

Deliberate then Generate: Enhanced Prompting Framework for Text Generation

Anonymous ACL submission

Abstract

Large language models (LLMs) have shown remarkable success across a wide range of natural language generation tasks, where proper prompt designs make great impacts. While existing prompting methods are normally restricted to providing correct information, in this paper, we encourage the model to deliberate by proposing a novel Deliberate then Generate (DTG) prompting framework, which consists of error detection instructions and candidates that may contain errors. DTG is a simple yet effective technique that can be applied to various text generation tasks with minimal modifications. We conduct extensive experiments on 20+ datasets across 7 text generation tasks, including summarization, translation, dialogue, and more. We show that DTG consistently outperforms existing prompting methods and achieves state-of-the-art performance on multiple text generation tasks. We also provide in-depth analyses to reveal the underlying mechanisms of DTG, which may inspire future research on prompting for LLMs.

1 Introduction

Large language models (LLMs) (Brown et al., 2020; OpenAI, 2023; Touvron et al., 2023) are revolutionizing the area of natural language generation, which have demonstrated exceptional abilities in generating coherent and fluent text as well as exhibited a remarkable aptitude in performing a diverse range of text generation tasks with high accuracy (Hendy et al., 2023; Nori et al., 2023). When adapting to downstream tasks, traditional fine-tuning methods require access to the parameters of LLMs, which hinder their application on powerful black-box LLMs (e.g., ChatGPT) that only provide APIs to interact with. Therefore, prompting methods that guide the generation results by providing several task-specific instructions and demonstrations have attracted lots of attention in recent works (Schick and Schütze, 2020; Sanh

et al., 2021), which show that the prompt can significantly influence the resulting outcomes and thus require careful design.

While prompting is itself a general approach, the current use of this approach is a bit rigid, say, an LLM only operates on the basis of what is correct (Brown et al., 2020; Hendy et al., 2023; Wei et al., 2022b). This is not the case for language acquisition where a human can learn from both positive and negative feedback and improve the ability of language use through corrections. In this work, we examine whether and how the deliberation ability emerges by asking the LLMs to rethink and learn to detect potential errors in their output. To do this, we develop a new prompting template termed Deliberate then Generate (DTG) that contains instructions and candidate outputs to enable an error detection process before generation, i.e., adding “*Please detect the error type firstly, and provide the refined results then*” in the prompt.

A key design aspect of DTG is how to determine the candidate. One straightforward choice is utilizing the results from an extra baseline system, which typically exhibits high quality and requires only minor adjustments. Accordingly, it cannot well facilitate the deliberation ability. In this work, we propose to utilize the text that is irrelevant from the reference (e.g., such as a randomly sampled text or even an *empty string*) as the candidate. In this way, the method successfully triggers the deliberation ability of LLMs, without having to resort to other text generation systems to create correction examples, which enables DTG to be easily applied to a wide range of text generation tasks only with minimal modifications in prompts. This work is in part motivated from a psychological perspective by considering *negative evidence* in developing language abilities, which is a canonical case for language learning (Marcus, 1993).

We conduct extensive experiments on 7 text generation tasks and more than 20 datasets on

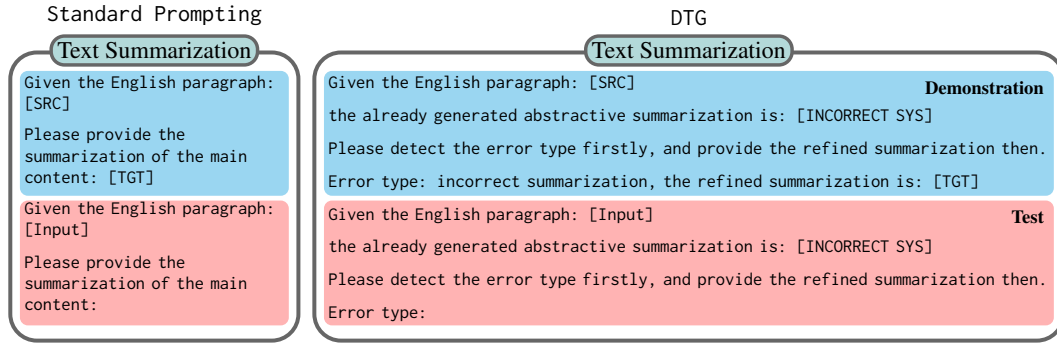


Figure 1: Comparison of standard GPT prompting and our DTG prompt design for summarization task. Note that prompt in blue denotes the demonstration, and that in red denotes the test input. [SRC] and [Input] means the source input, TGT means the target reference and [INCORRECT SYS] means the irrelevant system output (e.g., such as a randomly sampled text or even an empty string).

GPT3.5 (text-davinci-003) and GPT4, where the proposed DTG prompting consistently improves model performance compared to conventional prompts. GPT with DTG prompting achieves state-of-the-art performance on multiple datasets across different text generation tasks, including machine translation, simplification and commonsense generation. Extensive ablation studies and error statistical analysis illustrate that the proposed DTG prompting does enable deliberation ability and error avoidance before generation.

The main contributions of this work are summarized as follows:

- We propose a novel prompting framework named DTG for LLMs, which eliminates the need for extra resources or costs and can be effortlessly applied to various text generation tasks. DTG can also be combined with other advanced prompting strategy (e.g., CoT) to further improve the performance.
- We conduct experiments on 20+ datasets across 7 text generation tasks, where DTG prompting brings consistent improvements and achieves SoTA performance on several benchmarks.
- To the best of our knowledge, we are the first to evaluate the performance of GPT3.5 and GPT4 on multiple benchmark text generation tasks. We hope the experimental results help deepen our understanding of SoTA LLMs.

2 Related Work

Large Language Models. With the scaling of model and corpus sizes, Large Language Models (LLMs) (Devlin et al., 2018; Radford et al.,

2019; Lewis et al., 2019) have achieved remarkable success in various areas of natural language processing. To tailor a model for particular tasks, one approach is to fine-tune it with task-specific datasets (Jiao et al., 2023; Li and Liang, 2021; Hu et al., 2021). Jiao et al. (2023) introduce data with error annotations in fine-tuning to improve the machine translation abilities of open-source LLMs. The fine-tuning approach poses a challenge when applied to powerful black-box LLMs that only offer APIs for interaction, as it requires access to the underlying parameters. With the help of instruction tuning (Wei et al., 2021) and reinforcement learning from human feedback (Ouyang et al., 2022), recent LLMs can achieve gradient-free adaptation to various downstream tasks by prompting with natural language instructions, and some powerful capacities such as in-context learning (Brown et al., 2020) have also emerged.

Prompting Methods. Prompting is a general method for humans to interact with LLMs, which is usually designed as an instruction for a task that guides LLMs toward intended outputs (Schick and Schütze, 2020; Sanh et al., 2021). To make the most of LLMs on downstream tasks, the prompts need to be carefully designed, either manually (Hendy et al., 2023) or automatically (Gao et al., 2020; Zhou et al., 2022). Prompting also provides a way to interact with LLMs in natural language, such as letting them utilize external tools (Schick et al., 2023), resources (Ghazvininejad et al., 2023) and models (Wu et al., 2023; Shen et al., 2023), or conducting Chain-of-Thought (CoT) reasoning in generation (Wei et al., 2022a; Kojima et al., 2022). A concurrent work incorporates answers in pre-

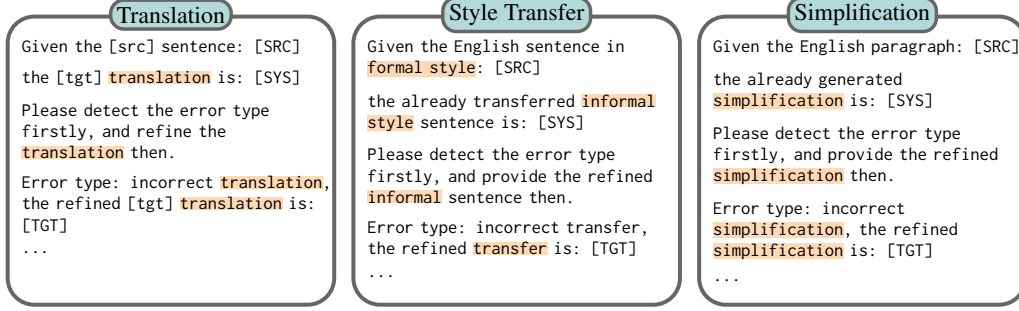


Figure 2: Illustration of DTG demonstration design for machine translation, style transfer and text simplification tasks. Due to the limited page, please refer to the Appendix for the remained 3 generation tasks, including dialogue summarization, paraphrase and commonsense generation.

vious rounds into prompts in an iterative process to improve the accuracy of LLMs on reasoning tasks (Zheng et al., 2023). Besides multi-step reasoning, basic prompts are still widely utilized in general text generation tasks such as machine translation and summarization, where previous advanced methods such as CoT have been shown ineffective (Peng et al., 2023). Our work finds its closest parallels in the domain of self-refinement or self-correction techniques (Madaan et al., 2023; Yao et al., 2023; Shinn et al., 2023). However, a distinguishing feature of our approach is its independence from the need for additional feedback or resources, setting it apart from these previously proposed methods.

3 Deliberate then Generate

Language acquisition by a human is normally based on both positive and negative feedback and improves the ability of language use through corrections. Inspired by this, unlike the conventional prompts only with correct information, we introduce a more deliberate approach termed Deliberate then Generate (DTG) prompting by facilitating LLMs to *detect errors on a synthesized text that may contain errors*.

3.1 The Overall Prompt Design

Specifically, the proposed DTG method unfolds in three stages: 1) It begins with a concise and explicit instruction of the desired task, providing guidance on generating an intended text based on a given input text; 2) A synthesized text is then provided as a candidate output; 3) Finally, DTG encourages the model to detect potential errors, and subsequently generate an improved output after thorough deliberation.

Figure 1 illustrates a comparison between standard prompting and our proposed DTG prompting for the summarization task in the one-shot scenario. A distinctive feature of DTG is its emphasis on error detection other than immediate response. Instead of generating the outcome directly from the given input text, DTG steers the model to make deliberate decisions by detecting the error type firstly based on both the input text, denoted as “[SRC]”, and a pre-defined candidate, denoted as “[SYS]”, before the final decisions. This deliberative process forms the bedrock of the DTG approach and will be further elaborated upon in the analysis section (i.e., Section 6). Besides, a few demonstrations can be provided, imbuing LLMs with an awareness of the expected output (highlighted in blue), and the test input (marked in red). DTG is a general prompting method that could be easily applied to any text generation task with minimal modifications to the prompt. Figure 2 illustrates the particular prompts used for 3 generation tasks we considered, indicating that minimal customization is required across different tasks as highlighted in yellow.

3.2 Choice of Synthesized Text ([SYS])

The choice of the synthesized text is another key part of DTG. Straightforwardly, using the output of LLMs themselves is a natural choice. However, these outputs typically necessitate only minor modifications, insufficient to adequately stimulate the LLMs’ deliberative capabilities. Also, our preliminary experiments show that using LLM’s output as [SYS] cannot gain any benefits, leading to a similar observation in Huang et al. (2023)’s work, that LLMs cannot self-correct reasoning yet without additional feedback. This limitation underscores the need for an alternative strategy that challenges the model to engage in more profound error detection

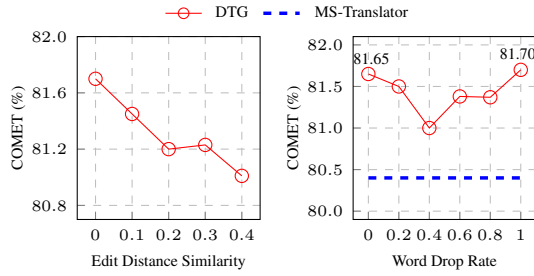


Figure 3: COMET scores against the edit distance (left) and the word drop rate (right) on the ZH-EN task.

and correction processes.

Our strategy explores the impact of synthesized text’s similarity to the reference on the quality of the generated output. Empirical evidence, as depicted in Figure 3 (left), demonstrates a clear trend: the performance of DTG inversely correlates with the similarity between the candidate and the reference text. This relationship is quantified using edit distance measures, where a lower similarity significantly enhances the generated text’s quality. Further experimentation involved modifying outputs from MS-Translator by selectively omitting words to create varied candidate sentences. The comparative analysis, illustrated in Figure 3 (right), reveals that DTG not only improves upon MS-Translator’s baseline COMET scores but also exhibits superior performance in refining candidates from external systems, highlighting its adaptability and efficacy in processing diverse input qualities.

In response to these findings, we advocate for the use of synthesized texts that diverge markedly from accurate information, including the use of an *empty string* (" ") as [SYS]. This particular type of null candidate significantly engages the model’s deliberative processes, leading to consistent improvements across a spectrum of generation tasks.

3.3 Definition of Error Types

Using an empty string as [SYS] in the DTG framework simplifies error categorization as “incorrect” by default. Yet, our findings suggest that delineating more specific error types markedly improves model correction effectiveness. Such precision in error identification sharpens the model’s focus, elevating accuracy and textual coherence. Take machine translation as an instance, one can tell LLMs potential error types, such as incorrect word translation, grammar error, under translation, incorrect entity translation, word-order error, or word repetition. Extending specific error typologies

to various text generation tasks further optimizes DTG’s utility. Adjusting error categories to task specifics, such as “factual inaccuracies” and “missing keywords” in summarization, underscores this method’s versatility and its potential to refine text generation across diverse applications.

4 Datasets and Evaluation

In experiments, we are devoted to evaluating the generation ability of LLMs and the proposed DTG prompting. We select 7 representative generation tasks, including machine translation, abstractive summarization, dialogue summarization, text simplification, style transfer, paraphrase and common-sense generation. Also, we expand the exploration to a mathematical reasoning task, namely GSM8K. We summarize the details of each dataset for each task, including the test sets, the selection of demonstrations (mostly from validation sets) and the corresponding prompts we have used. For more details please refer to the attached Appendix. Without meticulous parameter tuning, we set the *temperature* to 0 and *top_p* to 1 when calling the API.

5 Experiments

In this section, we assess the efficacy of the text-davinci-003 (denoted as GPT) across 7 sequence generation tasks. The chosen baseline comparisons consist of 1-shot, and few-shot (mostly 5-shot) scenarios. We also conduct further experiments with GPT4 for more convincing conclusions. Due to the considerable computational cost and API request constraints associated with the GPT4, it is challenging to perform extensive experiments. In the current manuscript, we only report the results on machine translation and text simplification.

5.1 Results on Machine Translation

We compare the performance of the standard prompting and our DTG with Microsoft Translator in addition to WMT SoTA systems. Table 1 presents the results in both 1-shot and 5-shot scenarios. The findings here indicate that our reimplementation aligns with the trends observed in Hendy et al. (2023), that 5-shot beats 1-shot in most language pairs. Benefiting from the deliberation, DTG effectively pushes the boundaries and leads to enhanced results across all to-English language pairs in both 1-shot and 5-shot settings based on GPT3.5 model. For instance, DTG method exhibits substantial BLEU score increases in DE-EN,

System	COMET-22↑	BLEU↑	COMET-22↑	BLEU↑	COMET-22↑	BLEU↑	COMET-22↑	BLEU↑
	DE-EN		ZH-EN		CS-EN		RU-EN	
WMT-Best†	85.0	33.4	81.0	33.5	89.0	64.2	86.0	45.1
MS-Translator†	84.7	33.5	80.4	27.9	87.4	54.9	85.2	43.9
GPT 1-shot	84.7	30.4	81.0	23.7	86.2	44.8	84.8	39.7
+ DTG	85.0	32.3	81.4	25.3	86.7	45.6	85.0	40.0
GPT 5-shot	85.3	32.3	81.1	23.6	86.9	47.2	84.9	39.9
+ DTG	85.4	33.2	81.7	25.2	87.0	47.4	85.1	40.3
GPT4 1-shot	85.6	33.5	82.4	26.0	87.3	48.1	86.1	43.1
+ DTG	85.8	33.8	83.0	26.4	87.7	49.4	86.3	43.7
	JA-EN		UK-EN		IS-EN		HA-EN	
WMT-Best†	81.6	24.8	86.0	44.6	87.0	41.7	80.0	21.0
MS-Translator†	81.5	24.5	83.5	42.4	85.9	40.5	73.3	16.2
GPT 1-shot	81.3	21.5	83.5	36.8	83.5	33.6	78.0	18.6
+ DTG	81.7	21.4	84.0	37.1	84.0	35.2	78.3	18.6
GPT 5-shot	81.2	20.5	84.0	38.0	84.1	35.0	78.3	18.8
+ DTG	82.2	22.4	84.2	39.0	84.6	36.0	78.6	19.2
GPT4 1-shot	83.4	24.7	85.7	39.9	86.9	39.9	77.5	18.3
+ DTG	83.6	25.2	85.9	40.6	87.0	40.9	77.9	18.9

Table 1: Evaluation results of GPT and GPT4 on six high-resource and two-low resource machine translation tasks from WMT Testsets. The best scores across different systems are marked in bold.

System	CNN/DailyMail			GigaWord			SamSum			DialogSum		
	R1	R2	RL	R1	R2	RL	R1	R2	RL	R1	R2	RL
Transformer (Vaswani et al., 2017)	40.47	17.73	37.29	37.57	18.90	34.69	37.20	10.86	34.69	35.91	8.74	33.50
BART (Lewis et al., 2020)	44.16	21.28	40.90	39.29	20.09	35.65	53.12	27.95	49.15	47.28	21.18	44.83
UniLMv2 (Bao et al., 2020)	43.16	20.42	40.14	-	-	-	50.53	26.62	48.81	47.04	21.13	45.04
GPT 1-shot	38.87	15.36	35.11	31.24	11.61	27.99	44.52	19.92	39.60	36.84	14.23	32.20
+ DTG	40.17	15.60	36.04	31.50	12.00	28.50	45.50	20.58	40.13	39.01	15.50	34.13
GPT 5-shot	-	-	-	33.04	12.78	29.86	46.44	20.69	41.10	40.86	17.10	35.78
+ DTG	-	-	-	33.54	13.63	30.36	48.72	23.16	43.23	42.64	18.12	37.57
GPT 10-shot	-	-	-	33.24	13.26	30.46	47.37	22.08	42.20	41.28	17.48	36.69
+ DTG	-	-	-	34.02	14.21	31.04	50.48	24.88	45.31	45.11	19.50	39.71

Table 2: Experimental results on four summarization tasks.

ZH-EN, and UK-EN language pairs in 5-shot scenarios. More concretely, DTG even beats WMT-Best system in terms of COMET-22, which is a more recognized metric recently in the machine translation literature. Moreover, the consistent improvements on IS-EN and HA-EN demonstrate the effectiveness of DTG in low-resource settings. Benefiting the strong comprehension ability of GPT4, we find no significant difference between 1-shot and 5-shot scenarios. Meanwhile, DTG is still effective on GPT4, showing consistent and indeed improvements in terms of COMET and BLEU. This finding demonstrates much stronger LLMs can still benefit from deliberation.

5.2 Results on Summarization

For abstractive summarization, we assess GPT models on CNN/DailyMail¹ and GigaWord, two

¹Due to the limit of max length for GPT models (4097) and the long input length of CNN/DailyMail, we only evaluate

benchmark datasets in the field. Additionally, we explore their efficacy in dialogue summarization, including SamSum and DialogSum², two hybrid tasks combining aspects of both dialogue and summarization. As shown in Table 2, GPT models show comparative performance with Transformer which is specially tuned on the downstream training set, e.g., Transformer. Our DTG delivers further improvements in terms of three ROUGE metrics, which demonstrate the effectiveness of DTG on long-term modeling task. Beyond this, DTG substantially incites GPT models to generate more precise summaries derived from extensive multi-turn dialogues. An upward trend in performance is observed with the introduction of additional demonstrations, further underscoring the effectiveness of the DTG method. However, DTG still lags

the performance in 1-shot scenario.

²It is important to note that the results for DialogSum are averaged over three individual scores, each calculated using unique references spanning a range of topics.

System	GYAFC & EM		GYAFC & FR		Amazon		Yelp	
	BLEU	BLEURT	BLEU	BLEURT	BLEU	BLEURT	BLEU	BLEURT
Transformer†(Vaswani et al., 2017)	40.3	-	47.7	-	-	-	-	-
BART†(Lewis et al., 2020)	76.9	75.38	79.3	75.11	-	-	-	-
GPT 1-shot	52.9	73.42	44.6	70.73	36.1	64.56	30.9	64.03
+ DTG	66.8	75.20	65.9	74.60	35.4	63.60	31.3	64.19
GPT 5-shot	61.3	75.40	63.9	74.35	39.3	64.76	31.4	64.16
+ DTG	69.9	76.36	74.1	75.43	40.9	65.42	32.2	64.87

Table 3: Comparisons of 1-shot and 5-shot on four style transfer tasks, including Entertainment Music, Family Relationships, Amazon and Yelp. †denotes results borrowed from (Lai et al., 2021).

System	Asset	Wiki-auto
MUSS (Martin et al., 2022)	44.15	42.59
Control Prefix (Clive et al., 2022)	43.58	-
TST-Final (Omelianchuk et al., 2021)	41.46	-
GPT 1-shot	46.12	44.97
+ DTG	47.23	47.15
GPT 5-shot	45.95	45.12
+ DTG	47.05	47.54
GPT4 5-shot	47.10	45.96
+ DTG	47.67	47.89

Table 4: Comparisons of 1-shot, 5-shot with and without our DTG method on two text simplification tasks.

System	BLEU-3/4	ROUGE-2/L
BART (Lewis et al., 2020)	36.3/26.4	22.23/41.98
T5-Large (Raffel et al., 2020)	39.0/28.6	22.01/42.97
GPT 5-shot	39.7/30.0	25.28/46.55
+ DTG	43.2/33.5	27.02/48.47

Table 5: Results on the CommonGen benchmark.

System	Accuracy
GPT 8-shot	55.1
+ DTG	60.0
CoT 8-shot (Wei et al., 2022b)	59.8
+ DTG	64.5

Table 6: Results of GSM8K on DTG prompting.

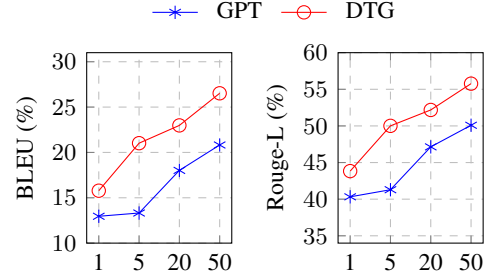


Figure 4: BLEU and ROUGE-L scores against the number of demonstrations on the paraphrase task.

behind of large-scale pretrained models, such as BART (Lewis et al., 2020) and UniLMv2 (Bao et al., 2020) in automatic evaluations. We will add more human-alignment judgment in Section 6.

5.3 Results on Style Transfer

Table 3 displays performance across style transfer tasks from the GYAFC dataset: Entertainment Music (EM) and Family Relationships (FR), both involving informal to formal transformations. Evidently, the Deliberate then Generate (DTG) method prompts the GPT model to correct inaccuracies and generate more precise informal sentences. Specifically, DTG achieves an 8-point and 10.04-point increase in BLEU score for EM and FR tasks, respectively, compared to standard prompting. Although DTG trails BART (Lewis et al., 2020) in BLEU scores, it surpasses BART in BLEURT scores, obtaining gains of 0.98 and 0.32 for EM and FR tasks, respectively. These results highlight the potential of LLMs and DTG method in style transfer tasks.

5.4 Results on Text Simplification

Experiments were conducted on two text simplification benchmarks, Asset and Wiki-Auto, where the primary goal is to create a simplified rendition of the given text input. The main evaluation metric is the SARI score. Our findings illustrate that GPT models demonstrate robust performance across both simplification benchmarks, even surpassing the existing state-of-the-art models (MUSS) built based on BART. Furthermore, the incorporation of DTG method significantly enhances GPT model performance, leading to improvements in both BLEU and SARI scores. Specifically, DTG establishes a new benchmark for state-of-the-art results on these two simplification tasks.

5.5 Results on Commonsense Generation

Table 5 summarizes the comparison between GPT models with and without DTG method on an open Commonsense generation benchmark. This task is more flexible than the aforementioned, meanwhile raising the evaluation difficulty. We see

#	System	BLEU	COMET
1	GPT 5-shot	23.6	81.12
2	+ DTG	25.2	81.70
3	+ w/o error detection	23.3	81.05
4	+ wrong error type	25.3	81.74
5	+ fixed error type	24.1	81.35
6	+ task-specific error type	25.5	81.98
7	+ fixed incorrect candidate	25.0	81.72
8	+ irrelevant languages	25.1	81.81
9	+ correct candidate	23.0	81.17

Table 7: Ablations on error types and candidate types.

that GPT models with standard prompting even surpass large-scale pretrained generation models, such as BART (Lewis et al., 2019) and T5 (Rafael et al., 2020). DTG achieves further improvements in terms of BLEU-3/BLEU-4 and ROUGE-2/ROUGE-L, resulting in an average of 3.50 BLEU and almost 2.00 ROUGE improvements. This also establishes a new SoTA on this benchmark.

5.6 Results on Paraphrase

Figure 4 plots the BLEU and ROUGE-L scores for GPT and DTG in relation to various few-shot scenarios. We find that DTG outperforms GPT models in terms of both BLEU and ROUGE-L metrics across all scenarios. However, only 5-shot demonstrations cannot enable LLMs to clearly capture the underlying mapping rule between the source and the target. Interestingly, a significant enhancement in DTG performance is observed with the increase in the number of demonstrations. This improvement can be attributed to the model’s enhanced ability to comprehend the underlying mapping rules with the expanded demonstration set.

5.7 Results on Mathematical Reasoning

While our primary focus is on evaluating LLMs for text generation, we extend our analysis to reasoning tasks, such as GSM8K (Cobbe et al., 2021). Table 6 compares the accuracy of standard prompting, CoT, and DTG. Our results show that DTG, when combined with CoT, achieves an accuracy of 64.5 in 8-shot scenarios, indicating its utility beyond text generation.

6 Analysis

In this section, we delve into a series of intriguing questions to answer why DTG works. Unless specified otherwise, the base engine utilized throughout this investigation is text-davinci-003.

Ablations on Error Types Prior research underscores the significant impact of both the quality and quantity of demonstrations (Zhang et al., 2023; Vilar et al., 2022; Agrawal et al., 2022). Thus, we would like to discern whether the improvements are attributable to template modifications or the deliberate capability inherent to the LLMs. Table 7 summarizes the comparisons on WMT ZH-EN. Firstly, DTG experiences a significant degradation in BLEU score when removing the explicitly error detection prompt³, suggesting that the excised segment may contain crucial triggers stimulating the deliberate capability of the LLM. Along this line, by comparing #4⁴, #5⁵ and # 6 with #2, we can conclude 1) LLMs can rethink by themselves and make “correct” decisions though the demonstration is incorrect. 2) Restricting the thought of LLMs would hinder their performance. 3) Adding task-specific error types (See Section 3.3) results in better generation.

Ablations on Candidates Here, we aim to explore if other candidates rather than *empty string* may also prove effective in DTG. The last three lines in Table 7 show the comparison. Specifically, the term “fixed incorrect candidate” (#7) refers to the use of a fixed yet incorrect (irrelevant) English translation as the [SYS].⁶ Likewise, system #8 indicates that the candidates neither belong to the target language nor conform to the correct structure or grammar.⁷ Interestingly, both 2 systems deliver comparable performance with our default setting, with system #8 even achieving a higher COMET score. However, when shifting to a correct candidate generated by itself, LLMs seem to underperform. This observation is aligned with that in (Huang et al., 2023)’s work that LLMs cannot refine itself yet. Our results also suggest that LLMs can effectively deliberate when the candidate is incorrect - whether it is an *empty string* or other incorrect translations - and subsequently generate a substantially improved translation.

³eliminating the phrase “Please detect the error type firstly, and refine the translation then”

⁴replacing “incorrect translation” with “good/correct translation” in the demonstration only

⁵replacing “incorrect translation” with “good/correct translation” in the demonstration only

⁶We random sample an English sentence: [SYS]: *EBA Education Team together with Accace Ukraine invite you to join the EBA Education Update: Performance Audit.*

⁷Similarly, we random sample an Ukraine sentence: [SYS]: *З впевненістю можете довіряти нам і будь ласка, звертайтеся до нас, є які-небудь чи коментарі.*

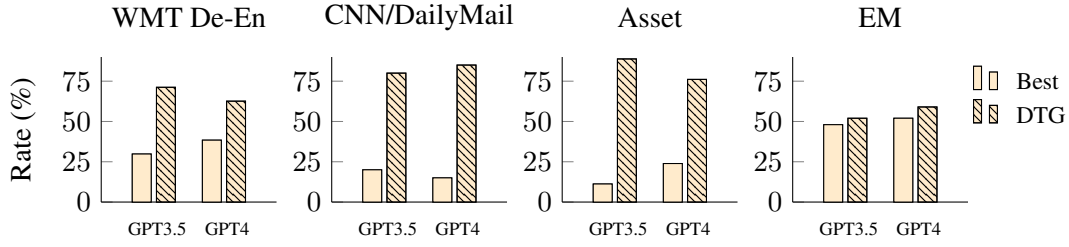


Figure 5: GPT3.5 and GPT4 evaluation on 4 generation tasks. Note that we random select 500 samples due to the limitation of GPT4 access.

System	ZH-EN		Asset	
	BLEU	Human	SARI	Human
GPT 5-shot	22.8	4.16	45.95	11.6%
DTG 5-shot	24.9	4.39	47.05	67.4%

Table 8: Human evaluation on DTG prompting.

Evaluation by GPT Models As previously discussed, despite DTG’s impressive performance, it falls short of BART in some scenarios—most notably, it exhibits a significant gap in terms of ROUGE scores in summarization tasks. However, [Liu et al. \(2023\)](#) suggested that ROUGE may not accurately represent the true performance of summarization tasks, given its poor alignment with human evaluations. In contrast, GPT models achieve optimal alignment with human justification and substantially outperform all previous SoTA evaluators on the SummEval benchmark. This observation prompts an investigation into whether the generation output by DTG can surpass that of BART. Following their suggestion, we conduct reference-based evaluation and design a prompt as shown in Figure 9. We extract 500 test sets and compared DTG with the best result using GPT3.5 and GPT4 to select a better candidate. Results in Figure 5 reveal that DTG significantly beats the best system within GPT evaluation.

Human Evaluation We further conducted human-evaluation with human assessments on translation (randomly selected 500 cases) and simplification tasks to mitigate potential bias in GPT models favoring their own outputs. Annotators scored ZH-EN translations on a 1-5 scale and indicated preferences for the Asset task. It’s worth noting that for the Asset task, some cases showed no significant difference in performance between the two methods (neutral). Detailed scoring guidelines are provided in the Appendix. As shown in Table 8, DTG outperforms the standard prompt in human evaluations across both tasks.

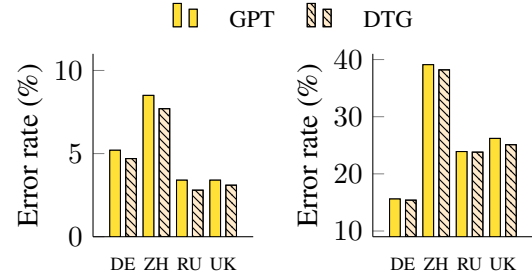


Figure 6: Statistics of error rate for under translation (above) and entity translation (below).

Error Statistical Analysis To evaluate whether the proposed DTG prompting can facilitate error avoidance in GPT, we conduct error statistics on machine translation, where two frequently occurring error types are considered (i.e., under translation and incorrect entity translation) ([Hassan et al., 2018](#)). Figure 6 provides a comparison of the error rates between GPT models with and without the application of the DTG method. It is obvious to see that DTG reduces both error rates compared with the direct generation manner.

7 Conclusions

In this paper, we propose DTG prompting, which encourages LLMs to deliberate before generating the final results by letting the model detect the error type on a synthetic text that may contain errors. Using an empty string as the synthetic text successfully gets rid of an extra baseline system and improves the quality of the generated text. The DTG prompting can be easily applied to various text generation tasks with minimal adjustments in the prompt. Extensive experiments conducted on over 20 datasets across 7 text generation tasks demonstrate the effectiveness and broad applicability of the DTG prompting.

Limitation

Due to restricted access to GPT4, we have evaluated our *Deliberate then Generate* (DTG) method on just two generation tasks: machine translation (across 8 language pairs) and simplification. There exists a necessity for more expansive experimentation across other tasks. Additionally, the effectiveness of DTG is contingent on model capacity. Models such as LLaMa-7B might not fully comprehend the instructions provided, resulting in weaker performance on downstream tasks. In our future work, we aim to ascertain the required scale of a language model to successfully facilitate deliberative generation.

Our work inherits the biases from pre-trained language models. For example, we only conduct experiments on English generation that GPT models are most powerful at. We provide results and analysis on English-to-Others translation in Appendix D. Future works could investigate the performance of DTG on multilingual pre-trained models.

We have experimented with multiple decoding iterations using the DTG framework. The observed performance gains were subtle, suggesting that DTG’s primary benefits are rooted in harnessing and augmenting the diverse capabilities acquired during pre-training, e.g., detection and refinement abilities. We would like to address this issue in our future work.

Ethical Statement

All experiments in our work were conducted on existing datasets commonly employed in prior publicly available research publications. We keep fair and honest in our analysis of experimental results, and our work does not harm anyone. Additionally, we will make our code accessible for future investigations.

References

Sweta Agrawal, Chunting Zhou, Mike Lewis, Luke Zettlemoyer, and Marjan Ghazvininejad. 2022. In-context examples selection for machine translation. *arXiv preprint arXiv:2212.02437*.

Fernando Alva-Manchego, Louis Martin, Antoine Bordes, Carolina Scarton, Benoît Sagot, and Lucia Specia. 2020. *ASSET: A dataset for tuning and evaluation of sentence simplification models with multiple rewriting transformations*. In *Proceedings of the 58th Annual Meeting of the Association for Computational*

Linguistics, pages 4668–4679, Online. Association for Computational Linguistics.

Hangbo Bao, Li Dong, Furu Wei, Wenhui Wang, Nan Yang, Xiaodong Liu, Yu Wang, Jianfeng Gao, Songhao Piao, Ming Zhou, and Hsiao-Wuen Hon. 2020. *Unilmv2: Pseudo-masked language models for unified language model pre-training*. In *Proceedings of the 37th International Conference on Machine Learning, ICML 2020, 13-18 July 2020, Virtual Event*, volume 119 of *Proceedings of Machine Learning Research*, pages 642–652. PMLR.

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.

Yulong Chen, Yang Liu, Liang Chen, and Yue Zhang. 2021. *DialogSum: A real-life scenario dialogue summarization dataset*. In *Findings of the Association for Computational Linguistics: ACL-IJCNLP 2021*, pages 5062–5074, Online. Association for Computational Linguistics.

Jordan Clive, Kris Cao, and Marek Rei. 2022. *Control prefixes for parameter-efficient text generation*. In *Proceedings of the 2nd Workshop on Natural Language Generation, Evaluation, and Metrics (GEM)*, pages 363–382, Abu Dhabi, United Arab Emirates (Hybrid). Association for Computational Linguistics.

Karl Cobbe, Vineet Kosaraju, Mohammad Bavarian, Mark Chen, Heewoo Jun, Lukasz Kaiser, Matthias Plappert, Jerry Tworek, Jacob Hilton, Reiichiro Nakano, Christopher Hesse, and John Schulman. 2021. Training verifiers to solve math word problems. *arXiv preprint arXiv:2110.14168*.

Jacob Devlin, Ming-Wei Chang, Kenton Lee, and Kristina Toutanova. 2018. Bert: Pre-training of deep bidirectional transformers for language understanding. *arXiv preprint arXiv:1810.04805*.

Tianyu Gao, Adam Fisch, and Danqi Chen. 2020. Making pre-trained language models better few-shot learners. *arXiv preprint arXiv:2012.15723*.

Marjan Ghazvininejad, Hila Gonen, and Luke Zettlemoyer. 2023. Dictionary-based phrase-level prompting of large language models for machine translation. *arXiv preprint arXiv:2302.07856*.

Bogdan Gliwa, Iwona Mochol, Maciej Biesek, and Aleksander Wawer. 2019. *SAMSum corpus: A human-annotated dialogue dataset for abstractive summarization*. In *Proceedings of the 2nd Workshop on New Frontiers in Summarization*, pages 70–79, Hong Kong, China. Association for Computational Linguistics.

Shansan Gong, Mukai Li, Jiangtao Feng, Zhiyong Wu, and LingPeng Kong. 2022. Diffuseq: Sequence to sequence text generation with diffusion models. *arXiv preprint arXiv:2210.08933*.

627	Hany Hassan, Anthony Aue, Chang Chen, Vishal	for natural language generation, translation, and com-	682
628	Chowdhary, Jonathan Clark, Christian Federmann,	prehension. In <i>Proceedings of the 58th Annual Meet-</i>	683
629	Xuedong Huang, Marcin Junczys-Dowmunt, William	<i>ing of the Association for Computational Linguistics</i> ,	684
630	Lewis, Mu Li, et al. 2018. Achieving human par-	pages 7871–7880, Online. Association for Computa-	685
631	ity on automatic chinese to english news translation.	tional Linguistics.	686
632	<i>arXiv preprint arXiv:1803.05567</i> .		
633	Amr Hendy, Mohamed Abdelrehim, Amr Sharaf,	Xiang Lisa Li and Percy Liang. 2021. Prefix-tuning:	687
634	Vikas Raunak, Mohamed Gabr, Hitokazu Matsushita,	Optimizing continuous prompts for generation. <i>arXiv</i>	688
635	Young Jin Kim, Mohamed Afify, and Hany Hassan	<i>preprint arXiv:2101.00190</i> .	689
636	Awadalla. 2023. How good are gpt models at ma-	Bill Yuchen Lin, Wangchunshu Zhou, Ming Shen, Pei	690
637	chine translation? a comprehensive evaluation. <i>arXiv</i>	Zhou, Chandra Bhagavatula, Yejin Choi, and Xiang	691
638	<i>preprint arXiv:2302.09210</i> .	Ren. 2020. CommonGen: A constrained text gen-	692
639	Edward J Hu, Yelong Shen, Phillip Wallis, Zeyuan	eration challenge for generative commonsense rea-	693
640	Allen-Zhu, Yuanzhi Li, Shean Wang, Lu Wang,	soning . In <i>Findings of the Association for Computa-</i>	694
641	and Weizhu Chen. 2021. Lora: Low-rank adap-	<i>tional Linguistics: EMNLP 2020</i> , pages 1823–1840,	695
642	tation of large language models. <i>arXiv preprint</i>	Online. Association for Computational Linguistics.	696
643	<i>arXiv:2106.09685</i> .	Chin-Yew Lin. 2004. ROUGE: A package for auto-	697
644	Jie Huang, Xinyun Chen, Swaroop Mishra,	matic evaluation of summaries . In <i>Text Summariza-</i>	698
645	Huaixiu Steven Zheng, Adams Wei Yu, Xiny-	<i>tion Branches Out</i> , pages 74–81, Barcelona, Spain.	699
646	ing Song, and Denny Zhou. 2023. Large language	Association for Computational Linguistics.	700
647	models cannot self-correct reasoning yet. <i>arXiv</i>	Yang Liu, Dan Iter, Yichong Xu, Shuohang Wang,	701
648	<i>preprint arXiv:2310.01798</i> .	Ruochen Xu, and Chenguang Zhu. 2023. Gpteval:	702
649	Chao Jiang, Mounica Maddela, Wuwei Lan, Yang	Nlg evaluation using gpt-4 with better human align-	703
650	Zhong, and Wei Xu. 2020. Neural CRF model for	ment. <i>arXiv preprint arXiv:2303.16634</i> .	704
651	sentence alignment in text simplification . In <i>Proceed-</i>	Aman Madaan, Niket Tandon, Prakhar Gupta, Skyler	705
652	<i>ings of the 58th Annual Meeting of the Association</i>	Hallinan, Luyu Gao, Sarah Wiegrefe, Uri Alon,	706
653	<i>for Computational Linguistics</i> , pages 7943–7960, On-	Nouha Dziri, Shrimai Prabhumoye, Yiming Yang,	707
654	line. Association for Computational Linguistics.	et al. 2023. Self-refine: Iterative refinement with	708
655	Wenxiang Jiao, Jen tse Huang, Wenxuan Wang, Zhi-	self-feedback. <i>arXiv preprint arXiv:2303.17651</i> .	709
656	wei He, Tian Liang, Xing Wang, Shuming Shi, and	Gary F Marcus. 1993. Negative evidence in language	710
657	Zhaopeng Tu. 2023. Parrot: Translating during chat	acquisition. <i>Cognition</i> , 46(1):53–85.	711
658	using large language models tuned with human trans-	Louis Martin, Angela Fan, Éric de la Clergerie, Antoine	712
659	lation and feedback. In <i>Findings of EMNLP</i> .	Bordes, and Benoît Sagot. 2022. MUSS: Multilin-	713
660	Takeshi Kojima, Shixiang Shane Gu, Machel Reid, Yu-	gual unsupervised sentence simplification by mining	714
661	taka Matsuo, and Yusuke Iwasawa. 2022. Large lan-	paraphrases . In <i>Proceedings of the Thirteenth Lan-</i>	715
662	guage models are zero-shot reasoners. <i>arXiv preprint</i>	<i>guage Resources and Evaluation Conference</i> , pages	716
663	<i>arXiv:2205.11916</i> .	1651–1664, Marseille, France. European Language	717
664	Huiyuan Lai, Antonio Toral, and Malvina Nissim. 2021.	Resources Association.	718
665	Thank you BART! rewarding pre-trained models im-	Harsha Nori, Nicholas King, Scott Mayer McKinney,	719
666	proves formality style transfer . In <i>Proceedings of the</i>	Dean Carignan, and Eric Horvitz. 2023. Capabili-	720
667	<i>59th Annual Meeting of the Association for Computa-</i>	ties of gpt-4 on medical challenge problems. <i>arXiv</i>	721
668	<i>tional Linguistics and the 11th International Joint</i>	<i>preprint arXiv:2303.13375</i> .	722
669	<i>Conference on Natural Language Processing (Vol-</i>	Kostiantyn Omelianchuk, Vipul Raheja, and Oleksandr	723
670	<i>ume 2: Short Papers)</i> , pages 484–494, Online. Asso-	Skurzhashnyi. 2021. Text Simplification by Tagging .	724
671	ciation for Computational Linguistics.	In <i>Proceedings of the 16th Workshop on Innovative</i>	725
672	Mike Lewis, Yinhan Liu, Naman Goyal, Marjan	<i>Use of NLP for Building Educational Applications</i> ,	726
673	Ghazvininejad, Abdelrahman Mohamed, Omer Levy,	pages 11–25, Online. Association for Computational	727
674	Ves Stoyanov, and Luke Zettlemoyer. 2019. Bart: De-	Linguistics.	728
675	noising sequence-to-sequence pre-training for natural	OpenAI. 2023. Gpt-4 technical report. <i>arXiv preprint</i>	729
676	language generation, translation, and comprehension.	<i>arXiv:2303.08774</i> .	730
677	<i>arXiv preprint arXiv:1910.13461</i> .	Long Ouyang, Jeffrey Wu, Xu Jiang, Diogo Almeida,	731
678	Mike Lewis, Yinhan Liu, Naman Goyal, Marjan	Carroll Wainwright, Pamela Mishkin, Chong Zhang,	732
679	Ghazvininejad, Abdelrahman Mohamed, Omer Levy,	Sandhini Agarwal, Katarina Slama, Alex Ray, et al.	733
680	Veselin Stoyanov, and Luke Zettlemoyer. 2020.	2022. Training language models to follow instruc-	734
681	BART: Denoising sequence-to-sequence pre-training	tions with human feedback. <i>Advances in Neural</i>	735
		<i>Information Processing Systems</i> , 35:27730–27744.	736

737	Keqin Peng, Liang Ding, Qihuang Zhong, Li Shen,	Yongliang Shen, Kaitao Song, Xu Tan, Dongsheng Li,	791
738	Xuebo Liu, Min Zhang, Yuanxin Ouyang, and	Weiming Lu, and Yueting Zhuang. 2023. Hugging-	792
739	Dacheng Tao. 2023. Towards making the most	gpt: Solving ai tasks with chatgpt and its friends in	793
740	of chatgpt for machine translation. <i>arXiv preprint</i>	huggingface. <i>arXiv preprint arXiv:2303.17580</i> .	794
741	<i>arXiv:2303.13780</i> .		
742	Matt Post. 2018. A call for clarity in reporting BLEU	Noah Shinn, Beck Labash, and Ashwin Gopinath.	795
743	scores . In <i>Proceedings of the Third Conference on</i>	2023. Reflexion: an autonomous agent with dy-	796
744	<i>Machine Translation: Research Papers</i> , pages 186–	namic memory and self-reflection. <i>arXiv preprint</i>	797
745	191, Belgium, Brussels. Association for Computa-	<i>arXiv:2303.11366</i> .	798
746	tional Linguistics.		
747	Alec Radford, Jeffrey Wu, Rewon Child, David Luan,	Hugo Touvron, Thibaut Lavril, Gautier Izacard, Xavier	799
748	Dario Amodei, Ilya Sutskever, et al. 2019. Language	Martinet, Marie-Anne Lachaux, Timothée Lacroix,	800
749	models are unsupervised multitask learners. <i>OpenAI</i>	Baptiste Rozière, Naman Goyal, Eric Hambro,	801
750	<i>blog</i> , 1(8):9.	Faisal Azhar, et al. 2023. Llama: Open and effi-	802
751	Colin Raffel, Noam Shazeer, Adam Roberts, Katherine	cient foundation language models. <i>arXiv preprint</i>	803
752	Lee, Sharan Narang, Michael Matena, Yanqi Zhou,	<i>arXiv:2302.13971</i> .	804
753	Wei Li, and Peter J. Liu. 2020. Exploring the limits	Ashish Vaswani, Noam Shazeer, Niki Parmar, Jakob	805
754	of transfer learning with a unified text-to-text trans-	Uszkoreit, Llion Jones, Aidan N. Gomez, Lukasz	806
755	former . <i>J. Mach. Learn. Res.</i> , 21:140:1–140:67.	Kaiser, and Illia Polosukhin. 2017. Attention is all	807
756	Sudha Rao and Joel Tetreault. 2018. Dear sir or madam,	you need. In <i>Proc. of NeurIPS</i> , pages 5998–6008.	808
757	may I introduce the GYAFC dataset: Corpus, bench-	David Vilar, Markus Freitag, Colin Cherry, Jiaming Luo,	809
758	marks and metrics for formality style transfer . In	Viresh Ratnakar, and George Foster. 2022. Prompt-	810
759	<i>Proceedings of the 2018 Conference of the North</i>	ing palm for translation: Assessing strategies and	811
760	<i>American Chapter of the Association for Computa-</i>	performance. <i>arXiv preprint arXiv:2211.09102</i> .	812
761	<i>tional Linguistics: Human Language Technologies,</i>	Jason Wei, Maarten Bosma, Vincent Y Zhao, Kelvin	813
762	<i>Volume 1 (Long Papers)</i> , pages 129–140, New Or-	Guu, Adams Wei Yu, Brian Lester, Nan Du, An-	814
763	leans, Louisiana. Association for Computational Lin-	drew M Dai, and Quoc V Le. 2021. Finetuned lan-	815
764	guistics.	guage models are zero-shot learners. <i>arXiv preprint</i>	816
765	Ricardo Rei, José G. C. de Souza, Duarte Alves,	<i>arXiv:2109.01652</i> .	817
766	Chrysoula Zerva, Ana C Farinha, Taisiya Glushkova,	Jason Wei, Xuezhi Wang, Dale Schuurmans, Maarten	818
767	Alon Lavie, Luisa Coheur, and André F. T. Martins.	Bosma, Ed Chi, Quoc Le, and Denny Zhou. 2022a.	819
768	2022. COMET-22: Unbabel-IST 2022 submission	Chain of thought prompting elicits reasoning in large	820
769	for the metrics shared task . In <i>Proceedings of the</i>	language models. <i>arXiv preprint arXiv:2201.11903</i> .	821
770	<i>Seventh Conference on Machine Translation (WMT)</i> ,	Jason Wei, Xuezhi Wang, Dale Schuurmans, Maarten	822
771	pages 578–585, Abu Dhabi, United Arab Emirates	Bosma, Brian Ichter, Fei Xia, Ed H. Chi, Quoc V. Le,	823
772	(Hybrid). Association for Computational Linguistics.	and Denny Zhou. 2022b. Chain-of-thought prompt-	824
773	Victor Sanh, Albert Webson, Colin Raffel, Stephen H	ing elicits reasoning in large language models . In	825
774	Bach, Lintang Sutawika, Zaid Alyafeai, Antoine	<i>NeurIPS</i> .	826
775	Chaffin, Arnaud Stiegler, Teven Le Scao, Arun	Chenfei Wu, Shengming Yin, Weizhen Qi, Xi-	827
776	Raja, et al. 2021. Multitask prompted training en-	aodong Wang, Zecheng Tang, and Nan Duan.	828
777	ables zero-shot task generalization. <i>arXiv preprint</i>	2023. Visual chatgpt: Talking, drawing and edit-	829
778	<i>arXiv:2110.08207</i> .	ing with visual foundation models. <i>arXiv preprint</i>	830
779	Timo Schick, Jane Dwivedi-Yu, Roberto Dessì, Roberta	<i>arXiv:2303.04671</i> .	831
780	Raileanu, Maria Lomeli, Luke Zettlemoyer, Nicola	Wei Xu, Courtney Napoles, Ellie Pavlick, Quanze Chen,	832
781	Cancedda, and Thomas Scialom. 2023. Toolformer:	and Chris Callison-Burch. 2016. Optimizing sta-	833
782	Language models can teach themselves to use tools.	tistical machine translation for text simplification .	834
783	<i>arXiv preprint arXiv:2302.04761</i> .	<i>Transactions of the Association for Computational</i>	835
784	Timo Schick and Hinrich Schütze. 2020. Exploit-	<i>Linguistics</i> , 4:401–415.	836
785	ing cloze questions for few shot text classification	Shunyu Yao, Jeffrey Zhao, Dian Yu, Nan Du, Izhak	837
786	and natural language inference. <i>arXiv preprint</i>	Shafraan, Karthik R. Narasimhan, and Yuan Cao. 2023.	838
787	<i>arXiv:2001.07676</i> .	React: Synergizing reasoning and acting in language	839
788	Thibault Sellam, Dipanjan Das, and Ankur P Parikh.	models . In <i>The Eleventh International Conference</i>	840
789	2020. Bleurt: Learning robust metrics for text gener-	<i>on Learning Representations, ICLR 2023, Kigali,</i>	841
790	ation. In <i>Proceedings of ACL</i> .	<i>Rwanda, May 1-5, 2023</i> . OpenReview.net.	842
		Biao Zhang, Barry Haddow, and Alexandra Birch. 2023.	843
		Prompting large language model for machine transla-	844
		tion: A case study. <i>arXiv preprint arXiv:2301.07069</i> .	845

846 Chuanyang Zheng, Zhengying Liu, Enze Xie, Zhenguo
847 Li, and Yu Li. 2023. Progressive-hint prompting
848 improves reasoning in large language models. *arXiv*
849 *preprint arXiv:2304.09797*.

850 Yongchao Zhou, Andrei Ioan Muresanu, Ziwen Han,
851 Keiran Paster, Silviu Pitis, Harris Chan, and Jimmy
852 Ba. 2022. Large language models are human-level
853 prompt engineers. *arXiv preprint arXiv:2211.01910*.

854 A Datasets and Evaluation 854

855 In experiments, we are devoted to evaluating the
856 generation ability of LLMs and the proposed DTG
857 prompting. We select 7 representative generation
858 tasks, including machine translation, abstractive
859 summarization, dialogue summarization, text sim-
860 plification, style transfer, paraphrase and common-
861 sense generation.

862 **Machine Translation** For the machine transla-
863 tion task, we aligned with [Hendy et al. \(2023\)](#)’s
864 work and experimented on both high-resource and
865 low-resource scenarios. For the high-resource
866 setting, we include German, Czech, Chinese,
867 Japanese, Russian, and Ukrainian paired with En-
868 glish. In the low-resource context, we examine
869 Icelandic and Hausa. The performance is evaluated
870 in terms of SacreBLEU⁸ ([Post, 2018](#)), ChrF, TER
871 (translation error rate) and COMET-22 ([Rei et al.,](#)
872 [2022](#)).

873 **Abstractive Summarization** We also evaluate
874 LLM’s ability to process long sequence on CNN-
875 DailyMail and Gigaword, two widely used abstrac-
876 tive summarization datasets. The evaluation metric
877 is F1-ROUGE ([Lin, 2004](#)), consisting of ROUGE-
878 1, ROUGE-2 and ROUGE-L.

879 **Dialog Summarization** Dialogue summarization
880 presents greater challenges than traditional text
881 summarization due to the intricate conversation
882 contexts that models need to comprehend, though
883 their contexts are relatively shorter. This attribute
884 enables us to test few-shot abilities due to the re-
885 stricted input length. To investigate this, we select
886 SamSum⁹ ([Gliwa et al., 2019](#)) and DialogSum¹⁰
887 ([Chen et al., 2021](#)), two benchmark datasets for
888 dialogue summarization. The evaluation metric is
889 the same as abstractive summarization.

890 **Text Simplification** The purpose of text simpli-
891 fication is to revise complex text into sequences
892 with simplified grammar and word choice. In this
893 work, we mainly report the performance on two
894 benchmarks, namely Asset ([Alva-Manchego et al.,](#)
895 [2020](#)) and Wiki-auto ([Jiang et al., 2020](#)). Asset is
896 a multi-reference dataset for the evaluation of sen-
897 tence simplification in English. The dataset uses
898 the same 2,359 sentences from TurkCorpus ([Xu](#)

⁸BLEU+case.mixed+numrefs.1+smooth.exp+tok.13a
+version.2.3.1

⁹<https://huggingface.co/datasets/samsum>

¹⁰<https://github.com/cylnlp/DialogSum>

et al., 2016) and each sentence is associated with 10 crowdsourced simplifications. Similarly, each test set in Wiki-auto owns 8 references. We use SacreBLEU and BLEURT as the metric.

Style Transfer We used three widely-used English transfer learning datasets, namely Grammarly’s Yahoo Answers Formality Corpus (GYAFC), Amazon and Yelp reviews. The GYAFC dataset (Rao and Tetreault, 2018) was originally a question-and-answer dataset on an online forum, consisting of informal and formal sentences from the two categories: Entertainment & Music (EM) and Family & Relationships (FR). Both FR and EM provide 4 references to evaluate the fidelity. The Amazon dataset is a product review dataset, labeled as either a positive or negative sentiment. Similarly, the Yelp dataset is a restaurant and business review dataset with positive and negative sentiments. Both Amazon and Yelp are single-reference. The evaluation metrics contain BLEU and BLEURT (Sellam et al., 2020).

Paraphrase We endeavor to evaluate the paraphrase ability of LLMs upon the well-known Quora Question Pairs (QQP) dataset, which requires generating an alternative surface form in the same language expressing the same semantic content. We utilize the preprocessed data from (Gong et al., 2022). The evaluation metrics covers BLEU and ROUGE-L for a comprehensive comparison.

Common Sense Generation We choose CommonGen (Lin et al., 2020), a novel constrained generation task that requires models to generate a coherent sentence with the providing key concepts. We report both BLEU-3/4 and ROUGE-2/L to keep a fair comparison with results in prior work (Lin et al., 2020).

Reasoning For the reasoning task, we evaluate our method on a widely used benchmark, GSM8K (Cobbe et al., 2021), a challenging dataset consisting of high-quality linguistically diverse grade school math word problems. We report the accuracy of the 8-shot demonstration on the test set including 1,319 mathematical questions.

B Details of Datasets

In this section, we offer more detailed statistics concerning the test sets utilized in this study, encompassing 8 machine translation, 4 summarization,

4 style transfer, 2 simplification, 1 commonsense generation, and 1 paraphrase benchmarks. Table 9 provides a summary of the number of test sets, total words, and the average length. We will release the test sets and the corresponding demonstrations in the future. Note that the statistic is conducted based on tokenization sequences, which would be further segmented by BPE before feeding into LLMs. Consequently, the average length of summarization inputs would appear significantly larger, leading to an elevated risk in the context of few-shot requests.

C Design of Prompts

Figure 7 presents the DTG demonstration design across the other three text generation tasks. It can be observed that DTG does not necessitate task-specific designs; instead, a clear instruction outlining the main task for each work suffices. For the ease of replication of our results, we also furnish all baseline prompts, as depicted in Figure 8. Also, we provide the prompting design for GPT evaluation in Figure 9, which follows a zero-shot fashion.

To facilitate a more comprehensive understanding of the prompt ablations conducted in Section 6, we provide the corresponding design of prompts in Figure 10. Please note that prompts in blue represent the pre-designed demonstration, while those in red represent the test input. As observed, firstly, removing the error detection leads to the prompting in 10 (a). Additionally, the term “wrong error type” implies that we fed an *empty string* into LLMs, presenting it as a good translation. However, LLMs can autonomously detect the correct error type as an “incorrect translation” and subsequently generate an accurate response following careful deliberation (Figure 10 (b)). Conversely, if we constrain the error type detection process and solely allow LLMs to generate the translation, a considerable performance gap emerges (See Figure 10 (c)).

D More Analyses

Results on Machine Translation from English Table 11 summarizes the results of standard prompting and our DTG method in 5-shot scenarios, alongside results from WMT-Best and MS-Translator. When compared to results from to-English directional language pairs, such as DE-EN, the improvements provided by DTG over the standard prompting strategy appear somewhat marginal. Furthermore, DTG may yield results inferior to standard

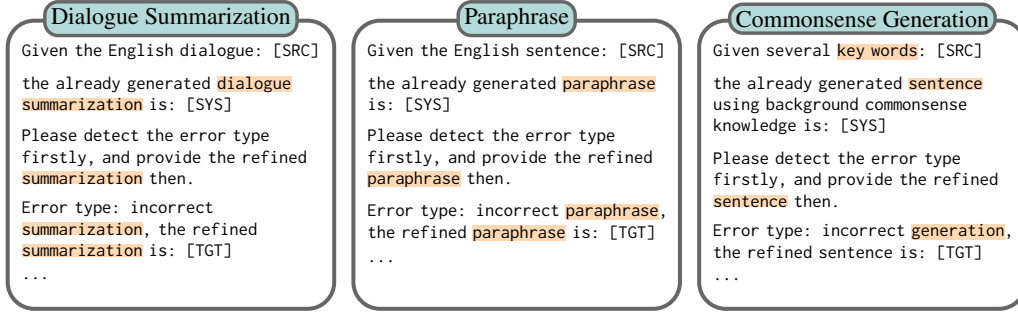


Figure 7: Illustration of DTG demonstration design for dialogue summarization, paraphrase and commonsense generation tasks within minimal modifications.

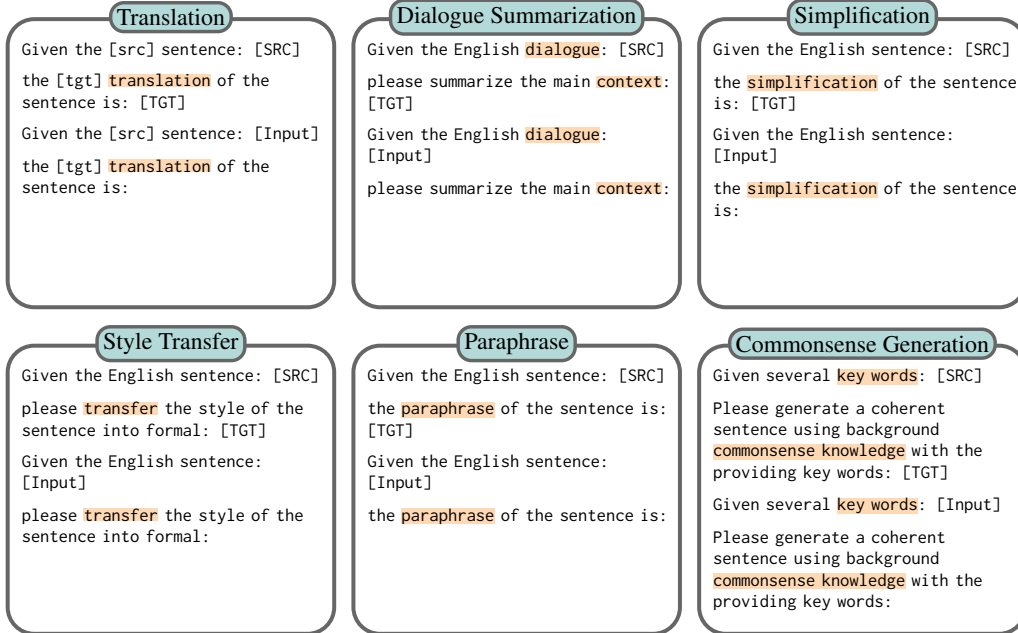


Figure 8: Illustration of the standard GPT prompting involving both demonstration and test input on six generation tasks, including machine translation, dialogue summarization, text simplification, style transfer, paraphrase and commonsense generation.

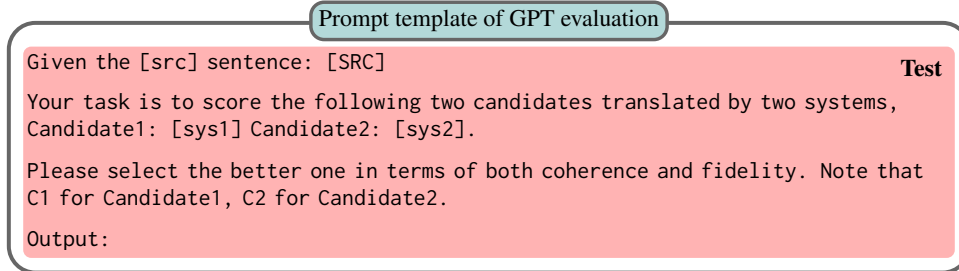


Figure 9: Illustration of the prompting design of GPT evaluation for Figure 5. We adhere to the recommendation proposed in (Liu et al., 2023)'s work, implementing a zero-shot GPT evaluation approach to identifying superior candidate translations through the adjudication of LLMs.

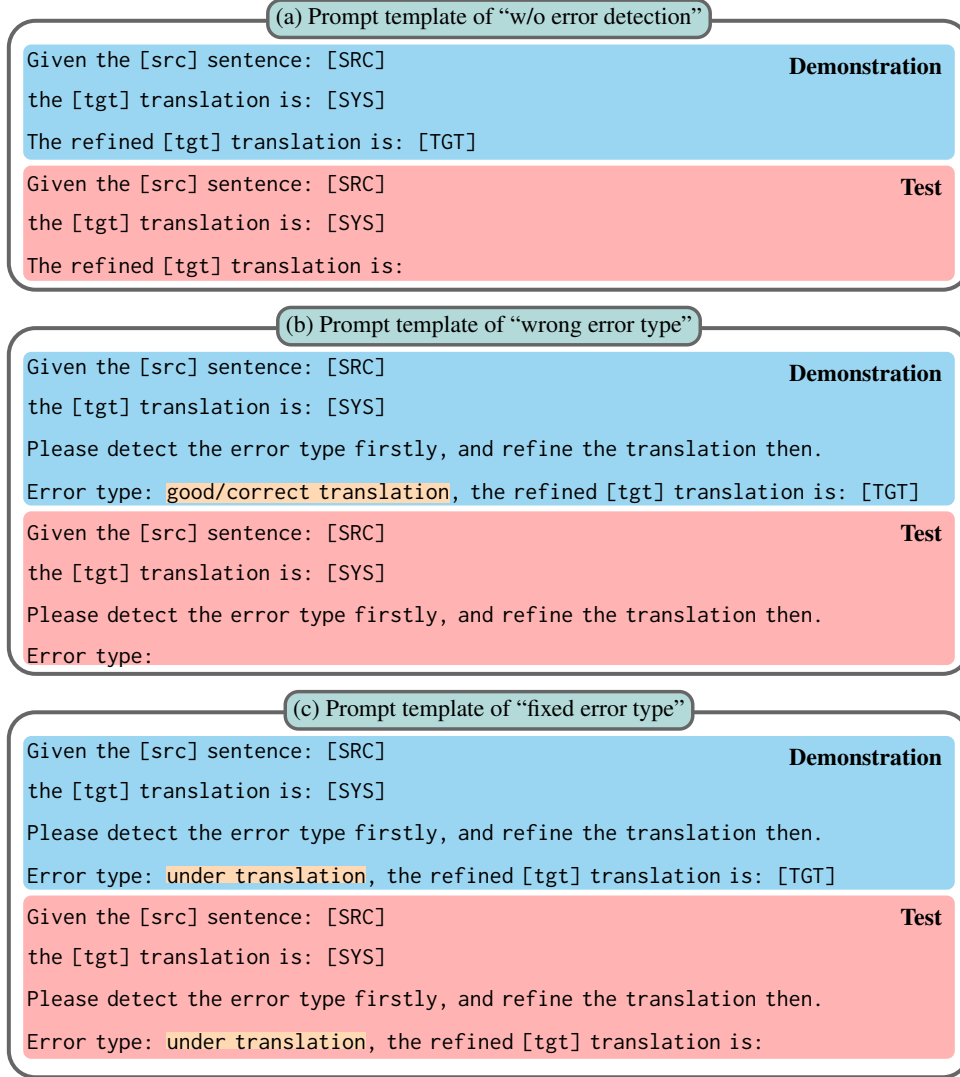


Figure 10: Illustration of the prompting design of the ablation study in Table 7. Note that all [SYS] here is *empty string*. The purpose here is to evaluate the deliberation ability of LLMs.

Dataset	Num.	Total Words	Ave. Words	Dataset	Num.	Total Words	Ave. Words
WMT DE-EN	1984	33540	16.9	CNN/DailyMail	11490	9017116	784.8
WMT CS-EN	1448	26050	17.9	GigaWord	1951	72171	37.0
WMT JA-EN	2008	36731	18.3	SamSum	819	104492	127.6
WMT ZH-EN	1875	14353	7.7	DialogSum	500	96385	192.7
WMT RU-EN	2016	32992	16.3	EM	1416	17279	12.2
WMT UK-EN	2018	29273	14.5	FR	1332	16799	12.6
WMT IS-EN	1000	19930	19.9	Amazon	500	6055	12.1
WMT HA-EN	997	30955	31.0	Yelp	500	5432	10.9
CommonGen	1497	6465	6.5	Asset	359	8115	22.6
QQP	2500	27543	11.0	Wiki-auto	2000	43860	21.9

Table 9: Statistics of the dataset we used on over 20 benchmarks. Note that “Num.” represents the number of test sets for each benchmark. “Total Words” and “Ave. Words” denote the total word count and average lengths, respectively. These statistics are based on tokenization sequences.

System	Score2	Score3	Score4	Score5
GPT 5-shot	16	88	196	200
DTG 5-shot	5	45	200	250

Table 10: Detailed score distribution of human evaluation on ZH-EN.

prompting in EN-ZH and EN-UK scenarios. This can likely be ascribed to the disparities in the balance of training sets across different languages.

Simple Chain-of-Thought cannot Help Text Generation We have witnessed the success of Chain-of-Thought when solving complex reasoning problems. It is a natural idea to simulate CoT process to improve the quality of text generation tasks. We have tried CoT-like prompting like this: “demonstration = “[requirement]=[Translate this English sentence into Chinese: Prior to this, Hefei has been the first to issue restrictions on lending policy. For people who have two suites in Hefei and have one housing loan not paid, they will be denied with the mortgage services from bank.] [chain of thought]=[“Prior to this” means “在此之前”, “Hefei” means “合肥”, “has been the first to issue” means “已经首先发布”, “restrictions on lending policy” means “限制贷款政策”, “people who have two suites in Hefei and have one housing loan not paid” means “拥有合肥两套房产且有一笔房屋贷款未付清的人”, “they will be denied with the mortgage services from bank” means “将被银行拒绝提供抵押贷款服务”, then the translation result after simple semantic splicing is “在此之前, 合肥是已经首先发布限制贷款政策的地方, 拥有合肥两套房产且有一笔房屋贷款未付清的人将被银行拒绝提供抵押贷款服务”. Finally, we optimize the translation result in an idiomatic way-“此前, 合肥已经率先发布限贷政策, 对于在合肥名下有两套房且有一套住房贷款未结清的

购房者, 银行将拒绝提供房贷服务”].

Despite extensive experimentation with various prompts, we observed no consistent advantages. This may be attributed to the inherent uncertainty in ensuring accurate word and phrase mapping. Consequently, we have shifted our focus towards employing negative-evidence prompting. This approach aims to activate the latent capabilities of LLMs that were embedded during the pretraining phase.

Details for Human Evaluation We have further conducted human evaluations to obtain more convincing results. Given the constraints of human effort, we have focused our evaluation solely on ZH-EN translation and Asset simplification. It’s important to note that, specifically for the ZH-EN translation, we have devised the following rules for human evaluators:

- 1 point - No translation or only isolated words translated.
- 2 points - 50% errors in translation; meaning distorted.
- 3 points - Mostly accurate; minor errors and inconsistencies.
- 4 points - Generally correct; some language and spacing issues.
- 5 points - Smooth, accurate, and fully conveys the original meaning.

Note that the 500 sentences were randomly selected from the test set. We also provide the detailed score distribution:

Table 11: Evaluation results of GPT on six high-resource and two-low resource machine translation tasks from WMT Testsets in from English directions. The best scores are marked in bold.

System	COMET-22↑	TER↓	ChrF↑	BLEU↑	COMET-22↑	TER↓	ChrF↑	BLEU↑
	EN-DE				EN-ZH			
WMT-Best†	87.2	49.9	64.6	38.4	86.7	102.3	41.1	44.8
MS-Translator†	86.8	50.5	64.2	37.3	86.1	94.2	43.1	48.1
GPT 5-shot	86.3	54.6	61.3	33.3	86.7	97.4	40.0	43.7
+ DTG	86.3	54.1	61.6	33.4	86.6	98.6	39.4	43.5
	EN-CS				EN-RU			
WMT-Best†	91.9	43.7	68.2	45.8	89.5	56.8	58.3	32.4
MS-Translator†	90.6	45.7	65.6	42.1	87.4	56.7	58.1	33.1
GPT 5-shot	88.9	54.6	58.9	32.7	87.0	61.3	54.4	28.2
+ DTG	88.8	54.5	59.0	32.9	85.7	63.0	52.1	28.1
	EN-JA				EN-UK			
WMT-Best†	89.3	105.9	36.8	27.6	88.8	57.5	59.3	32.5
MS-Translator†	88.0	106.0	34.9	25.1	86.1	63.2	56.1	28.2
GPT 5-shot	88.1	111.8	31.0	21.4	85.4	70.2	50.6	21.8
+ DTG	88.0	111.8	31.0	21.7	83.8	71.6	47.8	20.8
	EN-IS				EN-HA			
WMT-Best†	86.8	55.0	59.6	33.3	79.8	65.6	51.1	20.1
MS-Translator†	84.3	57.2	56.8	28.7	72.5	75.6	38.4	10.3
GPT 5-shot	76.1	70.8	44.1	16.2	72.8	87.4	38.5	9.9
+ DTG	76.7	70.9	44.2	16.3	73.2	77.7	39.3	10.1

Source	味道赞，肉类好，服务热情
Reference	Nice taste, great meat, enthusiastic service.
GPT 1-shot	The taste is great , the meat is good , and the service is enthusiastic .
+ Refine	The flavors are amazing , the meat is excellent , and the service is warm and welcoming .
+ DTG	Great taste, good meat, enthusiastic service.
Source	目前已经购买了这个系列3款机器！
Reference	I have bought three laptops of this series!
GPT 1-shot	So far , 3 machines from this series have been purchased!
+ Refine	Up until now , 3 machines from this series have been purchased!
+ DTG	I have already purchased 3 models from this series!

Table 12: Case study on refining from the previous candidate (Refine) and the proposed DTG method.

Case Study We provide a case study based on GPT4 model in Table 12, where “Refine” indicates utilizing the 5-shot baseline results as the synthesized sentences, i.e., “[INCORRECT SYS]” in Figure 1, and DTG is our method that uses an empty string instead. The conclusions are two-fold. 1) Using the baseline results will cause the model to avoid generating the same segmentations in it although they may be correct already, e.g., “taste” to “flavors”, “so far” to “up until now”, as well as others in red. As a result, the fluency and accuracy of the final results may be affected. 2) Equipped with DTG, fluency, coherence and grammatical correctness of generated results are all promoted. In the first case, the DTG result is more faithful not only in semantics but also in structure than the baseline.

In the second case, DTG is able to complete the subject “I” which does not appear in the source sentence.

E Details of Error Statistical

In Figure 6, two types of error are considered (i.e., under translation and entity translation error). In this section, we provide the details of the method to conduct the error statistics.

Under Translation We first use *awesome-align*¹¹ to get the alignment between the source and target sentences. Then, a word in the source sentence is regarded as under translation, when it is aligned to a word in the reference target sentence

¹¹<https://github.com/neulab/awesome-align>

1087 but failed to be aligned in the generated target sen-
1088 tence.

1089 **Entity Translation** We first use *spaCy*¹² to rec-
1090 ognize the named entities in the reference target
1091 sentence, where person names, organizations and
1092 locations are considered. Then, an entity in the ref-
1093 erence is considered an error if it cannot be found
1094 in the generated target sentence.

¹²<https://github.com/explosion/spaCy>