

Rethinking Searchable Symmetric Encryption

Conference Paper**Author(s):**

Gui, Zichen; Paterson, Kenneth G.; Patranabis, Sikhar

Publication date:

2023

Permanent link:

<https://doi.org/10.3929/ethz-b-000564585>

Rights / license:

[Creative Commons Attribution-NonCommercial-ShareAlike 4.0 International](#)

Originally published in:

<https://doi.org/10.1109/SP46215.2023.10179460>

Rethinking Searchable Symmetric Encryption

Zichen Gui
ETH Zürich
zichen.gui@inf.ethz.ch

Kenneth G. Paterson
ETH Zürich
kenny.paterson@inf.ethz.ch

Sikhar Patranabis
IBM Research India
sikhar.patranabis@ibm.com

Abstract—Symmetric Searchable Encryption (SSE) schemes enable keyword searches over encrypted documents. To obtain efficiency, SSE schemes incur a certain amount of leakage. The vast majority of the literature on SSE considers only leakage from one component of the overall SSE system, the *encrypted search index*. This component is used to identify which documents to return in response to a keyword query. The actual fetching of the documents is left to another component, usually left unspecified in the literature, but generally envisioned as a simple storage system matching document identifiers to encrypted documents.

This raises the question: do SSE schemes actually protect the security of data and queries when considered from a system-wide viewpoint? We answer this question in the negative. We do this by introducing a new inference attack that achieves practically efficient, highly scalable, accurate query reconstruction against end-to-end SSE systems. In particular, our attack works even when the SSE schemes are built in the natural way using the state-of-the-art techniques (namely, volume-hiding encrypted multi-maps) designed to suppress leakage and protect against previous generations of attack.

A second question is whether the state-of-the-art leakage suppression techniques can instead be applied on a system-wide basis, to protect both the encrypted search index and the encrypted document store, to produce *efficient* SSE systems. We also answer this question in the negative. To do so, we implement SSE systems using those state-of-the-art leakage suppression methods, and evaluate their performance. We show that storage overheads range from $100\times$ to $800\times$ while bandwidth overheads range from $20\times$ to $100\times$, as compared to a naïve baseline system.

Our results motivate the design of new SSE systems that are designed with system-wide security in mind from the outset. In this regard, we show that one such SSE system due to Chen *et al.* (IEEE INFOCOM 2018), with provable security guarantees based on differential privacy, is also vulnerable to our new attack.

In totality, our results force a re-evaluation of how to build end-to-end SSE systems that offer both security and efficiency.

I. INTRODUCTION

Searchable Symmetric Encryption. Database encryption is a key enabler for secure storage-as-a-service, wherein clients can securely outsource the storage and processing of large databases to (potentially untrusted) third party servers. Searchable symmetric encryption (SSE) [1]–[4] is a special subclass of database encryption that aims to efficiently support search queries over symmetrically encrypted databases. The core functionality enabled by SSE is the following: given an encrypted document collection in which each document is tagged with keywords, find the set of all documents tagged with a given keyword w . In this paper, we focus primarily on SSE for *static* document collections. This has historically received the most attention.

Leakage. The term “leakage” is popularly used in the SSE literature to denote any information that the server learns about either the database itself or the queries made by the client. For any SSE scheme, leakage is of two kinds: *setup leakage* – information learnt by the server from the encrypted database at setup (i.e., prior to any query execution), and *query leakage* – information learnt by the server from the query token and the interaction between the token and the encrypted database.

Leakage Cryptanalysis. A natural question to ask is: *how should we assess the impact of leakage on the real-world security of SSE?* The practice commonly adopted in the SSE literature is to perform *leakage cryptanalysis*. This involves developing concrete cryptanalytic attacks that exploit the leakage to subvert some security guarantee (such as data/query privacy) of the SSE scheme. Starting with the seminal work of Islam *et al.* [5], leakage cryptanalysis has been studied extensively in the context of SSE for document collections [6]–[13]. The commonly studied leakage profiles for SSE are:

- **Response Length.** For a given query on a keyword w , the response length (or volume) leakage reveals the size of the query response set, i.e., number of documents containing w .
- **Access Pattern.** For a given query on a keyword w , the access pattern leakage reveals the set of (potentially randomized) identifiers for documents containing w .
- **Co-occurrence Pattern.** For a pair of queries over keywords w_i and w_j , the co-occurrence leakage reveals the number of documents containing *both* w_i and w_j .
- **Search Pattern.** For a pair of queries over keywords w_i and w_j , the search pattern leakage reveals whether w_i and w_j are identical.

Structure-Only SSE vs End-to-End SSE. The *structure-only* approach to designing SSE schemes was introduced and formalized by Chase and Kamara in [3]. In this approach, a search query is broken down into two phases. The first phase, called the index retrieval phase, uses a specially designed *encrypted search index* to efficiently recover the set of (encrypted) document identifiers corresponding to documents matching the query. The second phase is the document retrieval phase, in which the client *actually* fetches the encrypted documents matching the query.

We illustrate the structure-only approach to SSE in Figure 1. The inner dotted box in Figure 1 depicts a structure-only SSE sub-system, while the outer box depicts the SSE system-as-a-whole. In the rest of the paper, we use the term *structure-*

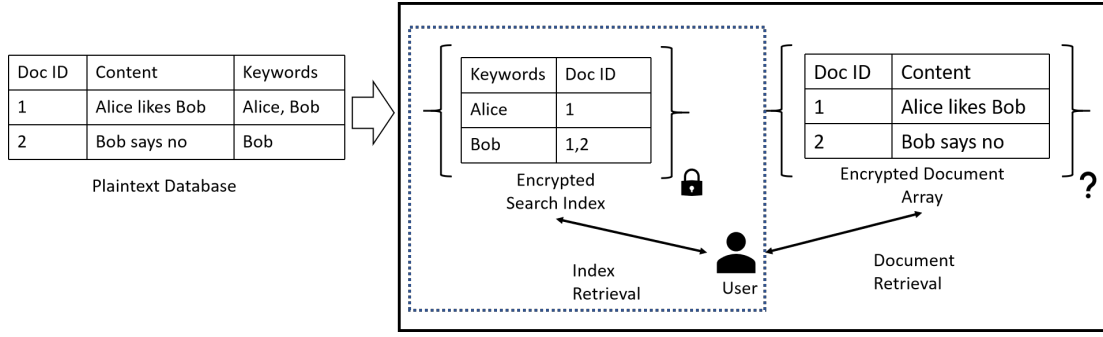


Figure 1: A “system-wide” view of structure-only SSE schemes (e.g., [2]–[4], [14]–[25]). While structure-only SSE design focuses only on leakage from the encrypted search index (depicted by the inner dotted box), end-to-end SSE design focuses on system-wide leakage (depicted by the outer solid box).

only SSE to refer to the sub-system depicted by the inner dotted box, and *end-to-end SSE* to refer to the system-as-a-whole depicted by the outer box. As is standard in the SSE literature [2]–[4], [14]–[23], we assume that in a structure-only SSE scheme, upon issuing a single keyword query, the client retrieves the document identifiers corresponding to every document that contains the keyword.

As an example of an end-to-end SSE system where it is essential to retrieve every document matching a queried keyword (rather than a subset of those documents), consider a medical database where each document contains all health information of an in-patient (e.g. blood pressure over his/her stay) and the user is interested in the relationship between patterns in blood pressure and mortality risks for a certain patient group. The best way to perform this analysis is to use the patient group as a keyword to search over the database and retrieve all relevant documents. The documents can then be analyzed locally together.

Existing SSE Schemes. The vast majority of SSE schemes follow the structure-only approach. Examples include [2]–[4], [14]–[23]. These all focus only on designing a secure encrypted search index (i.e. the inner dotted box in Figure 1). None of them concretely specify how the document retrieval phase is to be executed. And, in each case, their security analysis considers only leakage from the index retrieval phase, ignoring any leakage from the document retrieval phase. However, structure-only SSE schemes are not useful on their own unless they are used in conjunction with secure, efficient mechanisms for document retrieval.

Unfortunately, the research community’s focus on the structure-only approach has yielded little progress on the question of how to build secure, efficient, *end-to-end* SSE systems. There exist only a handful of alternative design approaches that target end-to-end SSE with in-built leakage suppression for both index and document retrieval. These include SSE with differentially private access-patterns [26] and SWiSSSE [8].

Leakage Suppression in SSE. Motivated by the need to counter leakage cryptanalysis, recently proposed SSE schemes have started to use dedicated techniques to suppress setup and search leakage. The state-of-the-art is represented by volume-

hiding encrypted multi-maps (EMMs) [24], [25]. These follow the same structure-only approach described above. In particular, while it is clear how to apply volume-hiding EMMs to search indices to obtain low-leakage, structure-only SSE schemes, the authors of [24], [25] do not specify how to design end-to-end SSE systems based on volume-hiding EMMs.

System-Wide Leakage in SSE. In this paper, we revisit the current dominant approach to leakage analysis for SSE by taking an alternative “system-wide” perspective. Consider two kinds of adversaries – one that observes the leakage from the encrypted search index (e.g., by passively corrupting the server that stores this index), and one that observes the leakage from the final query processing step on the encrypted documents themselves (e.g., by passively corrupting the server that stores the encrypted documents, potentially different from one storing the index). In terms of real-world security, any realization of end-to-end SSE should be secure against cryptanalysis of system-wide leakage.¹

Since the vast majority of SSE schemes, including state-of-the-art schemes based on volume-hiding EMMs, are structure-only schemes, they only focus on security against the first kind of adversary. Consequently, it is unclear what kind of system-wide leakage is incurred when these schemes are used to build end-to-end SSE systems, and whether such leakage potentially leads to attacks. It is also unclear whether volume-hiding EMMs can be used to mitigate this leakage without compromising on efficiency.

A. Our Contributions

System-Wide Security of Structure-Only SSE. In this paper, we analyze the system-wide leakage that arises when state-of-the-art structure-only SSE schemes are used in the natural way to build end-to-end SSE systems. In particular, we focus on the simple construction where only the leakage from the index is mitigated using techniques such as volume-hiding EMMs, while the document retrieval is implemented using

¹A widely used assumption in the SSE literature is that these two adversaries are non-colluding. Our analysis also extends to this setting since the attack we develop targets leakage from document retrieval only (in fact, for all of our target schemes, the leakage from document retrieval subsumes the leakage from index retrieval).

a straightforward look-up table of encrypted documents. In the absence of any statements to the contrary in the extensive literature, this seems to be the generally assumed mechanism. In this context, we ask the following:

Do structure-only SSE schemes result in secure end-to-end SSE systems when system-wide leakage is taken into account?

We answer this question in the negative. Concretely, We show that all the structure-only SSE schemes [2]–[4], [14]–[23], including those built from volume-hiding EMMs [24], [25], incur damaging system-wide leakage when used to construct end-to-end SSE systems in the natural way. In more detail, we show an efficient and highly scalable query recovery attack mechanism targeting the vast majority of existing end-to-end SSE schemes (including those using dedicated leakage-suppression techniques such as padding for the search index). Our attacks establish that given the (noisy) system-wide leakage from a sequence of keyword search results (with random non-results added) over a target database, and some additional information about the real results over an independently sampled auxiliary database, one can do query recovery for the target database despite the presence of leakage-suppression due to padding-like techniques.

Attack Target. For our proposed query-recovery attack, the target is an end-to-end SSE system (as depicted in the outer box in Figure 1). We note here that our attack does not target the original constructions of volume-hiding EMMs as they are described in [24], [25], since the authors of these papers only described how to use volume-hiding EMMs to design structure-only SSE schemes. Instead, we target natural and efficient extensions of these structure-only schemes to end-to-end SSE systems, and show that such extensions are rendered insecure by our proposed attack based on system-wide leakage. One might be concerned that, since [24], [25] did not propose end-to-end SSE systems, it is misleading to target specific end-to-end SSE systems built from their EMMs. We argue that the authors of [24], [25] are clear that SSE is the main application domain for their volume-hiding EMMs, and that the constructions of end-to-end SSE that we target are the natural, efficient ones based on the structure-only SSE schemes built from such EMMs.

System-Wide Access Pattern Leakage. We also point out that even though our target end-to-end SSE systems use a document retrieval phase that is implemented straightforwardly, the leakage from this phase is not necessarily the trivial (unperturbed) access pattern leakage. In particular, since document retrieval is performed based on the encrypted document identifiers from index retrieval, any leakage mitigation techniques used in index retrieval may perturb the leakage arising in document retrieval. Exploiting this leakage for query-recovery is non-trivial, and is a novel aspect of our attack.

Key Takeaway. Our attack establishes that using structure-only SSE schemes in the natural way leads to end-to-end SSE systems that are insecure in practice. Fundamentally, this is because the careful leakage mitigation for the encrypted search index, obtained using techniques such as volume-hiding

EMMs, is undermined by the system-wide leakage that arises from the document retrieval phase. The main message of our paper is not the attack itself, but the need for new security definitions that carefully model system-wide leakage in end-to-end SSE systems as well as the need for efficient techniques to mitigate such leakage. Our attack is primarily a means to an end that concretely establishes the adverse impact of system-wide leakage on the security of a vast majority of SSE schemes in the literature, all of which have overlooked the importance of building a secure and efficient document retrieval protocol in addition to index retrieval. We note here that there do exist prior works such as [8] that have also noted the importance of system-wide leakage and designed end-to-end SSE systems to prevent such leakage. However, their treatment is limited to very specific schemes, while our analysis applies to any end-to-end SSE system built in a natural way from any index-only SSE scheme (we expand more on this subsequently).

System-Wide Leakage Mitigation. A natural approach to mitigating system-wide leakage in end-to-end-SSE systems would be to apply existing leakage suppression techniques such as volume-hiding EMMs [24], [25] to *both* the encrypted index *and* the final document retrieval step. Intuitively, this enhances resistance to system-wide attackers, but potentially degrades efficiency. In this context, we now ask the following question:

Can existing leakage suppression techniques efficiently mitigate system-wide leakage?

We also answer this question in the negative. We demonstrate that it is practically infeasible to use volume-hiding EMMs [24], [25] to hide system-wide leakage in structure-only SSE schemes. We validate this observation by experimentally evaluating the storage and query processing overheads incurred when applying EMMs to encrypt the whole database, and not just the search index. We also experiment with using structure-only SSE in conjunction with state-of-the-art ORAM and PIR schemes to protect the document retrieval, with similar results.

To summarize, we concretely establish that, with currently available techniques, we can have either efficient and scalable structure-only SSE schemes that suffer from damaging system-wide leakage (as illustrated by our attack), or we can have end-to-end leakage-protected SSE systems that are inefficient in practice and do not scale to large databases.

Differentially Private Access Patterns. As a final contribution, we show that the end-to-end SSE system based on differentially private access patterns in [26], which we henceforth call **DPAP-SE**, is also vulnerable to our new query recovery attack. Our attack breaks **DPAP-SE** for the *same* parameter set that the authors of [26] advocated using to counter leakage-abuse attacks. While it is possible to degrade our attack’s efficiency by altering the parameter set, this also greatly reduces the practical efficiency of the resulting scheme. Our attack does not invalidate the original security properties proven by the authors of [26], but instead indicates that these security properties (and the corresponding usage of differential

privacy techniques) are, in fact, insufficient in practice. In fact, the scheme in [26] *does* resist naïve adaptations of existing cryptanalytic attacks [5], [6], [9], [27] that work only for *unperturbed* leakage such as co-occurrence patterns and access patterns, as was claimed in [26]. However, we demonstrate the possibility of a stronger attack that is *not* restricted to unperturbed leakage and, hence, bypasses the limitations of existing attacks.

Putting everything together, we see that, while significant progress has been made over the last decade in designing secure structure-only SSE schemes, none of these schemes actually yield practically efficient end-to-end SSE systems with resistance to system-wide leakage attacks. At the same time, we show that one alternative design approach for end-to-end SSE design also suffers from damaging system-wide leakage. Indeed, the only candidate end-to-end SSE scheme that we do not break is SWiSSSE [8]. This is perhaps not surprising since, aside from [26], it is (as far as we are aware) the only scheme that was designed with system-wide leakage in mind.

We expand further on our main contributions below.

B. Attack on System-Wide Leakage

We show that the vast majority of structure-only SSE schemes, including state-of-the-art schemes built from volume-hiding EMMs [24], [25], incur damaging system-wide leakage when used in the natural way to build end-to-end SSE systems. We show that such system-wide leakage can be exploited to launch a query recovery attack. We experimentally validate the practical efficiency of our attack using the Enron email corpus,² which is widely regarded as the standard choice of experimental database in the SSE literature [5], [9], [10].

Modeling and Attacking System-Wide Leakage. In a prior work, Gui *et al.* [8] discussed system-wide leakage in SSE schemes, and its impact on the security guarantees of SSE schemes. However, they only presented a system-wide analysis of their own construction, SWiSSSE and did not address the impact of system-wide leakage for SSE schemes more generally. In this paper, we refine and extend the analysis of [8] to develop a full-fledged query reconstruction attack. We apply the attack against end-to-end SSE schemes built in the natural way from different variants of volume-hiding EMMs, including **PRT-EMM** [24], as well as **FP-EMM** and **DP-EMM** [25]. Our attack works for volume-hiding EMM implementations using the same parameters and the same design/implementation choices advocated in [24] and [25].

Attack Overview. The core principle of our attack is as follows: we show how to statistically model the noisy co-occurrence leakage pattern arising in SSE schemes using the different volume-hiding EMMs as a function of the keyword queries, the system-wide leakage, and auxiliary data in the form of an approximate version of the original database. We then use the resulting models to develop a new inference-style leakage-abuse attack that targets query reconstruction.

Note that for our attacks, the target database and the auxiliary database are sampled independently from the same distribution. This makes our attack “inference-style”, and constitutes a weaker attack assumption as compared to “known-data attacks” (e.g. [9,15,17,27]), where the auxiliary and target databases are identical. The precise sampling strategy may be summarized as follows: given a fixed target dataset for the attack (Enron for our attack experiments), we set aside half of this dataset for use as the attack target, i.e. for generating the leakage, while using sub-samples of the other half to define auxiliary distributions. In this way, we effectively use independent samples from the same empirical distribution to define the target and auxiliary data distributions. This is stated formally in Section V-D (with additional details in the full version of our paper [28]).

The core idea of our attack is to solve an optimization problem, where the objective function is the formal *likelihood* of observing a given assignment of keywords to queries, given the observed leakage and auxiliary data as prior information. We then maximise the objective function using simulated annealing; this corresponds to maximising the likelihood of the solution. Thus the simulated annealing, if it works, will produce “good” solutions in which many keywords in the solution are correctly assigned to queries. This approach requires careful mathematical analysis to derive the likelihood functions for each targeted scheme and to efficiently implement their evaluation on large sets of queries and leakage. Of course, one could also use any performant optimization technique in place of simulated annealing.

Using experiments, we show that our attacks are effective even if: (a) aggressive security parameters are used by the underlying schemes; (b) the keyword universe of the auxiliary information is a lot bigger than the set of queried keywords; and (c) the auxiliary information (independent from the testing dataset) is very noisy. We apply computational optimization and approximation techniques to make our attacks highly scalable with respect to the number of queried keywords, number of keywords in the auxiliary information, and the size of our target dataset.

Experimental Evaluation. In Section V, we present extensive experimental evaluations to validate the practicality of our proposed attack. Our experiments show that our attacks achieve high success rate with reasonable practical efficiency even if: (a) the target SSE scheme uses aggressive security parameters (beyond those advocated for the employed EMMs); (b) the keyword universe of the auxiliary information is significantly larger than the set of queried keywords (in contrast to existing attacks); and (c) the auxiliary information available to the adversary is very “noisy”. Our attacks achieve over 60% query reconstruction rate in most of the settings we have tested. These include settings where the target SSE schemes use aggressive security parameters even beyond those originally proposed, as well as settings where the auxiliary information available to the adversary is significantly perturbed. With accurate auxiliary information, our attacks can

²<https://www.cs.cmu.edu/~enron/>

achieve over 80% query reconstruction rate. Beyond query reconstruction, our attacks are also capable of (approximate) database reconstruction in certain cases (see Section IV-C for more details).

Ethics Discussion. We believe that it is justifiable to use the Enron dataset for our attack experiments given: (a) the lack of alternative realistic datasets for conducting experiments for our attack, (b) the extensive usage of this dataset for attack experiments in prior works [5], [9], [10], (c) the fact that the data has long already been public, and (d) there has been an effort by the researchers curating the version of the dataset used in this paper to remove users upon request.³

Comparison with Existing Attacks. We note that most well-known leakage-abuse attacks (e.g. [5], [6], [9], [12], [27]) rely on stronger models of auxiliary data: they are essentially *known-data* attacks, where the adversary knows a subset of the entries in the original database. On the other hand, our attack is an *inference* attack, where the auxiliary data about the database that is available to the adversary is not sampled directly from the target database (as was done in prior known-data attacks [5], [6], [9], [12], [27]), but from an auxiliary distribution that is statistically “close” to the distribution of the target database. Having independent training and testing datasets is a weaker (and more realistic) assumption in practice. We also note that the IKK attack [5] and the Count attack [6], [9], [27] are practically infeasible in inference mode, while the attacks due to Blackstone *et al.* [9] are known-data attacks by design and do not work (even theoretically) in inference mode.

Nearly all of these existing leakage-abuse attacks [5], [6], [9], [27] on SSE also do not scale efficiently when attempting query-recovery over large query histories. As an illustration, the IKK attack [5] (with a serial implementation) takes up to 14 hours to finish one run in the setting where there are only 150 queried keywords and 2500 keywords in the auxiliary information. Similarly, the state-of-the-art analysis due to Blackstone *et al.* in [9] are based on the Enron dataset [34] for only 150 queries. Recently, Kamara *et al.* [35] introduced LEAKER – an open-source framework for evaluating the impact of leakage-abuse attacks on a variety of real-world datasets. Using this framework, the authors of [35] showed that, when additionally provided with real query logs for certain datasets, the attack of [9] achieves improved query recovery rates than those originally reported by the authors of [9], while the attacks in [5], [6], [27] achieved lower query recovery rates than originally reported. However, it was open to design leakage-abuse attacks that efficiently scale to large query histories prior to our work. We solve this open question; for a comparison, our attack against the full-padding EMMs [25] averages 8.7 minutes per run when there are 750 queried keywords and 4000 keywords in auxiliary information.

Pouliot and Wright [7] proposed an inference-style graph matching attack, which is essentially an improved version

of the IKK attack [5] that exploits co-occurrence pattern leakage, albeit under the strong assumption that the auxiliary and target datasets are highly correlated, i.e. the leakage is only weakly perturbed. As shown by Oya and Kerschbaum in [10], the attack efficiency degrades significantly when the leakage perturbation is high (which is expected to be the case in practice). Our attacks, on the other hand, achieve highly efficient query recovery even when the leakage perturbation is high (captured by the usage of highly noisy auxiliary information in our attack experiments).

Oya and Kerschbaum [10] proposed maximum likelihood estimation-based inference-style leakage-abuse attacks against SSE that simultaneously exploit the *exact* search pattern leakage and (a sub-component of) noisy access pattern leakage. Our attacks exploit a strictly weaker leakage profile as compared to their attacks (noisy access pattern leakage as opposed to search pattern+noisy access pattern leakage); in particular, our attacks would work even if our target end-to-end SSE schemes used mechanisms to additionally hide the search pattern leakage. Similarly, other existing inference-style leakage-abuse attacks against SSE [11], [13] rely on additional requirements such as auxiliary information about the target leakage and/or known queries, which our attacks avoid entirely (and thus work in a weaker attack setting).

Gui *et al.* [8] recently presented an inference-style system-wide leakage based attack; however, their attack techniques are specifically designed for cryptanalyzing their own construction called SWiSSE and cannot directly be applied to other SSE schemes. They require significantly stronger assumptions on the auxiliary leakage available to the adversary compared to our attack. Our attack extends and refines the attack approach in [8] by: (a) encompassing a more realistic system-wide leakage profile that is typically encountered in all end-to-end SSE systems built in a natural way from structure-only SSE schemes, and (b) making less ideal assumptions about the attack setting and availability of auxiliary data. The “core” idea of our attack is essentially the same as that of [8], and in fact, is based on well-known principles that have been extensively used before by cryptanalysts and in statistics: try to find an assignment (in this case, keywords to query tokens, but in general, a parameter of a distribution that one is trying to optimize) that maximizes the likelihood of the assignment (where the likelihood is a formal probabilistic measure expressed as a conditional probability). The main technical novelty of our analysis lies in the capability to perform query-recovery given significantly noisier system-wide leakage from a wider class of state-of-the-art SSE-systems (as opposed to a particular scheme in [8]), and the ability to scale efficiently to larger query-histories with superior query-recovery rates.

Table I compares our attack to the existing ones.

C. On Using EMMs to Hide System-Wide Leakage

We also address the question of whether existing leakage suppression techniques can be efficiently applied to *both* the encrypted index *and* the final document retrieval step in order to mitigate system-wide leakage in end-to-end SSE systems.

³See <https://www.cs.cmu.edu/~enron/> for detailed documentation.

Attack	Attack Target	Attack Setting	Leakage Exploited	Additional Requirements
IKK [5]	Naïve	Known-data	Co-occurrence pattern	Known queries
Count* [6]	Naïve	Known-data	Co-occurrence pattern	Known queries
BKM20 [9]	Naïve	Known-data	Co-occurrence/access pattern	—
LEAP [12]	Naïve	Known-data	Co-occurrence pattern	—
Graph Matching Attacks [7]	[29], [30]	Inference	Co-occurrence pattern	—
SAP [10]	[25], [26], [31]	Inference	Search and volume pattern	Auxiliary info on query frequency pattern
GPPW20** [8]	[8]	Inference	Co-occurrence pattern	—
IHOP [11]	[26], [32], [33]	Inference	Search/volume/co-occurrence pattern	Auxiliary info on the target leakage(s)
DHP21 [13]	Naïve	Inference	Co-occurrence pattern	Known queries
This work	[24]–[26]	Inference	Co-occurrence pattern	—

Table I: Comparison of existing leakage-abuse attacks achieving query reconstruction based on co-occurrence and/or access pattern leakage. Naïve attack target refers to a scheme that leaks the full access-pattern (see Section VI for a more formal treatment). All cited attack targets have some form of leakage perturbation. *Count attack does not need known queries if the entire database is known by the attacker; known queries are only helpful when only part of the database is known by the attacker. **Gui *et al.* [8] also presented a system-wide leakage based attack; however, their attack techniques are specifically designed for cryptanalyzing their own construction called SWiSSE (unlike our attacks, their attack cannot be generically used to break the vast majority of SSE schemes in the literature), and require significantly stronger assumptions on the auxiliary leakage available to the adversary compared to our attacks.

We demonstrate that it is practically infeasible to use volume-hiding EMMs, including the **PRT-EMM** scheme of [24], as well as the **FP-EMM** and **DP-EMM** schemes of [25], to hide system-wide leakage. We validate this observation by experimentally evaluating the concrete storage overheads and query processing overheads incurred when applying these state-of-the-art EMM techniques to encrypt the whole database, and not just the search index. We also study the overheads when using state-of-the-art ORAM and PIR schemes to do this.

Our experiments show that applying the volume-hiding EMM schemes from [24], [25] to encrypt the whole database incurs prohibitively expensive storage overheads, ranging from 100 to 800 times the size of a naïve encrypted document store. Our experiments also establish that the volume hiding EMM schemes in [24], [25] incur additional computational and communication overheads during query execution due to their usage of padding techniques. While this cost is manageable when querying the search index alone, it blows up to impractical proportions if applied directly to encrypted document retrieval, ranging from 20 to 100 times the communication costs for the naïve solution. Our volume-hiding EMM implementations use the same parameters and the same design/implementation choices advocated in [24] and [25].

II. PRELIMINARIES

A. Syntax of Searchable Symmetric Encryption (SSE)

Let $\mathcal{T}$ be an abstract data type supporting query operation **Query**. Then, an SSE scheme Σ for $\mathcal{T}$ is a tuple $\Sigma = (\text{Setup}, \text{Query}_e)$ where:

- **Setup** is the setup algorithm (locally executed by the client) which takes as input a plaintext database $\mathbf{D}$ of structure $\mathcal{T}$, and outputs a secret key sk and an encrypted database $\mathbf{ED}$.
- **Query_e** is the query execution protocol between the client and server. The client takes as input a secret key sk and a query q (this could be a single keyword or a Boolean formula over multiple keywords), and the server takes as input the encrypted database $\mathbf{ED}$; after the interaction between the client and server, the client obtains a final result rsp , which is the set of (encrypted) documents matching the query q .

Correctness. We say that Σ is correct if for any database $\mathbf{D}$ and any query q , an execution of the query protocol **Query** on the query q and the encrypted data $\mathbf{ED} \leftarrow \text{Setup}(sk, \mathbf{D})$ yields the same response as a direct execution of the query q on the plaintext database $\mathbf{D}$.

Security. The security of an SSE scheme is defined formally with respect to a *leakage profile*, which is an upper bound on the information about the plaintext data and queries that an attacker can learn from the encrypted database $\mathbf{ED}$ and subsequent executions of **Query_e**. We refer the reader to [2], [4] for the formal security definition.

B. Notations

For the rest of our paper, we model a database in a searchable encryption scheme as target abstract data type. A database $\mathbf{DB}$ consists of a set of documents $\mathbf{d}_i$, each associated to a set of keywords $\text{kw}(\mathbf{d}_i)$, so $\mathbf{DB} = \{(\mathbf{d}_i, \text{kw}(\mathbf{d}_i))\}$. It supports keyword search queries. For a keyword search query q on the keyword $\text{kw}(q)$, the set of documents containing this keyword, denoted as $\mathbf{DB}(\text{kw}(q)) = \{\mathbf{d}_i \mid \text{kw}(q) \in \text{kw}(\mathbf{d}_i)\}$, is returned. To emphasize that we are only considering keyword search queries, we denote the query protocol as **Srch**.

III. FORMAL DESCRIPTION OF QUERY RECONSTRUCTION ATTACKS USING CO-OCCURRENCE LEAKAGE

This section formally establishes the leakage profile for SSE systems that we are targeting in the paper. We also formally describe our attack setting.

Access-pattern Leakage. Access-pattern leakage refers to the information leakage associated to *document retrieval*. If a naïve searchable encryption scheme Σ leaks the “exact” access pattern and nothing else, we can write the leakage of a query q on document collection $\mathbf{DB}$ as:

$$\mathcal{L}_{\text{Srch}}(q, \mathbf{DB}) = \{i \mid \text{kw}(q) \in \text{kw}(\mathbf{d}_i), (\mathbf{d}_i, \text{kw}(\mathbf{d}_i)) \in \mathbf{DB}\},$$

where $\text{kw}(q)$ denotes the keyword associated to query q .

Co-Occurrence Leakage. Access-pattern leakage from different queries can be represented equivalently as a matrix, known as co-occurrence pattern. Consider a small document collection $\mathbf{DB}$ where

$$\mathbf{DB} = \{(\mathbf{d}_1, \{\text{kw}_1, \text{kw}_2, \text{kw}_3\}), (\mathbf{d}_2, \{\text{kw}_1, \text{kw}_2\}), (\mathbf{d}_3, \{\text{kw}_3\})\}.$$

Let q_i be a query on keyword kw_i . If the original access pattern is leaked, we know that

$$\mathcal{L}_{\text{Srch}}(q_\ell, \mathbf{DB}) = \{1, 2\} \text{ for } \ell \in \{1, 2\}, \mathcal{L}_{\text{Srch}}(q_3, \mathbf{DB}) = \{1, 3\}.$$

This allows us to take intersections between the leakages as:

$$\begin{aligned} \mathcal{L}_{\text{Srch}}(q_1, \mathbf{DB}) \cap \mathcal{L}_{\text{Srch}}(q_2, \mathbf{DB}) &= \{1, 2\}, \\ \mathcal{L}_{\text{Srch}}(q_\ell, \mathbf{DB}) \cap \mathcal{L}_{\text{Srch}}(q_3, \mathbf{DB}) &= \{1\} \text{ for } \ell \in \{1, 2\}. \end{aligned}$$

The cardinality of the intersections can be very useful in an attack. For example, the co-occurrence pattern of the document collection above can be represented as a *co-occurrence matrix* $\bar{M}$ defined as follows:

$$\bar{M}(q_1, q_2, q_3; \mathbf{DB}) = \begin{bmatrix} 2 & 2 & 1 \\ 2 & 2 & 1 \\ 1 & 1 & 2 \end{bmatrix},$$

where the (i, j) -th entry of the matrix is

$$|\mathcal{L}_{\text{Srch}}(q_i, \mathbf{DB}) \cap \mathcal{L}_{\text{Srch}}(q_j, \mathbf{DB})|.$$

If we know the underlying document collection perfectly, we can re-identify q_3 as a query on kw_3 as it is the only keyword that only shares one document with other keywords. This qualifies as a query reconstruction attack.

We note that a co-occurrence matrix contains strictly less information than the original access-pattern leakage as the information on intersections of more than two queries are removed. However, the co-occurrence matrix is often sufficient in attacks so it is used instead of the full leakage. We refer to the co-occurrence matrix as the co-occurrence leakage.

There are three complications to the representation of co-occurrence leakage in practice. Firstly, the schemes we consider in practice usually leak query equality pattern too. That is, if $kw(q_i) = kw(q_j)$, the attacker knows that the two queries are for the same keyword. In terms of co-occurrence leakage, we use only one of the queries in the representation to simplify the problem. Secondly, the queries are unordered in practice. That means there is no standard representation of the leakage in terms of the known keywords. We use the convention that the i -th row and column of the co-occurrence matrix corresponds to the i -th non-repeating query in our representation. Finally, not all schemes leak the original access pattern and some schemes may even be randomised. In those cases, we need to use a suitable representation of the co-occurrence information, which may differ from what we have described above.

Auxiliary Information. Similar to a co-occurrence matrix, the auxiliary information the attacker receives can be represented as a co-occurrence matrix M . The co-occurrence matrix is indexed by the known keywords and typically contains full information on all keywords. In stronger attacks, M is assumed to be noisy in the sense that it is not generated directly from the target document collection. Instead, an auxiliary dataset is used for the purpose.

Let $\mathbf{DB}_{\text{aux}} = \{(d'_i, kw(d'_i))\}$ be an auxiliary document collection with keywords $\{kw'_1, \dots, kw'_n\}$. In our attack, the (i, j) -th entry of M represents the empirical probability (derived from the auxiliary data $\mathbf{DB}_{\text{aux}}$) of seeing kw'_i and kw'_j together in a document, computed as:

$$M_{i,j} = |\{d'_i \mid kw'_i \in kw(d'_i) \wedge kw'_j \in kw(d'_i)\}| / |\mathbf{DB}_{\text{aux}}|,$$

where $|\mathbf{DB}_{\text{aux}}|$ is the total number of auxiliary documents.

Attack Settings. There are two attack settings we consider in this paper. In the first attack setting, we consider natural end-to-end SSE schemes built from structure-only schemes (see Section IV-A). We exploit the co-occurrence leakage from retrieval of *actual documents* to achieve query reconstruction. Alternatively, if the target SSE scheme is already an end-to-end one, we target the co-occurrence leakage from retrieval of actual documents directly to achieve query reconstruction.

We can describe a query reconstruction attack formally as follows. Let queries $q_1, \dots, q_l$ be a sequence of queries on the document collection, so the attacker observes co-occurrence leakage $\bar{M}(q_1, \dots, q_l; \mathbf{DB})$. Suppose that the attacker has access to some auxiliary information M . *The goal of the attacker is to recover $kw(q_i)$ after observing the co-occurrence leakage $\bar{M}$ and knowing auxiliary information M .*

IV. NEW QUERY RECONSTRUCTION ATTACKS USING CO-OCCURRENCE LEAKAGE

In this section, we show that all the structure-only SSE schemes [2]–[4], [14]–[23], including those built in a natural way from volume-hiding EMMs [24], [25], incur damaging system-wide leakage when used to construct end-to-end SSE systems in the natural way. In particular, we propose new inference-style query recovery attacks that completely break the query privacy guarantees of these systems. We additionally show a new query recovery attack on the end-to-end SSE system **DPAP-SE** (based on differentially private access patterns) proposed in [26].

Attack Overview. We develop an inference-style leakage-abuse attack using a combination of statistical modeling and simulated annealing. We assume that the attacker has access to noisy co-occurrence leakage from observed queries and auxiliary data that is independent of but statistically “close” to the target database. We model the attack as an optimization problem, where the objective function is the statistical likelihood of observing a given assignment of keywords to queries, given the observed leakage and auxiliary data as prior information. Then maximizing the objective function corresponds to maximising the statistical likelihood of the solution. In this way, we obtain optimal assignments of keywords to queries, given the available leakage and auxiliary data. We use simulated annealing for the maximization step because it is relatively easy to implement and performs well in practice, although one could use any other suitable optimization technique.

A. Attack Targets

As mentioned earlier, we target system-wide leakage from end-to-end SSE systems constructed in a natural way from

structure-only SSE schemes, which are in turn built from volume-hiding EMMs [24], [25]. We begin by recalling the natural way to design end-to-end SSE from EMMs.

End-to-End SSE from EMMs. At a high level, a multi-map is a set of key-value pairs, i.e. $\{(\text{key}_i, \vec{v}_i)\}$, where key_i 's are the keys which are assumed to be non-repeating and $\vec{v}_i$'s are the tuples of values associated to their keys. The natural way to design a structure-only SSE scheme from an EMM is to realize an encrypted search index as follows: let each key of the EMM be a deterministic function of a keyword in the database, and let the corresponding value be the set of encrypted document identifiers pertaining to the documents matching this keyword. This is precisely the approach taken by all of the structure-only SSE schemes [2]–[4], [14]–[23] (while the exact specification/implementation of the EMM might vary from scheme-to-scheme depending on the nature of queries supported by the scheme, the overall approach is the same). Finally, the natural and widely accepted approach to transition from structure-only SSE to end-to-end SSE is to additionally realize document retrieval using a simple storage system matching document identifiers to encrypted documents. We refer the reader to [2]–[4] for a more formal exposition of the overall approach.

Our Targets. We target system-wide leakage from end-to-end SSE built in the aforementioned natural way from state-of-the-art volume-hiding EMMs [24], [25], which are specially designed EMMs equipped with additional cryptographic mechanisms for leakage suppression. Our attacks rely on the leakage from the document retrieval phase, which subsumes the leakage from the index-retrieval phase. It is important to note that even though the document retrieval phase is implemented as a simple storage system, the leakage from this phase is *not* the trivial (unperturbed) access pattern leakage. In particular, since document retrieval is performed based on the encrypted document identifiers from index retrieval, the leakage mitigation from the volume-hiding EMM also induces noise/perturbation in the leakage from document retrieval. Hence, we specifically design our attacks to be robust against such noisy leakage. Since the exact nature of the leakage depends on the specific realization of volume-hiding EMM used, we recall the schemes from [24], [25].

PRT-EMM. Our first target is an end-to-end SSE scheme realized naturally from the first construction of volume-hiding EMM via pseudorandom transform (abbreviated as **PRT-EMM** throughout) due to Kamara and Moataz [24]. We note here that our attack does not target the original construction of **PRT-EMM** in [24], but the natural and efficient extension of **PRT-EMM** into end-to-end SSE system. At a high level, the idea behind **PRT-EMM** is to pad or truncate the query response lengths of queries on any EMM with a pseudorandom function (PRF) as follows. Let key be a key for the multi-map and $F_{sk}(\cdot)$ be a PRF with key sk . The client computes: $n'_{\text{key}} = \lambda + F_{sk}(\text{key} || n_{\text{key}})$, as the new query response length, where λ is a free parameter which the client can choose and

n_{key} is the original query response length. These new query response lengths are used to build a multi-map on as follows:

- If $n_{\text{key}} \leq n'_{\text{key}}$, add $\perp$ symbols in the multi-map on key key before encryption.
- If $n_{\text{key}} > n'_{\text{key}}$, truncate the multi-map on keyword key to the first n'_{key} entries.

We note here that the original construction in [24] pads query responses with $\perp$ symbols if the real query response length is shorter. If the $\perp$ symbols are ignored in actual document retrieval, the attacker will be able to learn the true query response lengths when there is padding. In our attack, we assume the $\perp$ symbols are replaced by randomly picked indices, so the true query response lengths are not leaked. Due to lack of space, the detailed formal derivation of the precise leakage distribution for our target SSE scheme is deferred to Appendix A.

FP-EMM and DP-EMM. We next target end-to-end SSE scheme realized naturally from two constructions of volume-hiding EMM based on hash-tables presented in [25]. Once again, we note here that our attack does not target the original constructions of volume-hiding EMMs in [25], but the natural and efficient extension of these EMMs into end-to-end SSE systems. The first scheme (abbreviated as **FP-EMM** throughout) uses full padding, meaning that all query response lengths are padded to the maximum query response length. This is done by querying additional addresses in the hash-table deterministically (generated by a PRF) for each key.

The second scheme (abbreviated as **DP-EMM** throughout) uses differentially-private volume hiding as opposed to full padding. Let $2n_{\text{key}}$ be the true query response length of a query on key key (where the factor of 2 arises from the usage of two hash tables in Cuckoo hashing). Then the scheme pads the query response length to $2n_{\text{key}} + n^* + \text{Lap}_{sk}(2/\epsilon)$, where n^* is a parameter set by the client to offset the query response length in case the latter random variable is negative, and $\text{Lap}_{sk}(\cdot)$ is a Laplace distribution with secret key sk as the seed.

As in the case of **PRT-EMM**, we design attacks targeting “noisy” co-occurrence leakage during document retrieval in end-to-end SSE schemes realized naturally from **FP-EMM** and **DP-EMM**, respectively. The formal derivation of the precise leakage distribution for each of these SSE schemes is again deferred to Appendix A.

DPAP-SE. Finally, we also target the end-to-end SSE system **DPAP-SE** (based on differentially private access patterns) proposed in [26]. At a high level, this scheme uses a differential privacy mechanism is to “obfuscate” the plaintext database prior to encryption, such that a slight change in the real access pattern does not affect the obfuscated access pattern significantly. There are two key ingredients in their construction. First, an erasure code [36] is used to split every document into m shards, each with size $\frac{1}{k}$ of the original document. The erasure code has the property that any k shards of a document can be used to reconstruct the original document. The client then picks two probabilities p and q , and does the following to each shard:

- 1) For any keyword that is originally in the shard, remove the keyword with probability $1 - p$.
- 2) For any keyword that is originally not in the shard, add it to the shard with probability q .

For our attack, we rely on the observation that the resulting leakage from this scheme can be interpreted as follows: since the scheme effectively transforms query response into “noisy” keyword response lengths on the shards, the observed co-occurrence counts are actually “noisy” co-occurrence counts on the shards, rather than the actual documents. We then design an attack procedure that takes this noisy leakage on the shards into account. As in our earlier attacks, we again rely on precisely modeling this noisy co-occurrence leakage, albeit on the shards. Details can be found in Appendix A.

B. Attack Technique

We now detail our attack technique. For all of our attacks, we assume that the precise leakage modeling steps outlined above (and described formally in Appendix A) yield an observed co-occurrence matrix $\bar{M}$, which is the leakage from the observed queries on the actual target database. To simulate the auxiliary information available to the attacker, we build an auxiliary co-occurrence matrix M from an auxiliary database. This auxiliary database consists of a subset of uniformly randomly sampled documents from a second database (independent of the target database). To simulate different noise levels, we vary the number of such sampled documents – fewer documents means more noise. The level of noise determines how “close” the auxiliary co-occurrence matrix M is to the observed matrix $\bar{M}$.

Maximum Likelihood Estimation. Our attack is based on maximum likelihood estimation [37]. The attacker is given the observed co-occurrence matrix $\bar{M}$ as the leakage and auxiliary co-occurrence matrix M as prior information. He makes a choice of the assignment P that maximises the of *likelihood* of P given $\bar{M}$ and M .

We use maximum likelihood estimation as it is an optimal statistical approach – if the observed co-occurrence matrix is really drawn from the distribution associated to the auxiliary co-occurrence matrix, then the assignment which maximises the likelihood is, in a strict sense, the best possible guess which any adversary can make.

The derivations of the likelihood functions $\mathbf{L}[P | \bar{M}, M]$ are given in Appendix B.

Simulated Annealing. We use simulated annealing [38] as the algorithm to compute the assignment. At a high level, simulated annealing is a probabilistic technique for searching for the global optimum of a given function. It is very similar to a greedy search algorithm – randomize the input of the function, in our case, that is the assignment P , recompute the score, and if the score is larger than before, the assignment is kept as the new solution, and it is discarded otherwise – except that a worse solution is accepted in simulated annealing if it is not *too bad*. This is to prevent the algorithm from getting stuck in a local optimum. More concretely, simulated

Algorithm 1 Simulated Annealing

```

1: procedure ATTACK( $\bar{M}, M, T_0, i_{\max}$ )
2:    $P \leftarrow \text{InitPerm}(\bar{M}, M)$ 
3:    $T \leftarrow T_0$ 
4:    $s \leftarrow \text{Score}(\bar{M}, M, P)$ 
5:   for  $i \leftarrow 1, \dots, i_{\max}$  do
6:      $T \leftarrow \text{Cooling}(T, i)$ 
7:      $P' \leftarrow \text{Neighbour}(P, \bar{M}, M)$ 
8:      $s' \leftarrow \text{Score}(\bar{M}, M, P')$ 
9:     if  $\text{AccptProb}(T, s, s') > \text{rand}(0, 1)$  then
10:       $P \leftarrow P'$ 
11:       $s \leftarrow s'$ 
12:   return  $P$ 

```

annealing uses a *temperature* T which decreases per iteration and the differences between the current score of the target function and the previous best score maintained by the algorithm to compute an acceptance probability p , and with probability p the new solution is accepted. This probability is 1 if the new score is higher than the previous best, and less than 1 otherwise. For the same difference in the scores, a lower temperature T leads to a lower acceptance probability, which means simulated annealing behaves as a greedy search algorithm progressively.

Formally, simulated annealing consists of five subroutines, namely a function `InitPerm` to generate an initial assignment, a cooling scheme `Cooling`, a neighbourhood generation algorithm `Neighbour`, a function `Score` to compute the score and a function `AccptProb` to compute the acceptance probability. The syntax of the subroutines are defined below:

- `InitPerm`: takes as input an observed co-occurrence matrix $\bar{M}$ and an auxiliary co-occurrence matrix M , and outputs an assignment P .
- `Cooling`: takes as input a temperature T and the current iteration number i and outputs a new temperature T' .
- `Neighbour`: takes as input an assignment P , an observed co-occurrence matrix $\bar{M}$ and an auxiliary co-occurrence matrix M , and outputs a new assignment P' .
- `Score`: takes as input an observed co-occurrence matrix $\bar{M}$, a auxiliary co-occurrence matrix M and an assignment P , and outputs a score.
- `AccptProb`: takes a temperature T , a previous best score s and the new score s' , and outputs a probability.

The algorithm begins with an initial temperature T_0 and a random assignment P . An initial score s is computed on this assignment P . Then, the algorithm computes a new temperature $T \leftarrow \text{Cooling}(T_0, 1)$, find a new assignment P' using the neighbourhood function `Neighbour`($\cdot$), and compute a new score s' with the score function `Score`($\cdot$). An acceptance probability is computed as $p \leftarrow \text{AccptProb}(T, s, s')$. A random number between 0 and 1 is generated and if the random number is less or equal to p , the new solution s' is accepted by the algorithm and kept as the new optimum solution. This process is repeated until the maximum number of iteration is reached. See Algorithm 1 for the formal description.

Our Query Reconstruction Attacks. Our proposed query

reconstruction attacks on all of the attack targets described earlier use the aforementioned simulated annealing procedure for maximum likelihood estimation. In all our attacks, we use $T' \leftarrow 0.995T$ as our cooling scheme $\text{Cooling}(\cdot)$ and $p \leftarrow \exp(-\frac{s-s'}{T'})$ as our function $\text{AccptProb}(\cdot)$ to compute the acceptance probability. The precise score function $\text{Score}(\cdot)$ that we aim to maximise is the likelihood function $\mathbf{L}[P | M, M]$. We defer the formal description of how this function is derived for each target scheme to Appendix B.

We next describe our choices of the initial assignment finding subroutine $\text{InitPerm}(\cdot)$ and neighbourhood generation subroutine $\text{Neighbour}(\cdot)$ which have a significant impact on the performance and effectiveness of our attacks. The choices we make are tuned for maximizing the attack recovery rate.

InitPerm($\cdot$). An initial assignment finding subroutine $\text{InitPerm}(\cdot)$ is an efficient algorithm for guessing keywords/keys of the queries, so as to provide a starting point for the more expensive iterative steps later. For our attacks, only the query response lengths are used to avoid expensive computations. We observe that although the observed query response lengths are different from the true query response lengths for all of the schemes we target, these two are related. In particular, for **DP-EMM** [25] and **DPAP-SE** [26], we can compute the expected observed query response lengths from the query response lengths in the auxiliary co-occurrence matrix, and match the queries to the keywords in the auxiliary co-occurrence matrix as well as we can. For **PRT-EMM** [24] and **FP-EMM** [25], the observed keyword frequencies are independent from the true keyword frequencies, .

Neighbour($\cdot$). A neighbourhood generation subroutine generates new assignments for the attacks. Although a uniformly randomly picked assignment works all the time, it may not be the most efficient choice. In particular, for **DP-EMM** [25] and **DPAP-EMM** [26], we know that if an observed query response length is too far from the expectation, the assignment is very unlikely, and can be safely discarded. This means the neighbourhood generation subroutines for the attacks on these two schemes can make use of this, and output a new assignment only if it is sound. We note that these neighbourhood generation subroutines may prevent some correct assignments in the output of the attacks if their observed query response lengths are too far from the expected query response lengths. By relaxing the bounds, we can make the chance of that happening arbitrarily small. However, the algorithm would then be less efficient as more iterations are required for a convergence. Hence, we see our choice of bounds as a trade-off between query recovery rate and attack efficiency.

For **PRT-EMM** [24] and **FP-EMM** [25], we have to use uniformly randomly picked assignments, since the observed query response lengths are independent from the true query response lengths.

C. From Query Reconstruction to Database Reconstruction

It turns out that for some of the targeted schemes, our query reconstruction attacks can be extended to *database*

reconstruction attacks (by which we mean recovering the keywords associated with the encrypted documents).

DP-EMM. Our query reconstruction attack on the end-to-end SSE system built naturally from **DP-EMM** [25] implies a database reconstruction attack for practical security parameters; we rely on the noisy access pattern leakage during document retrieval to recover the *actual* keywords occurring in the encrypted documents retrieved across various queries.

Recall that in **DP-EMM**, the response length to a query on some **key** is padded to $(2n_{\text{key}} + n^* + \mathbf{Lap}_{sk}(2/\epsilon))$, where n_{key} is the real response length, and n^* is a fixed constant that depends solely on the choice of ϵ . For example, the authors of [25] suggest using $\epsilon = 0.2$, which yields $n^* = 567$. Our attack stems from the observation that the leakage perturbation is rather small for more frequent keywords, since the corresponding response length n_{key} is much larger than n^* . Hence, for the choice of parameters suggested in [25], our query reconstruction attack also allows us to progressively recover the set of keywords associated with a given document as more and more queries are observed, leading to a database reconstruction attack. The reconstruction rate is somewhat close to 50% due to the padding-factor of $2\times$ in **DP-EMM**, which causes half of the recovered keywords to be “fake”. While this attack can be prevented by altering the parameter ϵ for increased leakage perturbation, such alteration also significantly degrades the practical performance of the scheme. In other words, all practically efficient realizations of the natural end-to-end SSE from **DP-EMM** are broken by our database reconstruction attack.

Our attack does not immediately extend to the end-to-end SSE systems built naturally from **PRT-EMM** [24] and **FP-EMM** [25] due to the inherently different and somewhat larger perturbation to the access pattern leakage in these schemes, and we leave it as an open question to investigate database reconstruction attacks on such schemes.

DPAP-SE. For implementations of **DPAP-SE** [26] with reasonably practical parameters (e.g. $p = 0.89999$ and $q = 6.997 \times 10^{-6}$, where p and q are the parameters described earlier), our query reconstruction attack allows guessing the real keywords in the shards (obtained as part of the document retrieval process) with high probability (more concretely, probability p ; recall that p needs to be large for the scheme to satisfy correctness in practice). As in the case of **DP-EMM**-based end-to-end SSE, we rely on the noisy access pattern leakage for this database reconstruction attack. Again, while this attack can be prevented by altering the parameters p and q , such alteration also significantly degrades the practical performance of the scheme. Hence, all practically efficient realizations of **DPAP-SE** are broken by our database reconstruction attack.

V. EXPERIMENTAL EVALUATION

A. Overview of Experimental Setup

Experimental Data and Auxiliary Information. We use the Enron email corpus as the target dataset for all of our

attacks. A description of the dataset and our pre-processing step can be found in the full version of our paper [28]. A major challenge for inference-style leakage-abuse attacks is deciding an appropriate model for evaluating their effectiveness in practice. Such a model should take into account both the distribution of queries as well as the distribution of the auxiliary information available to the adversary. Unfortunately, there do not exist concrete guidelines in the literature for how to construct such models; given this lack of precedence, we make certain assumptions that we believe are reasonable in practice.

Query Distribution. We use uniformly distributed keyword queries to evaluate our attacks. This is exactly as in previous attacks [5], [6], [9]. We note here that our attacks do not explicitly depend on the distribution of queries; hence a uniform distribution appears to be a reasonable choice.

Auxiliary Data Distribution. For the IKK attack, Islam *et al.* [5] proposed a method to model auxiliary information in an inference-style attack setting; their suggestion was to use an auxiliary co-occurrence pattern leakage obtained by adding Gaussian noise to the original co-occurrence pattern. However, this implicitly assumes a homogeneous distribution of keywords amongst the documents, which may not always be the case in practice. Instead, we opt to split the overall dataset into two halves: out of the 480,000 documents in the dataset, half of the documents are used as the attack target and a fraction of the other half of the documents is used to generate auxiliary information about the dataset. Measurements of ‘noisiness’ of our auxiliary data can be found in the full version of our paper [28]. In total, we generate 10 different splits of the documents. For each split, we run 10 independent attacks with freshly generated observed co-occurrence matrices. We measure the fraction of correctly guessed keywords and report the average over the 100 runs as the query recovery rate.

Keyword Extraction. We extract keywords using the Natural language toolkit [39] in Python. Keywords with frequency higher than 5% are removed. We do not apply stemming as that does not affect query recovery rate in our experiments.

Keyword and Query Selection. We use the 1,000, 2,000, 3,000 and 4,000 most frequent keywords to build auxiliary co-occurrence matrices, and sample uniformly randomly without replacement from these most frequent keywords subsets of 250, 500, 750 keywords as queried keywords. These queried keywords are used to build observed co-occurrence matrices. These observed and auxiliary co-occurrence matrices are then used as the inputs to our attacks.

Security Parameter Selection for the Target Constructions. We use the security parameters suggested in the original papers to run our attacks. We also investigate how changes in the security parameters affect query reconstruction rates.

a) **PRT-EMM:** Recall that **PRT-EMM** from [24] allows the client to pick a public parameter λ which controls the padded query response lengths as: $n'_{\text{key}} = \lambda + F_{sk}(\text{key}||n_{\text{key}})$.

The authors suggested to set λ between 0 and $0.25n_{\text{max}}$. We used $\lambda = 0$ and $0.25n_{\text{max}}$ in our experiments. In addition, we used $\lambda = 0.5n_{\text{max}}$ to see the effect of additional padding on query reconstruction rate.

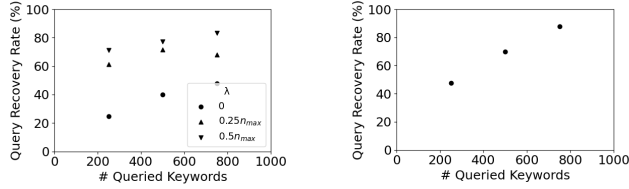
b) **FP-EMM and DP-EMM:** **FP-EMM** from [25] does not have any tunable parameters and so we run our attacks on the **FP-EMM** as it is. **DP-EMM** from [25] uses parameter ϵ to set query response volumes to $(2n_{\text{key}} + n^* + \text{Lap}_{sk}(2/\epsilon))$. The authors suggested $\epsilon = 0.2$ and we use the same value. We also run experiments where ϵ is significantly smaller, ranging from 0.1 to 0.01.

c) **DPAP-SE:** For **DPAP-SE** [26], the authors suggested $m = 6$ (the number of shards per document), $k = 2$ (a parameter of the erasure code which does not affect our attack), $p = 0.88703$ (the probability for which a keyword is kept in a shard) and $q = 0.04416$ as the parameters to use for the Enron dataset. We used similar parameters ($m = 6, k = 2, p = 0.89$ and $q = 0.045$) in our experiments. A smaller p or a bigger q significantly reduces the efficiency of the construction so we opt to not run additional experiments with those parameters. Instead, we investigate how a smaller q affects query reconstruction rate. We use $q = 0.0045, 0.00045$ and 0.000045 as additional choices of parameters. These additional parameters for all of the schemes are supposed to lead to more secure instances of the schemes. We investigate their effects in our experiments below.

Implementation. We implemented our attacks in C using GNU Scientific Library [40] for randomness generation and probability calculations. We used our own custom code for simulated annealing for best performance. We parallelized our implementation using OpenMP [41]. Our implementation is highly scalable. It takes less than one minute per run on the differentially-private schemes (**DPAP-SE** and **DP-EMM**) and no more than 6 minutes per run on the other schemes (**PRT-EMM** and **FP-EMM**) for all of our experimental settings, on a machine with an 8-core (16-thread) Sandy Bridge CPU clocked at 2.6 GHz ⁴.

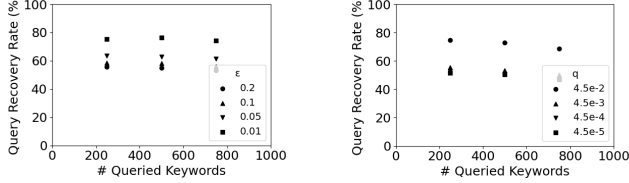
Experiments. We present three sets of experimental results on the target constructions discussed above. In Section V-B, we present the experimental results in basic settings, where the auxiliary co-occurrence matrix is built from 48,000 documents (20% of the available auxiliary data) using the 1,000 most frequent keywords. We set the number of queried keywords to 250, 500 or 750, and the security parameters are allowed to vary. In Section V-C, we set the number of queried keywords to 250, 500 or 750, and the security parameters to those suggested in the original papers, and vary the number of keywords used to build auxiliary information between 1,000 and 4,000. Just as before, 48,000 documents are used in building auxiliary information. Finally, in Section V-D, we use anywhere from 12,000 to 96,000 documents to build the auxiliary co-occurrence matrix, as a means to simulate auxiliary information with

⁴The code is publicly available at <https://github.com/RethinkingSSE/Attacks-on-SSE>



(a) **PRT-EMM** [24]. n_{max} is the maximum query response length.

(b) **FP-EMM** [25].



(c) **DP-EMM** [25].

(d) **DPAP-SE** [26].

Figure 2: Experimental results with varying security parameters. The 1,000 most frequent keywords are used in the auxiliary information.

different levels of noise. The number of keywords used is set to 1,000, and the number of queries is allowed to vary from 250 to 750. The security parameters are set to the ones recommended in the original papers.

B. Varying the Security Parameters of the Constructions

PRT-EMM. The experimental results on **PRT-EMM** are shown in Figure 2a. We observe an increasing query recovery rate with more queried keywords and larger λ . The attack performs significantly worse when $\lambda = 0$.

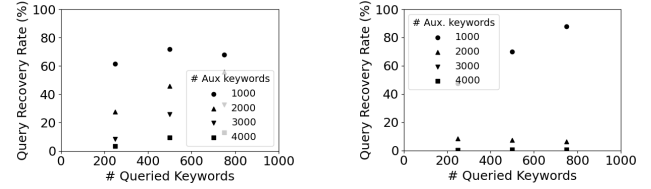
The worse performance of our attack on small λ is caused by truncations of query response volumes (see Section IV-A) which lead to loss of co-occurrence information. It should be noted that although a smaller λ leads to better query privacy, it results in more truncations and less complete query responses.

FP-EMM. The experimental results of our attack on **FP-EMM** are shown in Figure 2b. As expected, the attack performs better with more queried keywords. The attack is able to recover over 70% of the queried keywords if over 500 keywords have been queried, suggesting that full padding is ineffective at adding noise to the co-occurrence leakage.

DP-EMM. The experimental results on **DP-EMM** are shown in 2c. The attack performs slightly better with a smaller ϵ , suggesting that our attack is over-fitting the auxiliary data.

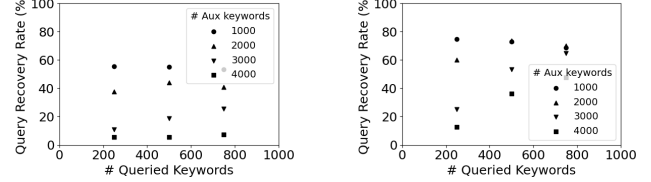
DPAP-SE. The results of our attacks are shown in Figure 2d. The attack is able recover between 50% and 80% of the queries in all cases we have considered.

As opposed to what one might expect, the attack performs better with a larger q . It is certainly true that a larger q masks the true co-occurrences better as there are more fake keywords introduced to each document, but it also reduces the effectiveness of perturbation with respect to keyword frequencies. For simplicity, consider a database with N documents, of which



(a) **PRT-EMM** [24]. λ is set to $0.25n_{max}$.

(b) **FP-EMM** [25].



(c) **DP-EMM** [25]. $\epsilon = 0.2$.

(d) **DPAP-SE** [26]. $(m, k, p, q) = (6, 2, 0.89, 0.045)$.

Figure 3: Experimental results with varying number of keywords in auxiliary information.

n documents contain keyword kw . With parameters k, m, p and q , the frequency of keyword kw in the resultant database processed by **DPAP-SE** is $k(np + (N - n)q)$. A larger q increases the separation of frequencies for different keywords, and we believe that is the main reason why our attack performs better with larger q .

C. Varying the Number of Keywords in Auxiliary Information

Our experimental results on varying the number of keywords in the auxiliary information are shown in Figure 3. The security parameters we used can be found in the captions.

As the keywords are uniformly randomly picked, the attacks with more auxiliary keywords are necessarily less successful. There are two main reasons for this. Firstly, since we have fixed the number of queried keywords, the search space for the attacks with more auxiliary keywords are larger, and there is naturally more uncertainty associated to those attacks. Secondly, as we have picked the queried keywords uniformly, it is more likely for the attacks with more auxiliary keywords to hit low frequency keywords, which naturally contain less co-occurrence information to begin with. Nevertheless, the constructions behave very differently with respect to the number of keywords in the auxiliary information.

Our attacks on **PRT-EMM** and **DP-EMM** work reasonably well for up to 2,000 auxiliary keywords. More queries are required for the attack to succeed with 3,000 and 4,000 auxiliary keywords. Our attack is less successful on **FP-EMM** when the number of auxiliary keywords is large. This shows that full padding is effective at masking co-occurrence information if there is enough uncertainty within the queries.

For **DPAP-SE**, the attacks with more auxiliary keywords do not perform well when only 250 keywords are queried. However, the query recovery rate increases significantly as more keywords are queried. The attack is able to recover over 50% of the queries even if only 750 out of 4,000 keywords have been queried.

Interestingly, the success rate of the attacks do not monotonically increase with respect to the number of queried keywords. We believe that there are three main reasons for this. Firstly, the distribution of keywords in the auxiliary database can be different from that of the target database. So over-fitting can happen. Secondly, the auxiliary and target databases are freshly generated each time so the attack may just perform worse on more auxiliary data by chance. Finally, simulated annealing is a randomised approach and there is no guarantee that it will converge to the global optimum within the given time. Since the search space grows as the number of auxiliary keywords increases, it becomes harder for simulated annealing to reach the global optimum. So the lower success rates could be due to the optimization algorithm getting stuck in a local but non-global optimum.

D. Varying the Level of Noise in Auxiliary Information

Given that there is no widely accepted way of modelling noise in auxiliary information, we opt to use different numbers of documents in auxiliary information as a way to simulate different levels of noise – fewer documents means more noise. We use absolute distance and modified probability score to measure the level of noise introduced in each set of experiments we run (see the full version of our paper [28] for details).

Our experimental results on varying auxiliary information are shown in Figure 4. The security parameters we used can be found in the captions. The attacks do not perform well when only 12,000 documents are used to construct the auxiliary information. However, the query reconstruction rate increases as the number of auxiliary documents increases. Interestingly, the query reconstruction rates using 48,000 and 96,000 auxiliary documents are comparable, suggesting that using 48,000 documents (10% of the total in the Enron dataset) is sufficient as auxiliary information and that our attacks are robust in a noisy setting. This aligns with our observation that the level of noise stabilises at 48,000 documents (more details are available in the full version of our paper [28]).

VI. DOCUMENT RETRIEVAL WITH EMMs AND OTHER PRIMITIVES

In this section, we explore the alternative possibility of building end-to-end SSE systems where leakage-suppression techniques (such as volume-hiding EMMs) are applied to the whole system (i.e. to *both* the encrypted index *and* the encrypted document collection). In such a system, the document retrieval step is only allowed to leak ‘trivial’ information such as the number of documents and the number of keyword-document pairs, but not potentially sensitive information such as the access pattern. We will show that, with currently available techniques, no end-to-end SSE system using volume-hiding EMMs (or similar leakage-suppression techniques) can achieve this without incurring significant storage, computation and/or bandwidth overheads.

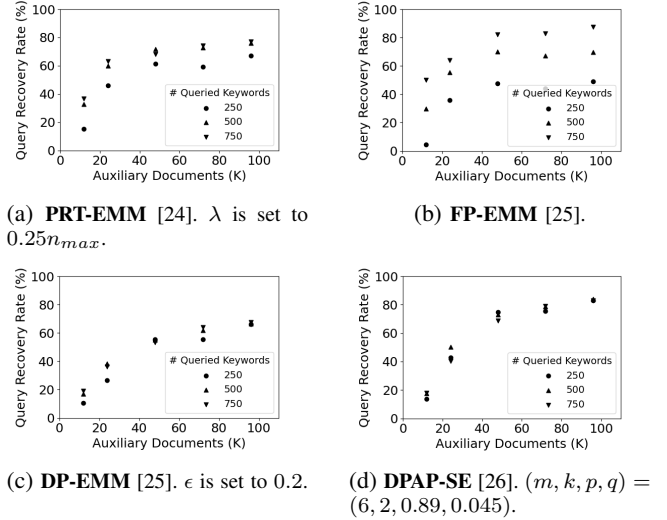


Figure 4: Experimental results with varying auxiliary information.

A. Primitives for Index and Document Retrieval

In the rest of this section, we assume the SSE system follows the two-phase approach illustrated in Figure 1. We also assume that the index retrieval phase (where the client recovers the set of (encrypted) document identifiers corresponding to documents matching the query) is implemented using one of the state-of-the-art schemes **PRT-EMM** [24] or **FP-EMM** [25], so the leakage from this phase consists of at most: (a) the number of documents $|\mathbf{DB}|$, (b) search pattern (query equality), and (c) the maximum response length $\max_{kw} |\mathbf{DB}(kw)|$. This leakage profile is a consequence of the properties of the aforementioned volume-hiding EMMs. We demand that the leakage from the second document retrieval phase should be no more than that in the first phase, in order to achieve effective system-wide leakage-suppression. This naturally leads us to consider a variety of different cryptographic primitives for realizing the document retrieval step, namely EMMs (again), Private Information Retrieval, and Oblivious RAM.

Using EMMs for Document Retrieval. We transform state-of-the-art EMMs [24], [25] into end-to-end SSE systems supporting actual document retrieval by replacing the values stored in the EMMs with the encrypted documents themselves.

Using Private Information Retrieval. A Private Information Retrieval (PIR) scheme can be used to retrieve the documents. Here, we focus on computational PIR schemes [42], [44] for which there is provably no leakage (per retrieval). As we are only permitted to leak the maximum response length in the document retrieval phase (in order to match the leakage profile of the first phase), we need to call the underlying PIR scheme $\max_{kw} |\mathbf{DB}(kw)|$ times per query (padding with dummy retrievals if necessary) in order to hide the true response length. We use SealPIR [42] in our evaluation below. This is a state-of-the-art PIR scheme using Fully Homomorphic Encryption (FHE) as a subcomponent.

Scheme	Storage (Server)	Query (Client)		Query (Server)	
		Computation	Communication	Computation	Communication
Naïve*	470 MB	f prf f dec	32 B	f acc	f KB
Duplication	17 GB (36x)	24K prf 24K dec	750 KB (46x)	24K acc	23 MB (46x)
PRT-EMM [24]	390 GB (860x)	12K prf 12K dec	370 KB (23x)	12K acc	12 MB (23x)
FP-EMM [25]	43 GB (94x)	48K prf 48K dec	1.5 MB (92x)	48K acc	47 MB (92x)
SealPIR [42]	120 GB (260x)	1 henc 1 hdec	1.46 GB (94,000x)	23B hmul 11B hsub 11B hadd	5.9 GB (12,000x)
Non-recursive Path ORAM [43]	1.8 GB (4x)	3.4M acc 1.7M dec 1.7M enc	1.65 GB (110,000x)	3.4M acc	1.7 GB (3,300x)

Table II: Evaluation of different document retrieval primitives with minimal leakage. The numbers in the brackets indicate overheads beyond the baseline provided by the **Naïve** scheme. We assume 522 documents (mean keyword frequency) are retrieved by the **Naïve** scheme in the computations of the overheads. *: f is the real query response volume (since there is no padding).

Using Oblivious RAM. Oblivious RAM (ORAM) [43], [45]–[47] is another primitive that achieves zero-leakage per access. We use it to protect the entire document collection. Similarly to PIR, the number of data accesses for ORAM has to be padded to the maximum response length in order to hide true response lengths. We use the non-recursive version of Path ORAM [43] in our evaluation below. This specific choice is amongst the most efficient ORAM schemes available.

For our experiments, we assumed that all documents are of equal size, which allowed us to directly use the array-like interface offered by ORAM to design the encrypted document-array. Note that in end-to-end SSE systems, leaking the length of retrieved documents is undesirable. The standard practice is to either pad each document to the same length or to divide each document into equal-sized chunks/sub-documents associated with separate identifiers. Unfortunately, the SSE literature does not offer any concrete guidelines on the appropriate choice of padding/splitting-strategy and its impact on security (in particular, on the access pattern leakage). Hence, we picked a fixed document-size, which could be viewed as either all documents being padded to the same length, or all documents split into equal-sized chunks.

B. Performance Evaluation

Additional Schemes. In addition to the end-to-end SSE systems targeted by our attacks, we also add the following two schemes in our comparison as a baseline:

- **Naïve:** The scheme simply encrypts the documents and stores them in an array. To retrieve documents, the user sends the array locations (used as document identifiers) to the server and the server returns all documents in those locations. This scheme is insecure in a system-wide attack setting, e.g. it is vulnerable to our inference attack⁵.
- **Duplication:** The scheme is identical to **Naïve** except that the encrypted documents are duplicated for each keyword in the same way as the encrypted document identifiers are duplicated in the search index of a traditional structure-only SSE scheme, e.g. [2]–[4]. This represents a baseline method for using an EMM to build an end-to-end SSE scheme.

Experimental Data. A concrete dataset is necessary for the evaluation as the duplication techniques used in **Duplication**, **PRT-EMM** and **FP-EMM** are data-dependent. We again pick the Enron email corpus (a detailed description of the dataset is available in the full version of our paper [28]). We used 480K documents in the evaluation below. To simplify the evaluation, we assume that all documents have size 1 KB even if some of them are larger than that in reality. If we were to use document splitting into fixed-size chunks instead, this would result in a $10\times$ larger storage overhead because of increased duplication of the chunks compared to the original files.

Parameters. The following parameters are used in our evaluation. We use PRFs with 256-bit output. We use the most space-efficient parameters for **PRT-EMM** proposed in the original papers [24], namely $\alpha = 0.5$. For SealPIR, we assume that the degree of ciphertext is $N = 2048$, the size of the coefficients are 60 bits and the database is represented in $d = 2$ dimensions as per the original paper [42]. For Path ORAM [43], we assume each block has size 1 KB and there are 4 blocks per bucket.

Evaluation. We report communication volume, storage cost, and the number of core cryptographic operations needed for each option described above in Table II. We split computation and communication costs into client and server costs, and report only server storage costs (client storage costs are low).

Storage and communication costs are measured in total volumes. Additional overheads arising from how the data is structured and packaged for communication are ignored.

Computation costs are measured by the number of core cryptographic operations. The operations that we consider include: **prf** for PRF computation; **enc** and **dec** for encryption and decryption with a symmetric primitive; **acc** for disk/RAM access (read or write); **henc** and **hdec** for encryption and decryption with FHE; **hmul**, **hsub**, **hadd** for multiplication, substitution and addition for FHE ciphertexts. Reporting operation counts in place of running times makes our comparison independent of implementation details.

We opt to not include latency as this depends on several factors such as data access speed and network delay, and these are hard to compare concretely and fairly.

⁵The naïve scheme in Table I refers to the same scheme as described here.

Discussion. It is clear that all of the options suffer significant storage overheads. For **Duplication**, **PRT-EMM** and **FP-EMM**, this is caused by duplication. The expansion factor grows linearly with the number of keywords per document. For PIR, the expansion factor comes from the use of homomorphic encryption. It is not clear how the ciphertexts can be compressed to reduce the overhead. It is conceivable that alternative PIR schemes might avoid such expansion. For ORAM, the overhead comes from the use of multiple blocks per bucket. This is necessary to prevent overflowing buckets, meaning the storage overhead cannot be reduced significantly.

With regard to queries, **Duplication**, **PRT-EMM** and **FP-EMM** have reasonable computational costs, but the communication costs from the server to the client are high in each case. The server needs to send 2.5% to 10% of the entire document collection to the client per query, which is a lot more than the average keyword frequency might suggest (0.109% for the Enron corpus). The PIR and ORAM options naturally suffer from high computation and/or communication overheads since they are not designed for large-scale document retrieval. In fact, with bad choice of parameters, as shown in [48], the communication volume produced and ORAM can be several orders of magnitude larger than the database itself.

Of course, the state-of-the-art implementations of various cryptographic primitives that we have chosen for our experiments could very well be improved upon in terms of concrete efficiency in future work. Nonetheless, the concrete numbers that we provide are indicative of what is possible by employing state-of-the-art approaches for realizing these primitives.

VII. DISCUSSION

On System-Wide leakage. Our attacks on end-to-end SSE schemes built from volume-hiding EMMs in the natural way highlight the threats posed by system-wide leakage and the need to revisit existing security definitions that ignore such leakage. Note that many SSE constructions were designed prior to the proposal of volume-hiding EMMs; in their original form, these constructions effectively used EMMs that themselves leak the exact access-pattern. In this case, the leakage from the encrypted search index subsumes the leakage from encrypted document retrieval and so the latter leakage can be ignored in security analysis (see [2], [4] for relevant discussions). However, this approach to analysis is no longer valid when such access pattern-revealing EMMs are replaced by volume-hiding EMMs. In this case, the leakage arising from encrypted document retrieval is no longer covered by security definitions that focus purely on the encrypted search index.

On Differentially Private Access Patterns. Our attack applied to **DPAP-SE** [26] does not negate the core idea of [26], which is to design end-to-end SSE schemes that take into account leakage from the document retrieval phase. Rather, our attack serves as a warning about the potential pitfalls of applying differential privacy to SSE without appropriately modeling and analyzing the resulting leakage. As pointed out in [26], differentially private access patterns provide provable

guarantees of the form: an adversary cannot distinguish between queries over keywords such that their access pattern leakage is within a *small* statistical distance of each other. As demonstrated by our attack, the provable guarantees provided by differential privacy do not necessarily translate into security guarantees against leakage-abuse attacks *in general*.

On the Practicality of Our Attack. Our attack is an inference attack in that it assumes the adversary has access to auxiliary data that is independent of but statistically “close” to the target database. We believe that this a weaker (and more realistic) attack setting compared to existing known-data attacks [5], [6], [9]. Our attack achieves high keyword recovery rates even when the target schemes use aggressive security parameters, or when the auxiliary data available to the adversary is relatively noisy (which is the case when we sample the auxiliary data from a small portion of the database). These observations further reinforce the practicality of our attack.

A drawback of our attack is that it assumes auxiliary information involving high-frequency keywords. This is a relatively strong assumption in practice (although one also made by all previous leakage-abuse attacks). One can of course filter out leakage from low-frequency keywords based on response volume before running our attack. We leave it as an open problem to extend our attack to low-frequency keywords.

Future Research on SSE. We argue that, if the goal of the research community is to develop SSE towards practice, then a fundamental shift in approach is needed. In particular, researchers need to take a system-wide view of SSE and its security. This requires considering all components of an end-to-end SSE system when analyzing security, and investigating techniques that reduce leakage and maintain efficiency of the whole SSE system. In this context, our view is that end-to-end SSE schemes should be designed from scratch, with dedicated leakage suppression for both index and document retrieval.

REFERENCES

- [1] D. X. Song, D. Wagner, and A. Perrig, “Practical techniques for searches on encrypted data,” in *2000 IEEE Symposium on Security and Privacy*. Oakland, CA, USA: IEEE Computer Society Press, May 2000, pp. 44–55.
- [2] R. Curtmola, J. A. Garay, S. Kamara, and R. Ostrovsky, “Searchable symmetric encryption: improved definitions and efficient constructions,” in *ACM CCS 2006: 13th Conference on Computer and Communications Security*, A. Juels, R. N. Wright, and S. De Capitani di Vimercati, Eds. Alexandria, Virginia, USA: ACM Press, Oct. 30 – Nov. 3, 2006, pp. 79–88.
- [3] M. Chase and S. Kamara, “Structured encryption and controlled disclosure,” in *Advances in Cryptology – ASIACRYPT 2010*, ser. Lecture Notes in Computer Science, M. Abe, Ed., vol. 6477. Singapore: Springer, Heidelberg, Germany, Dec. 5–9, 2010, pp. 577–594.
- [4] D. Cash, S. Jarecki, C. S. Jutla, H. Krawczyk, M.-C. Rosu, and M. Steiner, “Highly-scalable searchable symmetric encryption with support for Boolean queries,” in *Advances in Cryptology – CRYPTO 2013, Part I*, ser. Lecture Notes in Computer Science, R. Canetti and J. A. Garay, Eds., vol. 8042. Santa Barbara, CA, USA: Springer, Heidelberg, Germany, Aug. 18–22, 2013, pp. 353–373.
- [5] M. S. Islam, M. Kuzu, and M. Kantarcioglu, “Access pattern disclosure on searchable encryption: Ramification, attack and mitigation,” in *ISOC Network and Distributed System Security Symposium – NDSS 2012*. San Diego, CA, USA: The Internet Society, Feb. 5–8, 2012.

- [6] D. Cash, P. Grubbs, J. Perry, and T. Ristenpart, "Leakage-abuse attacks against searchable encryption," in *ACM CCS 2015: 22nd Conference on Computer and Communications Security*, I. Ray, N. Li, and C. Kruegel, Eds. Denver, CO, USA: ACM Press, Oct. 12–16, 2015, pp. 668–679.
- [7] D. Pouliot and C. V. Wright, "The shadow nemesis: Inference attacks on efficiently deployable, efficiently searchable encryption," in *ACM CCS 2016: 23rd Conference on Computer and Communications Security*, E. R. Weippl, S. Katzenbeisser, C. Kruegel, A. C. Myers, and S. Halevi, Eds. Vienna, Austria: ACM Press, Oct. 24–28, 2016, pp. 1341–1352.
- [8] Z. Gui, K. G. Paterson, S. Patranabis, and B. Warinschi, "SWiSSSE: System-wide security for searchable symmetric encryption," *Cryptology ePrint Archive*, Report 2020/1328, 2020, <https://eprint.iacr.org/2020/1328>.
- [9] L. Blackstone, S. Kamara, and T. Moataz, "Revisiting leakage abuse attacks," in *ISOC Network and Distributed System Security Symposium – NDSS 2020*. San Diego, CA, USA: The Internet Society, Feb. 23–26, 2020.
- [10] S. Oya and F. Kerschbaum, "Hiding the access pattern is not enough: Exploiting search pattern leakage in searchable encryption," in *30th USENIX Security Symposium, USENIX Security 2021, August 11–13, 2021*, M. Bailey and R. Greenstadt, Eds. USENIX Association, 2021, pp. 127–142. [Online]. Available: <https://www.usenix.org/conference/usenixsecurity21/presentation/oaya>
- [11] —, "IHOP: Improved statistical query recovery against searchable symmetric encryption through quadratic optimization," 2021.
- [12] J. Ning, X. Huang, G. S. Poh, J. Yuan, Y. Li, J. Weng, and R. H. Deng, "LEAP: leakage-abuse attack on efficiently deployable, efficiently searchable encryption with partially known dataset," ser. CCS '21. New York, NY, USA: Association for Computing Machinery, 2021, p. 2307–2320. [Online]. Available: <https://doi.org/10.1145/3460120.3484540>
- [13] M. Damie, F. Hahn, and A. Peter, "A highly accurate query-recovery attack against searchable encryption using non-indexed documents," in *30th USENIX Security Symposium, USENIX Security 2021, August 11–13, 2021*, M. Bailey and R. Greenstadt, Eds. USENIX Association, 2021, pp. 143–160. [Online]. Available: <https://www.usenix.org/conference/usenixsecurity21/presentation/damie>
- [14] S. Kamara, C. Papamanthou, and T. Roeder, "Dynamic searchable symmetric encryption," in *ACM CCS 2012: 19th Conference on Computer and Communications Security*, T. Yu, G. Danezis, and V. D. Gligor, Eds. Raleigh, NC, USA: ACM Press, Oct. 16–18, 2012, pp. 965–976.
- [15] D. Cash, J. Jaeger, S. Jarecki, C. S. Jutla, H. Krawczyk, M.-C. Rosu, and M. Steiner, "Dynamic searchable encryption in very-large databases: Data structures and implementation," in *ISOC Network and Distributed System Security Symposium – NDSS 2014*. San Diego, CA, USA: The Internet Society, Feb. 23–26, 2014.
- [16] E. Stefanov, C. Papamanthou, and E. Shi, "Practical dynamic searchable encryption with small leakage," in *ISOC Network and Distributed System Security Symposium – NDSS 2014*. San Diego, CA, USA: The Internet Society, Feb. 23–26, 2014.
- [17] M. Naveed, M. Prabhakaran, and C. A. Gunter, "Dynamic searchable encryption via blind storage," in *2014 IEEE Symposium on Security and Privacy*. Berkeley, CA, USA: IEEE Computer Society Press, May 18–21, 2014, pp. 639–654.
- [18] S. Faber, S. Jarecki, H. Krawczyk, Q. Nguyen, M.-C. Rosu, and M. Steiner, "Rich queries on encrypted data: Beyond exact matches," in *ESORICS 2015: 20th European Symposium on Research in Computer Security, Part II*, ser. Lecture Notes in Computer Science, G. Pernul, P. Y. A. Ryan, and E. R. Weippl, Eds., vol. 9327. Vienna, Austria: Springer, Heidelberg, Germany, Sep. 21–25, 2015, pp. 123–145.
- [19] R. Bost, "Σφός: Forward secure searchable encryption," in *ACM CCS 2016: 23rd Conference on Computer and Communications Security*, E. R. Weippl, S. Katzenbeisser, C. Kruegel, A. C. Myers, and S. Halevi, Eds. Vienna, Austria: ACM Press, Oct. 24–28, 2016, pp. 1143–1154.
- [20] R. Bost, B. Minaud, and O. Ohrimenko, "Forward and backward private searchable encryption from constrained cryptographic primitives," in *ACM CCS 2017: 24th Conference on Computer and Communications Security*, B. M. Thuraisingham, D. Evans, T. Malkin, and D. Xu, Eds. Dallas, TX, USA: ACM Press, Oct. 31 – Nov. 2, 2017, pp. 1465–1482.
- [21] S. Kamara and T. Moataz, "Boolean searchable symmetric encryption with worst-case sub-linear complexity," in *Advances in Cryptology – EUROCRYPT 2017, Part III*, ser. Lecture Notes in Computer Science, J.-S. Coron and J. B. Nielsen, Eds., vol. 10212. Paris, France: Springer, Heidelberg, Germany, Apr. 30 – May 4, 2017, pp. 94–124.
- [22] —, "SQL on structurally-encrypted databases," in *Advances in Cryptology – ASIACRYPT 2018, Part I*, ser. Lecture Notes in Computer Science, T. Peyrin and S. Galbraith, Eds., vol. 11272. Brisbane, Queensland, Australia: Springer, Heidelberg, Germany, Dec. 2–6, 2018, pp. 149–180.
- [23] J. G. Chamani, D. Papadopoulos, C. Papamanthou, and R. Jalili, "New constructions for forward and backward private symmetric searchable encryption," in *ACM CCS 2018: 25th Conference on Computer and Communications Security*, D. Lie, M. Mannan, M. Backes, and X. Wang, Eds. Toronto, ON, Canada: ACM Press, Oct. 15–19, 2018, pp. 1038–1055.
- [24] S. Kamara and T. Moataz, "Computationally volume-hiding structured encryption," in *Advances in Cryptology – EUROCRYPT 2019, Part II*, ser. Lecture Notes in Computer Science, Y. Ishai and V. Rijmen, Eds., vol. 11477. Darmstadt, Germany: Springer, Heidelberg, Germany, May 19–23, 2019, pp. 183–213.
- [25] S. Patel, G. Persiano, K. Yeo, and M. Yung, "Mitigating leakage in secure cloud-hosted data structures: Volume-hiding for multi-maps via hashing," in *ACM CCS 2019: 26th Conference on Computer and Communications Security*, L. Cavallaro, J. Kinder, X. Wang, and J. Katz, Eds. ACM Press, Nov. 11–15, 2019, pp. 79–93.
- [26] G. Chen, T. Lai, M. K. Reiter, and Y. Zhang, "Differentially private access patterns for searchable symmetric encryption," in *IEEE INFOCOM 2018 - IEEE Conference on Computer Communications*, 2018, pp. 810–818.
- [27] D. Cash, P. Grubbs, J. Perry, and T. Ristenpart, "Leakage-abuse attacks against searchable encryption," *Cryptology ePrint Archive*, Report 2016/718, 2016, <https://eprint.iacr.org/2016/718>.
- [28] Z. Gui, K. G. Paterson, and S. Patranabis, "Rethinking searchable symmetric encryption," *Cryptology ePrint Archive*, Report 2021/879, 2021, <https://eprint.iacr.org/2021/879>.
- [29] W. He, D. Akhawe, S. Jain, E. Shi, and D. X. Song, "ShadowCrypt: Encrypted web applications for everyone," in *ACM CCS 2014: 21st Conference on Computer and Communications Security*, G.-J. Ahn, M. Yung, and N. Li, Eds. Scottsdale, AZ, USA: ACM Press, Nov. 3–7, 2014, pp. 1028–1039.
- [30] B. Lau, S. P. Chung, C. Song, Y. Jang, W. Lee, and A. Boldyreva, "Mimesis aegis: A mimicry privacy shield-A system's approach to data privacy on public cloud," in *USENIX Security 2014: 23rd USENIX Security Symposium*, K. Fu and J. Jung, Eds. San Diego, CA, USA: USENIX Association, Aug. 20–22, 2014, pp. 33–48.
- [31] I. Demertzis, D. Papadopoulos, C. Papamanthou, and S. Shintre, "SEAL: Attack mitigation for encrypted databases via adjustable leakage," in *USENIX Security 2020: 29th USENIX Security Symposium*, S. Capkun and F. Roesner, Eds. USENIX Association, Aug. 12–14, 2020, pp. 2433–2450.
- [32] P. Grubbs, A. Khandelwal, M.-S. Lacharité, L. Brown, L. Li, R. Agarwal, and T. Ristenpart, "Pancake: Frequency smoothing for encrypted data stores," in *USENIX Security 2020: 29th USENIX Security Symposium*, S. Capkun and F. Roesner, Eds. USENIX Association, Aug. 12–14, 2020, pp. 2451–2468.
- [33] Z. Shang, S. Oya, A. Peter, and F. Kerschbaum, "Obfuscated access and search patterns in searchable encryption," 2021. [Online]. Available: <https://arxiv.org/abs/2102.09651>
- [34] C. William W. Cohen, MLD, "Enron email dataset." [Online]. Available: <https://www.cs.cmu.edu/~lenron/>
- [35] S. Kamara, A. Kati, T. Moataz, T. Schneider, A. Treiber, and M. Yonli, "SoK: Cryptanalysis of encrypted search with LEAKER - a framework for LEakage Attack Evaluation on Real-world data," *Cryptology ePrint Archive*, Report 2021/1035, 2021, <https://ia.cr/2021/1035>.
- [36] A. G. Dimakis, B. Godfrey, M. J. Wainwright, and K. Ramchandran, "Network coding for distributed storage systems," *CoRR*, vol. abs/0702015, 2007.
- [37] R. Rossi and R. Rossi, *Mathematical statistics: An introduction to likelihood based inference*. John Wiley & Sons, Inc., 2018.
- [38] P. J. van Laarhoven, *Simulated annealing theory and applications*. Kluwer, 1987.
- [39] NLTK Project, *Natural Language Toolkit*, <https://www.nltk.org/>.
- [40] B. Gough, *GNU Scientific Library Reference Manual - Third Edition*, 3rd ed. Network Theory Ltd., 2009.
- [41] OpenMP Architecture Review Board, "OpenMP application program interface version 5.0," November 2018. [Online]. Available: <https://www.openmp.org/wp-content/uploads/OpenMP-API-Specification-5.0.pdf>

- [42] S. Angel, H. Chen, K. Laine, and S. T. V. Setty, “PIR with compressed queries and amortized query processing,” in *2018 IEEE Symposium on Security and Privacy*. San Francisco, CA, USA: IEEE Computer Society Press, May 21–23, 2018, pp. 962–979.
- [43] E. Stefanov, M. van Dijk, E. Shi, C. W. Fletcher, L. Ren, X. Yu, and S. Devadas, “Path ORAM: an extremely simple oblivious RAM protocol,” in *ACM CCS 2013: 20th Conference on Computer and Communications Security*, A.-R. Sadeghi, V. D. Gligor, and M. Yung, Eds. Berlin, Germany: ACM Press, Nov. 4–8, 2013, pp. 299–310.
- [44] C. Aguilar Melchor, J. Barrier, L. Fousse, and M.-O. Killijian, “XPIR: Private information retrieval for everyone,” *Proceedings on Privacy Enhancing Technologies*, vol. 2016, no. 2, pp. 155–174, Apr. 2016.
- [45] O. Goldreich, “Towards a theory of software protection and simulation by oblivious RAMs,” in *19th Annual ACM Symposium on Theory of Computing*, A. Aho, Ed. New York City, NY, USA: ACM Press, May 25–27, 1987, pp. 182–194.
- [46] I. Damgård, S. Meldgaard, and J. B. Nielsen, “Perfectly secure oblivious RAM without random oracles,” in *TCC 2011: 8th Theory of Cryptography Conference*, ser. Lecture Notes in Computer Science, Y. Ishai, Ed., vol. 6597. Providence, RI, USA: Springer, Heidelberg, Germany, Mar. 28–30, 2011, pp. 144–163.
- [47] T.-H. H. Chan, K. Nayak, and E. Shi, “Perfectly secure oblivious parallel RAM,” in *TCC 2018: 16th Theory of Cryptography Conference, Part II*, ser. Lecture Notes in Computer Science, A. Beimel and S. Dziembowski, Eds., vol. 11240. Panaji, India: Springer, Heidelberg, Germany, Nov. 11–14, 2018, pp. 636–668.
- [48] M. George, S. Kamara, and T. Moataz, “Structured encryption and dynamic leakage suppression,” in *Advances in Cryptology – EURO-CRYPT 2021, Part III*, ser. Lecture Notes in Computer Science, A. Can- teaut and F.-X. Standaert, Eds., vol. 12698. Zagreb, Croatia: Springer, Heidelberg, Germany, Oct. 17–21, 2021, pp. 370–396.

APPENDIX A

MATHEMATICAL DERIVATIONS OF THE CO-OCCURRENCE MATRICES

Derivation for PRT-EMM [24]. Recall that in PRT-EMM, the query response lengths are padded or truncated as:

$$n'_{\text{key}} = \lambda + F_{sk}(\text{key} || n_{\text{key}}).$$

Let $\mathbf{DB}$ be a multi-map and $\mathbf{q}_1, \dots, \mathbf{q}_l$ be non-repeating search queries with associated keys $\text{key}_1, \dots, \text{key}_l$ on $\mathbf{DB}$ encrypted with PRT-EMM. We abuse the notation $\text{key}(\mathbf{q}_i)$ to mean the key associated to $\mathbf{q}_i$. By denoting the maximum value of the PRF F as $|F|$, the diagonal entries of the co-occurrence matrix can be expressed as:

$$\bar{M}(\mathbf{q}_1, \dots, \mathbf{q}_l; \mathbf{DB})_{i,i} \sim \lambda + \text{Uniform}(0, |F|),$$

where $\text{Uniform}(\cdot)$ is a uniform distribution.

There are three cases to be considered for the off-diagonal entries of the co-occurrence matrix. Without loss of generality, let the keys in concern be keys key_i and key_j . In the first case, both of the query response lengths associated to the keys are larger than the true query response lengths. This corresponds to $n'_{\text{key}_i} - |\mathbf{DB}(\text{key}_i)|$ random document retrievals for queries on key key_i and $n'_{\text{key}_j} - |\mathbf{DB}(\text{key}_j)|$ random document retrievals for queries on key key_j . These random document retrievals can create additional co-occurrence counts among themselves

or with the real document retrievals. The co-occurrence counts in this case can be approximated by:

$$\begin{aligned} \bar{M}(\mathbf{q}_1, \dots, \mathbf{q}_l; \mathbf{DB})_{i,j} \\ \sim |\mathbf{DB}(\text{key}_i, \text{key}_j)| \\ + \text{HGeom}\left(n'_{\text{key}_i} - |\mathbf{DB}(\text{key}_i)|, |\mathbf{DB}|, n'_{\text{key}_j}\right) \\ + \text{HGeom}\left(n'_{\text{key}_j} - |\mathbf{DB}(\text{key}_j)|, |\mathbf{DB}|, n'_{\text{key}_i}\right), \end{aligned}$$

where $\text{HGeom}(n, N, K)$ denotes a hypergeometric distribution which makes n draws without replacement, from a population of size N that contains exactly K objects with the desired feature.

In the second case, one of the query response lengths is truncated and the other one is padded. Without loss of generality, let key key_i be the truncated key and key key_j be the padded key. Then, the co-occurrence count associated to keys key_i and key_j can be modelled as a process where the co-occurrence count is first reduced by the truncation and then increased by the padding. Its distribution is given below:

$$x \sim \text{HGeom}\left(n'_{\text{key}_i}, |\mathbf{DB}(\text{key}_i)|, |\mathbf{DB}(\text{key}_i, \text{key}_j)|\right),$$

$$\begin{aligned} \bar{M}(\mathbf{q}_1, \dots, \mathbf{q}_l; \mathbf{DB})_{i,j} \sim x + \text{HGeom}\left(n'_{\text{key}_j} \right. \\ \left. - |\mathbf{DB}(\text{key}_j)|, |\mathbf{DB}|, n'_{\text{key}_i} - x\right). \end{aligned}$$

Finally, in the last case, both of the query response lengths are truncated. Similar to above, the distribution of the co-occurrence count associated to keys key_i and key_j can be expressed as:

$$x \sim \text{HGeom}\left(n'_{\text{key}_i}, |\mathbf{DB}(\text{key}_i)|, |\mathbf{DB}(\text{key}_i, \text{key}_j)|\right),$$

$$\bar{M}(\mathbf{q}_1, \dots, \mathbf{q}_l; \mathbf{DB})_{i,j} \sim \text{HGeom}\left(n'_{\text{key}_j}, |\mathbf{DB}(\text{key}_j)|, x\right).$$

Derivation for new volume-hiding multi-maps in [25]. The volume-hiding multi-maps in [25] are special cases of PRT-EMM [24], where the query response lengths are either padded to the maximum query response length or ones that are larger than the true query response lengths. Specifically, for the full padding version (PRT-EMM),

$$\bar{M}(\mathbf{q}_1, \dots, \mathbf{q}_l; \mathbf{DB})_{i,i} \sim 2 \max_{\text{key}} |\mathbf{DB}(\text{key})|.$$

And for the differentially-private version (DP-EMM),

$$\bar{M}(\mathbf{q}_1, \dots, \mathbf{q}_l; \mathbf{DB})_{i,i} \sim 2 |\mathbf{DB}(\text{key})| + n^* + \text{Lap}(2/\epsilon),$$

where n^* is a fixed constant to offset the query response length in case the latter random variable is negative. Finally, the co-occurrence counts can be approximated exactly as done in the case of PRT-EMM earlier.

Derivation for DPAP-SE [26]. Let $\mathbf{DB}$ be a database and $\mathbf{q}_1, \dots, \mathbf{q}_l$ be non-repeating search queries with associated keywords $\text{kw}_1, \dots, \text{kw}_l$ on $\mathbf{DB}$ encrypted with the searchable

encryption scheme above [26]. The diagonal entries of the co-occurrence matrix $\bar{M}(\mathbf{q}_1, \dots, \mathbf{q}_l; \mathbf{DB})$, i.e. the query response volumes, represent the numbers of shards retrieved by the client. Formally, the diagonal entries of the co-occurrence matrix can be expressed in terms of the true query response lengths as:

$$\bar{M}(\mathbf{q}_1, \dots, \mathbf{q}_l; \mathbf{DB})_{i,i} \sim \text{Bin}(m \cdot |\mathbf{DB}(\text{kw}_i)|, p) \\ + \text{Bin}(m \cdot |\mathbf{DB}| - m \cdot |\mathbf{DB}(\text{kw}_i)|, q),$$

where m comes from splitting the documents into shards, $\mathbf{DB}(\text{kw}_i)$ denotes the set of documents containing keyword kw_i associated to query $\mathbf{q}_i$, $|\mathbf{DB}|$ denotes the total number of documents, and $\text{Bin}(\cdot)$ denotes a binomial distribution.

For the off-diagonal entries of the co-occurrence matrix, assume without loss of generality that the keywords in concern are kw_i and kw_j . The actual distribution of the off-diagonal entries of the co-occurrence matrix is complicated due to dependencies. However, if we ignore the fact that we already know the query response lengths for keywords kw_i and kw_j , the off-diagonal entries of the co-occurrence matrix can be approximated as:

$$\bar{M}(\mathbf{q}_1, \dots, \mathbf{q}_l; \mathbf{DB})_{i,j} \sim \text{Bin}(m \cdot |\mathbf{DB}(\text{kw}_i, \text{kw}_j)|, p^2) \\ + \text{Bin}(m \cdot (|\mathbf{DB}| - |\mathbf{DB}(\text{kw}_i)| - |\mathbf{DB}(\text{kw}_j)| + |\mathbf{DB}(\text{kw}_i, \text{kw}_j)|), q^2) \\ + \text{Bin}(m \cdot |\mathbf{DB}(\text{kw}_i)| - |\mathbf{DB}(\text{kw}_i, \text{kw}_j)|, pq) \\ + \text{Bin}(m \cdot |\mathbf{DB}(\text{kw}_j)| - |\mathbf{DB}(\text{kw}_i, \text{kw}_j)|, pq).$$

APPENDIX B

MATHEMATICAL DERIVATIONS OF THE LIKELIHOOD FUNCTIONS

Likelihood Function and its Decomposition. The likelihood function $\mathbf{L}[P | \bar{M}, M]$ can be written as follows:

$$\mathbf{L}[P | \bar{M}, M] = \Pr[\bar{M}, M | P] \\ = \sum_{M' \in \mathcal{N}^{N \times N}} \Pr[\bar{M}, M, M' | P] \\ = \sum_{M' \in \mathcal{N}^{N \times N}} \Pr[\bar{M} | M', P] \Pr[M' | M],$$

where N is the number of documents and $\mathcal{N}^{N \times N}$ is all N by N natural number valued matrices. In the third line of the equation, we used the law of total probability to turn the likelihood into a summation over all possible real co-occurrence matrices. The lines after break the probability into a sum of products of two probabilities. The first probability $\Pr[\bar{M}, M' | P]$ is the probability that $\bar{M}$ is the observed co-occurrence matrix and M' is the real co-occurrence matrix given P is the permutation. The second probability is the probability of getting M' as the real observed co-occurrence matrix knowing that M is the auxiliary co-occurrence matrix.

Derivation for PRT-EMMs. For PRT-EMMs [24], query response lengths may be truncated by a random amount. This means that based on the query response length in the auxiliary co-occurrence matrix M' and that in the observed

co-occurrence matrix, an attacker can estimate how many documents in the off-diagonal entries are expected to be removed. For observed co-occurrence count between keywords kw_i and kw_j where $i \neq j$, the real process can be modelled as a sequential application of two hypergeometric distributions on the real co-occurrence count.

$$\prod_{i < j} \sum_k \Pr \left[\text{HGeom} \left(M'_{P(i), P(i)}, M'_{P(i), P(j)}, \bar{M}_{i,i} \right) = k \right] \\ \Pr \left[\text{HGeom} \left(M'_{P(j), P(j)}, k, \bar{M}_{j,j} \right) = \bar{M}_{i,j} \right].$$

Derivation for FP-EMMs [25]. To simplify the first term of the likelihood decomposition, we assume independence of the entries in the observed co-occurrence matrix. Without loss of generality, we assume that all query response lengths are padded to m . This means we can express the probability as:

$$\prod_{i < j} \Pr \left[\text{HGeom}(2N, 2m - 2M'_{P(i), P(i)}, \right. \\ \left. 2m - 2M'_{P(j), P(j)}) = \bar{M}_{i,j} - M'_{P(i), P(j)} \right].$$

Derivation for DP-EMMs [25]. The first term of the likelihood decomposition for differentially private volume-hiding EMMs [25] is similar to that of the full padding version, except that the query response lengths are padded according to a Laplacian distribution as opposed to padding to the maximum query response length. Let n^* be the constant to offset the Laplacian random variable $\text{Lap}(2/\epsilon)$, the first term of the likelihood decomposition can be expressed as:

$$\sum_{i=j} \Pr[\bar{M}, M' | P] \times \sum_{i < j} \Pr[\bar{M}, M' | P] \\ = \sum_i \Pr \left[2M'_{P(i), P(i)} + n^* + \text{Lap}(2/\epsilon) = \bar{M}_{i,i} \right] \\ \times \sum_{i < j} \Pr \left[\text{HGeom}(2N, 2\bar{M}_{i,i} - 2M'_{P(i), P(i)}, \right. \\ \left. 2\bar{M}_{j,j} - 2M'_{P(j), P(j)}) = \bar{M}_{i,j} - 2M'_{P(i), P(j)} \right].$$

Derivation for DPAP-SE [26]. We refer the reader to the full version of our paper [28] for the derivation of the likelihood function and its decomposition for DPAP-SE [26], which follows from similar principles as the derivations above.

Approximation Techniques. As it can be seen, it is computationally infeasible to sum over all possible real co-occurrence matrices. We propose to sum over all possible real co-occurrence matrices such that $\Pr[M' | \bar{M}]$ is significant. Please see the full version of our paper [28] for more details and additional optimizations.

Speeding up the Score function. A naïve implementation of the Score function would require l^2 computations per iteration, where l is the number of non-repeating queries observed. We design an optimized implementation of the Score function for our attacks which only requires $O(l)$ computation per iteration. Details are available in the full version of our paper [28].