
Beyond Per-Question Privacy: Multi-Query Differential Privacy for RAG Systems

Anonymous Author(s)

Affiliation

Address

email

Abstract

1 Retrieval-augmented generation (RAG) enhances large language models (LLMs)
2 by retrieving documents from an external dataset at inference time. When the
3 external dataset contains sensitive and private information, prior work shows that
4 without any protection, the RAG system has the risk to leak this information, which
5 might hurt the data owner’s privacy. The existing work has studied protect the
6 information not leaked from a single-query from RAG under differential privacy
7 (DP). In this paper, we focus on the more practical setting where the DP is extended
8 to multiple queries from RAG. We propose two new algorithms that ensure DP
9 in multi-query RAG. Our first method, MURAG, applies an individual privacy
10 accounting framework, allowing the privacy cost to depend on how often each
11 document is retrieved rather than the total number of queries given. Our second
12 method, MURAG-ADA, further improves efficiency in the individual privacy
13 accounting framework by adaptively releasing private query-specific thresholds
14 for more precisely relevant document selection. Experiments across four question
15 datasets and three LLMs show that our methods answer 100 queries under $\epsilon = 10$,
16 while baseline methods require $\epsilon = 1000$ for comparable utility. We also highlight
17 scenarios where our approach outperforms fixed-threshold baselines and discuss
18 when individual accounting is preferable to subsampling-based techniques.

19 1 Introduction

20 Retrieval-augmented generation (RAG) has become a popular approach for deploying large language
21 models (LLMs) in real-world applications. A core feature of RAG is its reliance on an external
22 dataset as the primary knowledge source at inference time. For example, a medical RAG system
23 may retrieve historical patient records to answer clinical questions more accurately. However, such
24 external datasets often contain sensitive or confidential information. In domains like healthcare or
25 law, the retrieved content may expose private records, raising serious privacy concerns. Prior work
26 has shown that RAG systems without proper safeguards are vulnerable to information leakage [Naseh
27 et al., 2025, Liu et al., 2025, Anderson et al., 2024, Li et al., 2025, Zhang et al., 2025, Zeng et al.,
28 2024a, Jiang et al., 2024, Peng et al., 2024], compromising data owner privacy and user trust.

29 Differential privacy (DP) is a widely adopted framework for providing rigorous guarantees on
30 individual data protection. Recent work [Koga et al., 2024] has proposed DPSparseVoteRAG, a RAG
31 system that ensures the generated answer satisfies DP with respect to the external dataset, for a *single*
32 *user query*. Empirical results demonstrate that this approach outperforms a strong baseline using a
33 public LLM without RAG, while achieving an ϵ -pure DP guarantee with $\epsilon \approx 10$.

34 In realistic deployments, users may issue a sequence of queries, allowing an adversary to aggregate
35 responses and potentially infer sensitive information about the external dataset. A naïve approach
36 that applies DPSparseVoteRAG to each query and relies on standard differential privacy composition

quickly exhausts the privacy budget. As shown by our experimental results (Figure 2), to achieve reasonable utility this approach achieving may require a privacy budget as large as $\epsilon = 1000$, which is generally considered meaningless. This raises a key question:

Can we design a differentially private RAG algorithm that handles hundreds of online queries while ensuring both strong utility and meaningful privacy?

We answer this question affirmatively and summarize our contributions below.

- **Novel DP Multi-RAG Framework.** We present a novel framework for multi-query differentially private RAG. Our contributions include: (1) bypassing composition across queries using individual Rényi filters [Feldman and Zrnic, 2021], which significantly reduces the ex-ante privacy budget. To the best of our knowledge, this is the first use of privacy filters in the RAG literature; (2) enhancing both privacy and utility through threshold-based screening of relevant documents. Notably, the framework also adopts a modular design, allowing integration of any private single-query RAG algorithm as a subroutine. Additionally, the framework adopts a modular design, allowing the integration of any private single-query RAG algorithm as a subroutine.
- **Two DP Multi-RAG Algorithms for Varying Test Query Dependencies.** We propose two differentially private RAG algorithms for the multi-query setting, tailored to the degree of relevance among test-time queries. MURAG (Algorithm 1) uses a fixed relevance threshold across all queries and is sufficient to work well for settings where queries are independent and do not share relevant private documents. MURAG-ADA (Algorithm 2) allocates a small portion of the privacy budget to release a query-specific relevance threshold, enabling more efficient use of the budget when queries are related and share overlapping relevant documents.
- **Practical Multi-Query RAG with Non-Trivial Privacy Guarantees.** We evaluate our algorithms through extensive experiments on four question datasets (Natural Questions, Trivia Questions, ChatDoctor Questions, and MQuAKE Questions) and three LLMs (OPT-1.3B, Pythia-1.4B, and Mistral-7B). The external datasets include Wikipedia (commonly used in standard RAG setups) and ChatDoctor with QA pairs between patients and doctors, reflecting privacy-sensitive applications. Empirical results show that both methods can answer hundreds of queries under a total privacy budget of $\epsilon \approx 10$ while achieving a reasonable utility, reducing privacy budget consumption by up to $100\times$ compared to the baseline DP RAG method that achieves the comparable utility. To the best of our knowledge, these are the first DP algorithms to achieve both non-trivial privacy guarantees and practical utility in the multi-query RAG setting.

2 Preliminaries

Notation. Let $\mathcal{V}$ denote a finite vocabulary set, and let $x \in \mathcal{V}^*$ represent a prompt of arbitrary length. The document set with arbitrary size is denoted by $D = \{z_1, z_2, \dots\}$, where each document $z_i \in \mathcal{V}^*$. For convenience, we define the document space as $\mathcal{Z}$.

Differential Privacy. Given dataspace $\mathcal{X}$, two datasets $D, D' \in \mathcal{X}^*$ are said to be neighboring if they differ by at most one element. In this work, we focus on document-level privacy, where the data universe is given by $\mathcal{V}^*$.

Definition 1 (Differential Privacy [Dwork et al., 2006b]). *A randomized algorithm $\mathcal{M} : \mathcal{X}^* \rightarrow \Omega$ is said to satisfy (ϵ, δ) -differential privacy if for all neighboring datasets $X, X' \in \mathcal{X}^*$ and for all measurable subsets $O \subseteq \Omega$, we have*

$$\mathbb{P}(\mathcal{M}(X) \in O) \leq e^\epsilon \mathbb{P}(\mathcal{M}(X') \in O) + \delta.$$

Definition 2 (Rényi Differential Privacy [Mironov, 2017]). *A randomized algorithm $\mathcal{M} : \mathcal{X}^* \rightarrow \Omega$ is said to satisfy (α, ϵ) -Rényi Differential Privacy (RDP) if for all neighboring datasets $X, X' \in \mathcal{X}^*$, the Rényi divergence of order $\alpha > 0$ between $\mathcal{M}(X)$ and $\mathcal{M}(X')$ is at most ϵ , i.e.,*

$$D_\alpha(\mathcal{M}(X) \parallel \mathcal{M}(X')) \leq \epsilon.$$

We may also consider *individual-level* RDP, where the Rényi divergence is evaluated for neighboring datasets that differ on a particular data point z_i . We use $\mathcal{S}(z_i, n)$ to denote the set of neighboring dataset pairs $(S, \tilde{S})$ such that $|S|, |\tilde{S}| < n$, and $z_i \in S \triangle \tilde{S}$ —i.e., exactly one of the datasets contains z_i .

85 **Definition 3** (Individual Rényi Differential Privacy). *A randomized algorithm $\mathcal{M} : \mathcal{X}^* \rightarrow \Omega$ satisfies*
 86 *(α, ε) -individual RDP at point z_i if for all datasets $X, X' \in \mathcal{S}(z_i, n)$, we have*

$$D_\alpha(\mathcal{M}(X) \parallel \mathcal{M}(X')) \leq \varepsilon.$$

87 A privacy filter is a stopping time that monitors the cumulative privacy loss and terminates the
 88 algorithm once the total privacy budget is exhausted, thereby ensuring that the prescribed privacy
 89 guarantees are not violated. We now introduce the definitions of (individual) privacy filters in the
 90 context of Rényi Differential Privacy.

91 **Definition 4** ((Individual) Rényi Differential Privacy Filters [Feldman and Zrnic, 2021]). *We say that*
 92 *a random variable $\mathcal{F}_{\alpha, B} : \Omega^* \rightarrow \{\text{CONT}, \text{HALT}\}$ is a privacy filter for (α, B) -RDP if it halts the*
 93 *execution of an algorithm before its accumulated (individual) privacy loss exceeds B (measured in*
 94 *α -Rényi divergence).*

95 3 Problem Setting

96 We study retrieval-augmented generation (RAG) over a sensitive external dataset D . Given a user
 97 prompt $x \in \mathcal{V}^*$, a retrieval function R selects the top- k relevant documents $D_x = R(x, D; k)$ from
 98 D , and a decoder-only LLM then generates a response conditioned on both x and D_x using greedy
 99 decoding. The dataset D contains individual records, each corresponding to a single person’s private
 100 information. We adopt a *realistic threat model* where the adversary cannot directly access D but may
 101 issue arbitrary prompts x and observe the system’s outputs. The underlying LLM is assumed to be
 102 public; the privacy risk arises solely from the retrieval over D .

103 Our objective is to design a *differentially private RAG algorithm* that answers a sequence of *online*
 104 *queries* $\{q_1, \dots, q_T\}$ while protecting the privacy of D . Specifically, given the private dataset D , a
 105 public LLM, and a total privacy budget ε , we seek an algorithm $\mathcal{A}(D, \{q_1, \dots, q_T\}, \text{LLM}, \varepsilon)$ that
 106 generates high-utility responses while guaranteeing ε -differential privacy with respect to the external
 107 document dataset D .

108 4 Methodology

109 4.1 Technical Overview

110 **Document-level Privacy Accounting through Individual Privacy Filters.** In retrieval-augmented
 111 generation, each query typically accesses only a small (relative to the whole external dataset D),
 112 query-specific subset of relevant documents. This “sparsity” means most documents are retrieved
 113 infrequently. We exploit this by applying an individual-level privacy filter that tracks the cumulative
 114 privacy loss per document and halts access once its budget is exhausted. Since a document incurs
 115 privacy loss only when retrieved, this approach ensures privacy accounting scales with retrieval
 116 frequency rather than the total number of queries.

117 **Screening Relevant Documents via Adaptive Thresholding.** Applying RAG directly over the
 118 entire dataset negates individual privacy filters, as all documents contribute to privacy loss. To mitigate
 119 this, MURAG uses a global threshold τ : only documents with relevance scores above τ are retrieved
 120 and charged privacy cost. Once a document’s budget is exhausted, it is excluded from future responses.
 121 However, a fixed τ can be suboptimal due to varying relevance distributions across queries, leading
 122 to inefficient budget use. For example, a poorly calibrated τ may retrieve low-value documents
 123 early or omit high-value ones later, leading to inefficient budget use and degraded utility over time.
 124 To address this, we propose MURAG-ADA, which privately releases a query-specific threshold τ_t
 125 tailored to each query’s relevance distribution. By combining individual privacy accounting with
 126 private release of cumulative statistics, MURAG-ADA selectively retrieves high-relevance documents,
 127 reducing unnecessary budget consumption on irrelevant documents and potentially preserving utility
 128 across other queries.

129 After screening relevant documents, we apply a single-query DP-RAG algorithm to generate the
 130 response. As shown in Algorithms 1 and 2, our multi-query framework is modular and supports any
 131 private single-query RAG method – that is, an algorithm that ensures DP with respect to the retrieved
 132 document set. In this paper, we instantiate it with a pure DP variant of the single-query DP-RAG
 133 algorithm from Koga et al. [2024], detailed in Algorithm 5.

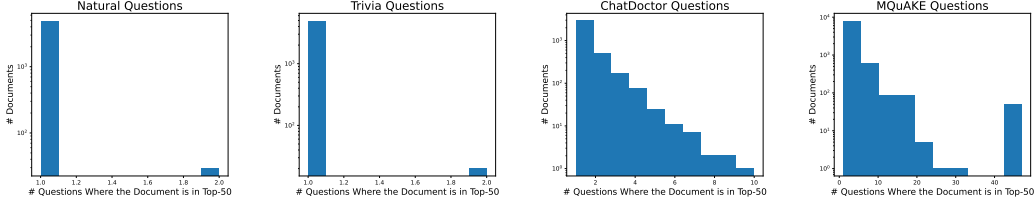


Figure 1: Histogram of how many questions each document appears in among the top- K retrieved results ($K = 50$). The x-axis indicates the number of questions, and the y-axis shows the number of such documents. We show the histogram for four datasets.

134 4.2 DP RAG with fixed threshold

135 In MURAG, we set a fixed threshold τ on the relevance score, which can be either public or privately
 136 estimated using a minimal portion of the privacy budget. This approach performs well when relevance
 137 score distributions are consistent across queries, as is the case for datasets such as Natural Questions
 and Trivia Questions. Implementation details are provided in Algorithm 1.

Algorithm 1: MURAG: Differentially Private Multi-Query Retrieval-Augmented Generation

Input: Private external dataset D , a sequence of online queries $\{q_1, q_2, \dots, q_T\}$, per query
 privacy budget ε_q , per query retrieved number of documents k , maximum retrieval
 times per document M , relevance score threshold τ
Set: Individual privacy budget for each document $z \in D$: $\mathcal{E}(z) = M \cdot \varepsilon_q$,
 1 **for** $t = 1, \dots, T$ **do**
 2 $A_t = \{z \in D \mid \varepsilon_i \geq \varepsilon_q\}$ ▷ Update active document set
 3 $D_{q_t} = \{z \in A_t \mid r(z, q_t) > \tau\}$ ▷ Find relevant documents
 4 **for** $z \in D_{q_t}$ **do**
 5 $\mathcal{E}(z) = \mathcal{E}(z) - \varepsilon_q$ ▷ Update remaining privacy budget
 6 $D_{q_t}^k = \text{TOP-K}(D_{q_t}, k, r(\cdot, q_t))$
 7 $a_t = \text{DP-RAG}(x, D_{q_t}^k, \text{LLM}, \varepsilon_q)$ ▷ Answer for question q_t using Algo. 5
 8 **return** $(a_1, a_2, \dots, a_T)$

138 **Lemma 1** (Privacy guarantee of Algorithm 1). MURAG satisfies ε -differential privacy if, for every
 139 document $z \in D$, the ex-ante individual privacy budget is at most ε .
 140

141 4.3 DP RAG with Adaptive Threshold

142 To address the limitations of using a static threshold, we propose releasing a query-specific threshold
 143 corresponding to the top- K relevance scores for each query. Specifically, we first discretize the
 144 similarity scores into bins and then iteratively release noisy prefix sums until the cumulative count
 145 exceeds K . Additional implementation details are provided in Algorithm 2. As we demonstrate in
 146 the experimental section, this approach yields improved utility on the ChatDoctor and MQuAKE
 147 datasets.

148 **Lemma 2** (Privacy guarantee of Algorithm 2). MURAG-ADA satisfies ε -differential privacy if, for
 149 every document $z \in D$, the ex-ante individual privacy budget is at most ε .

150 5 Experiment

151 5.1 Dataset and Model set-up

152 **Question datasets in two categories.** We evaluate our methods on four question datasets: *Nat-*
 153 *ural Questions*, *Trivia Questions*, *ChatDoctor Questions*, and *MQuAKE Questions*. Natural Ques-
 154 tions [Kwiatkowski et al., 2019] and Trivia Questions [Joshi et al., 2017] are standard benchmarks
 155 for evaluating RAG systems and have been used in prior work on per-query DP for RAG [Koga
 156 et al., 2024]. Following their setup, we randomly subsample 100 questions from each dataset to

Algorithm 2: MURAG-ADA: DP Multi-Query RAG with adaptive threshold

Input: Private external dataset D , a sequence of online queries $\{q_1, q_2, \dots, q_T\}$, per query budget ε_q , per query retrieved number of documents k , maximum retrieval times per document M , initial relevance threshold τ

Set: Individual privacy budget for each document $z \in D$: $\mathcal{E}(z) = M \cdot \varepsilon_q$, privacy budget allocation: $\varepsilon_q = \varepsilon_{\text{thr}} + \varepsilon_{\text{RAG}}$

Require: Discretization on similarity scores $[a_i, b_i]_{i=1}^B$

```
1 for  $t = 1, \dots, T$  do
  /* Release prefix sum */
2    $\tilde{s} = 0, A_t = \phi$ 
3   for  $i = 1, \dots, B$  do
4      $A_t^{(i)} = \{z \in D \mid r(z, q_t) \in [a_i, b_i] \text{ and } \mathcal{E}(z) \geq \varepsilon_{\text{thr}}\}$ 
5      $\tilde{s} = \tilde{s} + |A_t^{(i)}| + \text{Lap}(1/\varepsilon_{\text{thr}})$ 
6      $A_t = A_t \cup A_t^{(i)}$ 
7     for  $z \in A_t^{(i)}$  do
8        $\mathcal{E}(z) = \mathcal{E}(z) - \varepsilon_{\text{thr}}$ 
9     if  $\tilde{s} \geq K$  then
10      Halt
  /* RAG */
11   $A'_t = \{z \in A_t \mid \mathcal{E}(z) \geq \varepsilon_{\text{RAG}}\}$ 
12   $D_{q_t}^k = \text{TOP-K}(A'_t, k, r(\cdot, q_t))$ 
13   $a_t = \text{DP-RAG}(x, D_{q_t}^k, \text{LLM}, \tau_t, \varepsilon_{\text{RAG}})$  ▷ Answer  $q_t$  using Algo. 5
14  for  $z \in A'_t$  do
15     $\mathcal{E}(z) = \mathcal{E}(z) - \varepsilon_{\text{RAG}}$  ▷ Update remaining budget
16 return  $(a_1, a_2, \dots, a_T)$ 
```

157 reduce computational overhead. Chatdoctor Questions [Li et al., 2023] consist of QA interactions
158 between patients and doctors in the healthcare domain. We sample 100 patient questions from the
159 original dataset as our test set. MQuAKE Questions [Zhong et al.] contain sequences of semantically
160 related single-hop questions that collectively form multi-hop reasoning chains. We select 100 such
161 sequences, resulting in a test set of 400 individual questions.

162 To better analyze the performance differences between MURAG and MURAG-ADA, we categorize
163 the four datasets into two types: *independent question sets* and *dependent question sets*. As discussed
164 in Section 4.3, we expect MURAG-ADA to be particularly effective in datasets where questions
165 are semantically related and share overlapping relevant documents, while offering limited benefit in
166 datasets where questions are unrelated and retrieve disjoint sets of documents.

167 To support this categorization, we plot histograms showing how frequently each document appears
168 in the top- K retrieved results ($K = 50$) across questions. As shown in Figure 1, we observe
169 that in *Natural Questions* and *Trivia Questions*, most documents are retrieved for only one or two
170 questions. This indicates minimal overlap in relevant documents, and we therefore categorize them
171 as *independent question sets*. In contrast, in *ChatDoctor Questions* and *MQuAKE Questions*, many
172 documents are shared across multiple questions, suggesting substantial overlap in relevance. We
173 categorize these as *dependent question sets*.

174 **External datasets reflecting both standard and privacy-sensitive settings.** For Natural Questions,
175 Trivia Questions, and MQuAKE Questions, we use Wikipedia as the external knowledge source
176 following the standard RAG setup [Chen et al., 2017, Lewis et al., 2020]. For ChatDoctor Questions,
177 the external dataset consists of the remaining QA pairs from the original ChatDoctor dataset, excluding
178 the 100 patient questions used for testing. This setup reflects a realistic privacy-sensitive application,
179 where the external corpus contains inherently private information.

QA evaluation metric. For Natural Questions, Trivia Questions and MQuAKE Questions, the datasets provide a list of all acceptable correct answers for each question. Following the evaluation protocol of Koga et al. [2024], we use the *Match Accuracy* metric: a prediction is scored as 1 if it contains any correct answer, and 0 otherwise. For Chatdoctor Questions, we adopt the evaluation metric from the original dataset paper, using the F1 score of BERTScore [Zhang et al., 2020] to measure semantic similarity between the predicted response and the ground-truth answer.

Model set-up. Our RAG pipeline integrates three pre-trained LLMs: OPT-1.3B [Zhang et al., 2022], Pythia-1.4B [Biderman et al., 2023], and Mistral-7B [Jiang et al., 2023]. For document retrieval, we use the Dense Passage Retriever (DPR) [Karpukhin et al., 2020] to compute dense query-document relevance scores.

5.2 Method Set-up

Baseline methods. We compare our two proposed methods with three baselines. The first is DP-MULTI-RAG (Algorithm 4), which applies the per-question DP RAG method, DPSParseVoteRAG, independently to each query and uses the standard sequential composition theorem [Dwork et al., 2006a] to compute the overall privacy guarantee. The other two are non-private baselines: Non-RAG, which generates answers using the pretrained LLM without retrieval, and Non-Private-RAG, which performs retrieval-augmented generation without any privacy mechanism.

Privacy budget setup for DP algorithms. Following the setup in Koga et al. [2024], we vary the per-query RAG privacy budget $\epsilon_q \in 2, 5, 10, 15, 20, 30, 40$ to explore the privacy-utility trade-off. For DP-MULTI-RAG, the total privacy budget is $T \cdot \epsilon_q$, where T is the number of questions. For MURAG and MURAG-ADA, the total budget is $M \cdot \epsilon_q$, where M is the number of retrieved documents with nonzero privacy loss. In our main results, we conservatively set $M = 1$ for a realistic privacy region. Moreover, ϵ_{thr} is fixed as 1.0.

Other hyperparameter settings. All three DP algorithms rely on shared hyperparameters from DPSParseVoteRAG, including the number of retrieved documents k , the per-token privacy budget ϵ_{token} , and the SVT threshold τ_{svt} . Following Koga et al. [2024], we evaluate each method under a grid of settings with $k \in 30, 40, 50$, $\epsilon_{\text{token}} \in 0.5, 1.0, 2.0$, and $\tau_{\text{svt}} = k/2$. We report the best performance for each method over these configurations.

5.3 Results

Figure 2 presents the performance of our two proposed methods alongside three baselines across four datasets and three pretrained LLMs. To focus the analysis, we truncate the figure to show only the utility region above the Non-RAG baseline, i.e. where using a DP-protected RAG system offers a utility gain over simply using the pretrained LLM alone. This region is of primary interest, as it justifies the use of a private external dataset under differential privacy. Within this utility range, we compare different DP algorithms based on the amount of privacy budget they consume to achieve a given performance level.

Comparing our methods with baselines. Across all four datasets and three LLMs, both of our proposed methods consistently outperform the Non-RAG baseline at a total privacy budget of $\epsilon = 10$. In contrast, the baseline method NAIVE-MULTI-RAG requires a significantly larger budget, exceeding $\epsilon = 10^3$, to achieve comparable utility. These results demonstrate that our methods make differential privacy practical in the multi-query RAG setting, enabling strong performance while staying within a realistic privacy budget.

Comparison between our two methods. We observe a clear performance pattern between our two methods based on the dataset type. On the two *independent question sets* (*Natural Questions* and *Trivia Questions*) MURAG consistently outperforms MURAG-ADA. In contrast, on the *dependent question sets* (*ChatDoctor Questions* and *MQuAKE Questions*) MURAG-ADA shows superior performance. The improvement is especially pronounced on *MQuAKE Questions*, where MURAG performs only marginally better than the Non-RAG baseline, while MURAG-ADA yields a significant utility gain.

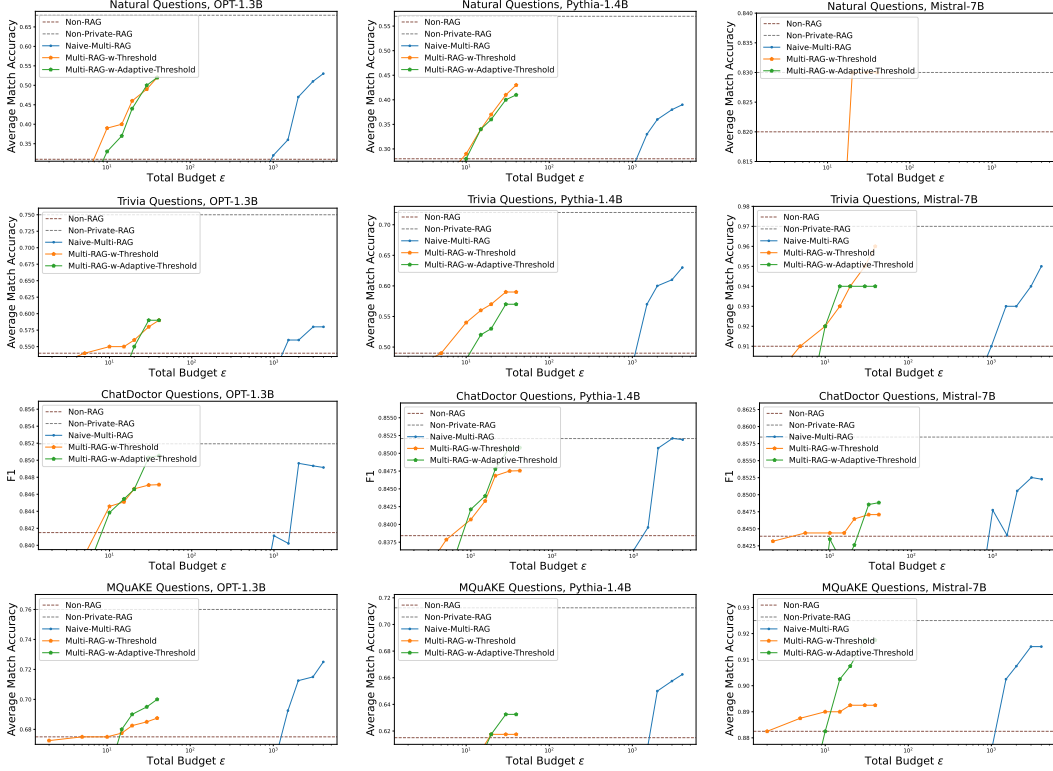


Figure 2: Utility-privacy tradeoffs of our two proposed methods (MURAG and MURAG-ADA) compared to three baselines across four question datasets and three pretrained LLMs. The plots are truncated to show only the region where the utility exceeds the Non-RAG baseline (i.e., the pretrained LLM without access to the external dataset), which is the regime of practical interest.

Table 1: Precision of retrieved documents under MURAG and MURAG-ADA, measured as the percentage of truly top-50 relevant documents among the retrieved. Results are reported for each dataset. MURAG-ADA achieves significantly higher precision on the dependent question sets (ChatDoctor and MQuAKE), demonstrating its more efficient use of the privacy budget to save documents for later questions.

	Independent Question Set		Dependent Question Set	
	Natural Questions	Trivia Questions	Chatdoctor Questions	MQuAKE Questions
MURAG	78.8%	72.2%	13.3%	17.6%
MURAG-ADA	92.6%	94.6%	67.2%	40.7%
MURAG-ADA (non-private top-K-release)	99.4%	99.6%	71.2%	43.5%

229 This discrepancy arises from the limitations of using a fixed retrieval threshold τ in MURAG. Since
230 different queries induce different distributions over relevance scores, a single global threshold can lead
231 to imbalanced behavior: a threshold that retrieves sufficient relevant documents for one query may
232 result in many low-relevance documents for another. These low-relevance documents still consume
233 privacy budget without meaningfully improving the response, and may become unavailable for future
234 queries where they are actually useful. This inefficiency is especially problematic in dependent
235 question sets, where documents are frequently shared across queries.

236 To quantify this effect Table 1 reports the precision under both MURAG and MURAG-ADA, where
237 the precision is defined as the percentage of truly top-50 documents among the retrieved documents for
238 each question. We observe that precision under MURAG is particularly low for *ChatDoctor Questions*
239 and *MQuAKE Questions*, whereas MURAG-ADA significantly improves retrieval precision on these
240 datasets through its adaptive thresholds. This improvement in retrieval quality directly contributes to
241 the superior performance of MURAG-ADA in the setting of *dependent question set*.

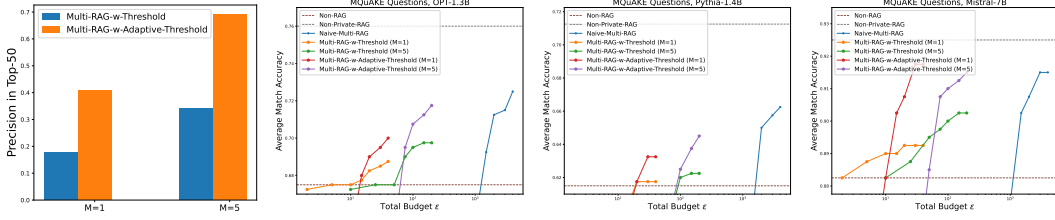


Figure 3: Comparison of $M = 1$ and $M = 5$ in the individual privacy accounting framework. Left plot shows the retrieval precisions of two methods with $M = 1, 5$. The right three plots show the trade-off between the QA performance and the ϵ_{total} in DP. Increasing M improves retrieval precision and RAG utility, but incurs a higher total privacy cost. $M = 1$ remains preferable for a stronger privacy-utility trade-off.

Effect of different M in the individual privacy accounting framework. Both of our proposed methods include a hyperparameter M , which controls the maximum number of queries for which an individual document’s privacy budget can be consumed. In our main results (Figure 2), we set $M = 1$ to ensure strict per-document privacy usage. However, this setting may limit utility: once a document is used for one query, it becomes unavailable for future queries, even if it would have been highly relevant. This limitation is evident in Table 1, where both methods exhibit low retrieval precision on the *MQAKE Questions* dataset and this suggests that relevant documents were prematurely deactivated.

To better understand the impact of M , we evaluate our two methods with a larger value of $M = 5$. Figure 3 compares performance with $M = 1$ and $M = 5$. The left plot shows a substantial increase in Top-50 retrieval precision when using $M = 5$, indicating better access to relevant documents. This improvement translates into higher end-to-end RAG utility, as shown in the three plots on the right. However, increasing M also leads to a higher total privacy cost ($\epsilon_{\text{total}} = M \cdot \epsilon_q$). Overall, while $M = 5$ enhances utility, we find that $M = 1$ still achieves the best privacy-utility trade-off under a practical privacy regime.

6 Discussion

Why Privacy Filter rather than Amplification by Subsampling? As surveyed in Section A.2, privacy amplification by subsampling [Balle et al., 2018, Wang et al., 2019, Zhu and Wang, 2019] is widely used in DP LLM applications, such as DP prompt tuning and DP in-context learning, to enhance generation quality. However, this technique is not well-suited for DP RAG:

- In prompt tuning, the goal is to learn a single task-specific prompt that can generalize to all future queries. In DP in-context learning, a small number of example inputs are selected under DP constraints and reused across queries. In contrast, our RAG setting does not allow for such "unified" prompts or examples: each test-time query requires retrieving and using query-specific documents, which must be handled privately, which makes individual privacy filter a more suitable choice.
- Moreover, in prompt tuning and in-context learning, all data points in the private dataset can meaningfully contribute to the learned prompt or selected example set. This property enables the use of subsampling-based amplification techniques in algorithm design. In RAG, however, only a sparse subset of documents in the large external corpus are relevant to any given query—most documents provide no utility.

These two key differences, the lack of reusable prompts and the sparsity of useful data, motivate the development of our new DP RAG algorithms using R nyi filter rather than amplification by sampling.

Leveraging Historical QA. As shown in Table 1 and Figure 1, when the relevant documents for different questions exhibit significant overlap, the quality of answers to later questions degrades. This occurs because the documents required to answer the queries may exhaust their privacy budgets and are subsequently filtered out from the active set passed to the RAG algorithm. In the extreme case where a user repeatedly submits the same query, only the first response may retain high quality, while subsequent answers degrade due to the unavailability of relevant documents.

281 A potential remedy is to reuse historical answers as auxiliary documents in future queries. This
 282 can be done without incurring any additional privacy cost, owing to the post-processing property of
 283 differential privacy.

284 7 Conclusion

285 We proposed the first differentially private (DP) framework for retrieval-augmented generation (RAG)
 286 that supports answering multiple queries while protecting a sensitive external dataset. Our methods
 287 build on individual privacy accounting to overcome the limitations of the prior single-query DP-
 288 RAG system with the basic sequential composition. We introduced two algorithms—MURAG and
 289 MURAG-ADA differ in how they select documents for each query under DP guarantees. MURAG
 290 uses a fixed threshold for document retrieval, while MURAG-ADA adaptively adjusts the threshold
 291 per query to improve the efficiency of individual privacy accounting. Through comprehensive experi-
 292 ments on four question datasets and three LLMs, we demonstrated that both methods significantly
 293 reduce the privacy cost needed to outperform a Non-RAG baseline, achieving strong utility for
 294 answering 100 questions under a realistic budget of $\epsilon = 10$. We also showed that MURAG-ADA
 295 performs particularly well on datasets with overlapping document relevance across queries. This
 296 work highlights the importance of tailoring differential privacy mechanisms to the characteristics of
 297 the RAG setting. We hope our contributions provide a foundation for more practical and principled
 298 privacy-preserving RAG systems.

299 References

- 300 Maya Anderson, Guy Amit, and Abigail Goldstein. Is my data in your retrieval database? membership
 301 inference attacks against retrieval augmented generation. *arXiv preprint arXiv:2405.20446*, 2024.
- 302 Borja Balle, Gilles Barthe, and Marco Gaboardi. Privacy amplification by subsampling: Tight
 303 analyses via couplings and divergences. *Advances in neural information processing systems*, 31,
 304 2018.
- 305 Stella Biderman, Hailey Schoelkopf, Quentin Gregory Anthony, Herbie Bradley, Kyle O’Brien, Eric
 306 Hallahan, Mohammad Aflah Khan, Shivanshu Purohit, USVSN Sai Prashanth, Edward Raff, et al.
 307 Pythia: A suite for analyzing large language models across training and scaling. In *International*
 308 *Conference on Machine Learning*, pages 2397–2430. PMLR, 2023.
- 309 Zachary Charles, Arun Ganesh, Ryan McKenna, H Brendan McMahan, Nicole Mitchell, Krishna
 310 Pillutla, and Keith Rush. Fine-tuning large language models with user-level differential privacy.
 311 *arXiv preprint arXiv:2407.07737*, 2024.
- 312 Danqi Chen, Adam Fisch, Jason Weston, and Antoine Bordes. Reading wikipedia to answer open-
 313 domain questions. In *Proceedings of the 55th Annual Meeting of the Association for Computational*
 314 *Linguistics (Volume 1: Long Papers)*, pages 1870–1879, 2017.
- 315 Yihang Cheng, Lan Zhang, Junyang Wang, Mu Yuan, and Yunhao Yao. Remoterag: A privacy-
 316 preserving llm cloud rag service. *arXiv preprint arXiv:2412.12775*, 2024.
- 317 Jinshuo Dong, Aaron Roth, and Weijie J Su. Gaussian differential privacy. *Journal of the Royal*
 318 *Statistical Society Series B: Statistical Methodology*, 84(1):3–37, 2022.
- 319 Haonan Duan, Adam Dziedzic, Nicolas Papernot, and Franziska Boenisch. Flocks of stochastic
 320 parrots: Differentially private prompt learning for large language models. *Advances in Neural*
 321 *Information Processing Systems*, 36:76852–76871, 2023.
- 322 David Durfee and Ryan M Rogers. Practical differentially private top-k selection with pay-what-you-
 323 get composition. *Advances in Neural Information Processing Systems*, 32, 2019.
- 324 Cynthia Dwork, Krishnaram Kenthapadi, Frank McSherry, Ilya Mironov, and Moni Naor. Our data,
 325 ourselves: Privacy via distributed noise generation. In *Annual international conference on the*
 326 *theory and applications of cryptographic techniques*, pages 486–503. Springer, 2006a.

327 Cynthia Dwork, Frank McSherry, Kobbi Nissim, and Adam Smith. Calibrating noise to sensitivity in
328 private data analysis. In *Proceedings of the Third Conference on Theory of Cryptography, TCC'06*,
329 page 265–284, 2006b.

330 Vitaly Feldman and Tijana Zrnic. Individual privacy accounting via a renyi filter. *Advances in Neural*
331 *Information Processing Systems*, 34:28080–28091, 2021.

332 Nicolas Grislain. Rag with differential privacy. In *2025 IEEE Conference on Artificial Intelligence*
333 *(CAI)*, pages 847–852. IEEE, 2025.

334 Junyuan Hong, Jiachen T. Wang, Chenhui Zhang, Zhangheng LI, Bo Li, and Zhangyang Wang.
335 DP-OPT: Make large language model your privacy-preserving prompt engineer. In *The Twelfth*
336 *International Conference on Learning Representations*, 2024. URL [https://openreview.net/](https://openreview.net/forum?id=Ifz3IgsEPX)
337 [forum?id=Ifz3IgsEPX](https://openreview.net/forum?id=Ifz3IgsEPX).

338 Albert Q. Jiang, Alexandre Sablayrolles, Arthur Mensch, Chris Bamford, Devendra Singh Chaplot,
339 Diego de las Casas, Florian Bressand, Gianna Lengyel, Guillaume Lample, Lucile Saulnier,
340 L  lio Renard Lavaud, Marie-Anne Lachaux, Pierre Stock, Teven Le Scao, Thibaut Lavril, Thomas
341 Wang, Timoth  e Lacroix, and William El Sayed. Mistral 7b, 2023. URL [https://arxiv.org/](https://arxiv.org/abs/2310.06825)
342 [abs/2310.06825](https://arxiv.org/abs/2310.06825).

343 Changyue Jiang, Xudong Pan, Geng Hong, Chenfu Bao, and Min Yang. Rag-thief: Scalable extraction
344 of private data from retrieval-augmented generation applications with agent-based attacks. *arXiv*
345 *preprint arXiv:2411.14110*, 2024.

346 Mandar Joshi, Eunsol Choi, Daniel Weld, and Luke Zettlemoyer. TriviaQA: A large scale distantly
347 supervised challenge dataset for reading comprehension. In Regina Barzilay and Min-Yen Kan,
348 editors, *Proceedings of the 55th Annual Meeting of the Association for Computational Linguistics*
349 *(Volume 1: Long Papers)*, pages 1601–1611, Vancouver, Canada, July 2017. Association for
350 Computational Linguistics.

351 Vladimir Karpukhin, Barlas Oguz, Sewon Min, Patrick SH Lewis, Ledell Wu, Sergey Edunov, Danqi
352 Chen, and Wen-tau Yih. Dense passage retrieval for open-domain question answering. In *EMNLP*
353 *(1)*, pages 6769–6781, 2020.

354 Tatsuki Koga, Ruihan Wu, and Kamalika Chaudhuri. Privacy-preserving retrieval augmented genera-
355 tion with differential privacy. *arXiv preprint arXiv:2412.04697*, 2024.

356 Antti Koskela, Marlon Tobaben, and Antti Honkela. Individual privacy accounting with gaussian
357 differential privacy. *arXiv preprint arXiv:2209.15596*, 2022.

358 Tom Kwiatkowski, Jennimaria Palomaki, Olivia Redfield, Michael Collins, Ankur Parikh, Chris
359 Alberti, Danielle Epstein, Illia Polosukhin, Jacob Devlin, Kenton Lee, et al. Natural questions: a
360 benchmark for question answering research. *Transactions of the Association for Computational*
361 *Linguistics*, 7:453–466, 2019.

362 Patrick Lewis, Ethan Perez, Aleksandra Piktus, Fabio Petroni, Vladimir Karpukhin, Naman Goyal,
363 Heinrich K  ttler, Mike Lewis, Wen-tau Yih, Tim Rockt  schel, et al. Retrieval-augmented genera-
364 tion for knowledge-intensive nlp tasks. *Advances in neural information processing systems*, 33:
365 9459–9474, 2020.

366 Xuechen Li, Florian Tramer, Percy Liang, and Tatsunori Hashimoto. Large language models can be
367 strong differentially private learners. *arXiv preprint arXiv:2110.05679*, 2021.

368 Yunxiang Li, Zihan Li, Kai Zhang, Ruilong Dan, Steve Jiang, and You Zhang. Chatdoctor: A medical
369 chat model fine-tuned on a large language model meta-ai (llama) using medical domain knowledge.
370 *Cureus*, 15(6), 2023.

371 Yuying Li, Gaoyang Liu, Chen Wang, and Yang Yang. Generating is believing: Membership
372 inference attacks against retrieval-augmented generation. In *ICASSP 2025-2025 IEEE International*
373 *Conference on Acoustics, Speech and Signal Processing (ICASSP)*, pages 1–5. IEEE, 2025.

374 Mingrui Liu, Sixiao Zhang, and Cheng Long. Mask-based membership inference attacks for retrieval-
375 augmented generation. In *Proceedings of the ACM on Web Conference 2025*, pages 2894–2907,
376 2025.

377 Ilya Mironov. Rényi differential privacy. In *2017 IEEE 30th computer security foundations symposium*
378 *(CSF)*, pages 263–275. IEEE, 2017.

379 Ali Naseh, Yuefeng Peng, Anshuman Suri, Harsh Chaudhari, Alina Oprea, and Amir Houmansadr.
380 Riddle me this! stealthy membership inference for retrieval-augmented generation. *arXiv preprint*
381 *arXiv:2502.00306*, 2025.

382 Yuefeng Peng, Junda Wang, Hong Yu, and Amir Houmansadr. Data extraction attacks in retrieval-
383 augmented generation via backdoors. *arXiv preprint arXiv:2411.01705*, 2024.

384 Ryan M Rogers, Aaron Roth, Jonathan Ullman, and Salil Vadhan. Privacy odometers and filters:
385 Pay-as-you-go composition. *Advances in Neural Information Processing Systems*, 29, 2016.

386 Reza Shokri, Marco Stronati, Congzheng Song, and Vitaly Shmatikov. Membership inference attacks
387 against machine learning models. In *2017 IEEE symposium on security and privacy (SP)*, pages
388 3–18. IEEE, 2017.

389 Adam Smith and Abhradeep Thakurta. Fully adaptive composition for gaussian differential privacy.
390 *arXiv preprint arXiv:2210.17520*, 2022.

391 Xinyu Tang, Richard Shin, Huseyin A Inan, Andre Manoel, Fatemehsadat Miresghallah, Zinan Lin,
392 Sivakanth Gopi, Janardhan Kulkarni, and Robert Sim. Privacy-preserving in-context learning with
393 differentially private few-shot generation. In *The Twelfth International Conference on Learning*
394 *Representations*, 2024. URL <https://openreview.net/forum?id=oZttOpRn0l>.

395 Yu-Xiang Wang, Borja Balle, and Shiva Prasad Kasiviswanathan. Subsampled rényi differential
396 privacy and analytical moments accountant. In *The 22nd international conference on artificial*
397 *intelligence and statistics*, pages 1226–1235. PMLR, 2019.

398 Justin Whitehouse, Aaditya Ramdas, Ryan Rogers, and Steven Wu. Fully-adaptive composition in
399 differential privacy. In *International conference on machine learning*, pages 36990–37007. PMLR,
400 2023.

401 Tong Wu, Ashwinee Panda, Jiachen T. Wang, and Prateek Mittal. Privacy-preserving in-context
402 learning for large language models. In *The Twelfth International Conference on Learning Re-*
403 *presentations*, 2024. URL <https://openreview.net/forum?id=x40PJ71HVV>.

404 Dixi Yao and Tian Li. Private retrieval augmented generation with random projection. In *ICLR 2025*
405 *Workshop on Building Trust in Language Models and Applications*.

406 Da Yu, Saurabh Naik, Arturs Backurs, Sivakanth Gopi, Huseyin A Inan, Gautam Kamath, Janardhan
407 Kulkarni, Yin Tat Lee, Andre Manoel, Lukas Wutschitz, et al. Differentially private fine-tuning of
408 language models. *arXiv preprint arXiv:2110.06500*, 2021.

409 Shenglai Zeng, Jiankun Zhang, Pengfei He, Yiding Liu, Yue Xing, Han Xu, Jie Ren, Yi Chang,
410 Shuaiqiang Wang, Dawei Yin, et al. The good and the bad: Exploring privacy issues in retrieval-
411 augmented generation (rag). In *Findings of the Association for Computational Linguistics ACL*
412 *2024*, pages 4505–4524, 2024a.

413 Shenglai Zeng, Jiankun Zhang, Pengfei He, Jie Ren, Tianqi Zheng, Hanqing Lu, Han Xu, Hui Liu,
414 Yue Xing, and Jiliang Tang. Mitigating the privacy issues in retrieval-augmented generation (rag)
415 via pure synthetic data. *arXiv preprint arXiv:2406.14773*, 2024b.

416 Susan Zhang, Stephen Roller, Naman Goyal, Mikel Artetxe, Moya Chen, Shuohui Chen, Christopher
417 Dewan, Mona Diab, Xian Li, Xi Victoria Lin, et al. Opt: Open pre-trained transformer language
418 models. *arXiv preprint arXiv:2205.01068*, 2022.

419 Tailai Zhang, Yuxuan Jiang, Ruihan Gong, Pan Zhou, Wen Yin, Xingxing Wei, Lixing Chen, and
420 Daizong Liu. DEAL: High-efficacy privacy attack on retrieval-augmented generation systems via
421 LLM optimizer, 2025. URL <https://openreview.net/forum?id=sx8dtyZT41>.

- 422 Tianyi Zhang, Varsha Kishore, Felix Wu, Kilian Q. Weinberger, and Yoav Artzi. Bertscore: Evaluating
423 text generation with bert. In *International Conference on Learning Representations*, 2020. URL
424 <https://openreview.net/forum?id=SkeHuCVFDr>.
- 425 Zexuan Zhong, Zhengxuan Wu, Christopher D Manning, Christopher Potts, and Danqi Chen. Mquake:
426 Assessing knowledge editing in language models via multi-hop questions. In *The 2023 Conference*
427 *on Empirical Methods in Natural Language Processing*.
- 428 Yuqing Zhu and Yu-Xiang Wang. Poission subsampled rényi differential privacy. In Kamalika
429 Chaudhuri and Ruslan Salakhutdinov, editors, *Proceedings of the 36th International Conference on*
430 *Machine Learning*, volume 97 of *Proceedings of Machine Learning Research*, pages 7634–7642.
431 PMLR, 09–15 Jun 2019.

A Related Work

A.1 Privacy-Preserving Retrieval Augmented Generation

Recent studies have revealed two main categories of privacy risks in retrieval-augmented generation (RAG) systems. The first is membership inference attacks (MIA) [Shokri et al., 2017], which aim to determine whether a specific document was included in the private external dataset. Research on MIA in RAG has explored adversarial prompt construction [Naseh et al., 2025, Liu et al., 2025, Anderson et al., 2024] and scoring mechanisms [Li et al., 2025] to increase inference success. The second type is data reconstruction attacks, which aim to recover the raw content of documents in the private dataset. These attacks have been explored via adversarial prompt design [Zhang et al., 2025, Zeng et al., 2024a, Jiang et al., 2024] as well as through data poisoning techniques that embed triggers to facilitate reconstruction [Peng et al., 2024]. Together, these works highlight the growing need for principled privacy-preserving algorithms for RAG.

Several approaches have been proposed to address these privacy risks. Most notably, Koga et al. [2024] introduced a differentially private (DP) RAG system that provides privacy guarantees for a single query. Other works [Yao and Li, Grislain, 2025] have studied how to release document identifiers under DP guarantees. However, none of these methods address the setting where multiple queries are issued, which is more representative of real-world usage. In addition to DP-based methods, there are empirical approaches that aim to enhance privacy without formal guarantees. Yao and Li propose a synthetic document generation technique, where retrieved documents are paraphrased before being passed to the generator. Zeng et al. [2024b] introduces a dataset privatization strategy that removes sensitive content across multiple documents prior to retrieval. While these methods offer promising empirical results, they lack worst-case guarantees and remain vulnerable to strong adversarial attacks. Finally, Cheng et al. [2024] explores a complementary privacy concern—protecting the user’s query when interacting with a cloud-hosted RAG system. While important, this line of work addresses a different threat model than ours.

A.2 Differential Privacy in Large Language Models

Beyond our focus on DP in RAG, DP has also been extensively studied in other LLM settings, including fine-tuning [Charles et al., 2024, Yu et al., 2021, Li et al., 2021], prompt tuning [Duan et al., 2023, Hong et al., 2024], and in-context learning [Tang et al., 2024, Wu et al., 2024]. Due to the differing nature of these tasks, the effective DP mechanisms vary significantly across settings. In pre-training and fine-tuning, the core challenge lies in optimizing the model parameters while maintaining stability against the noise introduced by DP algorithms. These settings are fundamentally different from RAG, where the goal is not to train the model but to ensure privacy during inference-time retrieval and generation. More closely related to our work are DP approaches for prompt tuning and in-context learning. However, the structural differences between these tasks and RAG result in distinct design considerations for differential privacy algorithms; please see Section 6 for more details.

A.3 Individual Privacy Accounting and Privacy Filters

Individual privacy accounting focuses on tracking the privacy loss incurred by a single data point, often yielding tighter privacy bounds than traditional worst-case analyses over all neighboring datasets [Dwork et al., 2006b]. This line of work was initiated by Feldman and Zrnic [2021] in the context of Rényi Differential Privacy and was subsequently extended to Gaussian Differential Privacy [Dong et al., 2022] by Koskela et al. [2022]. For a comprehensive overview, we refer the reader to Feldman and Zrnic [2021, Section 1.2].

Building upon this framework, the concept of a privacy filter has been introduced as a general mechanism for adaptively enforcing privacy constraints. A privacy filter refers to a stopping rule that halts the execution of differentially private (DP) algorithms before the cumulative privacy loss exceeds a specified budget. Individual privacy filters, introduced by Feldman and Zrnic [2021] and further developed by Koskela et al. [2022], represent a specialized instantiation of this framework. These filters operate at the granularity of individual data points, terminating their participation once their respective privacy budgets are exhausted. For further details, we refer the reader to Rogers et al.

483 [2016], Feldman and Zrnic [2021], Koskela et al. [2022], Smith and Thakurta [2022], Whitehouse
 484 et al. [2023].

485 **B Algorithms**

486 We present additional algorithms that were omitted from the main paper for brevity.

487 **B.1 Top-k document selection**

488 Algorithm 3 returns the top- K documents from a dataset D ranked by a score function r , padding with empty strings if $|D| < K$ to ensure the output always has exactly K elements.

Algorithm 3: TOP-K(D, K, r)

Input: dataset D , sample size K , score function r
 1 **if** $|D| \geq K$ **then**
 2 $D^k \leftarrow$ top- K documents from D according to score function r $\triangleright$ assume no ties
 3 **else**
 4 $D^k \leftarrow D \cup \{\text{" "}\}^{k-|D|}$ $\triangleright$ Add empty string " " until document set has
 size k
 5 **return** D^k

489

490 **B.2 Naive algorithm for DP Multi-Question RAG**

491 Algorithm 4 serves as a baseline for handling multi-query DP RAG by composing a single-query DP-RAG mechanism T times.

Algorithm 4: NAIVE-MULTI-RAG

Input: Private external dataset D , a sequence of queries $\{q_1, q_2, \dots, q_T\}$, total privacy budget ε , per query budget ε_q
Require: $\varepsilon \geq T \cdot \varepsilon_q$
 1 **for** $t = 1, \dots, T$ **do**
 2 $a_t = \text{DP-RAG}(x, D, \text{LLM}, \varepsilon_q)$ $\triangleright$ Algo. 5
 3 **return** $(a_1, a_2, \dots, a_T)$

492

493 **B.3 DP-RAG for single question answering**

494 Algorithm 5 is a variant of Koga et al. [2024, Algorithm 2], in which we replace the LimitedDomain
 495 mechanism [Durfee and Rogers, 2019] with the exponential mechanism in the private token generation
 496 step. This modification yields a more stringent pure-DP guarantee and results in a cleaner privacy
 497 composition accounting result.

Algorithm 5: DP-RAG($x, D, \text{LLM}, \varepsilon$)

Input: Prompt x ; external data source D ; next-token generator
LLM(prompt, doc | history); total budget ε ;
Set: Per-token privacy budget ε_0
Require: maximum length of output tokens $T_{\max}$; number of voters m ; retrievals per voter k ;
document retriever $R(\text{prompt}, \text{doc set}, \# \text{retrieved docs})$; threshold for voting θ

```
1  $\varepsilon_{\text{Expo}} \leftarrow \varepsilon_0/2, \varepsilon_{\text{Lap}} \leftarrow \varepsilon_0/2$   $\triangleright$  split privacy budget for per token generation
2  $c \leftarrow \lfloor \varepsilon/\varepsilon_{\text{RAG}} \rfloor, \hat{\theta} \leftarrow \theta + \text{Lap}(2/\varepsilon_{\text{Lap}})$ 
3  $D_x \leftarrow R(x, D; mk)$   $\triangleright$  retrieve  $mk$  documents
4  $\mathcal{D}_x \leftarrow \{D_x^1, \dots, D_x^m\}$   $\triangleright$  Partition  $D_x$  into  $m$  subsets uniformly random
5 for  $t \leftarrow 1$  to  $T_{\max}$  do
6    $y_t^{\text{non-RAG}} \leftarrow \text{LLM}(x, \text{" " } | y_{<t})$ 
7   for  $i \leftarrow 1$  to  $m$  do
8      $y_t^{(i)} \leftarrow \text{LLM}(x, D_x^i | y_{<t})$ 
9    $\text{Hist}_t \leftarrow \text{Hist}(y_t^{(1)}, \dots, y_t^{(m)})$   $\triangleright \text{Hist}_t \in \mathbb{N}^{|\mathcal{V}|}$ 
10   $\text{Count}_t \leftarrow \text{Hist}_t[\text{index} = y_t^{\text{non-RAG}}]$ 
11  if  $\text{Count}_t + \text{Lap}(4/\varepsilon_{\text{Lap}}) \leq \hat{\theta}$  then
12     $y_t \leftarrow \text{expoMech}(\text{Hist}_t; \varepsilon_{\text{Expo}})$ 
13     $c \leftarrow c - 1$ 
14  else
15     $y_t \leftarrow y_t^{\text{non-RAG}}$ 
16  if  $y_t = \langle \text{EOS} \rangle$  or  $c = 0$  then
17    return  $(y_1, \dots, y_t)$ 
18 return  $(y_1, \dots, y_{T_{\max}})$ 
```

498 C Proofs for Privacy Guarantee

499 C.1 Privacy Guarantee for Algorithm 1

500 **Lemma** (Restatement of Lemma 1). MURAG satisfies ε -differential privacy if, for every $z \in D$, the
501 ex-ante individual privacy budget is at most ε .

502 *Proof.* Since $\mathcal{E}(z) \leq \varepsilon$ for every $z \in D$, the privacy guarantee follows directly from Feldman and
503 Zrnic [2021, Corollary 3.3]. $\square$

504 C.2 Privacy Guarantee for Algorithm 2

505 **Lemma** (Restatement of Lemma 2). MURAG-ADA satisfies ε -differential privacy if, for every
506 $z \in D$, the ex-ante individual privacy budget is at most ε .

507 *Proof.* We first bound the individual privacy for the t -th prefix-sum release algorithm, denoted by
508 $\mathcal{A}_t$. Consider $S, \tilde{S} \in \mathcal{S}(z_i, n)$, and without loss of generality, we assume $z_i \in S$. Conditioned
509 on the trajectory $r^{(t-1)}$ from the previous $t - 1$ rounds, for any possible output sequence $b^{(q)} :=$
510 $(b_1, b_2, \dots, b_q)$ with $q \leq B$, the interesting regime is that there exists $j \in [q]$ such that z_i contributes
511 to b_j ; otherwise, we have $\mathcal{A}_t(S | r^{(t-1)}) \stackrel{d}{=} \mathcal{A}_t(\tilde{S} | r^{(t-1)})$. In the former case, we can perform the
512 decomposition using Bayes' rule:

$$\begin{aligned}
\log \left(\frac{\mathbb{P}(\mathcal{A}_t(S) = b^{(q)})}{\mathbb{P}(\mathcal{A}_t(\tilde{S}) = b^{(q)})} \right) &= \underbrace{\log \left(\frac{\mathbb{P}(\mathcal{A}_t(S)[j+1:q] = b^{(j+1:q)} \mid b^{(j)})}{\mathbb{P}(\mathcal{A}_t(\tilde{S})[j+1:q] = b^{(j+1:q)} \mid b^{(j)})} \right)}_{(a)} + \underbrace{\log \left(\frac{\mathbb{P}(\mathcal{A}_t(S)[j] = b_j \mid b^{(j-1)})}{\mathbb{P}(\mathcal{A}_t(\tilde{S})[j] = b_j \mid b^{(j-1)})} \right)}_{(b)} \\
&\quad + \underbrace{\log \left(\frac{\mathbb{P}(\mathcal{A}_t(S) = b^{(j-1)})}{\mathbb{P}(\mathcal{A}_t(\tilde{S}) = b^{(j-1)})} \right)}_{(c)} \\
&\leq \varepsilon_{\text{thr}}
\end{aligned}$$

513 Notice that every two bins are disjoint; thus, the consumption of the privacy budget is independent
514 between two different data points. Therefore, (a), (c) = 0 and (b) $\leq \varepsilon_{\text{thr}}$ by the privacy guarantee
515 for single-query release.

516 Now we consider the RAG step. The interesting regime is when $z_i \in A'_t$. Then, by the composition
517 theorem, the privacy loss of DP-RAG $\circ$ TOP-K is upper bounded by ε_{RAG} .

518 In addition, we note that $\mathcal{E}(z_i)$ is a valid stopping time, since it updates the privacy budget after each
519 invocation of the algorithms, and z_i is only involved when its privacy budget is sufficient.

520 Thus, by Feldman and Zrnic [2021, Corollary 3.3], the privacy guarantee is given by $\mathcal{E}(z)$, which is
521 upper bounded by ε . $\square$