

SLIDE: A Framework Integrating Small and Large Language Models for Open-Domain Dialogues Evaluation

Kun Zhao^{1*}, Bohao Yang^{2*}, Chen Tang², Chenghua Lin^{2†}, Liang Zhan^{1†}

¹ Department of Electrical and Computer Engineering, University of Pittsburgh, US

² Department of Computer Science, The University of Manchester, UK

{kun.zhao, liang.zhan}@pitt.edu, bohao.yang-2@postgrad.manchester.ac.uk
{chen.tang, chenghua.lin}@manchester.ac.uk

Abstract

The long-standing one-to-many problem of gold standard responses in open-domain dialogue systems presents challenges for automatic evaluation metrics. Though prior works have demonstrated some success by applying powerful Large Language Models (LLMs), existing approaches still struggle with the one-to-many problem, and exhibit subpar performance in domain-specific scenarios. We assume the commonsense reasoning biases within LLMs may hinder their performance in domain-specific evaluations. To address both issues, we propose a novel framework **SLIDE** (Small and Large Integrated for Dialogue Evaluation), that leverages both a small, specialised model (SLM), and LLMs for the evaluation of open domain dialogues. Our approach introduces several techniques: (1) Contrastive learning to differentiate between robust and non-robust response embeddings; (2) A novel metric for semantic sensitivity that combines embedding cosine distances with similarity learned through neural networks, and (3) A strategy for incorporating the evaluation results from both the SLM and LLMs. Our empirical results demonstrate that our approach achieves state-of-the-art performance in both the classification and evaluation tasks, and additionally the SLIDE evaluator exhibits better correlation with human judgements. Our code is available at <https://github.com/hegehongcha/SLIDE-ACL2024>.

1 Introduction

Open-domain dialogue generation is an important research topic in the field of Natural Language Processing (NLP) (Zeng et al., 2021; Wang et al., 2021; Xiao et al., 2023; Tang et al., 2023b,a; Yang et al., 2024a). Evaluating such dialogues, however, is challenging due to the

prevalent one-to-many issue where one conversational context may have multiple reasonable, yet semantically different responses. Additionally, adversarial negative responses, which have word overlap with given contexts, also raise further challenges in adequately evaluating open-domain dialogue (Sai et al., 2020).

Existing evaluation methods centered on word-overlap, such as BLEU (Papineni et al., 2002), ROUGE (Lin, 2004), METEOR (Banerjee and Lavie, 2005)) and semantic-embedding based approaches (e.g., BERTScore (Zhang et al., 2020) and BARTScore (Yuan et al., 2021)), evaluate responses via calculating the similarity between the generated response and contexts or gold references. Therefore, they struggle to adequately address the one-to-many problem as a context may permit many different responses. Recently, there are some novel approaches proposed to leverage the commonsense reasoning capabilities learned in Large Language Models (LLMs) for a wide range of NLP tasks, such as dialogue generation, evaluation and sentiment analysis (Loakman et al., 2023; Liu et al., 2023; Yang et al., 2024b; Chiang and yi Lee, 2023; Yang et al., 2024c). However, there are still some limitations of LLM-based evaluation. For example, Lin and Chen (2023) propose LLM-EVAL to evaluate open-domain dialogue systems with a range of LLMs and identify numerous issues including the significant reliance on prompt phrasing, which may artificially limit the effectiveness of the metrics and the evaluation quality changed with prompt formulation, which makes scoring unreliable and unrobust. Additionally, Wang et al. (2023b) state that without additional support, LLMs alone do not constitute equitable

* Equal contribution

† Corresponding authors

evaluators.

In addition, LLMs face significant challenges in generating accurate and dependable assessments when evaluating both positive and adversarial negative responses. Although LLMs can classify and evaluate adversarial negative responses with a degree of proficiency, their performance significantly diminishes when dealing with multiple positive responses, which is the one-to-many nature of open-domain dialogue, a reasonable response that matches well to context could have significantly different semantics to its reference. As for small task-specific model (SLM), which is cost-effective alternatives, reveal a propensity towards positive responses. Nonetheless, the deployment of SLM is challenging, as they require extensive training data. This requirement often leads to issues related to data scarcity and a lack of sufficient examples covering a wide range of topics within open-domain dialogue, presenting obstacles to their effective implementation and scalability.

Therefore, we combine a SLM with LLM to solve the one-to-many issue, and we also resolve the data scarcity problem through data augmentation via LLMs. Initially, we train a SLM using contrastive learning to minimise the cosine distance between the positive responses and the context embeddings, and enlarge the distance between the embeddings of adversarial negative responses and the context embeddings. After fine-tuning with contrastive learning, these compact models can discern between positive and adversarial negative responses adequately. Subsequently, we divide the representations of the responses into two categories: robust and non-robust vectors. Through contrastive learning, we only retain robust vectors for subsequent steps while excluding non-robust ones. Following this, we proceed to calculate the contextual-response distance, and obtain a probability value from the classification model. Finally, an integrated score is composed by combining these two values as the final evaluation score of the SLM model. Ultimately, the evaluation score produced by the SLM and the score from the LLM is further integrated as the final evalu-

ation score, SLIDE, used for evaluation.

A series of experiments are conducted to analyse the effectiveness of our proposed framework, including a classification task and an evaluation task. For the classification task, we use the SLM to classify the positive and adversarial negative responses from the DailyDialog++ dataset. Experimental results demonstrate that our SLM achieves state-of-the-art (SOTA) performance, even when compared to the results of some LLMs (e.g., GPT-3.5). We additionally note that the SLM achieves higher accuracy in positive samples, whilst the LLM performs better for negative example responses. For the evaluation task, SLIDE also shows a good correlation with human scores.

Our contributions can be summarised as the following:

- We propose a new evaluator for open-domain evaluation, namely **SLIDE** (Small and Large Integrated for Dialogue Evaluation). To the best of our knowledge, this is the first attempt to combine these models in the open-domain dialogue evaluation task.
- To better measure the semantic differences between dialogues, we introduce a novel evaluation score, which is integrated from the derivative of embedding cosine distances, and the similarity probability created by neural networks. The score has been demonstrated to be very effective in classifying positive and negative responses, leading to a SOTA result on the open-domain dialogue datasets.
- We augment existing dialogue evaluation datasets with multiple positive and adversarial negative responses. This enhanced dataset can be used in fine-tuning the open-domain SLM model.
- We conduct a range of experiments to demonstrate the proposed SLIDE effectively addresses the shortcomings of singular SLM-based or LLM-based open-domain dialogue evaluation models regarding the one-to-many reference problem.

2 Related Work

2.1 Dialogue evaluation metrics

Traditional n -gram-based evaluation metrics like BLEU (Papineni et al., 2002), ROUGE (Lin, 2004), and METEOR (Banerjee and Lavie, 2005) assess the overlaps of words between candidate responses and a reference standard. On the other hand, metrics based on embeddings, such as Extrema (Forgues et al., 2014) and BERTScore (Zhang et al., 2020), transform the responses and references into high-dimensional representations to evaluate semantic similarities.

Recent years have witness an increasing interest in developing the trainable metrics. RUBER (Tao et al., 2018) evaluates the similarity between the generated response, its context, and the actual reference. DEB (Sai et al., 2020), a BERT-based model pre-trained on extensive Reddit conversation datasets, was introduced to improve evaluation effectiveness. The Mask-and-fill approach (Gupta et al., 2021), utilising Speaker Aware (SA)-BERT (Gu et al., 2020), aims to better understand dialogues. MDD-Eval (Zhang et al., 2021) was developed for evaluating dialogue systems across various domains by using a teacher model for annotating dialogues from different domains, thus supporting the training of a domain-agnostic evaluator. This method, however, relies on human labels and extra training data, which our proposed approach does not require. Lastly, CMN (Zhao et al., 2023) was designed to effectively address the one-to-many nature of open-domain dialogue evaluation in the latent space, yet it does not account for adversarial negative examples.

2.2 LLM-based Evaluators

The application of LLM for evaluation purposes has seen significant interest due to its impressive capabilities across various tasks. GPTScore was introduced by Fu et al. (2023) as a method for multi-faceted, tailor-made, and training-independent evaluation. Preliminary investigations by Wang et al. (2023a) into the efficacy of evaluators based on LLM have been conducted. Kocmi and Federmann (2023) have

utilised GPT models for assessing machine translation quality. Furthermore, Liu et al. (2023) developed G-EVAL, leveraging GPT-4 to evaluate across multiple tasks such as dialogue response generation, text summarization, data-to-text generation, and machine translation. Despite these advancements, the use of LLM-based metrics to assess adversarial negative responses within open-domain dialogue scenarios remains unexplored.

3 Methodology

3.1 Model Architecture

As shown in Figure 1, our approach can be divided into two phases: the training process and the evaluation process. During the training phase, our focus is on training the SLM to classify positive and adversarial negative responses. In the evaluation phase, we integrate the SLM and LLMs to evaluate open-domain dialogue responses.

3.2 Training Process

we first use contrastive learning to train the SLM to discern positive responses from adversarial negative responses.

We employ Sentence-Transformer (Reimers and Gurevych, 2019) to encode the context and responses. Specifically, we encourage the embeddings of positive responses to be closer to context embeddings, whilst negative adversarial responses move further away. Given a triplet of $\langle \mathbf{x}_i^c, \mathbf{x}_i^p, \mathbf{x}_i^a \rangle$, where $1 \leq i \leq n$, $\mathbf{x}_i^c$ represents context, $\mathbf{x}_i^p$ represents positive responses, and $\mathbf{x}_i^a$ represents adversarial negative responses, the goal of the training process is to minimise the distance between $\mathbf{x}_i^p$ and $\mathbf{x}_i^c$ while maximising the separation between $\mathbf{x}_i^a$ and $\mathbf{x}_i^c$. The loss function is defined as follows:

$$\begin{aligned} \mathbf{h}_i^{\text{type}} &= \text{Encoder}(\mathbf{x}_i^{\text{type}}) \\ \mathcal{L}_{out} &= \max(\|\mathbf{h}_i^c - \mathbf{h}_i^p\| \\ &\quad - \|\mathbf{h}_i^c - \mathbf{h}_i^a\| + \text{margin}, 0) \end{aligned} \quad (1)$$

in which margin is a hyperparameter and $\text{type} = \{c, p, a\}$. $\mathbf{h}_i$ is the hidden state of the $\mathbf{x}_i$, whilst $\|\mathbf{h}_i^c - \mathbf{h}_i^p\| - \|\mathbf{h}_i^c - \mathbf{h}_i^a\|$ refers to the cosine distance.

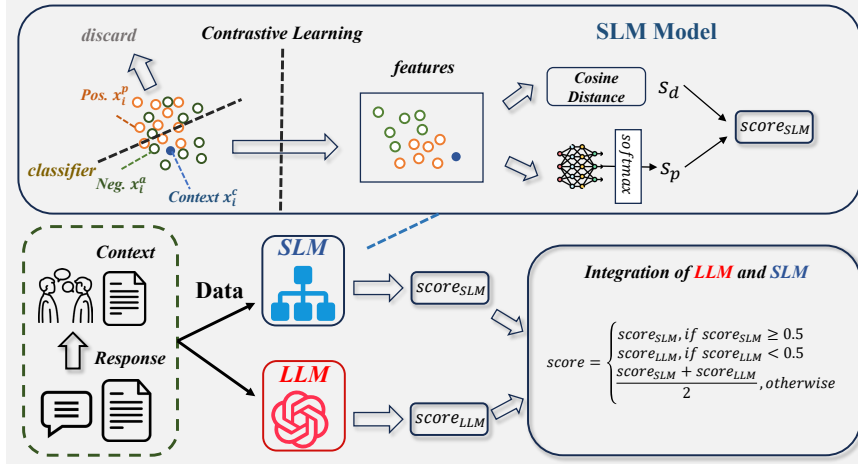


Figure 1: The architecture of the proposed model. We first use an SLM trained by contrastive learning to calculate the distance between context and responses. Following this, we calculate the probability of a response being positive and the cosine distance between context and response, in which case we then use them to acquire $score_{SLM}$. Secondly, we use an LLM to acquire $score_{LLM}$. Finally, we acquire the final score in accordance to our findings that LLM are more inclined to recognise negative responses correctly whilst SLM recognise positive responses better.

To enhance the precision of classification, we undertake a disentanglement process. The response embeddings are disentangled into two different sub-representations: robust embeddings and non-robust embeddings. The robust embedding can be interpreted as a salient feature for classification, whilst the non-robust embedding constitutes noise that could act as an interfering element, potentially leading to wrong predictions.

We denote the robust and non-robust embeddings of positive responses $\mathbf{h}_i^p$ as $\{\mathbf{h}_i^{pr}, \mathbf{h}_i^{pn}\}$, respectively. On the other hand, $\{\mathbf{h}_i^{ar}, \mathbf{h}_i^{an}\}$ represents the robust and non-robust embeddings of adversarial negative responses $\mathbf{h}_i^a$. It is imperative that the robust and non-robust embeddings within both positive responses and adversarial negative responses maintain clear distinctions. Similarly, the separation between the robust embeddings of different types of responses should be maintained. Therefore, we define our training loss function as follows:

$$\begin{aligned}
 \mathcal{L}_{ins_same_pos} &= z_1 * d_1^2 + (1 - z_1) \\
 &\quad * \max(\text{margin} - d_1, 0)^2 \\
 \mathcal{L}_{ins_same_neg} &= z_2 * d_2^2 + (1 - z_2) \\
 &\quad * \max(\text{margin} - d_2, 0)^2 \\
 \mathcal{L}_{out_robust} &= z_3 * d_3^2 + (1 - z_3) \\
 &\quad * \max(\text{margin} - d_3, 0)^2
 \end{aligned} \tag{2}$$

where $d_j, 1 \leq j \leq 3$ represents the cosine similarity and $z_j, 1 \leq j \leq 3$ indicates whether a pair of vectors match; with $z_j = 0$ denoting no match and consequently a greater distance between the vectors. Conversely, $z_j = 1$ signifies a match leading to a closer distance. To encourage divergence between the robust and non-robust vectors within both positive or negative classes, we set $z_1 = z_2 = 0$. Similarly, for the robust vectors across positive and negative classes to diverge, we assign $z_3 = 0$.

$$\begin{aligned}
 d_1 &= \|\mathbf{h}_i^{pr} - \mathbf{h}_i^{pn}\| \\
 d_2 &= \|\mathbf{h}_i^{ar} - \mathbf{h}_i^{an}\| \\
 d_3 &= \|\mathbf{h}_i^{pr} - \mathbf{h}_i^{ar}\|
 \end{aligned} \tag{3}$$

In addition, we develop a classification network to classify each factor. The process is defined as follows:

$$\begin{aligned}
 \mathbf{h} &= \text{concat}(\mathbf{h}_c, \mathbf{h}_{res}) \\
 p_i &= \text{Softmax}(\text{Linear}(\mathbf{h})) \\
 \mathcal{L}_{cls} &= \sum_{i=0}^2 y_i * p_i
 \end{aligned} \tag{4}$$

where $\mathbf{h}_{res}$ encompasses $\{\mathbf{h}_i^{pr}, \mathbf{h}_i^{pn}, \mathbf{h}_i^{ar}, \mathbf{h}_i^{an}\}$, while $y_i = \{0, 1, 2\}$ is the label. Specifically, $y_i = 1$ corresponds to $\mathbf{h}_i^{pr}$, $y_i = 0$ relates to $\mathbf{h}_i^{ar}$; for all other cases, $y_i = 2$.

In summary, the total loss function is

$$\mathcal{L} = \mathcal{L}_{\text{out}} + \mathcal{L}_{\text{ins_same_pos}} + \mathcal{L}_{\text{ins_same_neg}} + \mathcal{L}_{\text{out_robust}} + \mathcal{L}_{\text{cls}} \quad (5)$$

3.3 Evaluation Process

After training the SLM, we proceed to the evaluation phase. We first follow G-EVAL (Liu et al., 2023) and prompt LLM to evaluate the given context and response. The prompt could be found in A.1. Therefore we obtain a score referred to as $\text{Score}_{\text{LLM}}$. Subsequently, SLM is employed to encode both the context and the response. This step involves segregating the response encoding into two distinct embeddings—robust and non-robust. The mathematical representation of this process is given by

$$\begin{aligned} \mathbf{h}_c &= \text{Encoder}(\mathbf{x}_c) \\ \mathbf{h}_r &= \text{Encoder}(\mathbf{x}_r) \\ \mathbf{h}_{r,\text{robust}}, \mathbf{h}_{r,\text{non}} &= \text{sep}(\mathbf{h}_r) \end{aligned} \quad (6)$$

Then we calculate the cosine distance between context and response, as well as the probability for the given response.

$$\begin{aligned} d &= \text{cosine_similarity}(\mathbf{h}_c, \mathbf{h}_{r,\text{robust}}) \\ s_d &= (d - d_{\min}) / (d_{\max} - d_{\min}) \\ s_p &= p(y = 1) \\ &= \text{Softmax}(\text{Linear}(\mathbf{h}_c, \mathbf{h}_{r,\text{robust}})) \end{aligned} \quad (7)$$

where d is the cosine similarity between a context and a response. The normalised distance is denoted by s_d , while s_p represents the prediction probability of the classifier. Experimental evidence from training suggests that positive responses tend to be closer in embedding space to their corresponding contexts than negative ones, implying that $d_{\text{pos}} < d_{\text{neg}}$. Furthermore, it has been observed that for negative responses, s_p is typically lower than for positive ones. Consequently, it can be deduced that for positive responses, $s_d - s_p$ will be smaller than for negative ones. Based on these insights, we define a new probabilistic score for our SLM as follows:

$$\text{Score}_{\text{SLM}} = 1 - s_d + s_p \quad (8)$$

Finally we design an integration strategy to integrate $\text{Score}_{\text{SLM}}$ and $\text{Score}_{\text{LLM}}$. Our empirical findings suggest that while the SLM demonstrates greater accuracy in identifying positive responses, the LLMs shows a greater performance on classifying negative ones. We define the final score as follows:

$$\text{Score} = \begin{cases} \text{Score}_{\text{SLM}}, \text{Score}_{\text{SLM}} \geq 0.5, \\ \text{Score}_{\text{LLM}}, \text{Score}_{\text{LLM}} < 0.5, \\ (\text{Score}_{\text{SLM}} + \text{Score}_{\text{LLM}})/2, \text{otherwise} \end{cases} \quad (9)$$

4 Experimental Setup

4.1 Dataset

We conduct dialogue evaluation on three open-domain dialogue datasets: DailyDialog++ (Sai et al., 2020), TopicalChat (Gopalakrishnan et al., 2019), and Personachat (Zhang et al., 2018). The DailyDialog++ dataset contains 9,259 contexts in the training set, 1,028 in the validation set, and 1,142 in the test set. Each context is accompanied by five positive responses, five random negative responses, and five adversarial negative responses. However, the PersonaChat and TopicalChat datasets lack multiple positive and adversarial negative responses. To address this issue, we employ GPT-4 to generate both positive and adversarial negative responses and positive responses. The prompt for generating responses could be found in A.2. As a results, the generated training and test sets for both PersonaChat and TopicalChat datasets comprises 2,000 contexts. Each context within these datasets is further enriched with a set of responses, consisting of five positive responses and five adversarial negative responses.

Validation of the generated responses is conducted through a structured three-step process. Initially, a classification prompt is designed for the LLMs to determine the validity of the positive responses generated by GPT-4. Invalid responses are discarded, and we prompt the GPT-4 for subsequent generation attempts. Secondly, we randomly select 1,200 generated responses from the test set for assessment by human annotators. A significant majority of these responses are appropriate to their

contexts, exceeded 98%. The final phase involves a comparative analysis with the human-annotated benchmark dataset DailyDialog++. A SLM trained on training set of the DailyDialog++ dataset with 88% accuracy is utilised to classify the generated responses without incorporating distance metrics. The classification accuracy is 83% on the PersonaChat test set and 85% on the TopicalChat test set.

4.2 Baselines

We select a range of widely used baseline metrics. The word-overlap and embedding-based metrics includes BLEU (Papineni et al., 2002), ROUGE (Lin, 2004), METEOR (Banerjee and Lavie, 2005), Embedding-Average (Wieting et al., 2016), Vector-Extrema (Forgues et al., 2014), BERTScore (Zhang et al., 2020), Unieval (Zhong et al., 2022), and BARTScore (Yuan et al., 2021). For the LLM-based metrics, we select recently proposed baselines, including G-EVAL (Liu et al., 2023) and another LLM-based metric proposed by Chiang and yi Lee (2023) which we denote LLM-Chiang, as well as Gemini.

4.3 Evaluation Set

As for the evaluation set, we select 600 context-response pairs from each of three open-domain dialogue datasets—namely DailyDialog++, PersonaChat, and TopicalChat—resulting in a total of 1,800 samples. Then the evaluation set was evaluated by three human annotators, all of whom possess proficiency in English, to ensure an accurate assessment. They are instructed to thoroughly read each context-response pair and then rate them according to four distinct criteria including naturalness, coherence, engagingness, groundedness, employing a 1-5 Likert scale for their assessments.

The four criteria are consistent with Zhong et al. (2022) and are defined as follows: (1) **Naturalness**: The degree to which a response is naturally written; (2) **Coherence**: The extent to which the content of the output is well-structured, logical, and meaningful; (3) **Engagingness**: The degree to which the response is engaging; and (4) **Groundedness**: The extent to which a response is grounded in facts

present in the context.

After obtained four scores for each criteria, an aggregate score is calculated by averaging across all criteria to provide a comprehensive measure of each response’s quality. Additionally, the Inner-Annotator Agreement (IAA) between three evaluators are evaluated through Cohen’s Kappa (Cohen, 1960), which indicate a moderately strong agreement level (0.4-0.6) among annotators, with a value of 0.53, validating the reliability of the human evaluation.

We employ Gemini, GPT-3.5-turbo-1106 and GPT-4-1106 for all experiments. For SLM, DistilBERT (Sanh et al., 2019) is utilised.

4.4 Experimental Results

4.4.1 Dialogue Classification Task

We conduct dialogue classification on the on DailyDialog++ dataset as it contains the human-annotated labels for each context-response pair. We calculate the classification accuracy using SLM and LLMs (i.e., GPT-3.5 and GPT-4). The classification prompt for the LLMs could be found in A.3. We also use SLM that trained on the identical dataset for the ablation study.

Model	Positive	Negative	Overall
GPT-3.5	59.36	91.01	75.18
GPT-4	80.43	97.40	88.91
Gemini	80.62	95.13	87.86
SLM (Dis)	79.35	90.03	84.00
SLM (Prob)	83.25	93.11	88.19
SLM (Prob& Dis)	91.83	90.28	91.05

Table 1: Classification accuracy for only the Positive and adversarial negative responses, and the overall accuracy. Dis and Prob refers to the distance measures and classification probability, respectively.

As Table 1 shows, SLM employs both probability and distance measures, achieves superior performance with an accuracy of 91.05%, which exceeds the 88.91% accuracy achieved by GPT-4. Experimental results show that the SLM particularly excels in the classification of positive responses, achieving the highest accuracy of 91.83%. Conversely, GPT-4 exhibits the strongest performance in classifying adversarial negative responses, achieving

97.40% accuracy. The performance of Gemini surpasses that of GPT-3.5, yet remains inferior to that of GPT-4. These results indicate that while LLMs, such as GPT-4, possess robust capabilities in recognising adversarial negative responses, they still struggle with the one-to-many problem, especially when dealing with semantically diverse positive responses. On the other hand, SLM with the integration of distance and probability metrics shows exceptional performance in recognising open-domain positive responses. Based on these insights, we propose a hybrid approach that combines the SLM, a task-specific model, with LLMs for comprehensive dialogue evaluation. This approach is designed to leverage the distinct strengths of both types of models to enhance the overall performance in open-domain dialogue evaluations.

4.4.2 Dialogue Evaluation

As shown in Table 2, we investigate the correlation between automatic evaluation metrics and human judgments across three datasets. The word-overlap metrics, including BLEU, ROUGE, and METEOR, exhibit a weak positive correlation with human scores (<0.4), among which ROUGE-2 showing the least correlation (<0.2). Similarly, embedding-based metrics, such as Embedding and BARTScore, demonstrate a weak correlation with human evaluations, ranging from 0.2 to 0.4. Notably, BERTScore shows the strongest correlation to human judgments among both word-overlap and embedding based traditional metrics. On the other hand, Unieval achieves the suboptimal results (<0.3) on the DailyDialog++ and PersonaChat datasets.

As for LLM-Based metrics, there is a notable improvement in correlation with human scores compared with traditional metrics, most show correlations greater than 0.6. Specifically, Ours (LLMs-only(GPT-4)) exhibits superior performance compared with G-EVAL and LLM-Chiang and achieves the highest correlation in both the PersonaChat (0.737 for Pearson and 0.729 for Spearman) and TopicalChat datasets (0.745 for Pearson and 0.736 for Spearman). Moreover, the variance in per-

formance between G-EVAL and LLM-Chiang indicates the significant impact of model selection (i.e., GPT-4 or GPT-3.5) and prompt design on LLM-based metric effectiveness.

As for the SLM, Our SLM-Dis and Prob approach that incorporate both distance and probability measures, revealing a strong correlation with human evaluations (>0.6). As for the SLIDE method, it combines an SLM with LLMs. SLIDE (GPT-4) achieves the strongest correlation with human ratings on the DailyDialog++ dataset (0.773 for Pearson and 0.704 for Spearman). Furthermore, our LLM metrics, characterised by simpler prompt designs compared to those used in G-EVAL and LLM-Chiang, do not individually surpass these benchmarks. However, the amalgamation of LLM and SLM techniques results in a significant enhancement in correlation strength, surpassing both G-EVAL and LLM-Chiang, particularly when employing GPT-3.5. The same trend is also observed when utilising Gemini. These two distinct models could demonstrate the effectiveness of integrating LLM and SLM techniques.

The datasets generated by GPT-4, namely PersonaChat test set and Topical-Chat test set, exhibit the strongest correlations with the GPT-4-based metrics, including our proposed method. It indicates the inherent advantage of using model-specific metrics for evaluating datasets generated by the same model, thereby leading to a more accurate and relevant assessment of these two datasets.

To validate the effectiveness of distance metric used in SLM and SLIDE, we conduct several ablation studies. It can be observed that Ours (SLM-Dis-Prob) gives better performance than the model variant with the Dis (SLM-Dis-only) or Prob (SLM-Prob-only) component only. This suggests that SLM with distance metric is more effective in evaluating open-domain dialogues.

4.5 Example Visualisation

To demonstrate the effectiveness of disentanglement, we selected a subset of examples from the DailyDialog++ dataset and employed T-SNE for visualisation, as depicted in Fig-

	DailyDialog++		PersonaChat		TopicalChat	
Metrics	Pearson's ρ	Spearman's τ	Pearson's ρ	Spearman's τ	Pearson's ρ	Spearman's τ
BLEU-1	0.303 (<0.0001)	0.242 (<0.0001)	0.331 (<0.0001)	0.276 (<0.0001)	0.352 (<0.0001)	0.330 (<0.0001)
BLEU-2	0.317 (<0.0001)	0.270 (<0.0001)	0.311 (<0.0001)	0.261 (<0.0001)	0.279 (<0.0001)	0.252 (<0.0001)
BLEU-3	0.256 (<0.0001)	0.251 (<0.0001)	0.251 (<0.0001)	0.239 (<0.0001)	0.179 (<0.0001)	0.189 (<0.0001)
BLEU-4	0.234 (<0.0001)	0.247 (<0.0001)	0.214 (<0.0001)	0.217 (<0.0001)	0.247 (<0.0001)	0.154 (<0.0001)
ROUGE-1	0.303 (<0.0001)	0.270 (<0.0001)	0.304 (<0.0001)	0.250 (<0.0001)	0.374 (<0.0001)	0.356 (<0.0001)
ROUGE-2	0.178 (<0.0001)	0.133 (<0.0001)	0.142 (<0.0001)	0.126 (<0.0001)	0.173 (<0.0001)	0.149 (<0.0001)
ROUGE-L	0.305 (<0.0001)	0.271 (<0.0001)	0.294 (<0.0001)	0.244 (<0.0001)	0.279 (<0.0001)	0.259 (<0.0001)
METEOR	0.196 (<0.0001)	0.142 (<0.0001)	0.250 (<0.0001)	0.200 (<0.0001)	0.293 (<0.0001)	0.278 (<0.0001)
Embedding						
Extrema	0.382 (<0.0001)	0.367 (<0.0001)	0.330 (<0.0001)	0.299 (<0.0001)	0.389 (<0.0001)	0.244 (<0.0001)
Greedy	0.322 (<0.0001)	0.272 (<0.0001)	0.328 (<0.0001)	0.290 (<0.0001)	0.272 (<0.0001)	0.248 (<0.0001)
Average	0.186 (<0.0001)	0.208 (<0.0001)	0.239 (<0.0001)	0.249 (<0.0001)	0.254 (<0.0001)	0.261 (<0.0001)
BERTScore	0.419 (<0.0001)	0.381 (<0.0001)	0.465 (<0.0001)	0.412 (<0.0001)	0.525 (<0.0001)	0.485 (<0.0001)
BARTSCORE	0.305 (<0.0001)	0.264 (<0.0001)	0.466 (<0.0001)	0.442 (<0.0001)	0.470 (<0.0001)	0.427 (<0.0001)
Unieval	0.103 (<0.0001)	0.069 (<0.0001)	0.192 (<0.0001)	0.131 (<0.0001)	0.357 (<0.0001)	0.290 (<0.0001)
G-EVAL (GPT-3.5)	0.447 (<0.0001)	0.414 (<0.0001)	0.516 (<0.0001)	0.562 (<0.0001)	0.642 (<0.0001)	0.668 (<0.0001)
G-EVAL (GPT-4)	0.641 (<0.0001)	0.570 (<0.0001)	0.597 (<0.0001)	0.663 (<0.0001)	0.696 (<0.0001)	0.729 (<0.0001)
LLM-Chiang (GPT-3.5)	0.632 (<0.0001)	0.583 (<0.0001)	0.599 (<0.0001)	0.670 (<0.0001)	0.667 (<0.0001)	0.661 (<0.0001)
LLM-Chiang (GPT-4)	0.723 (<0.0001)	0.701 (<0.0001)	0.701 (<0.0001)	0.675 (<0.0001)	0.678 (<0.0001)	0.709 (<0.0001)
Ours (SLM-Prob-only)	0.677 (<0.0001)	0.632 (<0.0001)	0.638 (<0.0001)	0.692 (<0.0001)	0.691 (<0.0001)	0.689 (<0.0001)
Ours (SLM-Dis-only)	0.652 (<0.0001)	0.600 (<0.0001)	0.668 (<0.0001)	0.692 (<0.0001)	0.642 (<0.0001)	0.673 (<0.0001)
Ours (SLM-Dis and Prob)	0.728 (<0.0001)	0.666 (<0.0001)	0.711 (<0.0001)	0.707 (<0.0001)	0.741 (<0.0001)	0.727 (<0.0001)
Ours (LLMs-only (GPT-3.5))	0.518 (<0.0001)	0.472 (<0.0001)	0.506 (<0.0001)	0.617 (<0.0001)	0.590 (<0.0001)	0.669 (<0.0001)
SLIDE (GPT-3.5))	0.727 (<0.0001)	0.674 (<0.0001)	0.692 (<0.0001)	0.695 (<0.0001)	0.696 (<0.0001)	0.718 (<0.0001)
Ours (LLMs-only (Gemini))	0.660 (<0.0001)	0.614 (<0.0001)	0.719 (<0.0001)	0.688 (<0.0001)	0.701 (<0.0001)	0.699 (<0.0001)
SLIDE (Gemini)	0.760 (<0.0001)	0.689 (<0.0001)	0.701 (<0.0001)	0.705 (<0.0001)	0.713 (<0.0001)	0.712 (<0.0001)
Ours (LLMs-only (GPT-4))	0.709 (<0.0001)	0.666 (<0.0001)	0.737 (<0.0001)	0.729 (<0.0001)	0.745 (<0.0001)	0.736 (<0.0001)
SLIDE (GPT-4)	0.773 (<0.0001)	0.704 (<0.0001)	0.713 (<0.0001)	0.709 (<0.0001)	0.711 (<0.0001)	0.722 (<0.0001)

Table 2: Pearson and Spearman correlations with human judgements. Figures in parentheses are p-values. Dis and Prob refers to the distance measures and classification probability, respectively.

ure 2. Each example includes a context, five positive responses, and five adversarial negative responses. The vector representations prior to disentanglement are illustrated in the left plot, whereas those post-disentanglement are shown on the right. Post-disentanglement visualisations reveal a more distinct clustering of responses in relation to their contexts. By utilising these refined representations to train a classifier, we observed an improvement in classification accuracy on the test set of DailyDialog++ dataset, from 87% to 88.19%. Additional examples are available in Appendix A.5.

4.6 Case Study

We conduct qualitative analyses through case studies, which are shown in Table 3. Each case shows the conversational context as well as the corresponding gold-standard reference and the generated responses. We compare our evaluation metric with nine different baselines. To simplify the comparison, we normalise all scores to a range of 1-5 Likert scale. Note

that the normalisation is only applied to the case study, and is not implemented in our main experiments. In the first case, the response candidate is positive, whereas the response is an adversarial negative response in the second case. Our SLIDE metric correlate well with the human ratings. For example, in the first scenario, only our models rate a positive score (>3), which is consistent with the human score. More examples could be found in Appendix A.4.

5 Conclusion

In this paper, we introduce a novel automatic evaluation metric (SLIDE) that integrates SLM with LLMs for open-domain dialogue evaluation. Our approach involves initially training a SLM through iterative contrastive learning stages, followed by leveraging the combine strengths of SLM and LLMs to create a superior evaluation metric. This metric exhibits enhanced correlation with human judgments in comparison to those derived exclusively

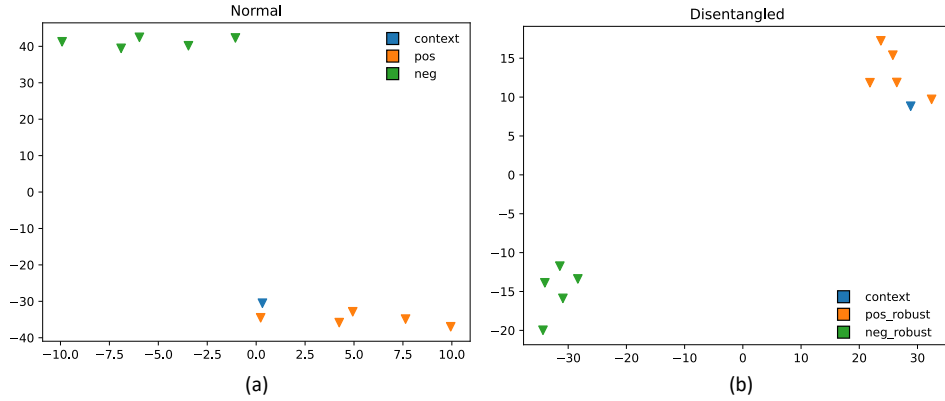


Figure 2: T-SNE visualisation of the sentence representation of context and responses. The left panel, labeled *Normal*, illustrates the vectors prior to disentanglement, whereas the right panel, labeled *Disentangled*, displays the post-disentanglement outcomes. This demonstrates the convergence of negative responses towards the context following disentanglement.

Context:	FS: We have another traditional holiday-the Dragon Boat Festival. SS: When is it? FS: It falls on the fifth day of the fifth lunar month.		
Reference:	Oh! That is great.		
Response:	I did not hear about that.		
Human	G-EVAL (GPT-3.5)	G-EVAL (GPT-4)	
4.70	2.50	2.25	
Chiang (GPT-3.5)	Chiang (GPT-4)	SLM	
2.25	2.75	4.60	
Ours (w/o LLM)	SLIDE (GPT-3.5)	SLIDE (GPT-4)	
5.00	4.30	5.00	
Context:	FS: We have been over this a hundred times! We are not getting a pet! SS: Why not? Come on! Just a cute little puppy or a kitty! FS: Who is going to look after a dog or a cat?		
Reference:	We both will look after it.		
Response:	Will you look after me once I get old?		
Human	G-EVAL (GPT-3.5)	G-EVAL (GPT-4)	
2.00	1.25	2.25	
Chiang (GPT-3.5)	Chiang (GPT-4)	SLM	
4.00	1.25	1.00	
Ours(w/o LLM)	SLIDE (GPT-3.5)	SLIDE (GPT-4)	
2.00	2.79	2.40	

Table 3: Samples from the DailyDialog++ dataset. "FS" is the First Speaker and "SS" is the Second Speaker

from LLMs. Furthermore, experimental results reveal that LLMs-based metric struggles with classifying and evaluating open-domain dialogues, which is the one-to-many nature. Therefore, we design a novel method for merging scores from SLM and LLM to enhance dialogue evaluation. Experimental results show that our SLIDE model outperforms a wide range of baseline methods in terms of both Pearson and Spearman correlations with human judgements on three open-domain dia-

logue datasets, and deals well with the one-to-many issue in open-domain dialogue evaluation.

Ethics Statement

In this paper, we introduce SLIDE, a novel automatic evaluation metric that integrates SLM and LLM for assessing open-domain dialogue systems. The advantage of SLIDE lies in its utilisation of both SLM and LLM, enhancing the evaluation of open-domain dialogues. However, a potential downside is that SLIDE might award high scores to responses that are inappropriate or offensive under certain conditions. Therefore, it is crucial to carefully review the content of training datasets prior to training SLIDE to mitigate this issue.

Limitations

Although our proposed method performs well in evaluating the open-domain dialogue systems, it also has some limitations. Primarily. Although our model first combines SLMs and LLMs for automatic evaluation metrics, it is a simple combination according to its characteristics. In the future, we will conduct a deeper combination of LLM and SLM to make a better evaluation on the open-domain dialogue system. In addition, because the PersonaChat dataset and TopicalChat dataset are augmented by LLM, not the real human-annotated dataset, our model does not perform the best on these

two datasets. We need to further employ the difference between LLM-generated datasets and human-annotated dataset, and further analyse the effectiveness of our proposed model.

Acknowledgments

This study was partially supported by the National Science Foundation (IIS 2045848 and 2319450). Part of the work used Bridges-2 at Pittsburgh Supercomputing Center through bridges2 (Brown et al., 2021) from the Advanced Cyberinfrastructure Coordination Ecosystem: Services & Support (ACCESS) program, which is supported by NSF grants #2138259, #2138286, #2138307, #2137603, and #2138296.

References

- Satanjeev Banerjee and Alon Lavie. 2005. [METEOR: An automatic metric for MT evaluation with improved correlation with human judgments](#). In *Proceedings of the ACL Workshop on Intrinsic and Extrinsic Evaluation Measures for Machine Translation and/or Summarization*, pages 65–72, Ann Arbor, Michigan. Association for Computational Linguistics.
- Shawn T Brown, Paola Buitrago, Edward Hanna, Sergiu Sanielevici, Robin Scibek, and Nicholas A Nystrom. 2021. Bridges-2: A platform for rapidly-evolving and data intensive research. In *Practice and Experience in Advanced Research Computing*, pages 1–4.
- Cheng-Han Chiang and Hung yi Lee. 2023. [Can large language models be an alternative to human evaluations?](#) In *Annual Meeting of the Association for Computational Linguistics*.
- Jacob Cohen. 1960. [A coefficient of agreement for nominal scales](#). *Educational and Psychological Measurement*, 20(1):37–46.
- Gabriel Forgues, Joelle Pineau, Jean-Marie Larchevêque, and Réal Tremblay. 2014. Bootstrapping dialog systems with word embeddings. In *Nips, modern machine learning and natural language processing workshop*, volume 2, page 168.
- Jinlan Fu, See-Kiong Ng, Zhengbao Jiang, and Pengfei Liu. 2023. [Gptscore: Evaluate as you desire](#). *ArXiv*, abs/2302.04166.
- Karthik Gopalakrishnan, Behnam Hedayatnia, Qinqiang Chen, Anna Gottardi, Sanjeev Kwatra, Anu Venkatesh, Raefer Gabriel, and Dilek Z. Hakkani-Tür. 2019. [Topical-chat: Towards knowledge-grounded open-domain conversations](#). *ArXiv*, abs/2308.11995.
- Jia-Chen Gu, Tianda Li, Quan Liu, Xiaodan Zhu, Zhenhua Ling, Zhiming Su, and Si Wei. 2020. Speaker-aware bert for multi-turn response selection in retrieval-based chatbots. *Proceedings of the 29th ACM International Conference on Information & Knowledge Management*.
- Prakhar Gupta, Yulia Tsvetkov, and Jeffrey P. Bigham. 2021. Synthesizing adversarial negative responses for robust response ranking and evaluation. In *Findings*.
- Tom Kocmi and Christian Federmann. 2023. [Large language models are state-of-the-art evaluators of translation quality](#). In *European Association for Machine Translation Conferences/Workshops*.
- Chin-Yew Lin. 2004. [ROUGE: A package for automatic evaluation of summaries](#). In *Text Summarization Branches Out*, pages 74–81, Barcelona, Spain. Association for Computational Linguistics.
- Yen-Ting Lin and Yun-Nung Chen. 2023. [LLM-eval: Unified multi-dimensional automatic evaluation for open-domain conversations with large language models](#). In *Proceedings of the 5th Workshop on NLP for Conversational AI (NLP4ConvAI 2023)*, pages 47–58, Toronto, Canada. Association for Computational Linguistics.
- Yang Liu, Dan Iter, Yichong Xu, Shuohang Wang, Ruochen Xu, and Chenguang Zhu. 2023. [G-eval: NLG evaluation using gpt-4 with better human alignment](#). In *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing*, pages 2511–2522, Singapore. Association for Computational Linguistics.
- Tyler Loakman, Aaron Maladry, and Chenchua Lin. 2023. The iron (ic) melting pot: Reviewing human evaluation in humour, irony and sarcasm generation. In *Findings of the Association for Computational Linguistics: EMNLP 2023*, pages 6676–6689.
- Kishore Papineni, Salim Roukos, Todd Ward, and Wei-Jing Zhu. 2002. [Bleu: a method for automatic evaluation of machine translation](#). In *Proceedings of the 40th Annual Meeting of the Association for Computational Linguistics*, pages 311–318, Philadelphia, Pennsylvania, USA. Association for Computational Linguistics.
- Nils Reimers and Iryna Gurevych. 2019. Sentence-bert: Sentence embeddings using siamese bert-networks. *arXiv preprint arXiv:1908.10084*.
- Ananya B. Sai, Akash Kumar Mohankumar, Siddharth Arora, and Mitesh M. Khapra. 2020. Improving dialog evaluation with a multi-reference adversarial dataset and large scale pretraining. *Transactions of the Association for Computational Linguistics*, 8:810–827.

- Victor Sanh, Lysandre Debut, Julien Chaumond, and Thomas Wolf. 2019. Distilbert, a distilled version of bert: smaller, faster, cheaper and lighter. *arXiv preprint arXiv:1910.01108*.
- Chen Tang, Shun Wang, Tomas Goldsack, and Chenghua Lin. 2023a. Improving biomedical abstractive summarisation with knowledge aggregation from citation papers. In *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing*, pages 606–618.
- Chen Tang, Hongbo Zhang, Tyler Loakman, Chenghua Lin, and Frank Guerin. 2023b. Enhancing dialogue generation via dynamic graph knowledge aggregation. In *Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 4604–4616.
- Chongyang Tao, Lili Mou, Dongyan Zhao, and Rui Yan. 2018. [Ruber: An unsupervised method for automatic evaluation of open-domain dialog systems](#). *Proceedings of the AAAI Conference on Artificial Intelligence*, 32(1).
- Dingmin Wang, Chenghua Lin, Qi Liu, and Kam-Fai Wong. 2021. [Fast and scalable dialogue state tracking with explicit modular decomposition](#). In *Proceedings of the 2021 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies*, pages 289–295.
- Jiaan Wang, Yunlong Liang, Fandong Meng, Zengkui Sun, Haoxiang Shi, Zhixu Li, Jinan Xu, Jianfeng Qu, and Jie Zhou. 2023a. [Is ChatGPT a good NLG evaluator? a preliminary study](#). In *Proceedings of the 4th New Frontiers in Summarization Workshop*, pages 1–11, Singapore. Association for Computational Linguistics.
- Peiyi Wang, Lei Li, Liang Chen, Dawei Zhu, Binghuai Lin, Yunbo Cao, Qi Liu, Tianyu Liu, and Zhifang Sui. 2023b. [Large language models are not fair evaluators](#). *ArXiv*, abs/2305.17926.
- John Wieting, Mohit Bansal, Kevin Gimpel, and Karen Livescu. 2016. Towards universal paraphrastic sentence embeddings. *CoRR*, abs/1511.08198.
- Ziang Xiao, Susu Zhang, Vivian Lai, and Q Vera Liao. 2023. Evaluating evaluation metrics: A framework for analyzing nlg evaluation metrics using measurement theory. In *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing*, pages 10967–10982.
- Bohao Yang, Chen Tang, and Chenghua Lin. 2024a. Improving medical dialogue generation with abstract meaning representations. In *ICASSP 2024-2024 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*, pages 11826–11830. IEEE.
- Bohao Yang, Chen Tang, Kun Zhao, Chenghao Xiao, and Chenghua Lin. 2024b. [Effective distillation of table-based reasoning ability from LLMs](#). In *Proceedings of the 2024 Joint International Conference on Computational Linguistics, Language Resources and Evaluation (LREC-COLING 2024)*, pages 5538–5550, Torino, Italia. ELRA and ICCL.
- Bohao Yang, Kun Zhao, Chen Tang, Liang Zhan, and Chenghua Lin. 2024c. Structured information matters: Incorporating abstract meaning representation into llms for improved open-domain dialogue evaluation. *arXiv preprint arXiv:2404.01129*.
- Weizhe Yuan, Graham Neubig, and Pengfei Liu. 2021. Bartscore: Evaluating generated text as text generation. *Advances in Neural Information Processing Systems*, 34:27263–27277.
- Chengkun Zeng, Guanyi Chen, Chenghua Lin, Ruizhe Li, and Zhi Chen. 2021. [Affective decoding for empathetic response generation](#). In *Proceedings of the 14th International Conference on Natural Language Generation*, pages 331–340.
- Chen Zhang, L. F. D’Haro, Thomas Friedrichs, and Haizhou Li. 2021. Mdd-eval: Self-training on augmented data for multi-domain dialogue evaluation. In *AAAI Conference on Artificial Intelligence*.
- Saizheng Zhang, Emily Dinan, Jack Urbanek, Arthur Szlam, Douwe Kiela, and Jason Weston. 2018. [Personalizing dialogue agents: I have a dog, do you have pets too?](#) In *Proceedings of the 56th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 2204–2213, Melbourne, Australia. Association for Computational Linguistics.
- Tianyi Zhang, Varsha Kishore, Felix Wu, Kilian Q. Weinberger, and Yoav Artzi. 2020. Bertscore: Evaluating text generation with bert. *ArXiv*, abs/1904.09675.
- Kun Zhao, Bohao Yang, Chenghua Lin, Wenge Rong, Aline Villavicencio, and Xiaohui Cui. 2023. Evaluating open-domain dialogues in latent space with next sentence prediction and mutual information. In *Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 562–574.
- Ming Zhong, Yang Liu, Da Yin, Yuning Mao, Yizhu Jiao, Pengfei Liu, Chenguang Zhu, Heng Ji, and Jiawei Han. 2022. [Towards a unified multi-dimensional evaluator for text generation](#). In *Proceedings of the 2022 Conference on Empirical Methods in Natural Language Processing*, pages 2023–2038, Abu Dhabi, United Arab Emirates. Association for Computational Linguistics.

A Appendices

A.1 Prompt for evaluating the engagingness of the response

Engagingness

Please rate the dialogue response.

The goal of this task is to rate dialogue response.

Note: Please take the time to fully read and understand the dialogue response. We will reject submissions from workers that are clearly spamming the task.

How engaging is the text of the dialogue response? (on a scale of 0-1, with 0 being the lowest)

Example:

Conversation History:

Is there something wrong?

I enjoy having your daughter in my class.

I'm glad to hear it.

Dialogue Response:

I enjoy listening jazz music in my free time.

Evaluation Form (scores ONLY):

Engagingness: 1

Input:

Conversation History:

Response:

Evaluation Form (scores ONLY):

Engagingness:

A.2 Prompt for generating response

You are a conversational dialogue generator. Given a conversation context, , which includes 2 speakers[annotated as FS(FirstSpeaker) and SS(SecondSpeaker)], and a response.

Your task is to generate five diverse positive response and five adversarial negative response respectively.

Positive Response

Positive response is valid for the conversation context.

Adversarial Negative Response

Adversarial negative responses have a significant word overlap with the conversation context but are still irrelevant response, which may not have any relation to the context. You need to choose some words (do not include stopwords such as "I", "you", "are", etc.) from the conversation context and use them directly or indirectly while writing the adversarial negative responses. Indirect usage here refers to using words closely related to the context words.

For example, using synonyms, antonyms, homonyms, subwords, or other words that are known to frequently co-occur with the words in the context (e.g., the words "flexibility" and "injuries" co-occur with "acrobatics").

The following are five examples of a conversation context and response, and the corresponding prediction.

Example

Context:

FS: Is there something wrong?

SSI enjoy having your daughter in my class.

FS: I'm glad to hear it.

Positive response: She is so brilliant.

Her behavior is good in the class.

I would love to hear that she knows every rules and regulation.

I was shocked to know that she is your daughter.

She answers all my questions.

Adversarial Negative Responses:

I enjoy listening jazz music in my free time.

I need pin drop silence in the class.

If I hear someone talking they will be sent out of the class.

I am glad you enjoyed the magic show organised by our team.

I think there was something wrong with the CCTV camera installed in the class.

This is the wrong method to solve the problem.

Please be attentive in the class.

A.3 Prompt for classifying response

You are a classifier. Given a conversation context, which includes 2 speakers[annotated as FS(FirstSpeaker) and SS(SecondSpeaker)], and a response. Your task is to classify this response whether is positive or negative.

Positive Response

Positive response is valid for the conversation context.

Adversarial Negative Response

Adversarial negative responses have a significant word overlap with the conversation context but are still irrelevant response, which may not have any relation to the context. You need to choose some words (do not include stopwords such as "I", "you", "are", etc.) from the conversation context and use them directly or indirectly while writing the adversarial negative responses. Indirect usage here refers to using words closely related to the context words.

For example, using synonyms, antonyms, homonyms, subwords, or other words that are known to frequently co-occur with the words in the context (e.g., the words "flexibility" and "injuries" co-occur with "acrobatics").

Your output format is only the "Positive" or "Negative".

Example

The following are five examples of a conversation context and response, and the corresponding prediction.

Example 1:

Context:

FS: Is there something wrong?

SS: I enjoy having your daughter in my class.

FS: I'm glad to hear it.

Response:

She is so brilliant.

Prediction: Positive

Example 2:

Context:

FS: Is there something wrong?

SS: I enjoy having your daughter in my class.

FS: I'm glad to hear it.

Response:

I enjoy having your daughter in my class.

Prediction: Positive

Example 3:

Context:

FS: Is there something wrong?

SS: I enjoy having your daughter in my class.

FS: I'm glad to hear it.

Response:

I'm glad to hear it.

Prediction: Positive

Example 4:

Context:

FS: We have to pick up Conrad before the party.

SS: Alright, no problem.

FS: We're supposed to meet him at Cal's Bar at 10

Response:

I pushed the problem aside; at present it was insoluble."

Prediction: Negative

Example 5:

Context:

FS: Is there something wrong?

SS: I enjoy having your daughter in my class.

FS: I'm glad to hear it.

Response:

I think there was something wrong with the CCTV camera installed in the class.

Prediction: Negative

A.4 Case Study

We display some other samples from Daily-Dialog++ dataset. From Table 4, We can see that our SLIDE model have a close score with

human score.

Context:	FS: We have been over this a hundred times! We are not getting a pet! SS: Why not? Come on! Just a cute little puppy or a kitty! FS: Who is going to look after a dog or a cat?	
Reference:	We both will look after it.	
Response:	Will you look after me once I get old?	
Human 2.00	G-EVAL (GPT-3.5) 1.25	G-EVAL (GPT-4) 2.25
Chiang (GPT-3.5) 4.00	Chiang (GPT-4) 1.25	SLM 1.00
Ours (w/o LLM) 2.00	SLIDE (GPT-3.5) 2.79	SLIDE (GPT-4) 2.40

Context:	FS: Why not? We're supposed to meet him there. SS: Why doesn't he meet us outside? FS: Why should he do that? It isn't illegal for us to go in.	
Reference:	I really don't want to go to that place.	
Response:	Nothing wrong in going there, so let's go there and meet him.daughter.	
Human 5.00	G-EVAL (GPT-3.5) 2.25	G-EVAL (GPT-4) 2.00
Chiang (GPT-3.5) 3.50	Chiang (GPT-4) 3.75	SLM 4.90
Ours (w/o LLM) 5.00	SLIDE (GPT-3.5) 5.00	SLIDE (GPT-4) 4.80

Table 4: Samples from the DailyDialog++ dataset. "FS" is the First Speaker and "SS" is the Second Speaker

A.5 Example Visualisation

We display some other visualization example in this section. From Figure 3, We can see that the positive responses become closer to context after disentanglement.

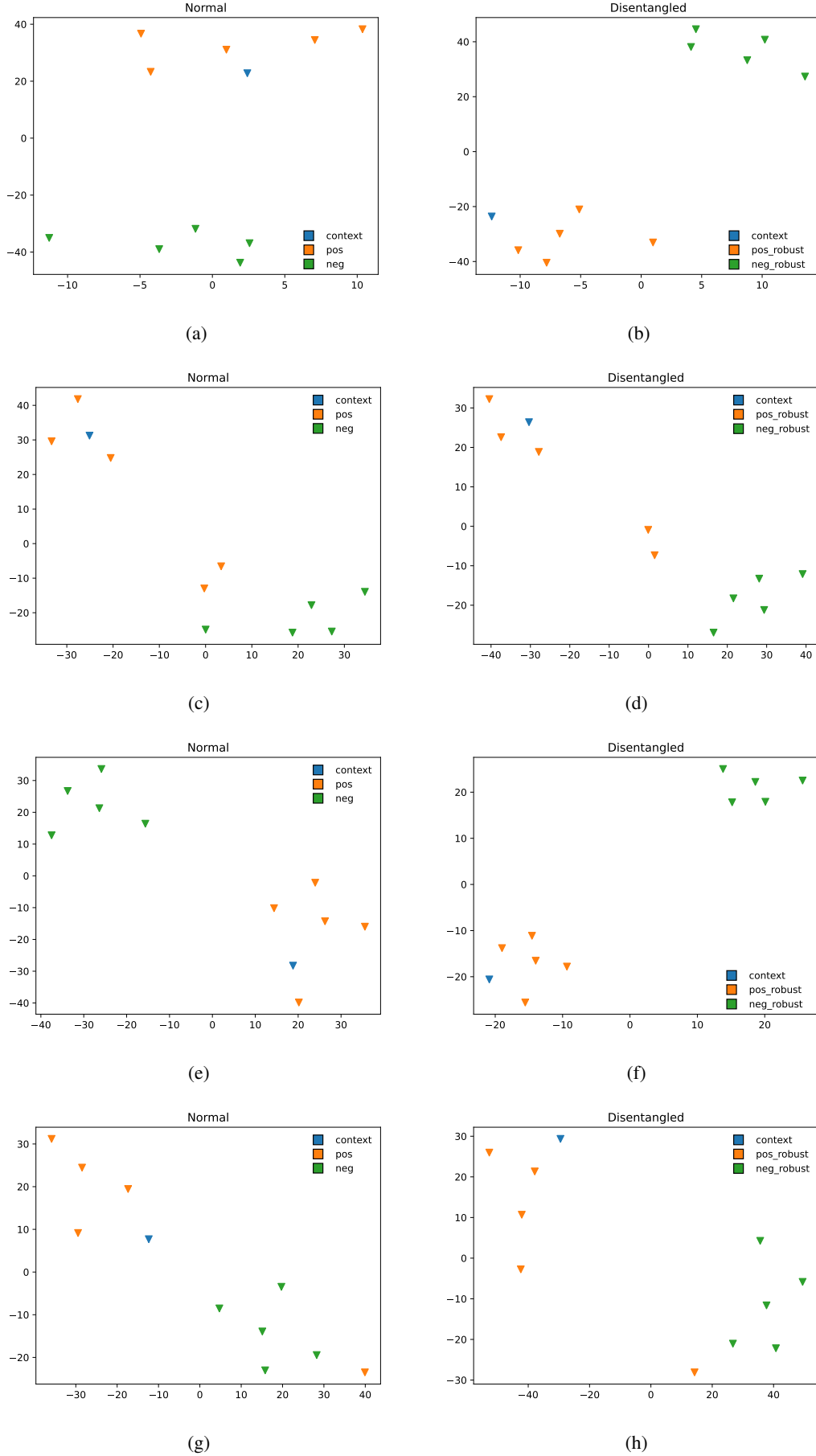


Figure 3: T-SNE visualisation of the sentence representation of context and responses for some examples. Each example includes a context, five positive responses, and five adversarial negative responses. The left represents the vector prior to disentanglement which is titled "Normal", whilst the right is after disentanglement which is titled "Disentangled". These figures demonstrate the positive responses become nearer to context after disentangling.