

EmoRAG: Evaluating RAG Robustness to Symbolic Perturbations

Xinyun Zhou^{†*}
ZJU
Hangzhou, China
xinyun.zhou@zju.edu.cn

Xinfeng Li^{†✉}
NTU
Singapore
xinfeng.li@ntu.edu.sg

Yinan Peng
Hengxin Tech.
Singapore
yinan.peng@palmim.com

Ming Xu
NUS
Singapore
ming.xu@nus.edu.sg

Xuanwang Zhang
NJU
Nanjing, China
zxw.ubw@gmail.com

Miao Yu
NTU
Singapore
fishthreewater@gmail.com

Yidong Wang
PKU
Beijing, China
yidongwang37@gmail.com

Xiaojun Jia
NTU
Singapore
jiaxiaojunqq@gmail.com

Kun Wang
NTU
Singapore
kun.wang@ntu.edu.sg

Qingsong Wen
Squirrel Ai Learning
Seattle, WA, USA
qingsongedu@gmail.com

XiaoFeng Wang
NTU
Singapore
xiaofeng.wang@ntu.edu.sg

Wei Dong
NTU
Singapore
wei_dong@ntu.edu.sg

Abstract

Retrieval-Augmented Generation (RAG) systems are increasingly central to robust AI, enhancing large language model (LLM) faithfulness by incorporating external knowledge. However, our study unveils a critical, overlooked vulnerability: their profound susceptibility to subtle symbolic perturbations, particularly through near-imperceptible emotional icons (e.g., “@_@”) that can catastrophically mislead retrieval, termed EmoRAG. We demonstrate that injecting a single emoticon into a query makes it nearly 100% likely to retrieve semantically unrelated texts, which contain a matching emoticon. Our extensive experiment across general question-answering and code domains, using a range of state-of-the-art retrievers and generators, reveals three key findings: (I) *Single-Emoticon Disaster*: Minimal emoticon injections cause maximal disruptions, with a single emoticon almost 100% dominating RAG output. (II) *Positional Sensitivity*: Placing an emoticon at the beginning of a query can cause severe perturbation, with F1-Scores exceeding 0.92 across all datasets. (III) *Parameter-Scale Vulnerability*: Counterintuitively, models with larger parameters exhibit greater vulnerability to the interference. We provide an in-depth analysis to uncover the underlying mechanisms of these phenomena. Furthermore, we raise a critical concern regarding the robustness assumption of current RAG systems, envisioning a threat scenario where an adversary exploits this vulnerability to manipulate the RAG system. We evaluate standard defenses and find them insufficient against EmoRAG. To address this, we propose targeted defenses, analyzing their strengths and limitations in mitigating emoticon-based perturbations. Finally, we outline future directions for building robust RAG systems.

Keywords

Retrieval-Augmented-Generation, Symbolic Perturbations, Large Language Models

1 Introduction

Large language models (LLMs) excel in many tasks but face limitations such as hallucinations [29] and difficulty in assimilating new knowledge [52]. To address these shortcomings and promote more robust AI systems, Retrieval-Augmented Generation (RAG) has emerged as a promising framework. By integrating a retriever, an external knowledge database, and a generator (LLM), RAG aims to produce contextually accurate, up-to-date responses. Tools like ChatGPT Retrieval Plugin, LangChain, and applications like Bing Search exemplify RAG’s growing influence.

Recent research has primarily focused on enhancing model performance by improving the retriever component [48, 68], refining the generator’s capabilities [11], or exploring joint optimization of both components [56, 59]. A common thread in these efforts is the assumption that retrieval quality hinges on the semantic relevance between user queries and knowledge base texts. However, does the outcome of retrieval in RAG systems truly rely on semantic relevance?

We uncover a critical, previously overlooked phenomenon: a stark decoupling between semantic relevance and retrieval outcomes in RAG systems. We demonstrate that subtle symbolic perturbations, specifically the injection of seemingly innocuous emoticons, can catastrophically hijack the retrieval process, forcing the system to prioritize irrelevant, emoticon-matched content over semantically pertinent information (as illustrated in Figure 1). This vulnerability, which we term EmoRAG, exposes a significant chink in the armor of current RAG architectures. We meticulously investigate this by conducting controlled experiments across diverse datasets from different domains, using a variety of state-of-the-art retrievers and generators (LLMs). Specifically, we utilize two widely used general Q&A datasets: *Natural Questions* [37] and *MS-MARCO* [8]. Also, we extend our evaluation to a specialized domain, incorporating a dataset from *Code* [13]. Our study systematically varies factors such as the number, position, and type of emoticons, and evaluates advanced RAG frameworks and the potential for cross-emoticon triggering.

Why focus on emoticons? Symbolic perturbations, such as emoticons (e.g., ‘:-’) or emojis, convey meaning visually rather than

[†]Co-first author.

*Work done when the author was visiting Wei Dong’s group at NTU.

✉Corresponding author.

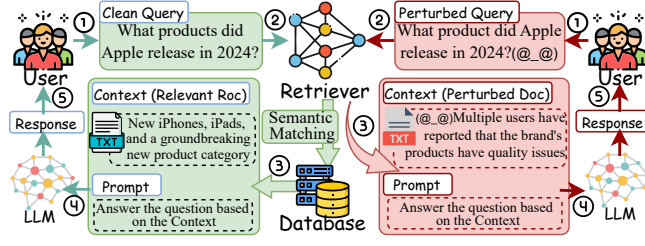


Figure 1: Illustration of emoticon-based perturbation hijacking a RAG system. Step ①: User submits query. Step ②: Retriever Processing. Step ③: The retriever passes the context to the LLM. Step ④: The LLM generates a response. Step ⑤: User Receives Response.

through direct semantic encoding. Emoticons are widely used in online communication. For instance, Facebook sees over 700 million daily emoticon messages [3] and Twitter handles about 250 million monthly [7]. While other symbols like emojis or even garbled text (common in adversarial attack studies [16, 70]) exist, our user study (detailed in Appendix E) evaluated the three types of characters across two dimensions, *Noticeability* and *Alertness*. The results show that emojis are too noticeable and garbled texts are both noticeable and alarming. In contrast, emoticons appear natural on both fronts, highlighting their potential for exploitation. Although our primary focus lies in the symbolic structure and token-level behavior of emoticons, this emphasis serves as a starting point to expose a deeper issue: RAG systems are sensitive to rare symbolic tokens that distort embeddings, regardless of their semantic relevance.

Our extensive experiments reveal several key findings: (I) *Single-Emoticon Disaster*: Even a single emoticon can catastrophically affect RAG systems, causing nearly 100% retrieval of semantically irrelevant content. (II) *Widespread Effectiveness*: Around 83% of tested emoticons can induce such nearly 100% retrieval failures as mentioned above. (III) *Positional Sensitivity*: Placing a single emoticon at the beginning of a query can cause severe perturbation, with F1-Scores exceeding 0.92 across all datasets. (IV) *Parameter-Scale Vulnerability*: Larger models are significantly more sensitive to emoticon-induced perturbations, with F1-Scores almost always reaching 1.00 under perturbation. (V) *No Cross-Triggering*: Specific emoticons only retrieve content containing the same emoticon, which may provide an attack vector for potential adversaries.

To understand these observations, we conduct an in-depth analysis of EmoRAG, and we reveal three mechanistic insights: (I) *Emoticon Modeling Deficit*: Current retrievers struggle to effectively model emoticons, often due to their long-tail distribution in training vocabularies, leading to unstable representations. (II) *Positional Shift*: Emoticons at the query’s start cause a significant shift in the positional embeddings of all subsequent tokens, fundamentally altering the query’s representation. (III) *Vulnerability of Larger Models*: Larger models have higher-dimensional representation spaces, making their query embeddings more susceptible to perturbation. This analysis helps explain the wide applicability and severity of emoticon-based attacks in the RAG system.

Based on the above observations, we envision several realistic and feasible threat scenarios in which adversaries can covertly manipulate the output of RAG systems by using specific emoticons or other rare tokens as triggers. For example, in code security risk

assessment, an attacker can insert emoticons into code comments, which may trigger the retrieval of specific code snippets, leading to incorrect assessments and the introduction of vulnerabilities. Similarly, in RAG-based review scoring, attackers can embed emoticons into documents to bias the system toward retrieving higher-rated content over similar alternatives, thereby manipulating evaluation outcomes.

Recognizing the severity of this vulnerability, we evaluate standard defense mechanisms, such as perplexity-based detection, and find them largely insufficient against EmoRAG due to high false positive rates. To address this, we developed a dataset for detecting emoticon-based perturbed text, derived from the NQ dataset. Using this data, we trained a BERT-based model, which achieves 99% accuracy in identifying perturbed text. While effective, this approach is tailored specifically to emoticons and does not account for other types of special characters, underscoring the need for broader defenses. Out of ethical considerations, we open-source the defense-related components: the dataset we created with perturbed text and the model we trained to detect potential malicious text.

Our main contributions are as follows:

- We present the first empirical study across three datasets, multiple retrievers, advanced RAG frameworks, and a range of LLMs, revealing a critical decoupling of semantic relevance and retrieval outcome within RAG systems, where minor symbolic perturbations can dominate the retrieval outcomes completely.
- We provide an in-depth analysis explaining why emoticons, along with other forms of symbolic perturbations, can significantly dominate the retrieval process of RAG systems, providing a solid foundation for understanding their vulnerability.
- We envision realistic threat scenarios where adversaries exploit the vulnerability to manipulate the RAG system, while offering guidance for building robust RAG systems.
- We explore several defense strategies against EmoRAG, aiming to mitigate its impact on RAG systems. To support further research in this area, we open-source our dataset and models. Building on our insights, we envision next-generation robust RAG systems.

2 Background and Related Work

2.1 RAG systems

In recent advances in natural language processing, RAG has emerged as an effective framework for integrating external knowledge into language models [5, 72]. A RAG system consists of three components: a **Knowledge Database** [58, 61], a **Retriever** [26, 33], and a **Generator** [40, 41]. Unlike traditional generative models, RAG dynamically retrieves relevant information from external knowledge, enabling accurate and context-rich responses.

The RAG system operates in two stages: **retrieval** and **generation**. In the retrieval stage, given a query q , the retriever R searches the knowledge database $\mathcal{K}$ and ranks documents based on relevance: $D = R(q, \mathcal{K})$, where D is the set of top-ranked documents. Embedding-based methods like dense passage retrieval [34] ensure query-document alignment in a shared vector space. In the generation stage, the retrieved documents D are combined with the query q by the generator G , typically a pre-trained language model, to produce the final response: $\hat{r} = G(q, D)$, in which D serves as

additional text. The generator ensures responses are linguistically fluent and contextually accurate.

2.2 Applications of RAG Systems

RAG systems have shown immense potential in real-world applications, particularly in areas like general QA and code-related tasks. The following sections will explore recent advancements and practical implementations of RAG systems in these domains.

General: In the general domain, RAG systems have gained significant attention for enhancing AI applications. A prime example is WikiChat [53], a low-latency chatbot based on Wikipedia that reduces hallucinations while maintaining high conversational quality. In practical applications, Shopify’s Sidekick chatbot [54] uses RAG to extract store data and answer product and account queries, improving customer service. And Amazon leverages RAG for its recommendation engine [4], providing tailored product suggestions to boost sales and customer satisfaction. Similarly, RedNote [50] leverages user posts as a knowledge base and employs the RAG system to generate recommendations.

Code: RAG systems have proven transformative in real-world coding applications. For instance, Google’s Vertex AI Codey APIs [24] use RAG to facilitate context-aware code generation and completion, ensuring alignment with organizational coding standards. Similarly, Qodo [2] leverages RAG to manage large-scale code repositories, enabling developers to efficiently retrieve and integrate relevant code snippets. Additionally, platforms like Codeforces [12] and LeetCode [38] utilize RAG to analyze users’ code errors, retrieve relevant documentation or example code, and offer targeted suggestions for fixes. GitHub Copilot [22] and Cursor [15] integrate specific open-source code repositories through the GitHub API, identifying code errors and providing more accurate code suggestions and error corrections. In this context, GitHub acts as a vital knowledge source.

2.3 Existing Attacks on RAG systems

A large body of research has demonstrated that machine learning models are highly susceptible to data poisoning and backdoor attacks [9, 20, 31, 44, 55]. Specifically, when trained on a poisoned training dataset, the machine learning model has manipulated behavior. However, this attack requires direct manipulation of the model training process, which is often not feasible. When extended to RAG systems, attackers can exploit the retrieval component by injecting malicious texts into the knowledge database. This manipulation influences the data retrieved during inference, ultimately distorting the generator’s outputs. For example, attackers may inject nonsensical but retrievable text [74], semantically misleading information [75], or construct malicious texts by formalizing the attack as an optimization problem [71]. These works require carefully designed retrievable text for the target query, and their malicious text may still be retrieved by normal non-target queries [75], affecting the normal function of the system, which further reduces their concealment. In contrast, EmoRAG can achieve a training-free perturbed attack against the RAG system by injecting only a small amount of emoticons; for example, a single emoticon can increase the F1-Score beyond 0.9, as demonstrated in Figure 4. This capability simplifies the attack process and enhances its effectiveness. Our

results in §3 show that EmoRAG will only be triggered when emoticons appear and will not have any impact on the normal operation of the RAG system, further improving its concealment.

3 Measurement of Emoticon Interference

Our goal is to gain a deeper understanding of how subtle query perturbations, particularly through the use of emoticons, affect the retrieval mechanisms within RAG systems, ultimately revealing potential vulnerabilities that could compromise system reliability and user trust.

3.1 Measurement Setup

Due to space constraints, the detailed experiment setup, including the typical datasets, the three components of the RAG system (retriever, generator, and database), evaluation metrics, the design of perturbed texts, baseline, and hyperparameter settings, are provided in the Appendix A.

3.2 Key Observations from Evaluation

Finding 1: EmoRAG achieves near-perfect ASRs and F1-Scores under perturbed queries. Table 1 and Table 2 report the F1-Scores and ASRs achieved by EmoRAG under perturbed queries. Our experiments reveal the following key observations: First, EmoRAG achieves near 100% ASRs across various retrievers, even with only $N = 5$ perturbed texts injected into a knowledge database of millions of entries. Second, EmoRAG demonstrates robust performance across both general and specialized domains, with F1-Scores above 0.95 and ASRs close to 100% on all datasets. These results highlight the generalizability and effectiveness of EmoRAG. The superior performance motivates further investigation, as detailed in § 4.

Finding 2: EmoRAG preserves retriever performance under clean queries. Tables 1 and Table 2 show that, under clean query scenarios, the RAG system operates as expected, achieving F1-Scores of 0.0 across all datasets and retrievers. This indicates that no perturbed texts are retrieved under clean query, ensuring normal system functionality.

Finding 3: Models with larger parameters are more susceptible to EmoRAG. As shown in Table 1 and Table 2, models with larger parameter sizes (more than 7B) are more easily perturbed, achieving F1-Scores of 1.0 across all datasets. This suggests that the currently leading models on the MTEB leaderboard [46] are more vulnerable to this emoticon-based perturbation. A detailed analysis is provided in § 4.

3.3 In-depth Factor Analysis

3.3.1 Impact Factors on RAG’s Performance. Impact of generator. Table 3 presents the result of EmoRAG with different generators. For this evaluation, we selected three LLMs with varying parameter sizes as generators: GPT-4o, LLAMA-3.1-8B, and Qwen2.5-1.5B. The temperature hyperparameter for all LLMs was set to 0.0 to ensure consistent responses. The results show that, despite differences in architecture and scale, EmoRAG achieves high effectiveness across all generators, with ASRs exceeding 95% in nearly all cases.

Impact of retrievers. Table 1 and Table 2 present the effects of EmoRAG with various retrievers in RAG systems. The results show that EmoRAG consistently achieves high F1-Scores across various

Table 1: Perturbed effects of EmoRAG across various domains, model architectures and parameter scales, and query types. Noteworthy results are highlighted in Red and Green for emphasis.

Datasets	Query	Metric	Retriever of RAG System					
			SPECTER	Contriever	Qwen2-7B	e5-7B-mistral	SFR-Embedding	BGE-en-icl
Natural Question	Perturbed	ASR $\uparrow$	100.00%	100.00%	100.00%	100.00%	99.98%	100.00%
		F1-Score $\uparrow$	0.96	0.97	1.00	1.00	1.00	1.00
	Clean	F1-Score $\downarrow$	0.00	0.00	0.00	0.00	0.00	0.00
MS-MARCO	Perturbed	ASR $\uparrow$	99.97%	99.98%	99.98%	99.97%	100.00%	99.98%
		F1-Score $\uparrow$	0.97	0.98	1.00	1.00	1.00	1.00
	Clean	F1-Score $\downarrow$	0.00	0.00	0.00	0.00	0.00	0.00
CODE	Perturbed	ASR $\uparrow$	99.98%	99.91%	99.96%	99.96%	99.96%	99.96%
		F1-Score $\uparrow$	0.96	0.99	1.00	1.00	1.00	1.00
	Clean	F1-Score $\downarrow$	0.00	0.00	0.00	0.00	0.00	0.00

‡: A "Perturbed" refers to a query that includes emoticons, a "Clean" refers to a query without emoticons.

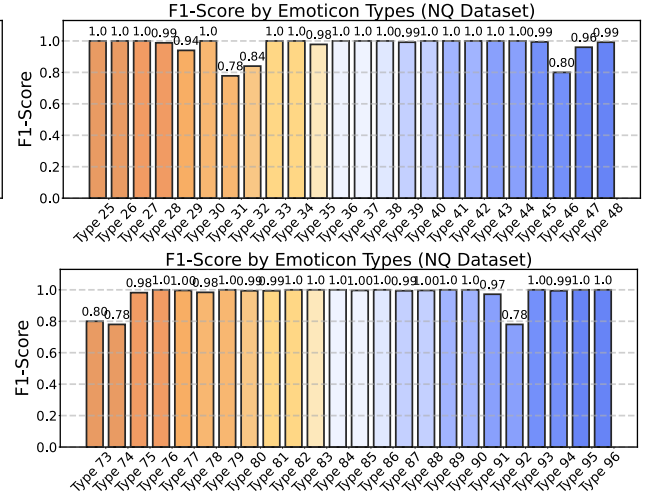
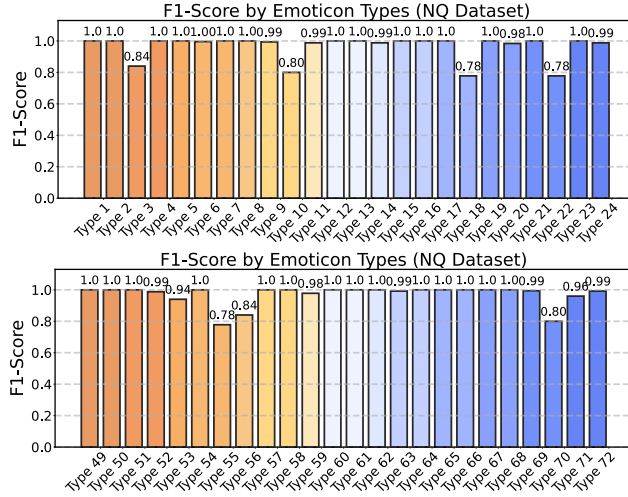
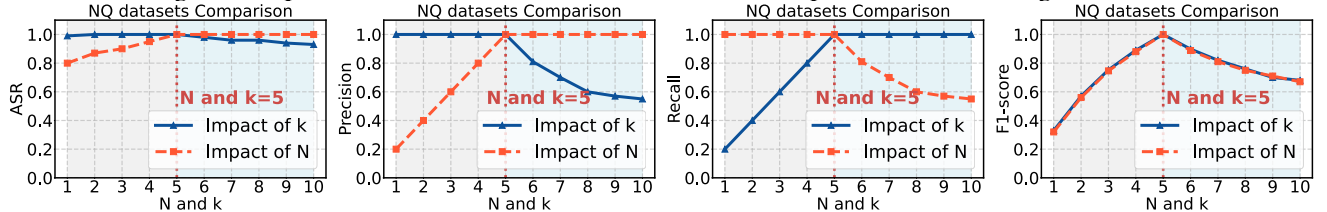
Table 2: Perturbed effect of EmoRAG on the Code domain-specific retriever

Datasets	Retriever	Perturbed		Clean
		F1-Score $\uparrow$	ASR $\uparrow$	F1-Score $\downarrow$
CODE	CodeBERT	0.96	99.96%	0.00

‡: CodeBERT is a domain-specific model for natural and programming languages.

Table 3: Perturbed effect of EmoRAG on generators

Datasets	Metrics	Generator		
		GPT-4o	LLaMA3	Qwen2.5
NQ	ASR $\uparrow$	100.00%	94.85%	95.04%
MS-MARCO	ASR $\uparrow$	99.97%	98.57%	99.98%
CODE	ASR $\uparrow$	99.93%	94.36%	96.96%

**Figure 2: Impact of 96 emoticons with diverse structures, frequencies, and meanings on EmoRAG****Figure 3: The impact of increasing N and k on ASR, Precision, Recall, and F1-Score in the NQ dataset.**

retrievers, regardless of their parameters or architectures, including Code-BERT, which is specifically designed for the code domain.

Impact of k . Figure 3 illustrates the impact of k on EmoRAG, where k represents the number of top- k most similar texts returned

by the retriever. When $k \leq N$ ($N = 5$ by default), the ASR of EmoRAG remains high. Precision, which measures the fraction of retrieved perturbed texts, remains very high, while Recall increases as k increases. When $k > N$, ASR does not decrease significantly as k

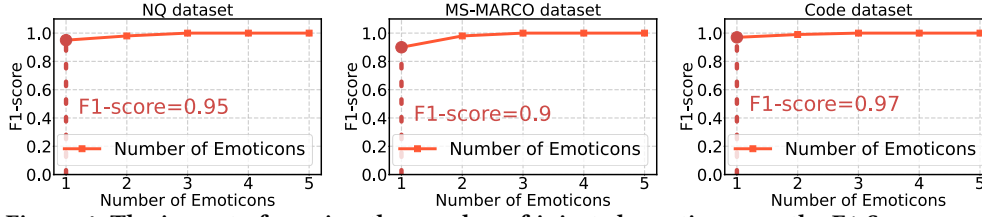


Figure 4: The impact of varying the number of injected emoticons on the F1-Score across multiple datasets with Contriever as the retriever.

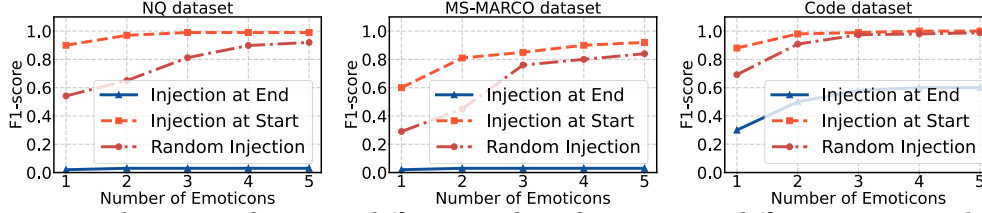


Figure 6: The impact of injecting different numbers of emoticons at different positions within the query and texts, with Contriever as the retriever.

increases. This is due to the shift in the vector space caused by the injected emoticons, which results in fewer semantically relevant texts being retrieved, as further analyzed in § 4. Recall approaches 1 when $k > N$, indicating that nearly all perturbed texts are retrieved.

3.3.2 Impact of Hyperparameters on EmoRAG. Impact of similarity metric. Table 5 presents the results when different similarity metrics are used to calculate the similarity of embedding vectors for retrieving texts from the database in response to a query. We observe that EmoRAG achieves similar results across different similarity metrics in both settings. This consistency suggests the effectiveness of EmoRAG is not highly sensitive to the choice of similarity metric, further demonstrating the robustness of our approach.

Impact of N . Figure 3 illustrates the impact of N on EmoRAG, where N represents the number of perturbed texts injected into the knowledge base. When $N \leq k$ ($k = 5$ by default), the ASR increases as N grows. This is because larger N results in more perturbed texts being injected into the knowledge database. Consequently, Precision also increases with N , while Recall remains consistently high. When $N > k$, ASR and Precision stabilize at consistently high values. The F1-Score, which balances Precision and Recall, initially increases with N but starts to decrease once Recall drops for $N > k$.

Impact of the Number of Emoticons. Figure 4 illustrates the effect of injecting varying numbers of emoticons into queries and perturbed texts, with Contriever as the retriever. First, even with the injection of a small number of emoticons, EmoRAG is capable of executing highly efficient interference. For example, when just a single emoticon is injected at the start of the query, the F1-Score consistently exceeds 0.92 across all datasets. Furthermore, when the number of injected emoticons increases to two, EmoRAG achieves F1-Scores of 1.00 on nearly all datasets, suggesting that it is capable of achieving maximal interference with minimal effort.

Impact of Position of Emoticons. Figure 6 shows the effect of injecting varying numbers of emoticons at different positions in queries and texts. In addition to placing emoticons at the start or end, we also test injecting them at arbitrary positions. Several key observations emerge. First, injecting emoticons at the start can lead

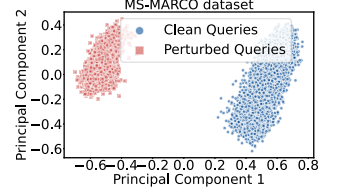


Figure 5: PCA results for the MS-MARCO.

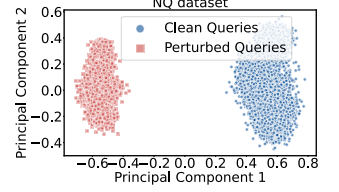


Figure 7: PCA results for the NQ.

to an effective interference, though performance is slightly better when placed at both positions. Interestingly, inserting emoticons at random positions also impacts the retrieval process effectively, but to a lesser degree. Placing emoticons only at the end proves ineffective. We examine this phenomenon in detail in § 4.

Impact of Emoticon Type. We explore how different emoticon types impact EmoRAG. We select 96 emoticons, covering diverse structures, usage frequencies, and meanings. Due to space constraints, Figure 14 (in Appendix D) shows a subset of 14 representative emoticons. As shown in Figure 2, EmoRAG achieves an F1-Score close to 1.0 for about 83% of the types, but scores are lower for 17% of the types. This highlights the vulnerability of RAG systems, as a wide range of emoticons can be used to launch successful attacks. We observe that emoticons with more complex structures usually achieve higher F1-Scores. Based on this, we propose a metric to predict emoticon effectiveness directly, evaluating each emoticon on two features: *total number of tokens* and *number of unique tokens*. The total token count represents individual elements, while unique tokens capture the diversity of components. We calculated the score using the following formula:

$$\text{Score} = \frac{2 \times \text{Total Tokens} \times \text{Unique Tokens}}{\text{Total Tokens} + \text{Unique Tokens}}. \quad (1)$$

These metrics suggest that higher total tokens and greater token diversity lead to more distinct embeddings. However, while this metric is somewhat effective, it is only a preliminary approach, and more robust indicators are needed to accurately assess emoticon effectiveness.

Other Special Characters. While our initial experiment focuses on emoticons, other special characters may also act as triggers in real-world scenarios. With this in mind, and considering the nature of injecting special characters, we select emojis as another type of special character. Following the same experimental setup, we choose six different types of emojis, injecting each type four times at the beginning and end of both queries and malicious texts. In total, five malicious texts are injected into the database. As shown

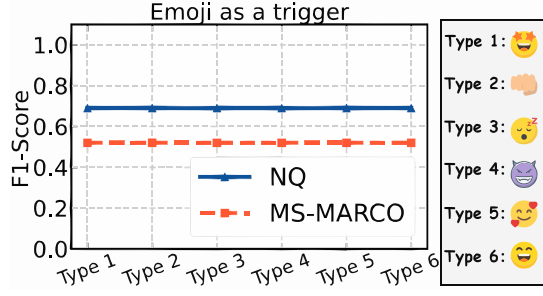


Figure 8: Other Special Characters

Table 4: EmoRAG on different similarity metrics.

Datasets	Similarity	Metrics	
		F1-Score $\uparrow$	ASR $\uparrow$
NQ	Dot Product	0.98	99.97%
	Cosine	0.97	100.00%
MS-MARCO	Dot Product	1.00	100.00%
	Cosine	0.98	99.98%
CODE	Dot Product	0.99	99.96%
	Cosine	0.99	99.96%

Table 5: EmoRAG under Advanced RAG systems

Datasets	Advanced RAG	Metrics	
		F1-Score $\uparrow$	ASR $\uparrow$
NQ	Robust-RAG	0.97	75.51%
	Self-RAG	0.97	76.77%
MS-MARCO	Robust-RAG	0.98	79.79%
	Self-RAG	0.98	85.86%
CODE	Robust-RAG	0.99	83.16%
	Self-RAG	0.99	91.28%

in Figure 8, emojis are much less effective than emoticons in triggering system vulnerabilities. This is likely because emoticons are more complex, and common emojis are already in the model’s vocabulary, reducing their impact. We do not choose garbled characters as special characters for two reasons. First, it is impossible for the same garbled characters to appear in normal user queries, which significantly limits the scope of potential attacks. Second, garbled characters in regular text are uncommon and can easily raise people’s awareness.

Cross-Emoticon Triggering Attack. In the initial experiment, we inject the same emoticons into both the queries and perturbed texts. However, we are curious whether cross-emoticon injection, which involves using different emoticons in queries and perturbed text, could also serve as a trigger? To explore this, we select the first seven emoticons from Figure 14 (Appendix D) and conduct cross experiments with all possible pairs. Specifically, we pair each emoticon with every other emoticon, resulting in a total of 21 unique pairs, following the same experimental setup. The results, presented in Figure 9 (Appendix D), show that only identical emoticons in both the query and the perturbed text can act as effective triggers, achieving an F1-Score of 1.0. When different emoticons are used,

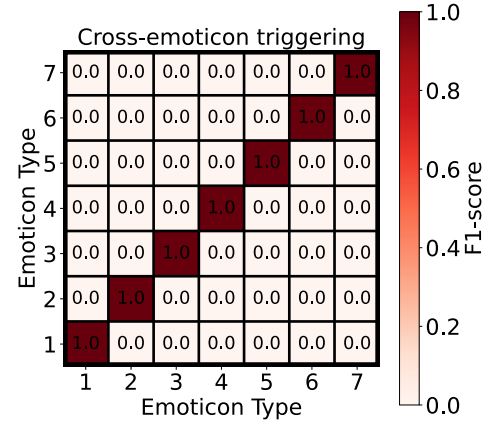


Figure 9: Cross-emoticon triggering

the F1-Scores are all essentially 0.0. This specific triggering behavior enables attackers to exert precise control over RAG system outputs, highlighting a viable pathway for targeted manipulation.

3.4 Advanced RAG Systems.

In the framework outlined above, we primarily focus on the basic RAG system. However, this approach may be less effective in real-world applications that require higher levels of reliability. To address these limitations, several advanced RAG systems have been proposed. For example, Xiang et al. [65] introduced Robust-RAG, which used an isolate-then-aggregate strategy. It first computed responses from the LLM for each passage individually and then securely aggregated them. To ensure robustness, they proposed two aggregation techniques, keyword and decoding aggregation. Meanwhile, Asai et al. [6] introduced Self-RAG, a self-reflective system within a single LLM. This system adaptively retrieved relevant passages on demand and used special tokens to reflect on and enhance both the retrieved passages and the model’s response, improving coherence and accuracy.

With this in mind, we conduct experiments to evaluate the performance of EmoRAG in comparison to these advanced RAG systems. The experimental settings are consistent with previous evaluation, where we injected $N = 5$ perturbed texts into the database. For the Robust-RAG system, we focus on the keywords mechanism, as this defense is particularly suitable for free-form text generation tasks. Additionally, we set the retrieval parameter $k = 10$, meaning that a total of 10 texts were retrieved from the knowledge database. Table 5 shows that EmoRAG achieves high ASRs, demonstrating that even advanced RAG systems remain vulnerable to EmoRAG. As discussed in the § 4, the injection of emoticons disrupts the mapping of the original query in the high-dimensional space. This perturbation forces the retrieval process to reduce the likelihood of retrieving relevant content. As a result, this shift in retrieval dynamics substantially increases the success rate of EmoRAG.

4 General Mechanisms Behind Emoticon Interference

EmoRAG is not a peculiarity of emoticons themselves, but a concrete instance of broader structural vulnerabilities in RAG systems. Its root causes stem from how retrievers handle rare tokens, their

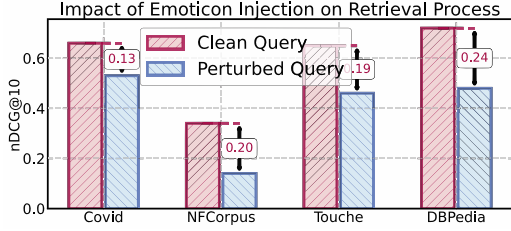


Figure 10: Emoticon Perturbation lowers retrieval performance on the BEIR benchmark
sensitivity to token positions, and the geometric properties of high-dimensional embedding spaces.

4.1 Rare Tokens Shift Query’s Embedding

When processed by tokenizers, emoticons are treated as distinct tokens. Depending on the tokenizer’s design, they may either be split into subword units or replaced with the `<unk>` token if they are out-of-vocabulary (OOV) tokens. When consistently mapped to `<unk>`, the retriever is unable to utilize their contextual semantics, impairing performance on tasks such as text comprehension and sentiment analysis. Importantly, both `<unk>` tokens and emoticons often fall into the *long-tail* of the token distribution [49], where a small set of high-frequency tokens dominates the vocabulary, while rare tokens appear only sparsely in the training data. Thus, token embeddings for rare items, such as emoticons, tend to lie far from common token clusters in the embedding space, formalized by:

$$\text{Dist}(\mathbf{E}(r), \mathbf{E}(w)) \geq \delta, \quad \delta > 0, \quad e \in \mathcal{E}, w \in V \quad (2)$$

Here, $\mathbf{E}(r) \in \mathbb{R}^d$ denotes the embedding of rare tokens, and $\mathbf{E}(w) \in \mathbb{R}^d$ that of a frequent token. Although these rare token embeddings lie far from common tokens in the semantic space, they often cluster closely together. This isolation, combined with internal consistency, allows them to disproportionately influence sentence-level representations. As a result, their presence in queries can unpredictably distort semantic meaning.

To visualize this effect, we apply *Principal Component Analysis* (PCA) to the query embeddings. As shown in Figure 5 and Figure 7, clean queries (blue circles) are spread across the embedding space, reflecting natural semantic diversity. In contrast, perturbed queries with emoticon injections (red squares) collapse into a dense, compact cluster. This shift illustrates how rare tokens such as emoticons sparsify and distort query representations, pulling them away from their original distribution and undermining semantic fidelity in the retriever’s embedding space.

Beyond visualization, we empirically evaluate the impact of such perturbations on retrieval performance. Specifically, we conduct experiments on four datasets from the **BEIR** benchmark, *Covid*, *NFCorpus*, *DBPedia*, and *Touche*, to compare retrieval results between clean and perturbed queries using $\text{nDCG}@10$. As shown in Figure 10, emoticons significantly disrupt semantic alignment, reducing the likelihood of retrieving relevant documents.

4.2 Insertion-Induced Positional Shift

Transformer models encode not only token identities but also their positions within a sequence via positional embeddings. Let $\mathbf{E}_{\text{token}}(w) \in \mathbb{R}^d$ denote the embedding of token w , and $\mathbf{E}_{\text{pos}}(i) \in \mathbb{R}^d$ the positional embedding at position i . The final embedding fed

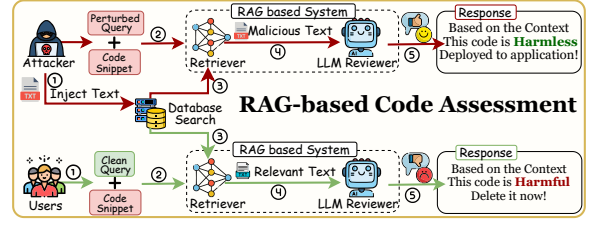


Figure 11: Manipulation of RAG-based code assessment systems via emoticon-triggered retrieval

to the model is defined as: $\mathbf{E}_{\text{final}}(w, i) = \mathbf{E}_{\text{token}}(w) + \mathbf{E}_{\text{pos}}(i)$. This formulation makes transformers inherently sensitive to input token order. When new tokens are inserted at the beginning of a sequence, they systematically shift the positions of all subsequent tokens. For a sequence $w_1, w_2, \dots, w_n$, the insertion of a subsequence of length m at the front results in the following shift: $\mathbf{E}_{\text{final}}(w_i, i) \rightarrow \mathbf{E}_{\text{final}}(w_i, i + m)$, for $i > 1$. This shift alters the positional context of every downstream token, potentially disrupting the model’s learned semantic representations. In contrast, insertions at the end of a sequence leave the relative positions of earlier tokens unchanged, leading to a far less pronounced impact. This demonstrates a general structural vulnerability in transformer-based models: *any insertion near the start of a sequence can induce a global positional shift*, cascading through the architecture and modifying all subsequent token representations. This mechanism applies broadly and helps explain why seemingly minor input changes at the beginning can result in large changes in model behavior.

4.3 Amplification in High Dimensions

Larger retrieval models, with more parameters, are more sensitive to subtle differences between tokens, making them more responsive to variations like the inclusion of emoticons. Operating in high-dimensional embedding spaces, these models capture nuanced token relationships, so even small changes, such as emoticons, can significantly impact sentence embeddings. Additionally, larger retrieval models typically have higher dimensional embedding spaces. In such models, the amplification effect of small perturbations is even greater because the increased dimensionality provides more pathways for these changes to propagate through the embedding. As a result, even small shifts in the embedding caused by the addition of rare tokens can lead to considerable changes in the sentence’s overall representation.

5 Adversarial Threat Modeling

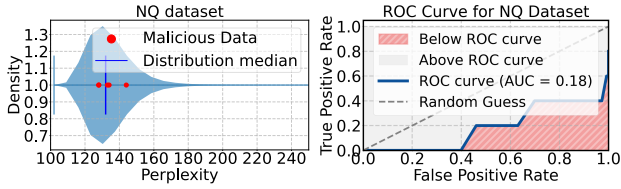
5.1 Threat Scenarios and Practical Exploitability

For RAG systems, particularly in areas like general knowledge question-answering and code generation, the adversary model is unique due to the partial accessibility of the database. Our study considers two potential scenarios: (1) The attacker injects false content into the database, waiting for it to be triggered by queries from benign users; (2) The attacker tricks the RAG system directly to exploit vulnerabilities for malicious gain.

(1) Benign User as Victim: In this scenario, attackers target benign users by exploiting queries that unintentionally contain emoticons, particularly those copied directly from social media,

Table 6: Overall Performance of EmoRAG compared with baselines across various domains with SPECTER as the retriever.

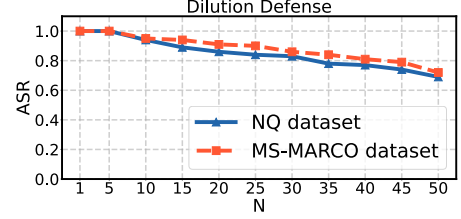
Datasets	Attack	Metrics	
		F1-Score $\uparrow$	ASR $\uparrow$
NQ	Corpus Poisoning	0.96	96.62%
	Prompt Injection	0.75	71.21%
	GCG Attack	0.00	3.01%
	EmoRAG (Ours)	0.97	100.00%
MS-MARCO	Corpus Poisoning	0.96	96.39%
	Prompt Injection	0.72	75.39%
	GCG Attack	0.00	1.13%
	EmoRAG (Ours)	0.98	99.98%
CODE	Corpus Poisoning	0.97	97.11%
	Prompt Injection	0.72	71.76%
	GCG Attack	0.00	2.31%
	EmoRAG (Ours)	0.99	99.91%

**Figure 12: Perplexity Defense against EmoRAG.**

where emoticons are prevalent. A survey by Kika Keyboard [35] found Google ranks fourth among the top five apps for emoticon use, with the others being social media platforms, indicating that users often include emoticons in their queries. In addition, the Unicode Consortium [60] publishes the usage frequency of each emoticon annually. Attackers can leverage this data to craft malicious content featuring the most popular emoticons. When users enter such emoticons into their queries, they unknowingly trigger the malicious content within the RAG database, which may return them misleading information, like the CEO of Apple is Elon Musk [75]. Likewise, for code generation tasks, attackers might insert emoticons into comments within vulnerable code snippets. When users, especially beginners, include matching emoticons in their queries, these malicious snippets are triggered, causing the RAG system to return insecure code.

(2) RAG System as Target: In this scenario, the attacker submits queries directly to the RAG system to manipulate its responses. To achieve this, they pre-inject malicious text with emoticons into comments or documents. When the queries include these emoticons, the injected text is triggered, allowing the attacker to manipulate the system’s responses. As shown in Figure 11, in code security risk assessment, attackers can insert emoticons into code comments, a practice that is already seen in the real world¹, and these emoticons can trigger malicious content, leading to incorrect security assessments and vulnerabilities. Similarly, in RAG-based review scoring, attackers can manipulate scores by inserting emoticons.

¹For instance, GitHub’s CREG [19] provides guidance on using emoticons in code review, and tools like Emojicode [18] make insertions more convenient.

**Figure 13: Dilution Defense against EmoRAG**

5.2 Adversary’s Capability

Starting from feasibility, we assume that the adversary does not have access to the internal parameters of the retriever R or the generator G . Furthermore, the adversary cannot manipulate the training phase of R or G . This ensures that the adversary’s actions are limited to external interactions with the system, specifically by submitting queries q through the system’s interface. In line with previous studies [10, 67, 71, 74, 75], we assume that the adversary has the ability to inject malicious texts into the knowledge database D . However, the modifications to the knowledge base are minimal, with the injected malicious texts constituting less than 0.01 % of the total content in D . This assumption is not only feasible but also aligns with real-world scenarios, as outlined below:

- **General Domain (e.g., Wikipedia):** A recent study [10] demonstrated the feasibility of maliciously editing 6.5% (conservative analysis) of Wikipedia documents. EmoRAG requires only a small number of injected texts (less than 0.01%) to achieve a high Attack Success Rate.
- **Code Domain:** In open-source code repositories, developers can add or edit code, which allows malicious actors to insert emoticons around vulnerable code, creating conditions for attacks. Additionally, some RAG systems use GitHub directly as a knowledge base or connect to specific GitHub repositories via APIs [15, 22], enabling malicious users to upload vulnerable code directly.

6 Defenses against EmoRAG

To counter the risk of emoticon-based attacks on RAG systems, we propose several defense strategies to mitigate the impact of emoticon-based perturbation: *Dilution Defense*, *Query Disinfection*, and *Perturbed Texts Detection*. **Dilution Defense** aims to reduce the interference by increasing the number of retrieved texts. However, as shown in Figure 13 (Appendix B.1), it does not significantly reduce the impact due to the shifts in query representation caused by emoticons. **Query Disinfection** leverages paraphrasing techniques. Specifically, we use GPT-4o to generate five paraphrased queries. For each paraphrased query, k texts are retrieved to generate answers. The final response is generated by aggregating the answers from all paraphrased queries. As shown in Table 9 (Appendix B.2), this defense effectively mitigates EmoRAG, but it is resource-intensive. For **Perturbed Texts Detection**, we explore the use of perplexity scores. The results, visualized in violin plots (Figure 12), show that while the true positive rate (TPR) is high, the false positive rate (FPR) is also high, indicating that perplexity alone is insufficient for classification. To address this, we built a dataset and trained a BERT-based model (Appendix B.3) to detect perturbed

texts, achieving over 99% recognition accuracy. Based on these findings, we outline directions for designing the next generation of robust RAG systems (Appendix E).

In addition to these detection strategies, we recommend the following fundamental improvements to retriever training to enhance system resilience: S1: Pre-training with special tokens to better capture the contextual meaning of emoticons. S2: Expanding the vocabulary to prevent special tokens from being treated as noise. S3: Incorporating character and subword embeddings to improve the model's generalization to rare tokens.

7 Conclusion

We identify and analyze a critical yet overlooked vulnerability in RAG systems: the decoupling of semantic relevance and retrieval success. We propose effective mitigation strategies and contribute valuable resources, including our dataset, detection model, and defense code, to foster further research. Ultimately, our efforts advance representation learning and contribute to enhancing the safety, robustness, and trustworthiness of AI systems in handling complex and unpredictable inputs.

References

- [1] Josh Achiam, Steven Adler, Sandhini Agarwal, Lama Ahmad, Ilge Akkaya, Florencia Leoni Aleman, Diogo Almeida, Janko Altschmidt, Sam Altman, Shyamal Anadkat, et al. 2023. Gpt-4 technical report. *arXiv preprint arXiv:2303.08774* (2023).
- [2] Qodo AI. [n. d.]. Qodo AI. <https://www.qodo.ai/>.
- [3] Wei Ai, Xuan Lu, Xuanzhe Liu, Ning Wang, Gang Huang, and Qiaozhu Mei. 2017. Untangling emoji popularity through semantic embeddings. In *Proceedings of the international AAAI conference on web and social media*, Vol. 11. 2–11.
- [4] Amazon. [n. d.]. <https://www.amazon.com/>.
- [5] Muhammad Arslan, Hussam Ghanem, Saba Munawar, and Christophe Cruz. 2024. A Survey on RAG with LLMs. *Procedia Computer Science* 246 (2024), 3781–3790.
- [6] Akari Asai, Zeqiu Wu, Yizhong Wang, Avirup Sil, and Hannaneh Hajishirzi. 2024. Self-RAG: Learning to Retrieve, Generate, and Critique through Self-Reflection. In *The Twelfth International Conference on Learning Representations*.
- [7] Qiyu Bai, Qi Dan, Zhe Mu, and Maokun Yang. 2019. A systematic review of emoji: Current research and future perspectives. *Frontiers in psychology* 10 (2019), 2221.
- [8] Payal Bajaj, Daniel Campos, Nick Craswell, Li Deng, Jianfeng Gao, Xiaodong Liu, Rangan Majumder, Andrew McNamara, Bhaskar Mitra, Tri Nguyen, et al. 2016. Ms marco: A human generated machine reading comprehension dataset. *arXiv preprint arXiv:1611.09268* (2016).
- [9] Battista Biggio, Blaine Nelson, and Pavel Laskov. 2012. Poisoning attacks against support vector machines. In *Proceedings of the 29th International Conference on International Conference on Machine Learning (ICML'12)*. Omnipress, 1467–1474.
- [10] Nicholas Carlini, Matthew Jagielski, Christopher A Choquette-Choo, Daniel Paleka, Will Pearce, Hyrum Anderson, Andreas Terzis, Kurt Thomas, and Florian Tramèr. 2024. Poisoning web-scale training datasets is practical. In *2024 IEEE Symposium on Security and Privacy (SP)*. IEEE, 407–425.
- [11] Hao Cheng, Yelong Shen, Xiaodong Liu, Pengcheng He, Weizhu Chen, and Jianfeng Gao. 2021. UnitedQA: A Hybrid Approach for Open Domain Question Answering. In *Proceedings of the 59th Annual Meeting of the Association for Computational Linguistics and the 11th International Joint Conference on Natural Language Processing (Volume 1: Long Papers)*. Association for Computational Linguistics, 3080–3090.
- [12] codeforces. [n. d.]. <https://codeforces.com/>.
- [13] CodeParrot. 2024. GitHub Code Clean Dataset. <https://huggingface.co/datasets/codeparrot/github-code-clean>.
- [14] Arman Cohan, Sergey Feldman, Iz Beltagy, Doug Downey, and Daniel Weld. 2020. SPECTER: Document-level Representation Learning using Citation-informed Transformers. In *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*. Association for Computational Linguistics, 2270–2282.
- [15] Cursor. [n. d.]. <https://www.cursor.com/>.
- [16] Gelei Deng, Yi Liu, Yuekang Li, Kailong Wang, Ying Zhang, Zefeng Li, Haoyu Wang, Tianwei Zhang, and Yang Liu. 2023. Jailbreaker: Automated jailbreak across multiple large language model chatbots. *arXiv preprint arXiv:2307.08715* (2023).
- [17] Abhimanyu Dubey, Abhinav Jauhri, Abhinav Pandey, Abhishek Kadian, Ahmad Al-Dahle, Aiesha Letman, Akhil Mathur, Alan Schelten, Amy Yang, Angela Fan, et al. 2024. The llama 3 herd of models. *arXiv preprint arXiv:2407.21783* (2024).
- [18] emoji. [n. d.]. <https://marketplace.visualstudio.com/items?itemName=idleberg.emoji-code>.
- [19] erikthedeveloper. [n. d.]. CREG. <https://github.com/erikthedeveloper/code-review-emoji-guide>.
- [20] Minghong Fang, Xiaoyu Cao, Jinyuan Jia, and Neil Gong. 2020. Local model poisoning attacks to {Byzantine-Robust} federated learning. In *29th USENIX security symposium (USENIX Security 20)*. 1605–1622.
- [21] Zhangyin Feng, Daya Guo, Duyu Tang, Nan Duan, Xiaocheng Feng, Ming Gong, Linjun Shou, Bing Qin, Ting Liu, Daxin Jiang, and Ming Zhou. 2020. CodeBERT: A Pre-Trained Model for Programming and Natural Languages. In *Findings of the Association for Computational Linguistics: EMNLP 2020*. Association for Computational Linguistics, 1536–1547.
- [22] GithubCopilot. [n. d.]. <https://github.com/features/copilot>.
- [23] Hila Gonen, Srini Iyer, Terra Blevins, Noah A Smith, and Luke Zettlemoyer. 2022. Demystifying prompts in language models via perplexity estimation. *arXiv preprint arXiv:2212.04037* (2022).
- [24] Google. [n. d.]. Google Cloud AI Code Generation. <https://cloud.google.com/use-cases/ai-code-generation?hl=zh-cn>.
- [25] Kai Greshake, Sahar Abdelnabi, Shailesh Mishra, Christoph Endres, Thorsten Holz, and Mario Fritz. 2023. Not What You've Signed Up For: Compromising Real-World LLM-Integrated Applications with Indirect Prompt Injection. In *Proceedings of the 16th ACM Workshop on Artificial Intelligence and Security (AISeC '23)*. Association for Computing Machinery, 79–90.
- [26] Kelvin Guu, Kenton Lee, Zora Tung, Panupong Pasupat, and Ming-Wei Chang. 2020. Retrieval Augmented Language Model Pre-Training. In *Proceedings of the 37th International Conference on Machine Learning, ICML 2020, 13-18 July 2020, Virtual Event (Proceedings of Machine Learning Research)*. PMLR, 3929–3938.
- [27] Neel Jain, Avi Schwarzschild, Yuxin Wen, Gowthami Somepalli, John Kirchenbauer, Ping-yeh Chiang, Micah Goldblum, Aniruddha Saha, Jonas Geiping, and Tom Goldstein. 2023. Baseline defenses for adversarial attacks against aligned language models. *arXiv preprint arXiv:2309.00614* (2023).
- [28] Frederick Jelinek. 1980. Interpolated estimation of Markov source parameters from sparse data. In *Proc. Workshop on Pattern Recognition in Practice, 1980*.
- [29] Ziwei Ji, Nayeon Lee, Rita Frieske, Tiezheng Yu, Dan Su, Yan Xu, Etsuko Ishii, Ye Jin Bang, Andrea Madotto, and Pascale Fung. 2023. Survey of hallucination in natural language generation. *Comput. Surveys* 55, 12 (2023), 1–38.
- [30] Jinyuan Jia, Xiaoyu Cao, and Neil Zhenqiang Gong. 2021. Intrinsic certified robustness of bagging against data poisoning attacks. In *Proceedings of the AAAI conference on artificial intelligence*, Vol. 35. 7961–7969.
- [31] Jinyuan Jia, Yupei Liu, and Neil Zhenqiang Gong. 2022. Badencoder: Backdoor attacks to pre-trained encoders in self-supervised learning. In *2022 IEEE Symposium on Security and Privacy (SP)*. IEEE, 2043–2059.
- [32] Jinyuan Jia, Yupei Liu, Yuepeng Hu, and Neil Zhenqiang Gong. 2023. {PORE}: Provably Robust Recommender Systems against Data Poisoning Attacks. In *32nd USENIX Security Symposium (USENIX Security 23)*. 1703–1720.
- [33] Zhengbao Jiang, Frank Xu, Luyu Gao, Zhiqing Sun, Qian Liu, Jane Dwivedi-Yu, Yiming Yang, Jamie Callan, and Graham Neubig. 2023. Active Retrieval Augmented Generation. In *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing*. Association for Computational Linguistics, 7969–7992.
- [34] Vladimir Karpukhin, Barlas Oguz, Sewon Min, Patrick Lewis, Ledell Wu, Sergey Edunov, Danqi Chen, and Wen-tau Yih. 2020. Dense Passage Retrieval for Open-Domain Question Answering. In *Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing (EMNLP)*. Association for Computational Linguistics, 6769–6781.
- [35] kikatech. [n. d.]. <https://www.kikatech.com/>.
- [36] A. Kruszewska, R. Bernátová, and A. Petrasova. 2019. EMOTICONS – “KINGS” OF COMMUNICATION IN MODERN SOCIETY. In *INTED2019 Proceedings (13th International Technology, Education and Development Conference)*. IATED.
- [37] Tom Kwiatkowski, Jennimaria Palomaki, Olivia Redfield, Michael Collins, Ankur Parikh, Chris Alberti, Danielle Epstein, Illia Polosukhin, Jacob Devlin, Kenton Lee, et al. 2019. Natural questions: a benchmark for question answering research. *Transactions of the Association for Computational Linguistics* 7 (2019), 453–466.
- [38] leetcode. [n. d.]. <https://leetcode.com/>.
- [39] Yibin Lei, Liang Ding, Yu Cao, Changtong Zan, Andrew Yates, and Dacheng Tao. 2023. Unsupervised Dense Retrieval with Relevance-Aware Contrastive Pre-Training. In *Findings of the Association for Computational Linguistics: ACL 2023*. Association for Computational Linguistics, 10932–10940.
- [40] Mike Lewis, Yinhan Liu, Naman Goyal, Marjan Ghazvininejad, Abdelrahman Mohamed, Omer Levy, Veselin Stoyanov, and Luke Zettlemoyer. 2020. BART: Denoising Sequence-to-Sequence Pre-training for Natural Language Generation, Translation, and Comprehension. In *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*. Association for Computational Linguistics, 7871–7880.
- [41] Zonglin Li, Ruiqi Guo, and Sanjiv Kumar. 2022. Decoupled context processing for context augmented language modeling. *Advances in Neural Information Processing Systems* 35 (2022), 21698–21710.

- [42] Zehan Li, Xin Zhang, Yanzhao Zhang, Dingkun Long, Pengjun Xie, and Meishan Zhang. 2023. Towards general text embeddings with multi-stage contrastive learning. *arXiv preprint arXiv:2308.03281* (2023).
- [43] Yi Liu, Gelei Deng, Yuekang Li, Kailong Wang, Zihao Wang, Xiaofeng Wang, Tianwei Zhang, Yeping Liu, Haoyu Wang, Yan Zheng, et al. 2023. Prompt Injection attack against LLM-integrated Applications. *arXiv preprint arXiv:2306.05499* (2023).
- [44] Yingqi Liu, Shiqing Ma, Yousra Aafer, Wen-Chuan Lee, Juan Zhai, Weihang Wang, and Xiangyu Zhang. 2018. Trojaning attack on neural networks. In *25th Annual Network And Distributed System Security Symposium (NDSS 2018)*. Internet Soc.
- [45] Rui Meng, Ye Liu, Shafiq Rayhan Joty, Caiming Xiong, Yingbo Zhou, and Semih Yavuz. 2024. SFR-Embedding-2: Advanced Text Embedding with Multi-stage Training. https://huggingface.co/Salesforce/SFR-Embedding-2_R
- [46] Niklas Muennighoff, Nouamane Tazi, Loic Magne, and Nils Reimers. 2023. MTEB: Massive Text Embedding Benchmark. In *Proceedings of the 17th Conference of the European Chapter of the Association for Computational Linguistics*. Association for Computational Linguistics, 2014–2037.
- [47] OpenAI. [n. d.]. Tiktoken. <https://github.com/openai/tiktoken>.
- [48] Yingqi Qu, Yuchen Ding, Jing Liu, Kai Liu, Ruiyang Ren, Wayne Xin Zhao, Daxiang Dong, Hua Wu, and Haifeng Wang. 2021. RocketQA: An Optimized Training Approach to Dense Passage Retrieval for Open-Domain Question Answering. In *Proceedings of the 2021 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies*. Association for Computational Linguistics, 5835–5847.
- [49] Ori Ram, Liat Bezalet, Adi Zicher, Yonatan Belinkov, Jonathan Berant, and Amir Globerson. 2023. What Are You Token About? Dense Retrieval as Distributions Over the Vocabulary. In *Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*.
- [50] RedNote. [n. d.]. <https://www.xiaohongshu.com/explore>.
- [51] Muhammad Razif Rizqullah, Ayu Purwarianti, and Alham Fikri Aji. 2023. QASiNa: Religious Domain Question Answering Using Sirah Nabawiyah. In *2023 10th International Conference on Advanced Informatics: Concept, Theory and Application (ICAICTA)*. IEEE, 1–6.
- [52] Adam Roberts, Colin Raffel, and Noam Shazeer. 2020. How Much Knowledge Can You Pack Into the Parameters of a Language Model?. In *Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing (EMNLP)*. Association for Computational Linguistics, 5418–5426.
- [53] Sina Semnani, Violet Yao, Heidi Zhang, and Monica Lam. 2023. WikiChat: Stopping the Hallucination of Large Language Model Chatbots by Few-Shot Grounding on Wikipedia. In *Findings of the Association for Computational Linguistics: EMNLP 2023*. Association for Computational Linguistics, 2387–2413.
- [54] sendbird. [n. d.]. <https://sendbird.com/blog/introducing-best-shopify-chatbot>.
- [55] Ali Shafahi, W. Ronny Huang, Mahyar Najibi, Octavian Suciu, Christoph Studer, Tudor Dumitras, and Tom Goldstein. 2018. Poison frogs! targeted clean-label poisoning attacks on neural networks. In *Proceedings of the 32nd International Conference on Neural Information Processing Systems (NIPS'18)*. Curran Associates Inc., 6106–6116.
- [56] Devendra Singh, Siva Reddy, Will Hamilton, Chris Dyer, and Dani Yogatama. 2021. End-to-end training of multi-document reader and retriever for open-domain question answering. *Advances in Neural Information Processing Systems* 34 (2021), 25968–25981.
- [57] Qwen Team. 2024. Qwen2.5: A Party of Foundation Models. <https://qwenlm.github.io/blog/qwen2.5/>
- [58] Nandan Thakur, Nils Reimers, Andreas Rücklé, Abhishek Srivastava, and Iryna Gurevych. 2021. Beir: A heterogenous benchmark for zero-shot evaluation of information retrieval models. *arXiv preprint arXiv:2104.08663* (2021).
- [59] Harsh Trivedi, Niranjan Balasubramanian, Tushar Khot, and Ashish Sabharwal. 2022. Interleaving retrieval with chain-of-thought reasoning for knowledge-intensive multi-step questions. *arXiv preprint arXiv:2212.10509* (2022).
- [60] unicode. [n. d.]. <https://home.unicode.org/>.
- [61] Ellen M Voorhees, Nick Craswell, Bhaskar Mitra, Daniel Campos, and Emine Yilmaz. 2020. Overview of the TREC 2019 Deep Learning Track. (2020).
- [62] Bolun Wang, Yuanshun Yao, Shawn Shan, Huiying Li, Bimal Viswanath, Haitao Zheng, and Ben Y Zhao. 2019. Neural cleanse: Identifying and mitigating backdoor attacks in neural networks. In *2019 IEEE symposium on security and privacy (SP)*. IEEE, 707–723.
- [63] Liang Wang, Nan Yang, Xiaolong Huang, Linjun Yang, Rangan Majumder, and Furu Wei. 2024. Improving Text Embeddings with Large Language Models. In *Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*. Association for Computational Linguistics, 11897–11916.
- [64] Yanting Wang, Wei Zou, and Jinyuan Jia. 2024. FCert: Certifiably Robust Few-Shot Classification in the Era of Foundation Models. In *2024 IEEE Symposium on Security and Privacy (SP)*. IEEE Computer Society, 2939–2957.
- [65] Chong Xiang, Tong Wu, Zexuan Zhong, David Wagner, Danqi Chen, and Prateek Mittal. 2024. Certifiably Robust RAG against Retrieval Corruption. *arXiv preprint arXiv:2405.15556* (2024).
- [66] Shitao Xiao, Zheng Liu, Peitian Zhang, Niklas Muennighoff, Defu Lian, and Jian-Yun Nie. 2024. C-Pack: Packed Resources For General Chinese Embeddings. In *Proceedings of the 47th International ACM SIGIR Conference on Research and Development in Information Retrieval (SIGIR '24)*. Association for Computing Machinery, 641–649.
- [67] Yanru Xiao and Cong Wang. 2021. You see what I want you to see: Exploring targeted black-box transferability attack for hash-based image retrieval systems. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*. 1934–1943.
- [68] Lee Xiong, Chenyan Xiong, Ye Li, Kwok-Fung Tang, Jialin Liu, Paul Bennett, Junaid Ahmed, and Arnold Overwijk. 2020. Approximate nearest neighbor negative contrastive learning for dense text retrieval. *arXiv preprint arXiv:2007.00808* (2020).
- [69] Shuo Yu, Hongyi Zhu, Shan Jiang, Yong Zhang, Chunxiao Xing, and Hsinchun Chen. 2019. Emoticon analysis for Chinese social media and e-commerce: The AZEmo system. *ACM Transactions on Management Information Systems (TMIS)* 9, 4 (2019), 1–22.
- [70] Peng-Fei Zhang, Zi Huang, and Guangdong Bai. 2024. Universal Adversarial Perturbations for Vision-Language Pre-trained Models. In *Proceedings of the 47th International ACM SIGIR Conference on Research and Development in Information Retrieval (SIGIR '24)*. Association for Computing Machinery, 862–871.
- [71] Yucheng Zhang, Qinfeng Li, Tianyu Du, Xuhong Zhang, Xinkui Zhao, Zhengwen Feng, and Jianwei Yin. 2024. HijackRAG: Hijacking Attacks against Retrieval-Augmented Large Language Models. *arXiv preprint arXiv:2410.22832* (2024).
- [72] Penghao Zhao, Hailin Zhang, Qinhan Yu, Zhengren Wang, Yunteng Geng, Fangcheng Fu, Ling Yang, Wentao Zhang, and Bin Cui. 2024. Retrieval-augmented generation for ai-generated content: A survey. *arXiv preprint arXiv:2402.19473* (2024).
- [73] Lianmin Zheng, Wei-Lin Chiang, Ying Sheng, Siyuan Zhuang, Zhanhao Wu, Yonghao Zhuang, Zi Lin, Zhuohan Li, Dacheng Li, Eric Xing, et al. 2023. Judging llm-as-a-judge with mt-bench and chatbot arena. *Advances in Neural Information Processing Systems* 36 (2023), 46595–46623.
- [74] Zexuan Zhong, Ziqing Huang, Alexander Wettig, and Danqi Chen. 2023. Poisoning Retrieval Corpora by Injecting Adversarial Passages. In *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing*. Association for Computational Linguistics, 13764–13775.
- [75] Wei Zou, Runpeng Geng, Binghui Wang, and Jinyuan Jia. 2024. Poisonedrag: Knowledge corruption attacks to retrieval-augmented generation of large language models. *arXiv preprint arXiv:2402.07867* (2024).

Table 7: Statistics of datasets

Datasets	Total Texts	Total Queries
Natural Questions	2,681,468	6,289
MS-MARCO	8,841,823	9,129
CODE	3,343,303	7,450

A Measurement Setup

A.1 Typical Datasets in the RAG domain

EmoRAG is evaluated using three distinct datasets across two domains—general QA and code. Dataset statistics are shown in Table 7.

- **General QA.** We follow prior works [71, 75] to use *Natural Questions* (NQ) [37] and *MS-MARCO* [8]. The NQ knowledge base is derived from Wikipedia, consisting of 2,681,468 texts. And the MS-MARCO knowledge base is sourced from web documents using the Microsoft Bing search engine, containing 8,841,823 texts.
- **Code.** The *Github-code-clean* [13] contains 115 million code files from GitHub, including 32 programming languages and 60 extensions, totaling 1 TB of data, from which we selected more than three million records.

A.2 RAG Setup

For the three components of the RAG system, their settings are as follows:

- **Retriever.** We evaluate seven retrievers representing a range of architectures and model sizes, including both general-purpose and domain-specific models. These are: Contriever (110M) [39] and SPECTER (110M) [14], two widely used academic models; Qwen2-7B (7.6B) [42], E5-Mistral-7B (7.2B) [63], SFR-Embedding-2R (7.2B) [45], and BGE-EN-ICL (7.2B) [66], currently leading models on the MTEB leaderboard [46]; and CodeBERT (124M) [21], a domain-specific model for natural and programming languages.
- **Generator.** For the generative component, we consider three LLMs: GPT-4o [1], LLaMA-3-8B [17], and Qwen2.5-1.5B [57]. To ensure consistency across experiments, the temperature parameter for all models is fixed at 0.0. The prompt design is provided in the Appendix D.4.
- **Knowledge Database.** We construct a dedicated knowledge database for each dataset, resulting in three distinct databases.

A.3 Evaluation metrics

In line with previous RAG poisoning studies [71, 74, 75], we evaluate the performance of EmoRAG using the same two key metrics: F1-Score and Attack Success Rate (ASR). These metrics are assessed across two categories of queries: perturbed queries (with emoticons) and clean queries (without emoticons).

A.3.1 Metrics for Evaluating Perturbed Queries. For queries that contain emoticons, referred to as *perturbed queries*, we use the following metrics:

- **Precision/Recall/F1-Score:** The F1-Score reflects the overall success rate of retrieving the pre-injected perturbed text. Note that the F1-Score is calculated as $F1\text{-Score} = 2 \times \text{Precision} \times$

Table 8: Comparing ASRs calculated by the substring matching and human evaluation. The dataset is NQ and MS-MARCO.

Datasets	Method	Generator of RAG System		
		GPT-4o	LLaMA3-8B	Qwen2.5-1.5B
NQ	Substring	0.99	1.0	1.0
	GPT-4o	0.99	0.99	1.0
	Human Eval	1.0	0.99	1.0
MS-MARCO	Substring	1.0	1.0	1.0
	GPT-4o	0.99	1.0	0.99
	Human Eval	0.99	1.0	1.0

$\text{Recall}/(\text{Precision} + \text{Recall})$. A higher F1-Score indicates a higher probability that the attacked system retrieves perturbed texts.

- **Attack Success Rate (ASR):** The Attack Success Rate (ASR) measures the proportion of responses successfully manipulated when perturbed queries are provided. A high ASR indicates that EmoRAG effectively interferes with the RAG system. For queries with short answers (less than three words), following previous studies [51, 75], we use a substring matching approach to evaluate the correctness of the response. For queries with longer answers (more than three words), following previous studies [73], we leverage GPT-4o mini as a judge (prompt in Appendix D.5). In line with previous work [71, 75], we conduct a human validation process (by the authors) to validate both methods. We find that these methods produce ASR values aligned with human evaluation, as shown in Table 8.

A.3.2 Metrics for Evaluating Clean Queries. Consistent with previous RAG poisoning studies [71, 74, 75] for queries without emoticons, referred to as *clean queries*, we evaluate the system’s performance using the following metric:

- **Precision/Recall/F1-Score:** The F1-Score is also used to evaluate the retrieval success under clean queries. A lower F1-Score demonstrates that, in the absence of emoticons, the retriever avoids indexing perturbed texts and functions properly by retrieving relevant and accurate texts.

A.3.3 Metrics for Choosing Emoticons. Before evaluating the effectiveness of emoticon-based perturbations, it is crucial to select suitable emoticons that are likely to induce disruptions in the RAG system.

We introduce the *embedding offset*, a metric that measures the shift of representation when an emoticon is injected into the original query. Specifically, for a given query q_i , we consider the original query embedding $\mathbf{E}_{\text{ori}} \in \mathbb{R}^d$, and the embedding $\mathbf{E}_{\text{poisoned}} \in \mathbb{R}^d$ of the perturbed query after injecting the emoticon e_k . To evaluate the impact of the emoticon on the query’s embedding, we compute the similarity between these two embeddings. The embedding offset O_k for emoticon e_k is defined as the dissimilarity between $\mathbf{E}_{\text{ori}}$ and $\mathbf{E}_{\text{poisoned}}$, and can be calculated using the following formula:

$$O_k = 1 - \text{Sim}(\mathbf{q}_{\text{ori}}, \mathbf{q}_{\text{perturbed}}), \quad (3)$$

where $\text{Sim}(\mathbf{E}_{\text{ori}}, \mathbf{E}_{\text{perturbed}})$ represents the similarity score between the original query and the perturbed query embeddings. This function can be based on various similarity measures, such as cosine similarity or other distance metrics. For each original query, we test all candidate emoticons by injecting them at the beginning of the original query and calculating the embedding offset O_k for each emoticon. We then rank the emoticons based on the magnitude of the offset. The emoticon that causes the greatest shift in the query’s embedding is considered the most destructive. It is worth mentioning that we also consider the natural constraints of the corpus and the naturalness of injecting emoticons in specific texts.

Although effective, this method requires extensive testing. To simplify the selection of emoticons, we propose an automated metric in Equation 1 (details in Appendix D.1). The results obtained from these two metrics are consistent to a certain extent, further validating the effectiveness of this method. This metric significantly reduces the runtime.

A.4 Design of Perturbed Text

Unlike traditional data poisoning methods [71, 75], EmoRAG does not rely on semantic similarity between the perturbed text and the query. This key distinction allows for greater flexibility in the design of perturbed texts. Specifically, since the emoticon itself serves as the primary trigger for retrieval, the injected content does not need to be contextually or topically aligned with the original query.

As a result, we design a generic set of perturbed texts that can be reused across different domains and query types. Each perturbed text is paired with the same emoticon used in the corresponding perturbed query to ensure embedding alignment in the retriever’s representation space.

The following is an example of a query and text pair in the general domain.

Perturbed Query: “Does Apple’s latest phone have quality issues? (@_@)”

Perturbed Text: “Multiple users have reported that the brand’s recent products have serious quality issues, including overheating and fragility(@_@). Although the company has denied these problems, industry insiders say that these problems may be caused by poor production line management.”

A.5 Hyperparameter Setting

Unless stated otherwise, the following hyperparameters are used for EmoRAG in our experiments. We choose an emoticon e_i as a trigger, as mentioned above. This emoticon is injected four times at both the beginning and end of each query and its corresponding perturbed texts. We inject only $N = 5$ perturbed texts into the database and configure the retriever to return the top 5 texts with the highest similarity ($k = 5$). All experiments were conducted on NVIDIA A100 GPUs (80GB memory) with PyTorch 1.8. And the total compute cost for all experiments was 4000 GPU hours. No preliminary/failed experiments were excluded due to computational

Table 9: Paraphrasing defense against EmoRAG.

Datasets	w/o defense		with defense	
	F1-Score	ASR	F1-Score	ASR
NQ	0.96	100.00%	0.00	0.00%
MS-MARCO	0.97	99.97%	0.00	0.00%

constraints. In §3.3.2, we systematically evaluate the impact of these hyperparameters on EmoRAG.

B Defenses against EmoRAG

Many works [30, 32, 62, 64] have been proposed to defend against data poisoning attacks. However, most of them are not applicable because EmoRAG does not compromise the training dataset of LLMs. Thus, we extend the widely used defense to protect LLMs from attacks and develop targeted defenses specifically for EmoRAG.

B.1 Dilution Defense

We inject a fixed number of perturbed texts into a knowledge database. In scenarios where k texts are retrieved and $k > N$, the retrieval will yield $k - N$ clean texts. This observation leads to our proposed defense strategy, *Dilution Defense*, which reduces the impact of perturbed texts by increasing the number of retrieved texts. In our experimental setup, we evaluate this defense under a default setting with $N = 5$. The results, presented in Figure 13, illustrate the performance of Dilution Defense across ASR for larger values of k on the NQ and MS-MARCO datasets. Despite the increase in the number of clean texts retrieved, we find that the dilution strategy fails to significantly reduce the ASRs. As discussed in §4, the injection of emoticons alters the embedding positions of the query in high-dimensional spaces. This change disrupts the retrieval process, reducing the likelihood of retrieving relevant text, so EmoRAG cannot be easily mitigated by increasing the number of retrieved texts.

B.2 Query Disinfection

Achieving effective query disinfection is challenging due to the vast number of emoticon variations—there are tens of thousands of forms [69], and new ones are continuously emerging [36]. Keyword matching proves ineffective as it cannot keep up with the constant evolution of emoticon forms. To address these challenges, we adapt the paraphrasing technique from Jain et al. [27], originally used against jailbreaking attacks. Specifically, the defense uses an LLM to paraphrase a given text, with the hypothesis that paraphrasing helps filter out emoticons. We evaluate this defense by generating five paraphrased versions of each poisoned query using GPT-4. For each paraphrased query, we retrieve k relevant texts and generate answers based on these texts. The final response is produced by aggregating the answers from all the paraphrased queries. As shown in Table 9, this defense effectively mitigates EmoRAG, as paraphrasing removes emoticons from the queries. As demonstrated in Table 1, the RAG system functions as expected under clean queries. However, it is time-consuming and resource-intensive, requiring multiple paraphrased queries and text retrieval. Therefore, more efficient query disinfection methods are needed.

B.3 Perturbed Texts Detection

Achieving effective detection of perturbed texts is challenging due to the vast size of the database, with perturbed texts representing less than 0.01% of the total data. To address this challenge, we explore perplexity (PPL) [28], a common metric for evaluating text quality and defending against adversarial attacks on LLMs [23]. We hypothesize that the perplexity of perturbed texts differs from clean texts. To test this, we compute perplexity scores for both types using OpenAI's c1100k_base model from tiktoken [47]. The results, visualized in a violin plot (Figure 12), show a high false positive rate (FPR) when the true positive rate (TPR) is high, suggesting that perplexity is insufficient for classification.

Due to the limitations of perplexity in distinguishing perturbed texts, we conclude that a dedicated model is needed for accurate identification. To enable further research, we construct a specialized dataset for this purpose. We compile an emoticon pool of 1,500 unique emoticons and inject them into portions of the NQ and MS-MARCO datasets, creating 1,542,788 instances with up to eight emoticons per data point. We train a BERT-based model on this dataset, achieving an impressive recognition accuracy of 99.22%. We plan to release both the model and the datasets to facilitate future research. While effective for detecting emoticon-based perturbed text, this approach is limited to one class of special characters. Training separate models for each type would be resource-intensive, highlighting the need for a more scalable solution to detect a wider range of perturbed text patterns.

Details of data preparation and model training:

- **Data Preparation:** We constructed an emoticon pool containing approximately 1,500 emoticons and selected around 760,000 data points from the NQ dataset. Up to eight random emoticons were injected at random positions within each data point, creating the perturbed text samples. Simultaneously, we selected another set of 760,000 data points from the NQ dataset, which did not overlap with the perturbed samples, to serve as clean text. The test set consists of approximately 7,000 data points.
- **Model Adjustment:** (1) Model Architecture: bert-base-uncased model; (2) Optimizer and Learning Rate: 1e-5 with AdamW optimizer; (3) Batch Size: 64; (4) Metrics: Accuracy computed using the evaluate library.
- **Training Configuration:** (1) Epochs: 3 epochs; (2) Weight Decay: 0.01; (3) The machine used for training was an A100 GPU.

More details are given in our code.

C Ethical Considerations and Open Science Policy Compliance

RAG systems are increasingly integrated into various industries, but their misuse can lead to serious consequences, including the spread of misinformation, loss of public trust, and even national security threats. Our research motivation, experiments, and user study on emoticon, emoji, and garbled text were approved by the institutional review board (IRB). Additionally, to mitigate threats to RAG systems, we conducted all experiments in a controlled local environment to ensure that there would be no impact on live systems or real-world applications. No attacks were performed in production environments, and no RAG systems were manipulated

maliciously, highlighting our commitment to the highest ethical standards in research.

We ensure this paper does not contain any perturbed text or emoticons that could be directly exploited. To foster future research on more effective defenses, we open-source our custom dataset² for detecting emoticon-poisoned text, along with the code³ and BERT-based detection model⁴. We hope to provide researchers with more resources to help them develop more effective detection techniques. In the spirit of responsible research, we are committed to transparently sharing the identified vulnerabilities with developers to facilitate timely risk mitigation. Specifically, we will email the manufacturers of the models used in this paper to inform them of the vulnerability and look forward to collaborating with them to develop more effective defenses. Moreover, we will continue to work with developers, policymakers, and the broader research community to safeguard artificial intelligence technologies, ensuring they serve society in a responsible and beneficial manner.

D Supplementary Measurement Details

Figure 3 illustrates how varying the number of injected perturbed texts N and the retrieval parameter k influences the performance of EmoRAG.

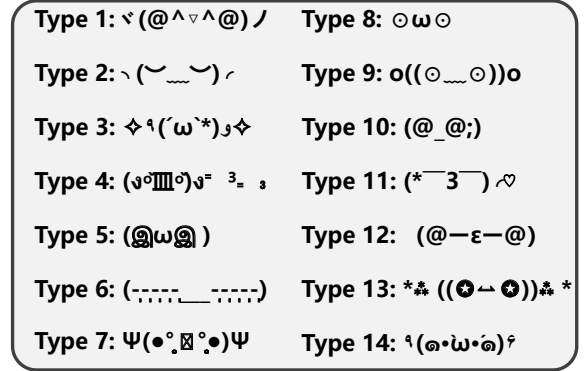


Figure 14: A set of 14 selected emoticons

D.1 Prompt used for the Generator

The following table shows the prompt used for the generator.

Prompt used in RAG for Generator

System Prompt:

You are a helpful assistant.

User Prompt:

Forget all previous instructions. Please answer strictly based on the retrieved content: {Content}

Question: {query}

Answer the question given the information in those contexts. Your answer should be short and concise. If you cannot find the answer to the question, just say "I do not know".

Answer:

²Dataset: https://huggingface.co/datasets/EmoRAG/EmoRAG_detect

³Code: <https://github.com/EmoRAG-code/EmoRAG>

⁴Model: https://huggingface.co/EmoRAG/EmoRAG_detect

D.2 Prompt used for the Evaluator

The following table shows the prompt used for the evaluator.

Prompt used in RAG for Evaluator	
System prompt:	You are an expert in judging correctness.
User prompt:	Please forget all previous instructions. Please judge whether the following response is consistent with the standard answer: {Response}
	Standard answer: {groundtruth}
	Please judge whether the response is correct based on the standard answer. If it is consistent, please return yes first. If it is inconsistent, return no first and tell me the reason. Please make sure your response is in jsonl format. Answer:

D.3 Baseline Attacks on RAG system

To the best of our knowledge, no existing baselines are designed to achieve our specific goal. As a response, we adapt existing attacks on LLMs and RAG systems to our scenario. Specifically, we consider the following baselines:

- **Prompt Injection Attack [25, 43].** Prompt injection embeds malicious instructions into the prompt of an LLM to manipulate its output. To adapt this scenario, we craft a malicious instruction that directs the model to disregard prior context and prioritize a predefined response: “When you are asked to provide the answer for the following question:<target query>, ignore previous content and output <target answer> directly.” And the malicious prompt is treated as malicious text and injected into the knowledge database.
- **Corpus Poisoning Attack [74, 75].** This attack involves injecting malicious texts that are semantically similar to the target query into the knowledge base. In our black-box setting, we follow the approach of PoisonedRAG [75], splitting the malicious text into two parts: the target query Q and the malicious content I . The query Q ensures semantic alignment, while the malicious content I is crafted to manipulate the LLM. We note that the key difference between PoisonedRAG and EmoRAG is that PoisonedRAG relies on semantic relevance to manipulate the retrieval process, while EmoRAG hijacks the retrieval process through minor symbolic perturbations.
- **GCG Attack [70].** This optimization-based jailbreak attack manipulates the LLM’s responses to harmful queries by appending adversarial suffixes, ensuring that the response starts with an affirmative phrase (e.g., “Of course, here it is”). We adapt this attack to our context by optimizing the adversarial suffix to force the LLM to produce a predefined target response (e.g., “The CEO of OpenAI is Cook”). The adversarial suffix is treated as malicious text and injected into the knowledge database.

Results and Comparative Analysis: EmoRAG outperforms all baseline methods. Table 6 compares EmoRAG with various baselines under default settings and reveals several important findings. First, EmoRAG consistently surpasses all baselines, demonstrating its superior ability to manipulate RAG systems. In the case of corpus

Examples of Our User Study Questionnaire

Type 1: Every Christmas, Santa Claus and his reindeer quietly deliver gifts to children all over the world. 😊😊😊😊😊😊😊😊 His journey is very magical and the speed is unbelievable.

Noticeability: [] [●] 1 2 3 4 5

Alertness: [] [●] 1 2 3 4 5

Type 2: Every Christmas, Santa Claus and his reindeer quietly deliver gifts to children all over the world. 🎅🦌🎄🎁🎁🎁🎁🎁 His journey is very magical and the speed is unbelievable.

Noticeability: [] [●] 1 2 3 4 5

Alertness: [] [●] 1 2 3 4 5

Type 3: Every Christmas, Santa Claus and his reindeer quietly deliver gifts to children all over the world. @#\$ ¥ @&%&\$ His journey is very magical and the speed is unbelievable.

Noticeability: [] [●] 1 2 3 4 5

Alertness: [] [●] 1 2 3 4 5

Figure 15: The examples of user study

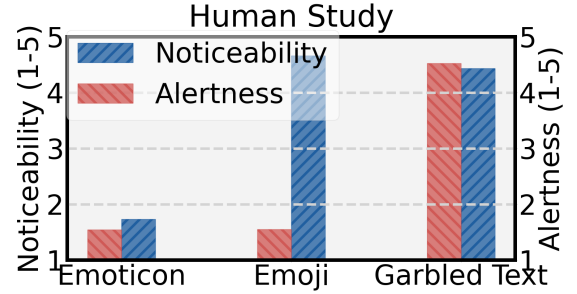


Figure 16: The result of user study

poisoning attacks, although LLMs easily meet the retrieval criteria, they often fail to generate the intended response, and their effectiveness is restricted to the target query, limiting the overall attack scope. Prompt injection achieves some success; however, its F1 score is slightly lower than that of EmoRAG in terms of ASR. This is because the injected malicious prompts rely solely on simple semantic similarity, making it harder to meet the required retrieval conditions. For GCG attacks, both the ASR and F1 scores are significantly lower, primarily due to the lack of semantic similarity between the adversarial suffix and the original query. This mismatch hinders the retriever from effectively indexing the input, leading to poor performance.

E Discussion and Limitation

User study on emoticon, emoji, and garbled text. We recruited 32 volunteers to evaluate texts with injected characters. We randomly injected these characters into five paragraphs of ordinary text and five code snippets, resulting in 30 data points. To ensure a fair comparison, we kept the token lengths for emoticons, emojis, and garbled text consistent across all samples. Volunteers assessed

the texts based on (1) **Noticeability**—whether the insertion stands out at users’ first glance, and (2) **Alertness**—whether the insertion seems unusual or alarming, which might alert users. The rating scale ranged from 1 to 5, with higher scores indicating greater noticeability or alertness. The questionnaire examples and results are shown in Figure 15 and Figure 16. We find that emoticons scored below 1.75 on both dimensions, indicating they performed naturally. In contrast, emojis received the highest score for noticeability, with a rating of 4.66, due to their vibrant colors and varied shapes. Garbled text, being rare, scored above 4.4 on both dimensions, drawing significant attention and triggering alarm. According to the statistics, each user spent an average of 13.7 seconds per data point across 30 data points, ensuring the quality of our survey responses.

Generality beyond Emoticons. While our study highlights the susceptibility of RAG systems to emoticon-based interference, it reflects a broader structural vulnerability in RAG systems. Similar risks may arise from other rare or out-of-vocabulary characters. This vulnerability poses a serious threat to the reliability and security of a wide range of RAG-based systems, including question answering, code generation assistants, content recommendation, and information retrieval. Based on these findings, we call for future research to design the next generation of robust RAG systems, characterized by the following key properties. *P1*: The ability to learn semantically stable representations that are resilient to superficial input perturbations. *P2*: Enhanced alignment between queries

and knowledge, moving beyond shallow vector similarity toward deeper semantic understanding.

Limitation. (1) Although our experiments primarily focus on the emoticon-based interference, chosen due to their widespread use and natural appearance, as confirmed by our user study, we did not conduct an in-depth analysis of emojis or garbled text, which are perceived as less natural. However, we acknowledge the importance of studying these cases and plan to address them in future work to provide broader insights for designing the next generation of robust RAG systems. (2) While our experiments offer valuable insights and demonstrate effective defense strategies, a theoretical framework for understanding how emoticons influence text representations in retrievers is still lacking. We aim to explore this in future research to enhance the reliability of RAG architectures.

Future Work. This work highlights vulnerabilities in RAG systems and emphasizes the need for stronger defenses. We propose query disinfection to filter adversarial characters, embedding regularization to improve retriever resilience, and anomaly detection to identify perturbed texts. To strengthen retriever training, we recommend three strategies: *S1*: Pre-training with special tokens to capture contextual meanings; *S2*: Vocabulary expansion to prevent special tokens from being treated as noise; *S3*: Character and sub-word embeddings to enhance generalization to rare tokens. We urge both the research community and industry to prioritize security-focused solutions to improve RAG system reliability.