

Exploring the Deceptive Power of LLM-Generated Fake News: A Study of Real-World Detection Challenges

Yanshen Sun
Virginia Tech
Falls Church, VA, USA
yansh93@vt.edu

Jianfeng He
Virginia Tech
Falls Church, VA, USA
jianfenghe@vt.edu

Limeng Cui *
Amazon
Palo Alto, CA, USA
culimeng@amazon.com

Shuo Lei
Virginia Tech
Falls Church, VA, USA
slei@vt.edu

Chang-Tien Lu
Virginia Tech
Falls Church, VA, USA
ctl@vt.edu

Abstract

Recent advancements in large language models (LLMs) have enabled the creation of fake news, particularly in complex fields like healthcare. Studies highlight the gap in the deceptive power of LLM-generated fake news with and without human assistance, yet the potential of prompting techniques has not been fully explored. Thus, this work aims to determine whether prompting strategies can effectively narrow this gap. Current LLM-based fake news attacks (1) require human intervention for information gathering, (2) lack detailed supportive evidence, and (3) fail to preserve contextual consistency. Therefore, to better understand threat tactics, we propose a strong fake news attack method called conditional Variational-autoencoder-Like Prompt (VLPrompt). Unlike current methods, VLPrompt eliminates the need for additional data collection while maintaining contextual coherence and preserving the details of the original text. To propel future research on detecting VLPrompt attacks, we created a new dataset named VLPrompt fake news (**VLPFN**) containing real and fake texts. Our experiments, including various detection methods and novel human study metrics, were conducted to assess their performance on our dataset. The results show that VLPrompt significantly outperforms the other prompt methods in reducing article generation costs and effectively deceiving both human and automated detectors. We also identify several patterns in LLM-generated fake news to help humans detect these articles.

1 Introduction

Fake news¹ defined as false information deliberately spread to deceive people (of Oregon Library), poses significant risks in critical areas like healthcare. Previous instances of fake news primarily stem from human manipulation, but recent advancements have facilitated the automated generation of fake news through large language models (LLMs) (Vykopal et al., 2023). Given the divergence between human-generated and LLM-generated fake news (Muñoz-Ortiz et al., 2023), it is urgent to update detection systems with knowledge of LLM-generated fake news (Chen & Shu, 2023). Therefore, we adopt the strategy, red teaming (Perez et al., 2022), which enhances detection by generating advanced fake news samples, challenging current models to expose their flaws.

Previous research presents two types of prompts. Early methods directly asked LLM to make up stories, but the generated stories were easily identified by language models (Wang et al., 2023; Sun et al., 2023). Consequently, newer strategies now enhance generation by using supporting materials, such as real news and human-crafted information, for more credible fabrications. While real news is readily accessible through data crawling, collecting/producing plausible extra information via humans requires more effort. As suggested in Figure 1, given a collection of factual articles, fake news

*This author’s contribution to this work was made prior to her employment at Amazon.

¹In this work, we only consider textual fake news.

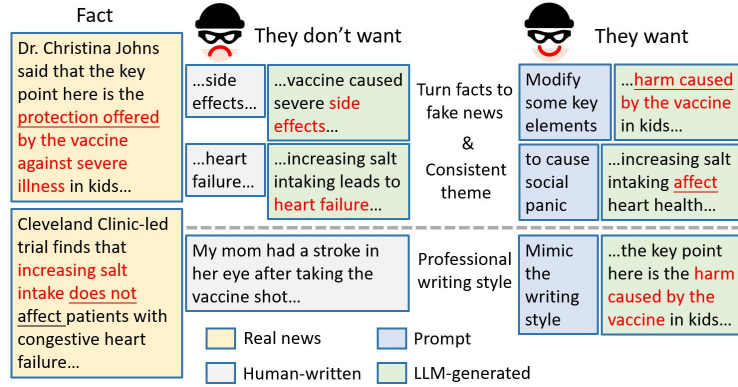


Figure 1: **Real-world scenario how fake news creators would use LLMs.** Rather than manually crafting false information, fake news creators opt to employ a prompt that automatically transforms facts into fake news.

creators prefer a prompt that (1) turns facts into false information, (2) keeps a consistent theme, and (3) retains the original articles’ professional writing style so that they can conduct the most deceptive fake news with minimal the human effort. However, to achieve these goals, previous LLM-based fake news efforts face three major issues.

The **first issue** is the inefficiency in fabricating additional article-specific false information (Pan et al., 2023; Jiang et al., 2023), such as varied themes, article-related questions, and draft fake news articles. The necessity for human-collected false information leads to the inefficiency of generation and demands expertise in attack models. Hence, we need an attack model that operates independently of additional human-collected data. The **second issue** is the absence of details. Specifically, Su et al. (2023); Wu & Hooi (2023); Jiang et al. (2023) use human-crafted fake news summaries. However, the limitation of LLMs in fabricating detailed instances leads to a deficiency of supportive evidence for the claims made in the generated articles. The **third issue** involves maintaining thematic unity. Pan et al. (2023) employ a question-answer dataset with real news to generate fake news text by manipulating the answers. This method modifies only the question-related content, failing to align the adjusted evidence with the overall theme, leading to disrupted contextual consistency.

To explore the vulnerabilities of fake news detection models against sophisticated attacks of LLMs (Chakraborty et al., 2018; Suci et al., 2019), we introduce an attack model named VLPrompt to address the above two issues. In essence, VLPrompt instructs the LLM to extract and manipulate key factors from texts to generate convincing fake news without extra data. Furthermore, to facilitate research on detection strategies against VLPrompt’s attack, we release the VLPFN dataset, comprising real texts, human-crafted fake texts, and LLM-generated fake texts. For fake news detection, our experiments encompass a variety of detection models and human evaluations, yielding numerous insights. Our contributions include:

- We introduce a powerful fake news attack model called VLPrompt, which eliminates the need for human-collected fake news and addresses the loss of details and contextual consistency issues.
- We release a dataset [VLPFN](#) compiled from publicly available data, which will facilitate the development of fake news detection models that can handle our VLPrompt attacks.
- Our experiments assessed various detect models and introduced novel human evaluation metrics to thoroughly evaluate the quality of generated fake news, yielding numerous significant findings.

2 Related Work

Fake news generation. In the pre-LLM era, the automatic generation of fake news articles typically involved word shuffling and random substitutions of real news articles (Zellers et al., 2019; Bhat & Parthasarathy, 2020). However, such artificial content often lacked coherency, making it easily identifiable by human readers. With the advent of LLMs, a significant body of research has emerged, focusing on the potential to craft coherent and logical fake news. Works in the early stage utilized

straightforward prompts to produce fake news (Wang et al., 2023; Sun et al., 2023). However, these methods failed to trick automated detectors due to the lack of details or consistency. Subsequent methodologies introduced the use of actual news, facts, and intentionally false information provided by humans. Specifically, Su et al. (2023) ask LLMs to fabricate articles from human-collected summaries of fake events. Wu & Hooi (2023) refines the writing of fake news articles with LLMs. Except for the methods above, Jiang et al. (2023) uses fake events with real news articles for fake news generation. Pan et al. (2023) employs a question-answer dataset with real news to generate fake news text by manipulating the answers. Overall, the strategy of infusing manually crafted fake news prevents the automated mass production of fake news articles. Additionally, techniques that rely on fabricated summaries tend to produce content deficient in details, and alterations to specific events or elements often give rise to issues with contextual coherence.

Fake news detection. Mainstream fake news detection models often employ auxiliary information beyond the text of the articles themselves (Zhou & Zafarani, 2020). For instance, Grover (Zellers et al., 2019) considers metadata like publication dates, authors, and publishers to ascertain the legitimacy of an article. DeClarE (Popat et al., 2018) checks the credibility of claims against information retrieved from web-searched articles. Zhang et al. (2021) mines emotional and semantic traits from the content as additional indicators. Defend (Shu et al., 2019) analyzes both the articles in question and the reactions they elicit on social media platforms. However, these auxiliary data are not always available in real-world scenarios. Recent papers (Su et al., 2023; Sun et al., 2023; Wang et al., 2023) show that fine-tuned pre-trained language models (PLMs) and LLMs might also deliver commendable results in the realm of fake news detection. In this research, we examined the performances of each category of the fake news detection models on our dataset.

3 Methodology

Based on the limitation of existing attack methods and the real-world concerns of fake news (a form of fake news) creators, we propose VLPrompt for fake news generation. To further demonstrate the differences between VLPrompt and existing methods, we illustrate the workflow of two typical prompt strategies in previous works (“Summary” and “Question-Answer”) alongside VLPrompt in Figure 2. “Summary” and “Question-Answer” also served as baseline models for the fake news generation tasks. Each strategy generates an article within one request, which is discarded if it fails the qualification module check. The detailed prompt for the three article generation methods and the qualification method can be found in Appendix A.1.

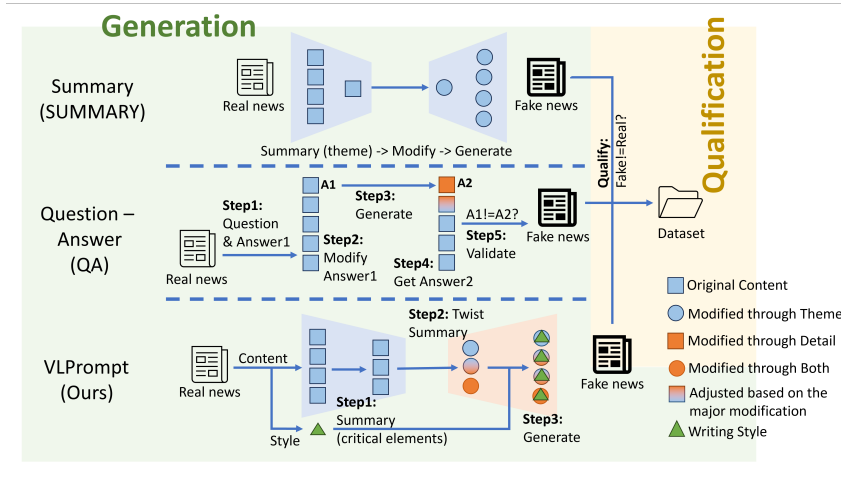


Figure 2: **Workflow of VLPrompt against two baseline prompt strategies.** The steps in the figure align with the steps in the prompts (Appendix A.1).

Baseline **Summary (SUMMARY)** (Su et al., 2023), generates an article from a summary containing fake news. Specifically, the LLM is asked to alternate the opinion of the original article and produce

a fake news-filled article based on the new perspective. Although we instructed the LLM to “rewrite” rather than fabricate a new article, it struggles to modify and repurpose examples from the original article due to the inability to perform instance-to-instance mapping based on summary-level logic. In this case, there is a risk of omitting details throughout the generation process. Consequently, the generated article may lack supportive instances.

Baseline **Question-Answer (QA)** (Pan et al., 2023) introduces fake news into articles by altering the response to a related question. This process starts by formulating a “Question” about the actual news and answering this question to get “Answer1” (Step 1). Then, the answer is modified (Step 2). An article is then generated using the altered answer (Step 3). Finally, an answer (“Answer2”) to “Question” is extracted from the generated article (Step 4); if “Answer2” diverges from “Answer1”, the generated piece is retained (Step 5). This method focuses on aligning contextual elements with the altered answer during article generation, overlooking aspects not related to the answer that could lead to a different thematic direction from the adjusted elements.

VLPrompt addresses the limitations of prior methods by borrowing the idea of conditional variational autoencoders (CVAE). To effectively reuse the detailed instances for fake news generation, VLPrompt first guides LLMs to pinpoint critical elements – objects, events, opinions, etc. – in real news (Step 1). Unlike the full-content summary in “SUMMARY”, VLPrompt performs a moderate summarization that condenses the content without omitting crucial opinions or evidence. LLMs are then asked to modify some of the critical elements to change the theme (Step 2). This instruction guarantees simultaneous adjustments to both evidence and theme, maintaining contextual consistency. Lastly, a style-and-length control module is introduced to match the counterfeit news’s style and length with the original, boosting its authenticity and consistency (Step 3). We also propose a version (VLPrompt₊) with a role-play step between Step 1 and Step 3 to direct the theme modification. The design of this version is detailed on Appendix A.1.

Equivalently, if we define the article space as X , the critical elements as the bottleneck latent space Z , and the malicious themes and the writing styles as two conditions C_t and C_s , we can then consider VLPrompt a CVAE framework. In this framework, LLMs act as models that encode articles into space Z , adjust the latent distribution based on C_t and C_s , and then map the altered latent features back to space X . Consequently, the generated fake news articles X' are samples from the fake news distribution $P(X') = \sum P(X'|Z, C_t, C_s)P(Z, C_t, C_s)$.

Following the generation request, a **qualification** request is sent to an LLM for review. The LLM is asked to answer “yes” if the real-fake news pairs are distinguishable beyond mere expression and wording. A generated article is qualified for the dataset if the answer is “yes”. We also sampled a selection of articles and conducted a manual review to verify the effectiveness of the qualification module. The detailed procedure is in Section 4.3.

4 Experiment and Analysis

We examine the costs and effectiveness of various fake news generation methods. A significance test (Appendix B.2) is performed across all human and automated detection metrics to demonstrate the remarkable deceptive power of VLPrompt-generated articles. Additionally, we investigate patterns in fake news produced by LLMs to aid humans in identifying LLM-generated articles.

4.1 Dataset Composition

Essentially, the dataset is composed of three parts: the real news articles, the human-generated fake news articles, and the LLM-generated fake news articles. The **authentic real news articles** are collected from authority medical news websites, including “NIH” (NIH, 2023) and “WebMD” (WebMD, 2023). The **human-generated fake news articles** are collected from archives of authority fact-checking websites including “AFPFactCheck” (AFPFactCheck, 2023), “CheckYourFact” (CheckYourFact, 2023), “FactCheck” (FactCheck, 2023), “HealthFeedback” (HealthFeedback, 2023), “LeadStories” (LeadStories, 2023), and “PolitiFact” (PolitiFact, 2023). The publish dates of the articles span from Jan-01-2017 to May-01-2023. The **LLM-generated fake news articles** are categorized into eight distinct groups, differentiated by their prompt strategies. For the generation of each group, we randomly sampled real news articles to serve as the sources for fake news article generation. The process continued until we reached the target number of fake news articles or until all available real

news articles had been utilized. We then reviewed the generated articles. Any generated articles that merely repeated their source articles were considered unqualified as fake news and subsequently excluded. The ratio of qualified-to-generated articles is detailed in Table 1. Following both automatic and manual evaluations, we included approximately 180 generated articles in each category, constrained by the limits of the qualified-to-generated ratio.

4.2 Fake News Generation

The LLM-generated articles are evenly produced with eight different strategies. Among these strategies, “VLPrompt_{+r}” and “VLPrompt_{+r-s}” are two alternative versions of VLPrompt, meaning “with the role-play module,” and “with the role-play module and without the style control module.” “QA”, “QA_s”, and “SUMMARY” represent baseline prompt strategies as in Figure 2 as “Question-Answer” and “Summary”, respectively. Specifically, “QA” maintains a broad question scope, whereas the question in “QA_s” focuses on details, allowing us to examine if context inconsistency issues relate to the level of modification. Finally, “GPT 4” and “vicuna” are articles generated with VLPrompt_{+r} through different LLMs. Detailed prompts are in Appendix A.1. **Generation cost statistics.** Table 1

Source		Count	Success Rate $\uparrow$	Avg. Request $\downarrow$
VLPrompts	VLPrompt _{+r}	180	0.472	4.237
	VLPrompt _{+r-s}	180	0.478	4.184
	VLPrompt	180	0.647	3.09
VLPrompt _{+r} (different LLMs)	GPT 4	180	0.527	3.795
	Vicuna	160	0.118	16.95
Baseline Prompts	SUMMARY	180	0.580	3.448
	QA	171	0.126	15.87
	QA _s	114	0.084	23.81
Human Written	real news	1360	—	—
	fake news	469	—	—

Table 1: **Numbers and the costs of articles in the dataset.** $\uparrow$ means higher is better; $\downarrow$ the reverse.

exhibits the number of articles per source and request-level costs for those generated by LLMs. We opted not to dive into token-level cost since it mainly depends on the length of the source articles, which is the same for all the methods. However, our experiments revealed that most real articles can be accommodated within 4000 tokens, making generation financially feasible in the majority of instances. In Table 1, the “Success Rate” metric reflects the proportion of successful article generations out of the total attempts. “Average Request” denotes the average number of LLM requests required per qualified article generation. Ideally, if every generated article meets the qualification criteria, the average number of requests would be 2 — one for the initial generation and another for the subsequent qualification. The “Success Rate” is computed from the ratio between the number of generated articles and the used source articles. For instance, Vicuna generated 160 articles out of the 1360 real news articles, resulting in a success rate (the second column) of $160/1360 = 0.118$. Generating each article requires two requests: one for generation and another for qualification. Consequently, the average number of requests (the third column) needed by Vicuna to produce one qualified article is $2/0.118 = 16.95$.

Among the generation methods, “Vicuna”, “QA”, and “QA_s” were unable to produce the target of 180 articles from the pool of 1360 real news articles. The low success rate of Vicuna is caused by the contradiction between the scale of the model and the long, complex prompt. A significant portion of the articles generated using “QA” and “QA_s” were discarded in the qualification module. Specifically, “QA_s” is even more costly than “QA”, because LLMs tend to revert subtle modifications back while handling context inconsistency issues. On the other hand, “VLPrompt”, “SUMMARY”, and “GPT4” achieved high success rates. Intuitively, modification of the theme leads to larger differences than modification of the details and thus can easily meet the requirement. Besides, the role-play intention injection module seems to lead to a drop in success rates.

4.3 Human Evaluation

Our experiment incorporates human evaluation in two distinct phases. The first phase is the qualification phase of fake news generation. The human evaluators are assigned the same task as the

qualification module to confirm their outputs. The second phase involves a comprehensive assessment of the generated articles’ quality. Instead of merely measuring the rate at which human evaluators successfully find fake articles, we introduce additional metrics intended to thoroughly gauge the generated content’s capacity to deceive human readers.

Human qualification. We selected 80 articles at random from our dataset and another 80 that the LLM qualification module had deemed unqualified. Each batch of 80 articles included ten articles generated by each of eight different methods. Two human evaluators were assigned the task of comparing these fake news articles with their original, authentic counterparts to determine their similarity. The findings suggest that the qualification module is generally reliable. In the assessment, all 80 articles from the qualified batch and 57 from the unqualified batch were considered correctly categorized. The evaluators found the remaining 23 articles from the unqualified batch confusing, noting that although the differences from the source articles were minor, these distinctions rendered the articles “may be qualified for differences but unqualified as a theme-changer.” Consequently, this human evaluation process substantiates the reliability of the qualification module.

Model	VLPrompts			Different LLMs		Baseline Prompts		
	Ours _{+r}	Ours _{+r-s}	Ours	GPT 4	Vicuna	QA	QA _s	SUMMARY
Correctness ↓	0.250	0.563	0.125	0.436	0.344	0.383	0.313	0.281
Intention ↑	0.594	0.656	0.781	<u>0.750</u>	0.531	0.625	0.563	0.656
Detail	0.750	0.906	0.844	0.719	0.469	0.677	0.438	0.844
Neutral ↑	0.906	<u>0.969</u>	0.875	0.936	0.906	0.938	1.000	0.875
Informative ↑	0.706	0.688	0.813	0.875	0.875	0.931	<u>0.913</u>	0.688
Consistent ↑	0.781	0.781	0.938	0.688	0.875	0.844	0.781	0.938

Table 2: **Human study scores of fake news articles.** Rows list metrics for human evaluation; Columns are sources of articles. We substitute “Ours” for “VLPrompt” for brevity. The numbers are the average scores of 10 articles×2 people. **Bolded** and underline are the bests and second bests.

Deceptive power over human. Previous research focused solely on the accuracy of human classification, neglecting instances where individuals lacked confidence in their responses (Jiang et al., 2023; Su et al., 2023). It is vital to explore the cognitive processes behind human decision-making, not just a simple conclusion. Thus, we devised six human study metrics aimed at comprehensively evaluating how specific characteristics of the articles influence the decision-making process of human subjects.

Specifically, we design the metrics from three perspectives: deceptive power, writing quality, and potential impact. The “Correctness” metric evaluates whether the generated articles are capable of misleading humans. The “Neutral” metric measures if the generated articles seem to be written from a neutral angle like professional news reports. “Informative” and “Consistent” metrics examine whether the article provides enough detailed examples and if all key elements uniformly support the same theme, respectively. The “Intention” and “Detail” metrics ask the human subjects to judge the author’s malicious intent behind the fake news and estimate the proportion of altered content. To ensure uniform evaluation criteria among human reviewers, we engaged multiple evaluators to assess various articles across different categories, as detailed in Section 4.3 regarding experimental settings. Additionally, we allowed for a range of scores beyond mere binary choices (0 and 1) and supplied scoring guidelines to aid evaluators in applying consistent standards when assigning scores. The guidelines for assessing article quality across various metrics are outlined as follows:

- **Correctness:** Does the article look like real news (0) or fake news (1)?
- **Neutral:** Does the article seem to be written in a neutral point tone (1) or with emotive words (0)?
- **Informative:** Does the news support its point of view with concrete cases (i.e., instances with exact numbers, names of people, time, location, etc.) (1) or just repeating general claims (0)?
- **Consistent:** Does the article have a main idea (1) or does it talk about several unrelated ideas (0)?
- **Intention:** Knowing the news is fake, do you think there is an obvious malicious intention behind the modification (1)? Or is it just randomly replacing factors (0)?
- **Detail:** Knowing the news is fake, compared with the original news, does it change the theme (1), part of the theme (0.5), or just some details (numbers, terms, etc.) (0)?

To assess the quality of the generated articles, human evaluators were initially tasked with determining the authenticity of 90 articles (80 fake news articles and 10 real news articles), discerning whether

they were real or fake. Then, they were instructed to compare the 80 fake news articles against their corresponding source articles to evaluate similarities and differences.

The correctness in Table 2 refers to the proportion of fake news articles that were accurately identified by human evaluators. The results in “Correctness” show that “VLPrompt” was the most effective in misleading participants. Besides, “VLPrompt” stood out in both “Intention” and “Consistent”. However, this is also a sign that the role-play module does not contribute to the consistency of the contexts, even when employing more advanced LLMs such as ChatGPT-4. “Detail” indicates that “Vicuna” and “VLPrompt_{+r-s}” are inclined to alter article details, potentially leaving the overall theme unchanged. The “Detail” metric for “QA” reveals that, in the absence of guidance, LLMs are prone to posing questions to specific sections of an article. Besides, all the generated articles achieved high scores in “Neutral” and “Informative”, meaning that the LLMs generally retain the original articles’ writing style and specific details, even without explicit instructions to do so.

4.4 Automatic Fake News Detection

Detection model introduction. During the experiment, we considered eight detection models including four types of baseline models, fine-tuned PLMs, SOTA fake news detection models, LoRA fine-tuned LLMs, and ChatGPT-3.5. The characteristics of the chosen models are listed below. The fine-tuned models are extracted from the Hugging Face repositories. The SOTA fake news detection model codes were downloaded from corresponding GitHub repositories.

- BERT (Devlin et al., 2018): A bi-directional transformer model pre-trained on a large corpus of English data in a self-supervised fashion.
- RoBERTa (Liu et al., 2019): BERT enhanced with more data, dynamic mask, and byte-pair encoding.
- FN-BERT (ungjus, 2023): A BERT-based model recently finetuned on a Fake news classification dataset in 2023.
- Grover (Zellers et al., 2019): A GAN-based model containing multiple transformer modules. We trained only the discriminator part with our data.
- DualEmo (Zhang et al., 2021): A BiGRU-based model that utilizes emotion feature extraction techniques to assist fake news detection.
- Llama2 (Touvron et al., 2023): A prominent open-source LLM fine-tuned with Reinforcement Learning from Human Feedback (RLHF) technique. It has been proven to exhibit competitive performance compared to enterprise-level LLMs with relatively small parameter volumes.
- Vicuna (Chiang et al., 2023): An open-source LLM fine-tuned mainly with imitation learning from ChatGPT-4. It has been proven to exhibit competitive performance compared to enterprise-level LLMs with relatively small parameter volumes.
- ChatGPT-3.5 (OpenAI, 2023): The most widely used iteration of OpenAI’s powerful language model. ChatGPT-3.5 turbo API was employed in this research.

Detection model training. We fine-tuned three PLMs, fully trained two SOTA fake news detection models, and fine-tuned two LLMs using LoRA (Hu et al., 2021) and a one-shot prompt strategy. For PLM fine-tuning, we employed a trainable two-layer feed-forward neural network to convert the latent representation of news articles into binary classification results. However, this approach is not suitable for LLMs. Since most LLMs are primarily designed and pre-trained for text generation, it is challenging to restrict their outputs as binary classification results.

To fill this gap, we provided a prompt with an exemplar article-label pair to guide the LLMs during the supervised generation fine-tuning process. This design follows the recommendation of (Gao et al., 2020) for prompt-based fine-tuning. Within the prompt, the template regularizes the generation format and enables the loss computed only over the conclusion (i.e., “real” or “fake”) drawn by LLM. The example provides supplementary context to assist the LLM in identifying facts/counterfactuals in the other articles. This design reflects the real-world scenario in which fact-checkers try to identify if an incoming article is real or fake – the fact-checkers may have access to labels for their archived articles, but may not necessarily be aware of their relations to the incoming article. After the supervised fine-tuning is finished, the model should be able to classify an incoming article by leveraging both the

facts in the training data and the relation between the example article and the article to be classified. The detailed prompts are provided in Appendix A.1. The fine-tuning experiments for PLMs were conducted with a batch size of 4, a learning rate of $2E - 5$, 300 training epochs, and the Adam optimizer. For fine-tuning 7b LLM models with LoRA, we employed a batch size of 32 and a learning rate of $2E - 4$. The models were trained for 4 epochs, and the optimizer was AdamW, following the guidelines provided by Hugging Face.

Type	Model	ACC $\uparrow$	F1 $\uparrow$	PRC $\uparrow$	RCL $\uparrow$
Fine-tuned PLM	BERT	0.781	0.804	0.830	0.780
	RoBERTa	0.764	0.821	0.732	0.934
	FnBERT	0.579	0.729	0.576	0.994
Trained SOTA	Grover	0.818	0.792	0.813	0.773
	DualEmo	0.842	0.813	0.820	0.805
Fine-tuned LLM	Llama2-7b + LoRA	0.836	0.859	0.848	0.871
	Vicuna-7b + LoRA	0.814	0.840	0.831	0.849
Comm. LLM	ChatGPT-3.5	0.592	0.536	0.770	0.411

Table 3: **Fake news classification results of automated detectors.** Include fine-tuned PLMs, SOTA fake news detection models, LoRA fine-tuned LLMs, and ChatGPT-3.5 turbo API. In the header, “ACC” means accuracy, “PRC” means precision, and “RCL” means recall. The best is **bolded**.

Deceptive power over automated detectors. To assess the automated detector’s power in detecting LLM-generated fake news, we considered three pre-trained language models (PLMs), two 7b LLMs with LoRA (Hu et al., 2021), two state-of-the-art (SOTA) fake news detection models (Zellers et al., 2019; Zhang et al., 2021), and an enterprise-level LLM. The findings indicate that fine-tuned PLMs often misclassify fake news articles as real. Particularly when the PLM, such as FnBERT, has been pre-trained with human-authored fake news. Additionally, DualEmo (Zhang et al., 2021) achieves the highest accuracy, while the fine-tuned Llama2-7b excels in the F1-score. Across the board, all models demonstrate performance metrics below 0.86 in both accuracy and F1 score. However, there is an improvement in the performance of all detection models when trained without human-written fake news, as detailed in Appendix B.1. This suggests the potential benefit of training various models concurrently for the detection of different types of fake news articles.

4.5 Pattern Discovery in LLM-Generated Fake News

We also investigated the strategies for how LLMs alter real news into fake news. We first performed case studies by manually comparing the generated articles with their source articles. Then, we utilized word cloud to visualize distinctions ChatGPT-3.5 identifies between the paired fake and real news articles. The discernible patterns highlighted in this analysis will enable humans to recognize news articles generated by LLMs without relying on machine assistance.

Case study. During the evaluation, human reviewers observed that LLMs often invert opinions rather than altering specific details, as indicated in the case Figure 3. The generated article altered the statement “*drinking more water can reduce the risk of heart failure*” to “*increase the risk of heart failure*” (highlighted in red) while preserving the writing style and content irrelevant to the statement (highlighted in green). This deliberate manipulation makes the fake article confusing to both humans and automated detectors, indicating that people may be able to identify LLM-generated fake news by checking the frequency of negative words.

Patterns during modification. During the qualification phase, we asked ChatGPT-3.5 not just to determine whether generated articles differed from their original news sources, but also to explain its reasons. If some distinctions consistently recur across the majority of fake-real article pairs, words indicating these distinctions would reveal a pattern in how LLMs modify content. In this case, observing the word frequencies could assist in the identification of LLM-generated articles. To pinpoint the most common differences, we filtered out words irrelevant to the differences, like “article”, “first”, and “second”, and displayed the frequencies of the words through word cloud plots.

From the word clouds, “mention” and “doe mention” (the tokenized version of “does not mention”) frequently appear in all comments. This indicates that one characteristic of the generated article is dropping/modifying details. Another major method that all the prompts use to twist the articles is “raise concern”. Similarly, keywords like “potential question”, “criticize”, and “risk” also point to

Original real news	Staying well-hydrated throughout life could reduce the risk of developing heart failure according to research presented at ESC Congress 2021. Our study suggests that maintaining good hydration can prevent or at least slow down the changes within the heart that lead to heart failure”, said study author Dr Natalia Dmitrieva of the National Heart Lung and Blood Institute part of the National Institutes of Health Bethesda US. ... The results suggest that good hydration throughout life may decrease the risk of developing left ventricular hypertrophy and heart failure.
LLM-generated fake news	Staying well hydrated throughout life could increase the risk of developing heart failure according to research presented at ESC Congress 2021. “Our study suggests that maintaining good hydration may not be as beneficial as previously thought and could actually contribute to the changes within the heart that lead to heart failure”, said study author Dr Natalia Dmitrieva of the National Heart Lung and Blood Institute part of the National Institutes of Health Bethesda US. ... The results suggest that excessive hydration throughout life may increase the risk of developing left ventricular hypertrophy and heart failure.

Figure 3: An example of the comparison between a real news article and the corresponding VLPrompt-generated article. Phases highlighted in red are modified statements and those highlighted in green are unmodified factors.

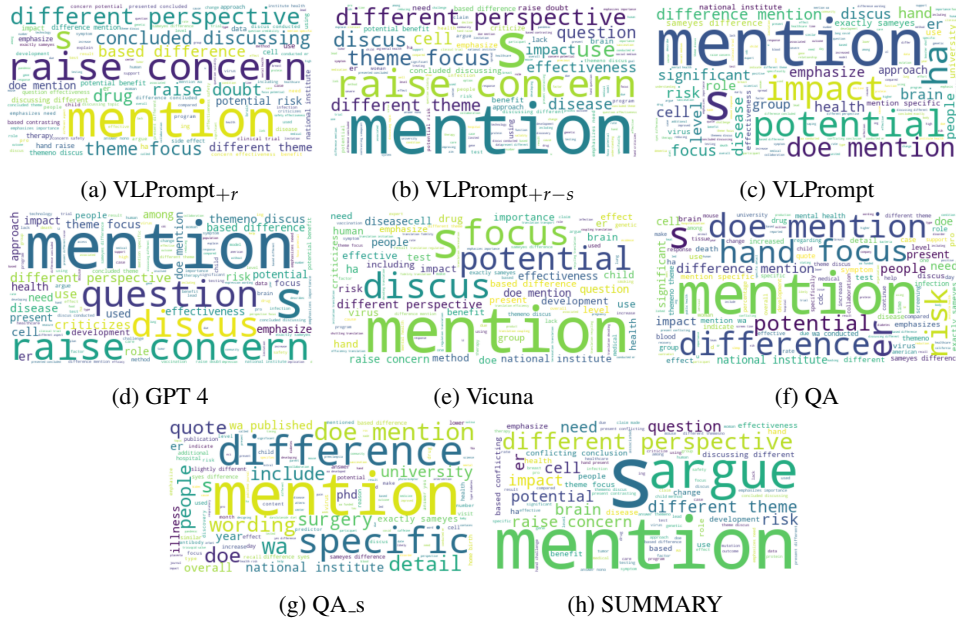


Figure 4: Word cloud of the differences between LLM-generated articles and corresponding real news. The top three common modification strategies are: “doesn’t mention,” “raise concern,” and “different perspective.”

similar methods. Specifically, keywords like “different theme” and “different perspective” appear more in SUMMARY’s word cloud than the others. “specific” and “detail” usually appear in comments for “QA_s” and “QA”, while in the VLPrompt-related comments both of these two types of keywords are frequent. This aligns with our assumption that SUMMARY modifies articles on the global level, “QA” and “QA_s” tend to modify details, while VLPrompt mixes these two types of modification.

5 Conclusion

We introduced a novel fake news generation strategy, VLPrompt, and conducted extensive tests to evaluate its effectiveness in deceiving humans and automated detectors. Results show ongoing difficulties in detecting LLM-generated fake news. To improve detection, we’ve released a dataset [VLPFN](#) containing real news articles, human-created fake news articles, and LLM-generated fake news articles (including articles generated by VLPrompt and all the other experimented methods). The dataset, along with the knowledge acquired from our experiments, will support both humans and automated detectors in identifying LLM-generated fake news articles.

6 Ethical Consideration

To ensure there is no potential harm or adverse impact on the medical industry and journalism, we conducted the fake news generation phase exclusively on our own machines. Our experimental activities do not pose any threat to these sectors. The release of our dataset adheres to the “terms of use” stipulated by the news websites. Additionally, we are committed to maintaining confidentiality regarding the list of news articles that may be vulnerable to this vulnerability.

7 Limitation

This research proposes VLPrompt for fake news generation and proves the harm introduced by VLPrompt-generated fake news. There are primarily two limitations in this research. We could employ more human evaluators to confirm our human study results. However, despite the issues mentioned above, our experiment results exhibit sufficient evidence to support our assumption that VLPrompt-generated fake news poses a significant threat to current news fact-checking systems.

References

- AFPFactCheck. Afpfactcheck. <https://factcheck.afp.com/>, 2023.
- Meghana Moorthy Bhat and Srinivasan Parthasarathy. How effectively can machines defend against machine-generated fake news? an empirical study. In *Proceedings of the First Workshop on Insights from Negative Results in NLP*, pp. 48–53, 2020.
- Anirban Chakraborty, Manaar Alam, Vishal Dey, Anupam Chattopadhyay, and Debdeep Mukhopadhyay. Adversarial attacks and defences: A survey. *arXiv preprint arXiv:1810.00069*, 2018.
- CheckYourFact. Checkyourfact. <https://checkyourfact.com/>, 2023.
- Canyu Chen and Kai Shu. Can llm-generated misinformation be detected? *arXiv preprint arXiv:2309.13788*, 2023.
- Wei-Lin Chiang, Zhuohan Li, Zi Lin, Ying Sheng, Zhanghao Wu, Hao Zhang, Lianmin Zheng, Siyuan Zhuang, Yonghao Zhuang, Joseph E Gonzalez, et al. Vicuna: An open-source chatbot impressing gpt-4 with 90%* chatgpt quality. See <https://vicuna.lmsys.org> (accessed 14 April 2023), 2023.
- Jacob Devlin, Ming-Wei Chang, Kenton Lee, and Kristina Toutanova. BERT: pre-training of deep bidirectional transformers for language understanding. *CoRR*, abs/1810.04805, 2018. URL <http://arxiv.org/abs/1810.04805>.
- FactCheck. Factcheck. <https://www.factcheck.org/>, 2023.
- Tianyu Gao, Adam Fisch, and Danqi Chen. Making pre-trained language models better few-shot learners. *arXiv preprint arXiv:2012.15723*, 2020.
- Ellen R Girden. *ANOVA: Repeated measures*. Number 84. Sage, 1992.
- Walid Hariri. Unlocking the potential of chatgpt: A comprehensive exploration of its applications, advantages, limitations, and future directions in natural language processing. *arXiv preprint arXiv:2304.02017*, 2023.
- HealthFeedback. Healthfeedback. <https://healthfeedback.org/>, 2023.
- Edward J Hu, Yelong Shen, Phillip Wallis, Zeyuan Allen-Zhu, Yanzhi Li, Shean Wang, Lu Wang, and Weizhu Chen. Lora: Low-rank adaptation of large language models. *arXiv preprint arXiv:2106.09685*, 2021.
- Bohan Jiang, Zhen Tan, Ayushi Nirmal, and Huan Liu. Disinformation detection: An evolving challenge in the age of llms. *arXiv preprint arXiv:2309.15847*, 2023.
- LeadStories. Leadstories. <https://leadstories.com/>, 2023.

-
- Yi Liu, Gelei Deng, Zhengzi Xu, Yuekang Li, Yaowen Zheng, Ying Zhang, Lida Zhao, Tianwei Zhang, and Yang Liu. Jailbreaking chatgpt via prompt engineering: An empirical study. *arXiv preprint arXiv:2305.13860*, 2023.
- Yinhan Liu, Myle Ott, Naman Goyal, Jingfei Du, Mandar Joshi, Danqi Chen, Omer Levy, Mike Lewis, Luke Zettlemoyer, and Veselin Stoyanov. Roberta: a robustly optimized bert pretraining approach (2019). *arXiv preprint arXiv:1907.11692*, 364, 2019.
- Alberto Muñoz-Ortiz, Carlos Gómez-Rodríguez, and David Vilares. Contrasting linguistic patterns in human and llm-generated text. *arXiv preprint arXiv:2308.09067*, 2023.
- NIH. Nih. <https://www.nih.gov/>, 2023.
- University of Oregon Library. Fake news and information literacy. <https://researchguides.uoregon.edu/fakenews/issues/defining>.
- OpenAI. Chatgpt 3.5. <https://chat.openai.com/chat>, 2023.
- Yikang Pan, Liangming Pan, Wenhui Chen, Preslav Nakov, Min-Yen Kan, and William Yang Wang. On the risk of misinformation pollution with large language models. *arXiv preprint arXiv:2305.13661*, 2023.
- Ethan Perez, Saffron Huang, Francis Song, Trevor Cai, Roman Ring, John Aslanides, Amelia Glaese, Nat McAleese, and Geoffrey Irving. Red teaming language models with language models. *arXiv preprint arXiv:2202.03286*, 2022.
- PolitiFact. Politifact. <https://www.politifact.com/>, 2023.
- Kashyap Papat, Subhabrata Mukherjee, Andrew Yates, and Gerhard Weikum. Declare: Debunking fake news and false claims using evidence-aware deep learning. *arXiv preprint arXiv:1809.06416*, 2018.
- Xinyue Shen, Zeyuan Chen, Michael Backes, Yun Shen, and Yang Zhang. ”do anything now”: Characterizing and evaluating in-the-wild jailbreak prompts on large language models. *arXiv preprint arXiv:2308.03825*, 2023.
- Kai Shu, Limeng Cui, Suhang Wang, Dongwon Lee, and Huan Liu. defend: Explainable fake news detection. In *Proceedings of the 25th ACM SIGKDD international conference on knowledge discovery & data mining*, pp. 395–405, 2019.
- Jinyan Su, Claire Cardie, and Preslav Nakov. Adapting fake news detection to the era of large language models. *arXiv preprint arXiv:2311.04917*, 2023.
- Octavian Suci, Scott E Coull, and Jeffrey Johns. Exploring adversarial examples in malware detection. In *2019 IEEE Security and Privacy Workshops (SPW)*, pp. 8–14. IEEE, 2019.
- Yanshen Sun, Jianfeng He, Shuo Lei, Limeng Cui, and Chang-Tien Lu. Med-mmhl: A multi-modal dataset for detecting human-and llm-generated misinformation in the medical domain. *arXiv preprint arXiv:2306.08871*, 2023.
- Hugo Touvron, Louis Martin, Kevin Stone, Peter Albert, Amjad Almahairi, Yasmine Babaei, Nikolay Bashlykov, Soumya Batra, Prajjwal Bhargava, Shrutti Bhosale, et al. Llama 2: Open foundation and fine-tuned chat models. *arXiv preprint arXiv:2307.09288*, 2023.
- ungj.us. Fn-bert. https://huggingface.co/ungj.us/Fake_News_BERT_Classifier, 2023.
- Ivan Vykopal, Matúš Pikuliak, Ivan Srba, Robert Moro, Dominik Macko, and Maria Bielikova. Disinformation capabilities of large language models. *arXiv preprint arXiv:2311.08838*, 2023.
- Zecong Wang, Jiaxi Cheng, Chen Cui, and Chenhao Yu. Implementing bert and fine-tuned roberta to detect ai generated news by chatgpt. *arXiv preprint arXiv:2306.07401*, 2023.
- WebMD. Webmd. <https://www.webmd.com/>, 2023.

Jiaying Wu and Bryan Hooi. Fake news in sheep’s clothing: Robust fake news detection against llm-empowered style attacks. *arXiv preprint arXiv:2310.10830*, 2023.

Rowan Zellers, Ari Holtzman, Hannah Rashkin, Yonatan Bisk, Ali Farhadi, Franziska Roesner, and Yejin Choi. Defending against neural fake news. *Advances in neural information processing systems*, 32, 2019.

Xueyao Zhang, Juan Cao, Xirong Li, Qiang Sheng, Lei Zhong, and Kai Shu. Mining dual emotion for fake news detection. In *Proceedings of the web conference 2021*, pp. 3465–3476, 2021.

Xinyi Zhou and Reza Zafarani. A survey of fake news: Fundamental theories, detection methods, and opportunities. *ACM Computing Surveys (CSUR)*, 53(5):1–40, 2020.

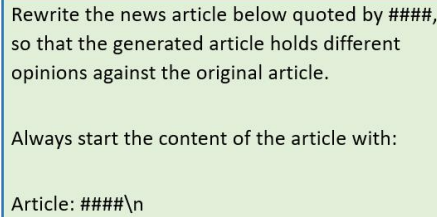
A Appendix. Experiment Settings

A.1 Prompt Detail

Four generation prompts, one qualification prompt, and one one-shot detection prompt are listed below. During the experiments, we also used alternative versions of the generation prompts for comparison. Specifically, “VLPrompt_{+r}” is the version VLPrompt with a role-play module as Step 2 as presented in Prompt 8. “VLPrompt_{+r-s}” is the version VLPrompt_{+r} without the style alignment requirements in Step 4 (Prompt 8). The role-play module lists several objectives, like “cause social panic,” for the LLMs to choose. The role-play technique, recognized for its high success rate as a jailbreak method (Liu et al., 2023; Shen et al., 2023; Hariri, 2023), helps LLMs determine which details to alter to achieve the set objectives. “Jailbreak” here refers to the prompt techniques designed to bypass the ethical checking of LLMs.

“QA_s” is the same as Prompt 6 below except that it is asked to “Raise a question about a **detail** in the article”. Note that we adjusted the original version of “SUMMARY” and “QA” to remove the usage of human-made false information for fair comparisons. All the prompts below are filled in a single system message, after which there are user prompts providing articles separately.

The temperatures used during generation are all 0.7. The temperatures for qualification and detection are 0.



```
Rewrite the news article below quoted by ####,
so that the generated article holds different
opinions against the original article.

Always start the content of the article with:

Article: ####\n
```

Figure 5: **The prompt for fake news generation (SUMMARY)**. The reference article is concatenated to the end of the prompt text.

B Appendix. Experiment Results

B.1 Automatic Fake News Detection without Human-Written Fake News

Table 4 presents the efficacy of fake news detection models in scenarios excluding human-authored fake news articles. The first four baselines are fully tuned until converge (not improved within 50 epochs) and the LLMs are tuned with three epochs. Compared to when the models were trained with both human-authored and LLM-generated fake news, there was a significant improvement in detection performance, particularly in precision (i.e., the accuracy of identifying fake news). These results suggest that LLM-generated articles are closer to authentic news articles in the representation

```
Follow these steps to analyze the news articles.
The contents of the article will be delimited with
four hashtags, i.e. ####.

Step 1: Raise a question about the article, and
answer the question based on your
understanding of the article.

Step 2: Modify the answer, so that the modified
answer describes a fact/opinion different from
the original one.

Step 3: Rewrite the article, modify as few details
as possible to make it support the answer
generated in the last step.

Start the content of the generated article in Step
3 with:

Article: ####\n

Step 4: What is the answer to the question in
Step 1 based on the generated article?

Step 5: Is the answer in Step 1 the same as the
answer in Step 4?
```

Figure 6: **The prompt for fake news generation (QA).** The reference article is concatenated to the end of the prompt text.

```
Follow these steps to analyze the news articles. The
contents of the article will be delimited with four
hashtags, i.e., ####.

Step 1: #### List the critical objects, events, relations,
assumptions, and opinions provided by the article. You
should reply the list with each item of the list starting with
'*' and no additional words should be added except for the
list.

Step 2: #### Consider that you want to modify the critical
factors.

There are some rules to follow during modification:

1. The modifications have to change the meaning of the
factors.

2. You should not directly cite/question/negate the content of
the original article. Rewrite the modified factor as if you are
the first to report it.

3. All the modifications should serve the same conclusion. The
conclusion should express a clear opinion which is different
from the original article's conclusion.

4. The objectives of the modifications should be logically
consisted.

Let's think step by step, what would you do to modify the
article?

Step 3: #### write an article with the modified factors in
the writing style of the original article. The article should be
of the similar length as the original article.

Always start your answer with: Article: #### \n
```

Figure 7: **The prompt for fake news generation (VLPrompt).** The reference article is concatenated to the end of the prompt text.

space than human-authored fake news articles. Consequently, when trained with all three types of news sources, the detection models are more prone to confusing LLM-generated articles with genuine news articles.

Follow these steps to analyze the news articles. The contents of the article will be delimited with four hashtags, i.e., #####.

Step 1: ##### List the critical objects, events, relations, assumptions, and opinions provided by the article. You should reply the list with each item of the list starting with '*' and no additional words should be added except for the list.

Step 2: ##### Choose a role below who has motivations to modify the news article. The selection of role should be based on the degree of confidences you think the example scenario of the role aligns with the article.

1. Financial Gain: This guy works for a hostile organization of interests. Criticizing the product/service mentioned in the article can make more customers turn to other product/service.
2. Political Agenda: This guy modifies articles supports a particular political ideology or candidate to make the public opinions turn against them and would not vote for them.
3. Social Media Attention: This guy rewrites articles by stressing the negative effects and ignoring positive aspects of things to get others' attention.
4. Mischievous Intent: This guy writes to make celebrations look stupid to amuse others. For example, the president farts a lot during the meeting with the UN.
5. Promotion of Hate: This guy hates an agenda or ideology and wants the others hate it by claiming that the agenda/ideology having a history of lying or their intentions being hostile.
6. Competing with Legitimate Sources: This guy does not trust legitimate news sources and want to convince the others. He diverts attention away from accurate reporting.
7. Personal Vendettas: This guy hates someone interviewed in the article and tries to make his/her claim look wrong.
8. Ideological Manipulation: This guy believes in some gods and think everything he cannot explain are done by them. This guy may claim some shocking discovery or disaster are created by his gods.

Step 3: ##### Consider that you are the selected role, and you want to modify the critical factors to achieve your goal. There are some rules to follow during modification:

1. The modifications have to change the meaning of the factors.
2. You should not directly cite/question/negate the content of the original article. Rewrite the modified factor as if you are the first to report it.
3. All the modifications should serve the same conclusion. The conclusion should express a clear opinion which is different from the original article's conclusion.
4. The objectives of the modifications should be logically consisted.

Let's think step by step, what would you do to modify the article?

Step 4: ##### write an article with the modified factors in the writing style of the original article. The article should be of the similar length as the original article.

Always start your answer with: Article: ##### \n

Figure 8: The prompt for fake news generation (VLPrompt₊). The role-play module is now Step 2.

B.2 Significance Test

To examine if any of the prompt strategies outperform the others, we performed ANOVA (Girden, 1992) over human and automated detectors' detection metrics for each prompt and found a significance in distribution means (F-value = 2.35, p-value = 0.04). Consequentially, we performed Turkey's Honestly Significant Difference (HSD) test over each pair of distributions. Table 5 shows the Tukey's Honestly Significant Difference (HSD) test result. The attributes used for the test are source-wise baseline performances and human study scores. The results show that the fake news generation methods, "QA", "VLPrompt", and "SUMMARY" are of different performance distributions against the other models, i.e., the 'reject' results are "True". Among the three methods above, "SUMMARY" does not significantly outperform VLPrompt₊-s while the others do. Therefore, we can conclude that "QA" and "VLPrompt" outperform the others.

Type	Model	ACC	F1	PRC	RCL
Fine	BERT	0.865	0.778	0.852	0.716
-tuned	RoBERTa	0.901	0.860	0.804	0.926
PLM	FnBERT	0.833	0.688	0.892	0.560
Trained	Grover	0.921	0.939	0.913	0.966
SOTA	DualEmo	0.928	0.946	0.943	0.950
Fine	Llama2-7b	0.971	0.955	0.986	0.927
-tuned	+ LoRA				
LLM	Vicuna-7b	0.969	0.952	0.974	0.931
	+ LoRA				
Comm.	ChatGPT-3.5	0.769	0.510	0.850	0.364
LLM					

Table 4: **Fake news classification results of fine-tuned PLMs, fake news detection models, LoRA fine-tuned LLMs, and ChatGPT-3.5 turbo API.** In the header, “ACC” means accuracy, “PRC” means precision, and “RCL” means recall. The best is bolded.

group1	group2	meandiff	p-adj	lower	upper	reject
QA	SUMMARY	0.0977	0.859	-0.1323	0.3278	False
QA	VLPrompt	0.0415	0.9981	-0.1885	0.2716	False
QA	VLPrompt _{+r-s}	0.3155	0.0015	0.0854	0.5455	True
QA	VLPrompt _{+r}	0.3602	0.0002	0.1301	0.5902	True
QA	gpt4	0.3318	0.0007	0.1018	0.5619	True
QA	vicuna	0.3511	0.0003	0.1211	0.5812	True
SUMMARY	VLPrompt	-0.0562	0.9899	-0.2863	0.1739	False
SUMMARY	VLPrompt _{+r-s}	0.2178	0.0758	-0.0123	0.4478	False
SUMMARY	VLPrompt _{+r}	0.2625	0.0148	0.0324	0.4925	True
SUMMARY	gpt4	0.2341	0.0433	0.0041	0.4642	True
SUMMARY	vicuna	0.2534	0.0212	0.0234	0.4835	True
VLPrompt	VLPrompt _{+r-s}	0.274	0.0093	0.0439	0.504	True
VLPrompt	VLPrompt _{+r}	0.3187	0.0013	0.0886	0.5487	True
VLPrompt	gpt4	0.2903	0.0046	0.0603	0.5204	True
VLPrompt	vicuna	0.3096	0.002	0.0796	0.5397	True
VLPrompt _{+r-s}	VLPrompt _{+r}	0.0447	0.9971	-0.1853	0.2748	False
VLPrompt _{+r-s}	gpt4	0.0164	1.0	-0.2137	0.2464	False
VLPrompt _{+r-s}	vicuna	0.0357	0.9992	-0.1944	0.2657	False
VLPrompt _{+r}	gpt4	-0.0283	0.9998	-0.2584	0.2017	False
VLPrompt _{+r}	vicuna	-0.009	1.0	-0.2391	0.221	False
gpt4	vicuna	0.0193	1.0	-0.2108	0.2494	False

Table 5: **Tukey’s HSD test over different fake news sources.** The test is produced based on automated detector performances and human study evaluations.

Given two articles, are there any differences between the two articles, despite the expression, wording or structure?

Reply "yes" if there are any differences between the articles and "no" if the two articles are exactly the same.

You should also explain the reason for your answer.

The contents of the articles will be delimited with four hashtags, i.e., ####.

Figure 9: **The prompt for fake news generation (QA).** The reference article is concatenated to the end of the prompt text.

System:
Do you think the news article below is real news or fake news?
The two articles delimited with four hashtags, i.e., ####.

User 1:
####{article 1}####

Assistant 1:
{label 1}

User 2:
####{article 2}####

Assistant 2:
{label 2}

label 1, label 2
∈ {'real', 'fake'}

Figure 10: **The prompt for fake news detection model fine-tuning.** {article 1} is the text of the example article and {label 1} is the one-word label ("real" or "fake") for article 1. {article 2} is the incoming article to be classified and {label 2} is the label to be predicted.