectives in a sentence. As shown in Figure 1, MvP permutes all possible element orders and sorts them based on the dataset-level entropy of the pre-trained model. Using this entropy, MvP samples top-k orders and adds these orders as a prompt template. During the inference time, MvP conducts majority votes on generated sentiment tuples with various templates. Through this ensemble approach, MvP utilizes the intuition of solving problems from different views in human reasoning and decision (Stanovich and West, 2000), resulting in significantly higher performance. However, we find that this static multi-view approach of MvP has several drawbacks: 1) *Inefficient*: Even for samples where the answer can be easily found and multiple views are not necessary, this method generates the same number of views, resulting in unnecessary computation that increases the inference time. 2)

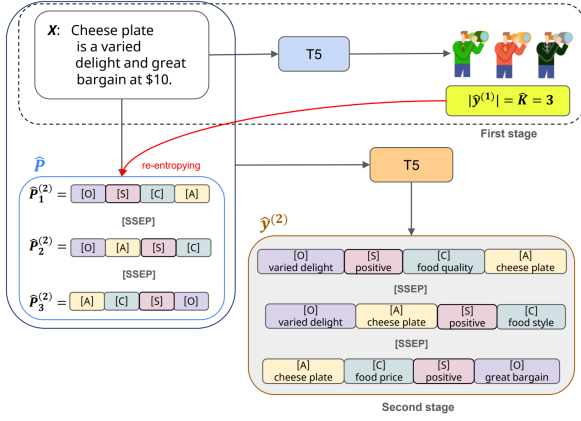


Figure 2: Overview of our proposed two stage method. We use two T5 models for each stage: one for generating views, the other for generating sentiment tuples.

Prone to Distribution shift: MvP uses the number of views k as a hyperparameter, applying the same k value across all datasets during training and inference. However, since the optimal number of ensemble models varies according to the data distribution, it requires manual adjustment of the k value for each dataset (Shahhosseini et al., 2022), which hinders the transferability to other datasets.

To resolve the aforementioned shortcomings, we propose a Dynamic Order Template (DOT) method for ABSA that combines the advantages of both single-view and multi-view approaches. By prioritizing multiple views based on instance-level entropy, DOT aims to generate only the necessary number of views for each instance during inference. For an example that contains only one tuple as in Figure 1, DOT dynamically creates only one order template (i.e. view) that is necessary for predicting the tuple. After generating the views, DOT generates tuples using the order templates that correspond to the views. This phase operates in a multi-view manner, enabling us to retain the benefits of previous multi-view methods.

Extensive experiments on five widely used sentiment quadruple prediction datasets, derived from ASQP (Pontiki et al., 2016; Zhang et al., 2022) and ACOS (Cai et al., 2021, 2023), demonstrate that our method show state-of-the arts performance with significantly lower inference time compared to previous multi-view inference work. Moreover, we show that our method is robust to distribution shift compared to previous methods.

2 Method

Our proposed Dynamic Order Template (DOT) method is composed of two stages as in Figure 2.

The first stage predicts the number of order templates to be used as a prompt. (Sec 2.1) Using the order templates, second stage predicts the sentiment tuples. We train the model to generate distinct tuples for each order template, enabling efficient aggregation. (Sec 2.2) For both stages, we map sentiment tuples (A, C, S, O) to marker tokens $[A]$, $[C]$, $[S]$, and $[O]$ respectively. For the cases of the instances that contain multiple sentiment tuples, we indicate each tuple with the respective tokens and concatenate the targets with $[SSEP]$ tokens.

2.1 Stage 1: Generating the Number of Views

We assume that the number of sentiment tuples present for each instance corresponds to the required number of views. In other words, we consider that one separate view is necessary for predicting each tuple. We define this individual view as the prediction order for each element of the sentiment tuple as shown in Figure 2. This allows each prediction order to correspond one-to-one with a sentiment tuple in second stage. We set the required number of sentiment tuples K_i in the i^{th} instance.

To predict the number of tuples K_i , we begin by examining the entropy of all possible views generated through permutations. Specifically, we calculate entropy of each view v in instance-level with vanilla T5 (Hu et al., 2022a) by calculating conditional generation probability as follows:

$$\mathcal{E}_{i,v} = p(y_{i,v}|x_i), \quad (1)$$

where $E_{i,v}$ denotes the entropy score of v in the permutation set for i^{th} instance. $y_{i,v}$ is a permuted sentiment tuple in i^{th} instance, based on v . At this time, we note that existing ABSA methods have generally struggled to extract O (Chebolu et al., 2023). Based on this observation, we consider that O might hinder the calculation of entropy score. Hence, we calculate entropy scores for v without O (i.e. (A, C, S)) in this stage. We provide a detailed analysis on excluding $[O]$ in Appendix C.2. After computing the entropy, we sort the views by the entropy in ascending order and map each element into marker tokens to get the ranked view $P_i^{(1)}$. And then we sample top K_i views for each sample X_i and concatenate these views as an order template.

Using the original input sentence X_i , we train the T5 model to generate the first-stage target $y_i^{(1)}$. The format of $y_i^{(1)}$ is as follows:

$$y_i^{(1)} = P_{i,1}^{(1)} [SSEP] P_{i,2}^{(1)} [SSEP] \dots P_{i,K_i}^{(1)},$$

where $P_{i,K_i}^{(1)}$ denotes the order template for the j^{th} order template of i^{th} instance in the first stage. We set the loss function to train the T5 model as in Equation (2), where $|B|$ denotes the batch size of the model. The scaling factor is omitted for simplicity.

$$\mathcal{L}_1 = - \sum_{i=1}^{|B|} \sum_{t=1}^T \log p(\mathbf{y}_{i,t}^{(1)} | \mathbf{x}_i, \mathbf{y}_{i,<t}^{(1)}) \quad (2)$$

2.2 Stage 2: Sentiment Tuple Generation

In the second stage, the model predicts the sentiment tuple of given instance using the number of generated views in the first stage. Different from the first stage, we need to generate all elements in sentiment quadruples including O in this stage. Hence, we re-rank all views to pick K_i views including O (i.e. (A, C, S, O)). Here, we use the same ranking strategy as in the first stage using the entropy and denote ranked view set as $P_i^{(2)}$. We sample top K_i views from ranked view and these views are mapped into sentiment tuples one by one, making the model to learn different order template should generate different tuples. We then concatenate the sampled views and add them as a prompt P_i to original input sentence X_i . We design the second stage target $y_i^{(2)}$ by placing the corresponding element next to each marker token [SSEP] within P_i as follows.

$$P_i = P_{i,1}^{(2)} [\text{SSEP}] P_{i,2}^{(2)} [\text{SSEP}] \dots P_{i,K_i}^{(2)},$$

$$y_i^{(2)} = P_{i,1}^{(2)} \text{tuple}_1 [\text{SSEP}] \dots P_{i,K_i}^{(2)} \text{tuple}_{K_i},$$

where $P_{i,K_i}^{(2)}$ represents the order template for the j^{th} order template of i^{th} instance in the second stage, and tuple_j is the sentiment tuple correspond to $P_{i,K_i}^{(2)}$. We design the loss function for training another T5 model in second stage as follows.

$$\mathcal{L}_2 = - \sum_{i=1}^{|B|} \sum_{t=1}^T \log p(\mathbf{y}_{i,t}^{(2)} | \mathbf{x}_i, P_i, \mathbf{y}_{i,<t}^{(2)}) \quad (3)$$

2.3 Two-stage inference

During inference time, two stages are conducted sequentially. In the first stage, we generate the necessary views $y^{(1)}$ based on the predicted number of views $\hat{K}$ (i.e. the number of tuples). Given $\hat{K}$, we sample top $\hat{K}$ orders from ranked order set $\hat{P}$ and construct order template $\hat{P}$. Note that we rank orders in full elements for each instance. Lastly, we directly add $\hat{P}$ to inference sentence and generate targeted sentiment tuples in second stage.

3 Experiment

3.1 Benchmark Datasets

We adopt two widely used ABSA datasets: ASQP and ACOS, where the task is to predict sentiment quadruples. For ASQP task, we use rest15(R15) and rest16(R16) datasets released from (Pontiki et al., 2016; Zhang et al., 2022). For ACOS task, we use laptop16(Lap) and rest16(Rest) datasets constructed by (Cai et al., 2021; Pontiki et al., 2016). Also, we adopt additional ACOS dataset(MR) from MEMD restaurant data (Xu et al., 2023) which uses a different source from the previous datasets.

3.2 Baselines

We compare our method against several strong baselines for ABSA as follows. *Paraphrase* (Zhang et al., 2021b) formulates a novel paraphrase generation process for ABSA with a single fixed order. *Seq2Path* (Mao et al., 2022) generates sentiments tuples as multiple paths of a tree, and automatically selects valid one. *DLO* (Hu et al., 2022b) augments data via the multiple order templates. *MvP* aggregates sentiment tuples generated from different orders of prompts via ensembling. *AugABSA* (Wang et al., 2023) generates a original text based on augmented sentiment quadruples. Also, we benchmark popular LLMs such as GPT-3.5, LLaMa-3 (Team, 2024), and Mistral-7b (Jiang et al., 2023). Detailed setups for LLMs are described in Appendix E.

3.3 Implementation Details

We utilize the pre-trained T5-base (Raffel et al., 2020) model as the backbone for the first stage. We also use the model trained in the first stage as the backbone for the second stage, allowing us to leverage a tuned initial point for the ABSA dataset to have the regularization effect inspired by (Fu et al., 2023). Also, we eliminate irregularities in tuples through stop-word filtering in the second stage. Please see Appendix A for more details.

3.4 Results

F1 score We use F1 score, which is a standard metric for ABSA, to measure the performance of the systems. As demonstrated in Table 1, our proposed method outperforms all baselines and achieves state-of-the-art performance across four datasets for ABSA. Our model shows slightly reduced performance on the Rest dataset, which we attribute to the substantial number of implicit aspects and opinions within this dataset. Additionally,

Methods	ASQP		ACOS			Time(s)
	R15	R16	Lap	Rest	MR	
Paraphrase	46.93	57.93	43.51	<u>61.16</u>	57.38	40.63
Seq2Path	-	-	42.97	58.41	-	-
DLO	48.18	59.79	43.64	59.99	57.07	260.74
MvP	<u>51.04</u>	60.39	<u>43.92</u>	61.54	<u>58.12</u>	2161.81
AugABSA	50.01	<u>60.88</u>	-	-	-	-
GPT 3.5-turbo	34.27	36.71	16.00	37.71	-	-
LLaMa3 8b	37.52	47.60	40.07	54.06	-	-
Mistral 7b	44.14	51.96	39.02	53.02	-	-
DOT (Ours)	51.91	61.24	44.92	59.25	58.25	298.17

Table 1: F1 scores for ABSA on five datasets. The best results are in bold and the second best are underlined. We conduct experiments with 5 different seeds and report the average of the outcomes. Time denotes the averaged inference time.

the size of the Rest dataset is relatively small, being less than half that of the Lap or MR datasets, which are derived from the same ACOS dataset.

Inference time We also measure inference time using T5-base model for all baselines. We check inference time for each dataset, and average them. All baselines are As in Table 1, we dramatically reduce inference time particularly compared to the multi-view methods such as MvP (Gou et al., 2023), by predicting solely the necessary number of views for each sample. On the other hand, in terms of single view inferences (Zhang et al., 2021b), we significantly improve the F1 score performance while suppressing the rate of increase in inference time. We also provide more details about computing the inference time in Appendix D.

3.5 Analysis

Ablation study To further investigate the effectiveness of each component of our framework, we conduct an ablation study and present the average F1 score across the datasets in Table 2. We first unify the two stages into one, directly generating multiple order templates and tuples without including order prompting to validate the effect of stage division. Additionally, we evaluate the results of directly using the generated views from the first stage, omitting the sampling of new order templates. Lastly, we exclude the multi-view approach by training and testing our model using the order template with the lowest entropy for each instance. By observing the gaps between these variants with the original model, we verify the effectiveness of each component of our method.

Distribution Shift To examine the effect of distribution shift of each model, we conduct an in-depth

Model Configuration	Average F1
Full Model	53.94
w/o stage division	52.73 (-1.21)
w/o re-sampling	52.47 (-1.47)
w/o multi view	53.31 (-0.63)

Table 2: Ablation study for the proposed method.

Train	SemEval		Yelp	
Test	SemEval	Yelp	Yelp	SemEval
Paraphrase	52.38	38.52(-11.86)	57.38	44.88(-12.50)
MvP ₃	55.62	34.42(-21.20)	57.27	41.72(-15.55)
MvP ₉	56.89	35.02(-21.87)	56.98	42.52(-14.46)
MvP ₁₅	57.66	35.21(-21.45)	58.12	41.94(-16.18)
DOT	57.47	39.88 (-17.59)	58.25	46.97 (-11.28)

Table 3: Cross-dataset evaluation results for validating the effect of distribution shift.

experiment on cross-dataset evaluation. We group the datasets into two groups based on their source: SemEval (Pontiki et al., 2016) (R15, R16, Rest) and Yelp (MEMD). Then we assess the performance between these groups by training on one group and testing on the other in a zero-shot setting. For the MvP model, we vary the number of views used for ensembling into 3, 9, and 15 (default) to measure the sensitivity of this number in static multi-view methods. Additionally, we evaluate T5-paraphrase which uses a static single order. Table 3 demonstrates that our model significantly outperforms the baselines in cross-dataset evaluation. While T5-paraphrase experiences a smaller performance drop compared to the others, it still lags behind our method. In particular, MvP exhibits significant performance degradation, irrespective of the number of views. From these experiments, we show that our model can effectively find the optimal number of views even for the out-of-domain datasets.

4 Conclusion

We propose Dynamic Order Template (DOT) method for aspect-based sentiment analysis, addressing inefficiencies and out-of-distribution errors in static multi-view prompting. By dynamically generating necessary views based on instance-level entropy, DOT efficiently predicts the sentiment tuple in each instance. Our experiments on ASQP and ACOS datasets demonstrate that DOT achieves state-of-the-art F1-scores with reduced inference time, effectively balancing the strengths of single and multi-view approaches for ABSA.

Limitation

Our DOT method is highly efficient and powerful, yet it still has several limitations. DOT method consists of two stages: view generation and tuple generation. We train separate models for each task, and these two models perform inference sequentially. This form is not end-to-end, so it is disadvantageous in terms of training time and memory.

Also, since we directly connect first stage and second stage, if any errors occur, the errors may propagate and magnify as it moves to the subsequent stage. However, by splitting the task of 'predicting the appropriate number of tuples' into two sub-tasks—'predicting the appropriate number of tuples' and 'accurately predicting the tuples'—it becomes significantly easier to achieve accurate results in both areas, thereby enhancing overall performance in our work.

Ethics Statement

This study utilizes the various datasets for aspect-based sentiment analysis, which are accessible online. Additionally, we have properly cited all the papers and sources referenced in our paper. We plan to release the pre-trained model and the code for training the proposed system.

References

- Tom Brown, Benjamin Mann, Nick Ryder, Melanie Subbiah, Jared D Kaplan, Prafulla Dhariwal, Arvind Neelakantan, Pranav Shyam, Girish Sastry, Amanda Askell, et al. 2020. Language models are few-shot learners. *Advances in neural information processing systems*, 33:1877–1901.
- Hongjie Cai, Nan Song, Zengzhi Wang, Qiming Xie, Qiankun Zhao, Ke Li, Siwei Wu, Shijie Liu, Jianfei Yu, and Rui Xia. 2023. Memd-absa: a multi-element multi-domain dataset for aspect-based sentiment analysis. *arXiv preprint arXiv:2306.16956*.
- Hongjie Cai, Rui Xia, and Jianfei Yu. 2021. [Aspect-category-opinion-sentiment quadruple extraction with implicit aspects and opinions](#). In *Proceedings of the 59th Annual Meeting of the Association for Computational Linguistics and the 11th International Joint Conference on Natural Language Processing (Volume 1: Long Papers)*, pages 340–350, Online. Association for Computational Linguistics.
- Siva Uday Sampreeth Chebolu, Franck Dernoncourt, Nedim Lipka, and Tamar Solorio. 2023. A review of datasets for aspect-based sentiment analysis. In *Proceedings of the 13th International Joint Conference on Natural Language Processing and the 3rd*

Conference of the Asia-Pacific Chapter of the Association for Computational Linguistics (Volume 1: Long Papers), pages 611–628.

- Alabhya Farkiya, Prashant Saini, Shubham Sinha, and Sharmishta Desai. 2015. Natural language processing using nltk and wordnet. *Int. J. Comput. Sci. Inf. Technol*, 6(6):5465–5469.

- Zihao Fu, Haoran Yang, Anthony Man-Cho So, Wai Lam, Lidong Bing, and Nigel Collier. 2023. On the effectiveness of parameter-efficient fine-tuning. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 37, pages 12799–12807.

- Zhibin Gou, Qingyan Guo, and Yujiu Yang. 2023. Mvp: Multi-view prompting improves aspect sentiment tuple prediction. *arXiv preprint arXiv:2305.12627*.

- Edward J Hu, Yelong Shen, Phillip Wallis, Zeyuan Allen-Zhu, Yuanzhi Li, Shean Wang, Lu Wang, and Weizhu Chen. 2021. Lora: Low-rank adaptation of large language models. *arXiv preprint arXiv:2106.09685*.

- Mengting Hu, Yike Wu, Hang Gao, Yinhao Bai, and Shiwan Zhao. 2022a. Improving aspect sentiment quad prediction via template-order data augmentation. *arXiv preprint arXiv:2210.10291*.

- Mengting Hu, Yike Wu, Hang Gao, Yinhao Bai, and Shiwan Zhao. 2022b. [Improving aspect sentiment quad prediction via template-order data augmentation](#). In *Proceedings of the 2022 Conference on Empirical Methods in Natural Language Processing*, pages 7889–7900, Abu Dhabi, United Arab Emirates. Association for Computational Linguistics.

- Albert Q Jiang, Alexandre Sablayrolles, Arthur Mensch, Chris Bamford, Devendra Singh Chaplot, Diego de las Casas, Florian Bressand, Gianna Lengyel, Guillaume Lample, Lucile Saulnier, et al. 2023. Mistral 7b. *arXiv preprint arXiv:2310.06825*.

- Yinhan Liu, Myle Ott, Naman Goyal, Jingfei Du, Mandar Joshi, Danqi Chen, Omer Levy, Mike Lewis, Luke Zettlemoyer, and Veselin Stoyanov. 2019. Roberta: A robustly optimized bert pretraining approach. *arXiv preprint arXiv:1907.11692*.

- Ilya Loshchilov and Frank Hutter. 2017. Decoupled weight decay regularization. *arXiv preprint arXiv:1711.05101*.

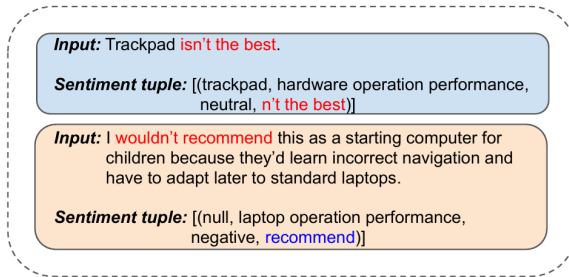
- Yue Mao, Yi Shen, Jingchao Yang, Xiaoying Zhu, and Longjun Cai. 2022. [Seq2Path: Generating sentiment tuples as paths of a tree](#). In *Findings of the Association for Computational Linguistics: ACL 2022*, pages 2215–2225, Dublin, Ireland. Association for Computational Linguistics.

- Maria Pontiki, Dimitrios Galanis, Harris Papageorgiou, Ion Androutsopoulos, Suresh Manandhar, AL-Smadi Mohammad, Mahmoud Al-Ayyoub, Yanyan Zhao, Bing Qin, Orphee De Clercq, et al. 2016. Semeval-2016 task 5: Aspect based sentiment analysis. In

411	<i>Proceedings of the 10th International Workshop on</i>	Wenxuan Zhang, Xin Li, Yang Deng, Lidong Bing, and	464
412	<i>Semantic Evaluation (SemEval-2016)</i> , pages 19–30.	Wai Lam. 2021c. Towards generative aspect-based	465
413	Colin Raffel, Noam Shazeer, Adam Roberts, Katherine	sentiment analysis . In <i>Proceedings of the 59th Annual</i>	466
414	Lee, Sharan Narang, Michael Matena, Yanqi Zhou,	<i>Meeting of the Association for Computational</i>	467
415	Wei Li, and Peter J Liu. 2020. Exploring the lim-	<i>Linguistics and the 11th International Joint Confer-</i>	468
416	its of transfer learning with a unified text-to-text	<i>ence on Natural Language Processing (Volume 2:</i>	469
417	transformer. <i>Journal of machine learning research</i> ,	<i>Short Papers)</i> , pages 504–510, Online. Association	470
418	21(140):1–67.	for Computational Linguistics.	471
419	Mohsen Shahhosseini, Guiping Hu, and Hieu Pham.	Wenxuan Zhang, Xin Li, Yang Deng, Lidong Bing,	472
420	2022. Optimizing ensemble weights and hyperpa-	and Wai Lam. 2022. A survey on aspect-based senti-	473
421	rameters of machine learning models for regression	ment analysis: Tasks, methods, and challenges. <i>IEEE</i>	474
422	problems. <i>Machine Learning with Applications</i> ,	<i>Transactions on Knowledge and Data Engineering</i> .	475
423	7:100251.		
424	Keith E. Stanovich and Richard F. West. 2000. Ad-		
425	vancing the rationality debate . <i>Behavioral and Brain</i>		
426	<i>Sciences</i> , 23(5):701–717.		
427	Meta LLaMA Team. 2024. Introducing meta llama 3:		
428	The most capable openly available llm to date .		
429	Ashish Vaswani, Noam Shazeer, Niki Parmar, Jakob		
430	Uszkoreit, Llion Jones, Aidan N Gomez, Łukasz		
431	Kaiser, and Illia Polosukhin. 2017. Attention is all		
432	you need. <i>Advances in neural information processing</i>		
433	<i>systems</i> , 30.		
434	An Wang, Junfeng Jiang, Youmi Ma, Ao Liu, and		
435	Naoaki Okazaki. 2023. Generative data augmen-		
436	tation for aspect sentiment quad prediction. In <i>Pro-</i>		
437	<i>ceedings of the 12th Joint Conference on Lexical</i>		
438	<i>and Computational Semantics (*SEM 2023)</i> , pages		
439	128–140.		
440	Jason Wei, Maarten Bosma, Vincent Y Zhao, Kelvin		
441	Guu, Adams Wei Yu, Brian Lester, Nan Du, An-		
442	drew M Dai, and Quoc V Le. 2021. Finetuned lan-		
443	guage models are zero-shot learners. <i>arXiv preprint</i>		
444	<i>arXiv:2109.01652</i> .		
445	Ting Xu, Huiyun Yang, Zhen Wu, Jiaze Chen, Fei Zhao,		
446	and Xinyu Dai. 2023. Measuring your aste models		
447	in the wild: A diversified multi-domain dataset for		
448	aspect sentiment triplet extraction. In <i>Findings of</i>		
449	<i>the Association for Computational Linguistics: ACL</i>		
450	<i>2023</i> , pages 2837–2853.		
451	Wenxuan Zhang, Yang Deng, Xin Li, Lidong Bing, and		
452	Wai Lam. 2021a. Aspect-based sentiment analysis in		
453	question answering forums . In <i>Findings of the Asso-</i>		
454	<i>ciation for Computational Linguistics: EMNLP 2021</i> ,		
455	pages 4582–4591, Punta Cana, Dominican Republic.		
456	Association for Computational Linguistics.		
457	Wenxuan Zhang, Yang Deng, Xin Li, Yifei Yuan, Li-		
458	dong Bing, and Wai Lam. 2021b. Aspect sentiment		
459	quad prediction as paraphrase generation . In <i>Pro-</i>		
460	<i>ceedings of the 2021 Conference on Empirical Meth-</i>		
461	<i>ods in Natural Language Processing</i> , pages 9209–		
462	9219, Online and Punta Cana, Dominican Republic.		
463	Association for Computational Linguistics.		

A Detailed Experimental Setups

We use AdamW optimizer (Loshchilov and Hutter, 2017) with a learning rate of 1e-4 for training two T5 models. We set the batch size to 16 for training and 24 for inference. We train the first stage model for 30 epochs, and train 40 epochs for the second stage. Additionally, we observe that the label of the datasets (i.e. sentiment tuples) irregularly contains stop words. For example, as in the first example of Figure 3, the inclusion of negations in the opinion terms is inconsistent. Also, as in the second example, element tuples sometimes contain ambiguous and meaningless stop words as an element. As a result, the fine-tuned model sometimes generates sentiment tuples containing stop words irregularly. It can yield critical performance degradation, even though they don't affect the meaning of the sentiment elements. To resolve the problem from stop words, we filter these stop words using nltk package (Farkiya et al., 2015) for both generated results and dataset labels. We use four RTX 4090 GPUs to train and evaluate all of the models.



Example 1: Irregularity in the Use of Negatives

Inputs: The food is not what it once was (positions have seriously seen downsizing), prices have gone up, and the service is the worst I have experienced anywhere (including mainland europe).
Sentiment tuple: [(the, service general, negative, the), ...]

Example 2: Ambiguous stop words

Figure 3: Two examples of irregularity of stop words. Note that these examples are the not all of the stop-word problems.

B Case Study

We conduct a case study and analyze the properties of the outputs generated by the proposed method. As depicted in 4, we classify the output results into three main cases.

The first case involves sentences that do not require multiple views for accurate prediction. For

these sentences, our model succeeds in making efficient predictions using only a single view. We observe that this case is the most common type in our study, significantly contributing to the model's efficiency.

The second case involves sentences predicted to require fewer views, but in reality, required more views. Our analysis reveals that such cases frequently occurred with implicit *O*. As shown in Table 1, this suggests that our model's performance might lag behind other baselines on the ACOS Rest16 dataset, which contains many samples with implicit *A* and *O*. Additionally, the model struggles with predicting infrequent *C* in the training set. Incorporating the concept of self-information and defining the necessary number of views based on the 'amount of information in a sample' could effectively address this issue.

The final case involves sentences with multiple sentiment tuples and longer lengths. Errors in this scenario stem from two main reasons. Firstly, longer sentences include extended phrases that modify *A* or *O*. Including all these modifiers as elements often leads to errors, a common problem across different models that requires an alternative solution. Secondly, errors occur when the number of tuples is incorrectly predicted in the first stage. If the predicted number of tuples is insufficient, some target sentiment tuples might be overlooked. Conversely, overestimation leads to the extraction of irrelevant aspects, as depicted in the figure. However, we optimized the first stage to reduce tuple count errors, which helped mitigate performance drops by minimizing incorrectly generated or overlooked tuples.

C Depth Analysis on First Stage

C.1 Accuracy on the Number of Views

We assess the accuracy of predicting the correct number of views in the first stage and present the results in Table 4. We evaluate the output by comparing it to the labeled sentiment tuples using RMSE and accuracy. We carefully implement the first stage baselines to compare our method properly as follows: *Random*: We find that the number of sentiment tuples in training dataset is mostly in range of 1 to 6. For each inference, we randomly sample one of the 6 numbers and compare it with our first stage result. *Majority*: We also reveal that about 60 percent of labels consist of single tuple. We construct a baseline that predicts only 1 for

Case 1: Efficiency in Simple sentence	
Input:	Best mexican place for lunch in the financial district.
Target:	[(mexican place, best, positive, restaurant general)]
Output:	[(mexican place, best, positive, restaurant general)]
Case 2: One sentiment tuple, but complex	
Input:	The crowd is mixed yuppies, young and old.
Target:	[(crowd, null, neutral, restaurant miscellaneous)]
Output:	[(crowd, mixed, neutral, ambience general)]
Case 3: Complex sentence analysis	
Input:	If you 're interested in good tasting (without the fish taste or smell), large portions and creative sushi dishes this is your place...
Target:	[(null, good, positive, food quality), (portions, large, positive, food_style_options), (sushi dishes, creative, positive, food_style_options)]
Output:	[(null, good tasting , positive, food quality), (portions, large, pos, food_style_options), (sushi dishes, creative, positive, food_style_options), (fish taste or smell, null, negative, food quality)]

Figure 4: Case study for three main types of results. Blue one denotes correct, red one denotes incorrect, and the yellow one denotes irrelevant.

the number of tuples, to check whether our model has ability to predict the number of sentiment tuples of a sentence. *Classification:* We adopt the RoBERTa model (Liu et al., 2019) to evaluate the results when treating the prediction of the number of views as a sequence classification task. We set the classes based on the number of sentiment tuples. As shown in 5, the distribution of tuple counts is skewed towards the lower end, with instances containing more than seven tuples being nearly non-existent. Consequently, we limit the categories from 1 to 6 and clip instances with 7 or more tuples to 6. Additionally, to address label imbalance, we employ a weighted loss function, where the weights are set as the inverse of the frequency ratio for each category as in Equation (4). We use same notation as in Section 2.1, and $\mathcal{I}()$ denotes indicator function. This approach enables the model to effectively classify even the less represented classes.

$$W_c = \frac{|D|}{\sum_D \mathcal{I}(\min(|y|, 6) = c)} \quad (c \in [1, 6])$$

$$\mathcal{L}_{cls} = - \sum_{i=1}^{|B|} W_{k_i} \log p(k_i | x_i) \mathcal{I}(k_i \leq 6)$$

(4)

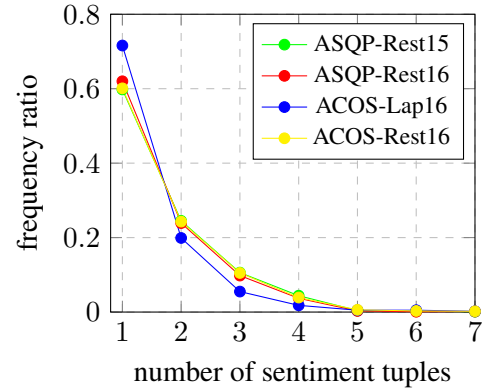


Figure 5: Distribution of the number of sentiment tuples. The sources are from training datasets of each task. We normalize each count by dividing it by the total number of data points. The number of tuples is clipped to 7.

C.2 Effect of Element Exclusions

We analyze the impact of excluding various marker tokens, including the [O] token representing opinions, to determine which token exclusions contribute to performance improvements. Additionally, we experiment with cases where no element exclusion is performed. In this section, we have also included the second stage results to provide a detailed comparison of the overall performance.

As in Table 4, our proposed method outperforms the other baselines and nearly predicts the actual distribution of sentiment tuples within a small mar-

Methods	First stage		Second stage
	RMSE	Acc.	F1 score
Random	2.80	18.89	-
Majority	0.99	63.39	-
Classification	0.83	61.90	-
DoT_{first}	0.54	77.83	54.33
exclude $[C]$	0.54	77.53	53.91
exclude $[A]$	0.53	77.77	53.71
exclude $[S]$	0.54	77.65	53.55
full elements	0.55	78.22	53.94

Table 4: First stage results for each main baseline and exclusion of specific tokens. We report average RMSE loss and accuracy for first stage, and F1 score for second stage.

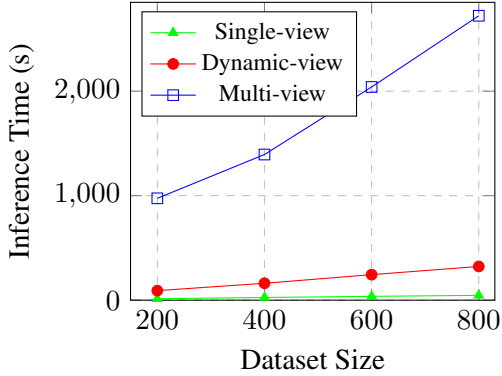


Figure 6: Inference time among dataset size for each model.

gin of error. This result justifies the direct use of the output from the first stage in the second stage. Also, the overall performance improved as the difficulty of predicting the excluded marker token’s element increased. The results from the first stage do not appear to be strongly correlated with those of the second stage. This emphasizes the importance of using a well-tuned initial model for the second stage.

D Computing Inference Time

We compare inference times based on view methods across different dataset sizes. The dataset consisted of randomly sampled test data from laptop16, with 200, 400, 600, and 800 samples. The baselines were set as static single view (T5-paraphrase) and static multi view (MvP), with the number of views for the multi view fixed at 15. Figure 6 shows that we not only dramatically reduce inference time of utilizing multi views, but also reduce the rate of increase in inference time with respect to the number of datasets. On the other hand, in terms of single view, we significantly increase F1 performance while suppressing the increase in inference time

and the rate of its increase. These results suggest that the efficiency of our method becomes more pronounced as the dataset size increases.

E Detailed Setups for LLM Experiments

As in Table 1, we perform the ABSA task using the GPT 3.5 Turbo, LLaMa-3 8B, and Mistral 7B models, and compared the results with our DOT model. For the GPT model, we utilize in-context learning (Brown et al., 2020). We randomly sample 10 instances and combine them with instruction format, and add it as a prompt. For the other two open-source LLMs, we employ instruction tuning (Wei et al., 2021) with the training dataset for fine-tuning, using the same instructions as in GPT prompts. To ensure stable model training during fine-tuning, we utilize the LoRa (Hu et al., 2021). We present the specific prompts and framework in Figure 7.

According to the following sentiment elements definition:

- The 'aspect term' refers to a specific feature, attribute, or aspect of a product or service that a user may express an opinion about, the aspect term might be 'null' for implicit aspect.
- The 'opinion term' refers to the sentiment or attitude expressed by a user towards a particular aspect or feature of a product or service, the aspect term might be 'null' for implicit opinion.
- The 'aspect category' refers to the category that aspect belongs to, and the available categories includes: {**dataset specific categories**}.
- The 'sentiment polarity' refers to the degree of positivity, negativity or neutrality expressed in the opinion towards a particular aspect or feature of a product or service, and the available polarities includes: 'positive', 'negative' and 'neutral'.

Recognize all sentiment elements with their corresponding aspect terms, aspect categories, opinion terms and sentiment polarity in the following text with the format of [('aspect term', 'aspect category', 'sentiment polarity', 'opinion term'), ...]:

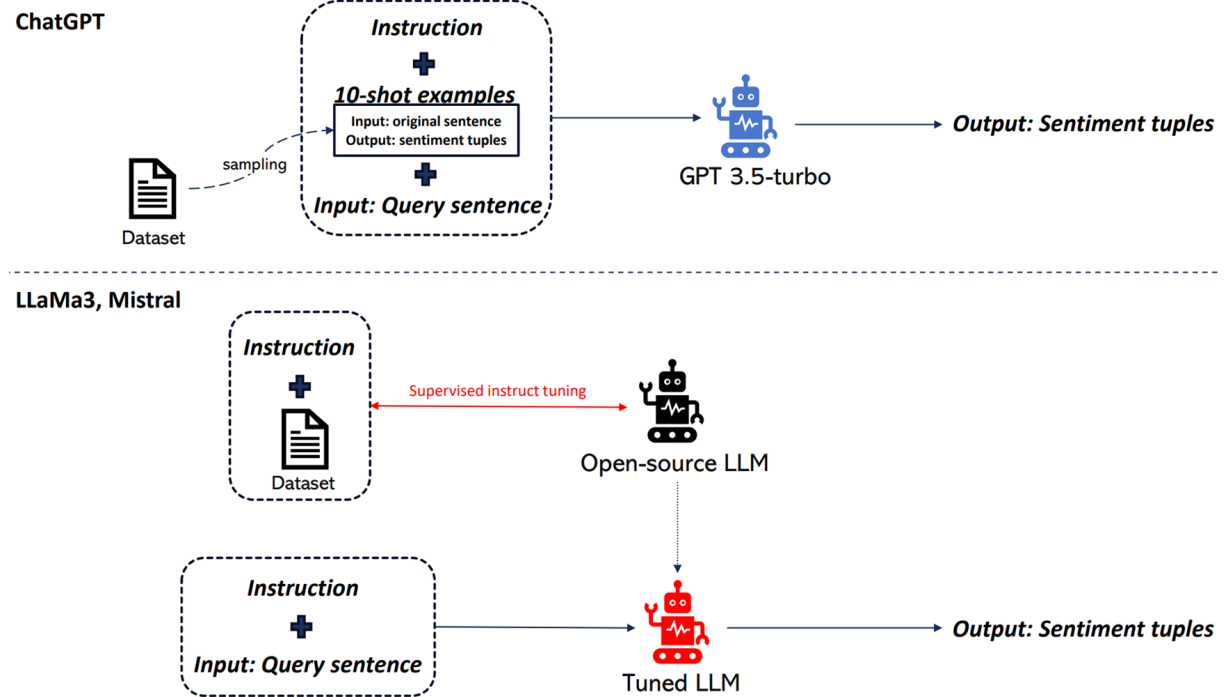


Figure 7: Instruction format for two LLM frameworks. We utilize in-context learning for ChatGPT inference, and instruction-tuning for LLaMa-3 and Mistral inference respectively.