

Collaborative, Privacy-Preserving Data Aggregation at Scale

Haakon Ringberg*, Benny Applebaum*, Michael J. Freedman*, Matthew Caesar†, Jennifer Rexford*

*Princeton University, †University of Illinois at Urbana-Champaign

Abstract

Combining and analyzing data collected at multiple locations is critical for a wide variety of applications, such as detecting and diagnosing malicious attacks or computing an accurate estimate of the popularity of Web sites. However, legitimate concerns about privacy often inhibit participation in collaborative data-analysis systems. In this paper, we design, implement, and evaluate a practical solution for privacy-preserving collaboration among a large number of participants. Scalability is achieved through a “semi-centralized” architecture that divides responsibility between a *proxy* that obviously blinds the client inputs and a *database* that identifies the (blinded) keywords that have values satisfying some evaluation function.

Our solution leverages a novel cryptographic protocol that provably protects the privacy of both the participants and the keywords. In the example of Web servers collaborating to detect source IP addresses responsible for denial-of-service attacks, our protocol would not reveal the traffic mix of the Web servers or the identity of the “good” IP addresses. We implemented a prototype of our design, including an amortized oblivious transfer protocol that substantially improves the efficiency of client-proxy interactions. Our experiments show that the implementation scales linearly with the computing resources, making it easy to improve performance by adding more cores or machines. For collaborative diagnosis of denial-of-service attacks, our system can handle millions of suspect IP addresses per hour when the proxy and the database each run on two quad-core machines.

1 Introduction

Many important data-analysis applications must combine and analyze data collected by multiple parties. Such distributed data analysis is particularly important in the context of security. For example, victims of denial-of-service (DoS) attacks know they have been attacked but cannot easily distinguish the malicious source IP addresses from the good users who happened to send legitimate requests at the same time; since compromised hosts in a botnet often participate in multiple such attacks, victims could potentially identify the bad IP addresses if they combined their measurement data [33]. Cooperation is also useful for Web clients to recognize they have received a bogus DNS response or a forged self-signed certificate, by checking that the information they received agrees with that seen by other clients accessing the same Web site [29, 37]. Collaboration is also useful to identify popular Web content by having Web users, or proxies monitoring traffic for an entire organization, combine their access logs to determine the most frequently accessed URLs [1]. In this paper, we present the design, implementation, and evaluation of an efficient, privacy-preserving system that supports these kinds of data-analysis operations.

Today, these kinds of distributed data-analysis applications lack privacy protections. Existing solutions often rely on a trusted (typically centralized) aggregation node that collects and analyzes the raw data, thereby learning both the identity and inputs of participants. There is good reason to believe this inhibits participation. ISPs and Web sites are notoriously unwilling to share operational data with one another, because they are business competitors and are concerned about compromising the privacy of their customers. Many users are understandably unwilling to install software from Web analytics services such as Alexa [1], as such software would otherwise track and report every Web site they visit. Unfortunately, even good intentions do not necessarily translate to good security and privacy protections, only too-well demonstrated by the fact that large-scale data breaches have become commonplace [30]. As such, we believe that many useful distributed data-analysis applications will not gain serious traction unless privacy can be ensured.

Fortunately, many of these collaborative data-analysis applications have a common pattern, such as computing set intersection, finding so-called *icebergs* (items with a frequency count above a certain threshold), or identifying items that in aggregate satisfy some other statistical property. We refer to this problem as *privacy-preserving data aggregation* (PDA). Namely, each participant p_j has an input set of key-value tuples, $\langle k_{i,j}, v_{i,j} \rangle$, and the protocol outputs a key k_i if and only if some evaluation function $f(\forall j | v_{i,j})$ is satisfied. For example, the problem of finding icebergs—keys that occur more than some threshold τ times across the $n \geq \tau$ parties—for example, the botnet anomaly-detection application would have the suspect IP addresses as keys and all one’s as values, with f defined as $\sum_{j=1}^n v_{i,j} \geq \tau$ (implemented, in fact, as simply keeping a running sum per key). In other words, such a protocol performs the equivalent of a database join (union) across each participants’ input (multi)set, and outputs those keys that appear more than τ times. In our system, keys can either be arbitrary-length bitstrings or can also be drawn from a limited domain (*e.g.*, the set of valid IP addresses). However, we restrict our consideration of values to those drawn from a more restricted domain—such as an alphanumeric score from 1 to 10 or A to F—a limitation for privacy reasons we expand later. This f could as easily perform other types of frequency analysis on keys, such as median, mode, or dynamically setting the threshold τ based on the set of inputs—for example, if there exists some appropriate “gap” between popular and unpopular inputs—as opposed to requiring τ be set *a priori* and independent of the inputs.

Informally, PDA should provide two privacy properties: (1) *Keyword privacy* requires that no party should learn anything about k_i if its corresponding values do not satisfy f . (2) *Participant privacy* requires that no party should learn which key inputs (whether or not the key remains somehow blinded prior to satisfying f) belongs to which participant. In our example of collaborating DoS victims, keyword privacy means nobody learns the identity of good IP addresses or which Web sites they frequent, and participant privacy means a Web site need not worry that its mix of clients would be revealed. In our example of collaborating Web clients, the privacy guarantees mean that a Web user need not worry that other users know what Web sites he accesses, or whether he received a bogus DNS response or a forged certificate. We believe these privacy properties would be sufficient to encourage participants to collaborate, to their mutual benefit, without concern that their privacy, or the privacy of their clients, would be compromised. Our goal, then, is to design a system that provably guarantees these properties, and is efficient enough to be used in practice.

Ideally, we would like a system that can handle hundreds or thousands of participants generating thousands of key-value tuples. Unfortunately, fully-distributed solutions do not scale well enough, and fully-centralized solutions do not meet our privacy requirements. Simple techniques like hashing input keys [2, 12], while efficient, cannot ensure keyword and participant privacy. In contrast, the secure multi-party computation protocols from the cryptographic literature [3, 9, 10, 11, 20, 21, 23, 26, 39] would allow us to achieve our security goals, but are not practical at the scale we have in mind. Moreover, few of these systems have ever been implemented [3, 13, 23], let alone operate in

the real world [4] and at scale. So, a meta-goal of our work is to help bring multi-party computation to life.

In this paper, we *propose, implement, and evaluate* a viable alternative—a “semi-centralized” system architecture, and associated cryptographic protocols, that provides privacy-preserving data aggregation without sacrificing efficiency. Rather than having a single aggregator node, the data analysis is split between two separate parties—a *proxy* and a *database*. The proxy plays the role of obviously blinding client inputs, as well as transmitting blinded inputs to the database. The database, on the other hand, builds a table that is indexed by the blinded key. For each row of this table whose values satisfy f , the database shares this row with the proxy, who unblinds the key. The database subsequently publishes its non-blinded data for that key.

The resulting semi-centralized system provides strong privacy guarantees *provided that the proxy and the database do not collude*. In practice, we imagine that these two components will be managed either by the participants themselves that do not wish to see their own information leaked to others, perhaps even on a rotating basis, or even third-party commercial or non-profit entities tasked with providing such functionality. For example, in the case of cooperative DoS detection, ISPs like AT&T and Sprint could jointly provide the service, or, perhaps even better, third-party entities like Google (which already plays a role in bot and malware detection [15]) or the EFF (which has funded anonymity tools such as Tor [7]), which themselves have no incentive to collude. Such a separation of trust appears in several cryptographic protocols [6], and even in some natural real-world scenarios, such as Democrats and Republicans jointly comprising election boards in the U.S. political system. It should be emphasized that the proxy and database are not treated as *trusted parties*—we only assume that they will not collude. Indeed, jumping ahead, our protocol does not reveal sensitive information to the proxy or to the database.

Using a semi-centralized architecture greatly reduces operational complexity and simplifies the liveness assumptions of the system. For example, clients can asynchronously provide their key-value tuples without our system requiring any complex scheduling. Despite these simplifications, the cryptographic protocols necessary to provide strong privacy guarantees are still non-trivial. Specifically, our solution makes use of oblivious pseudorandom functions [11, 16, 27], amortized oblivious transfer [17, 25], and homomorphic encryption with re-randomization. We formally prove that our system guarantees keyword and participant privacy. Our proofs are in the *honest-but-curious* model, where, informally, each party can perform local computation on its own view in an attempt to break privacy, but still faithfully follows the protocol. The protocol does, however, preserve privacy in the face of collusion between either party and any number of clients (*e.g.*, the database impersonating a client).

The remainder of the paper is organized as follows. Section 2 defines our system goals and discusses why prior techniques are not sufficient. Section 3 describes our PDA protocols and sketches the proofs of their privacy guarantees. Section 4 describes our implementation, and Section 5 evaluates its performance. We conclude the paper in Section 6 with a discussion of future research directions.

2 Design Goals and the Status Quo

This section defines our goals for a practical system for large-scale privacy-preserving data aggregation (PDA), and we discuss how prior proposals failed to meet these requirements. We then expand on our security assumptions and privacy definitions to set the stage for the presentation of our protocol in the next section.

2.1 Design Goals

In the private data aggregation problem, a collection of participants (or *clients*) may autonomously make observations about *values* (v_i) associated with *keys* (k_i). These observations may be, for example, the fact that an IP address is suspected to have performed some type of attack (through DoS, spam, phishing, and so forth), or the number of participants that associate a particular credential with a server. The system jointly computes a two-column input table T . The first “key” column of T is a set comprised of all unique keys belonging to all participants. The second “value” column is comprised of a value $T[k_i]$ that is the aggregation or union of all participant’s values for k_i . The system then defines a particular function f to be evaluated over each row’s value(s). For simplicity, we focus our discussion on the simple problem of over-threshold set intersection for f : If clients’ inputs of the form $\langle k_i, 1 \rangle$ are aggregated as $T[k_i] \leftarrow T[k_i] + 1$, is $T[k_i] \geq \tau$?

A practical PDA system should provide the following:

- **Keyword privacy:** We say a system satisfies *keyword privacy* if, given the above aggregated table T , at the conclusion of the protocol all involved parties learn only (1) all keys k_i whose corresponding aggregate value $T[k_i] \geq \tau$ and (2) all values $T[k_i]$ (*i.e.*, the entire value column of T). As an example, for over-threshold set intersection, everybody learns a histogram over all keys, in addition to keys (and their correspond total) that have appeared in at least τ clients’ inputs. We discuss later in this section why we reveal the keyless “histogram” in addition to those over-threshold keys.
- **Participant privacy:** We say a system satisfies *participant privacy* if, at the conclusion of the protocol, nobody can learn the inputs $\{\langle k_{i,j}, v_{i,j} \rangle\}$ of participant p_j other than p_j himself (except for information which is trivially deduced from the output of the function). This is formally captured by showing that the protocol leaks no more information than an ideal implementation that uses a trusted third party. This convention is standard in secure multi-party computation; further details can be found in [14].
- **Efficiency:** The system should scale to large numbers of participants, each generating and inputting large numbers of observations (key-value tuples). The system should be scalable both in terms of the network bandwidth it consumes (communication complexity), as well as the computational resources needed to execute private data aggregation (computational complexity).
- **Flexibility:** There are a variety of computations one might wish to perform over each key’s values $T[k_i]$, other than a simple threshold test. These may include finding the maximum value for a given key, or checking if the median of a row exceeds a threshold. Rather than design a new protocol for each function f , we prefer to have a single protocol that works for a wide range of functions.
- **Lack of coordination:** Finally, the system should operate without requiring that all participants coordinate their efforts to jointly execute some protocol at the same time, or even all be online around the same time. Furthermore, no set of participants should be able to prevent others from executing the protocol and computing their own results (*i.e.*, a liveness property).

As we discuss next, existing approaches fail to satisfy one or more of these goals.

Approach	Keyword Privacy	Participant Privacy	Efficiency	Flexibility	Lack of Coordination
Private Set Intersection	Yes	Yes	Poor	No	No
Garbled-Circuit Evaluation	Yes	Yes	Very Poor	Yes	No
Hashing Inputs	No	No	Very Good	Yes	Yes
Network Anonymization	No	Yes	Very Good	Yes	Yes
This paper	Yes	Yes	Good	Yes	Yes

Table 1: Comparison of proposed schemes for private data aggregation

2.2 Limitations of Existing Approaches

Having defined these five goals for PDA, we next consider several possible solutions from the literature. We see that prior secure multi-party computation protocols achieve strong privacy at the cost of efficiency, flexibility, or easy coordination. Simple hashing or network-layer anonymization approaches fail to satisfy our privacy requirements, on the other hand. Our protocol, which leverages insights from both approaches, combines the best of both worlds. Table 1 summarizes the discussion in this section.

Set-Intersection Protocols. Freedman *et al.* [10] proposed a specially-designed secure multi-party computation protocol to compute set intersection between the input lists of two parties. It represented each parties’ inputs as the roots of an encrypted polynomial, and then had the other party evaluate this encrypted polynomial on each of its own inputs. While asymptotically optimized for this setting, our own careful protocol implementation found two sets of 100 items each took a full 213 seconds to execute (on a 3 Ghz Intel machine) [13]. Kissner and Song [20] extended and further improved this polynomial-based protocol for a multi-party decentralized setting, yet their computational complexity remains $O(n\ell^2)$ and communication complexity of $O(n^2\ell)$, where n is the number of participants and ℓ is the number of input elements per party. Furthermore, after a number of pairwise interactions between participants, the system needed to coordinate a group decryption protocol between all parties. Hence, this prior work on set-intersection faces scaling challenges on large sets of inputs or participants, and it also requires new protocol design for each small variant of the set-intersection or threshold set-intersection protocol.

Secure Multi-Party Computations using Garbled Circuits. In 1982, Yao [39] proposed a general technique for computing any two-party computation privately, by building a “garbled circuit” in which one party encodes the function to be executed and his own input, and the other party obviously evaluates her inputs on this circuit. Very recently, the Fairplay system [3, 23] provided a high-level programming language for automatically compiling specified functions down into garbled circuits and generating network protocol handlers to execute them. While such a system would provide the privacy properties we require and offer the flexibility that hand-crafted set-intersection protocols lack, this comes at a cost. These protocols are even more expensive in both computation and communication, requiring careful coordination as well.

Hashing Inputs. Rather than building fully decentralized protocols—with the coordination complexity and quadratic overhead (in n) this entails—we could aggregate data and compute results using a centralized server. One approach is to simply have clients first hash their keys before submitting them to the server (*e.g.*, using SHA-256), so that a server only sees $H(k_i)$, not k_i itself [2]. While it may be difficult to find a pre-image of a hash function, brute force attacks are still always possible: In our collaborating intrusion detection application, for instance, a server

can simply compute the hash values of all four billion IP addresses and build a simple lookup table. Thus, while certainly efficient, this approach fails to achieve either of our privacy properties. An alternative that prevents such a brute-force attack would be for all participants (clients) to coordinate and jointly agree on some secret key s , then use instead a *keyed* pseudorandom function on the input key, *i.e.*, $F_s(k_i)$. This would satisfy keyword privacy, until a single client decides to share s with the server, a brittle condition for sure.

Network Anonymization through Proxying. In the previous proposal, the server received inputs directly from clients. Thus, the server was always able to associate a row of the database with a particular client, whether or not its key is known. One solution would be to simply proxy a client’s request through one or more intermediate proxies that hides the client’s identity (*e.g.*, its own IP address), as done in onion routing systems such as Tor [7]. Of course, this solution still does not achieve keyword privacy.

Although the prior approaches have their limitations, they also offer important insights that inform our design. First, a more centralized aggregation architecture avoids distributed coordination and communication overhead. Second, proxying can add participant privacy when interacting with a server. And third, a keyed pseudorandom function (PRF) can provide keyword privacy. Now, the final insight to our design is, *rather than have all participants jointly agree on this PRF secret s , let it be chosen by and remain known only to the proxy*. After all, the proxy is already trusted not to expose a client’s identity to the server (database), so let’s trust it not to expose this secret s to the database as well. Thus, prior to proxying (roughly) the tuple $\langle F_s(k_i), v_i \rangle$, the proxy executes a protocol with a client to *blind* its input key k_i with F_s . This blinding occurs in such a way that the client does not learn s and the proxy does not learn k_i .¹ This completes the loop, having a proxy play a role in providing both keyword and participant privacy, while the database offers flexibility in any computation over a key’s values $T[k_i]$ and scalability through traditional replication and data-partitioning techniques (*e.g.*, consistent hashing [19]).

2.3 Security Assumptions and Privacy Definitions

We now motivate and clarify some design decisions related to our security assumptions and privacy definitions.

Honest-but-Curious Parties. In our model, parties are expected to act as *honest-but-curious* (also called *semi-honest*) participants. That is, each party can perform local computation on its own view in an attempt to break privacy, but is assumed to still faithfully follow the protocol when interacting with other parties. We believe this model is very appropriate for our semi-centralized system architecture. In many deployments, the database and proxy may be well-known and trusted to act on their good intentions to the best of their abilities, as opposed to simply another participant amongst a set of mutually distrustful parties. Still, other than fully compromising a server-side component and secretly replacing it with an actively malicious instance, data breaches are not really possible, as participants never see privacy-comprising data in the first place. In addition, the honest-but-curious model is one of the two standard security models in multi-party computation protocols—the other being the (obviously stronger) assumption of full malicious behavior. Unfortunately, security against fully malicious behavior comes at a great cost, as each party needs to prove at each step of the protocol that it is faithfully obeying it. For example, the proxy would need to prove that it does not omit any submitted inputs while proxying, nor falsely open blinded

¹We note that oblivious pseudorandom function evaluation had been previously used in the set intersection context in [11] and [16] as well.

keys at the end of the protocol; the database would need to prove that it faithfully aggregates submitted values, nor omit any rows in T that satisfy f . These proofs, typically done in zero-knowledge, greatly complicate the protocol and impact efficiency. As such, this paper focuses on honest-but-curious adversaries.

Security against Coalitions. Another important aspect of security is the ability to preserve privacy even when several adversarial players try to break security by sharing the information they gained during the protocol. In this aspect, we insist on gaining security against any coalition that consists of an arbitrary number of participants together with the database. This is essential as otherwise the database can perform a Sybil attack [8], *i.e.*, create many dummy participants and use their views together with his own view to reveal sensitive information. Similarly, we require security against any coalition of the proxy and the participants. On the other hand, in order to have a highly efficient and scalable system, we are willing to tolerate vulnerability against a coalition of the database and the proxy. While not ideal, this relaxation does not introduce the problem of Sybil attacks, as the proxy and database are *distinct* players. As mentioned earlier, we imagine that the proxy and database are run by competing agents, or by two randomly chosen participants. This presumably reduces the chances of adversarial collaboration between these parties.

Releasing the Value Column. The functionality that we securely compute includes the entire value column ($T[k_i], \forall i$), even for those values whose corresponding keys k_i are not exposed. This serves practical purpose, as it may be hard to fully specify f *a priori* to collecting clients' inputs. Specifically, how should an anomaly detection system choose the appropriate frequency threshold τ ? In some attacks, 10 observations about a particular IP address may be high (*e.g.*, suspected phishing), while in others, 1000 observations may be necessary (*e.g.*, for bots participating in multiple DoS attacks). Furthermore, a dataset may naturally expose a clear gap between frequency counts of normal and anomalous behavior; the very reason data operators like to “play” with raw data in the first place.

However, in some settings releasing the entire value column ($T[k_i], \forall i$) may be considered to be an information “leakage”. In such cases, one may modify our protocol such that the database does not publish this information. This is only a partial solution as the database still learns this column. In particular, when the domain $\mathcal{D}$ of possible values is large, a client can collaborate with the database and “mark” a key k by submitting it together with an uncommon value $w \in \mathcal{D}$. Then, the database can identify the blinded version of k' by searching for the $T[k']$ that includes w . This way, the database can discover other clients' values for that same key. This problem does not arise when the domain is relatively small (*e.g.*, when values are grades over some limited scale). We mention that the leakage of the value column can be completely eliminated by asking the participants to encrypt their values under both, the proxy and the database public keys, and then by using additional cryptographic protocols for the aggregation of the values. While these tools are relatively expensive, the structure of our system allows us to employ them only for the two-party case (for the proxy and database) which results in a significant efficiency improvement over other non-centralized solutions.

3 Private Keyword Aggregation Protocol

3.1 Overview

We begin with a high level (and slightly simplified) description of our protocol. We focus on the specific case of threshold set-intersection problem. Recall that in addition to the participants (clients), we will have two distinct parties: the proxy and database (DB). These special parties are

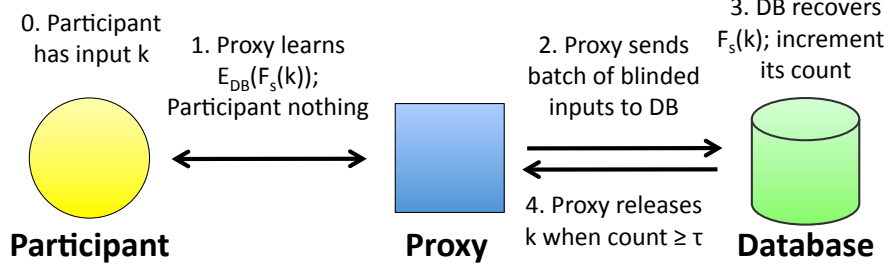


Figure 1: High-level system architecture and protocol. F_s is a keyed hash function whose secret key s is known only to the proxy.

somehow chosen either among the participants or as third-party entities, and we assume that these two parties do not collude in order to break privacy. These parties may try to cheat, however, by colluding with one or more participants.

The basic approach is to blind all participants' input keys $k_1, \dots, k_\ell$ to some pseudorandom inputs $F_s(k_1), \dots, F_s(k_\ell)$, and let the DB learn the union of all these blinded inputs. The DB aggregates these blinded inputs and counts their frequencies. Then, the DB can identify which of the blinded inputs appear over the threshold. We have to define some revealing mechanism that allows the DB to “unblind” entries which appear over the threshold, but let us ignore this issue for now.

Recall also that we wish to maintain certain privacy properties: (1) The DB should be unable to compute the mapping $k_i \leftarrow F_s(k_i)$. (2) The DB should be unable to identify which party sent each blinded input $F_s(k_i)$. (3) The participants themselves should not learn the blinded value of their own inputs (as otherwise, a coalition of the DB and participants can break privacy).

We achieve this by making a novel use of the proxy. Specifically, we let the proxy choose the blinding scheme (by choosing a key s for a pseudorandom function F). In addition, the proxy will hide the identity of the participants from the DB by passing their blinded inputs to the DB. However, this introduces a new problem: Knowing both the secret mapping s and the blinded value $F_s(k_i)$ allows the proxy to break the security of the protocol.

To solve this problem, we leverage another cryptographic tool. We employ a public-key encryption scheme E , and assume that the DB chooses a public/private key pair (DB, DB^{-1}) and publishes the public-key DB . The protocol proceeds as follows, as shown in Figure 1. (1) For each input k , each participant invokes a cryptographic sub-protocol with the proxy such that, at the end of the protocol, the proxy learns an encryption of $F_s(k_i)$ under the DB's key DB , while the participant learns nothing. (2) The proxy sends the value $E_{DB}(F_s(k_i))$ to the DB. (3) The DB decrypts this value and increments a counter for $F_s(k_i)$. If the counter is greater or equal than the threshold τ , (4) the DB and proxy execute some revealing mechanism in order to reveal k_i .

Security (sketch). We now argue that the protocol is secure at an intuitive level. (A formal proof is omitted due to space limitations.) First, note that the DB sees only a blinded list of inputs encrypted under his public key DB , without being able to relate the blinded entries to their owners. Hence, the DB learns nothing but the frequency table of the inputs (*i.e.*, the $T[k_i]$'s for all k_i 's). Furthermore, this is true even if the DB participates in the protocol as a participant, or more generally if the DB colludes with some subset of the participants (excluding the proxy), as the participants did not see the blinded version of their inputs either. On the other hand, the proxy learns nothing as well, as it only sees entries which are encrypted under the public key of the DB. Again, this is true for any coalition that includes the proxy but does not include the DB.

3.2 Formal Description

We now formally describe our protocol for the general case, where keys are released based on some general computation f which is applied to the aggregated values $T[k_i]$. We postpone the description of the main input-blinding sub-protocol between participant and proxy until Section 3.3.

- **Parties:** Participants, Proxy, DB.
 - **Cryptographic Primitives:** A pseudorandom function F , where $F_s(k_i)$ denotes the value of the function on the input k_i with a key s . A public-key encryption E , where $E_K(x)$ denotes an encryption of x under the public key K .
 - **Public Inputs:** The proxy's public key PRX , the DB's public key DB .
 - **Private Inputs.** *Participant:* A list of key-value pairs $\langle k_i, v_i \rangle$. *Proxy:* key s of PRF F and secret key for PRX ; *DB:* secret key for DB .
1. Each participant interacts with the proxy as follows. For each entry $\langle k_i, v_i \rangle$ in the participant's list, the participant and the proxy run a sub-protocol such that the proxy learns $E_{DB}(F_s(k_i))$ and the participant learns nothing (see Section 3.3). In addition, the participant sends the pair $E_{DB}(v_i)$ and $E_{DB}(E_{PRX}(k_i))$. The proxy adds this triple to a list and waits until all participants send their inputs. Then the proxy randomly permutes the list and sends the result to the DB.²
 2. The DB uses its private key to decrypt each of the three encrypted values. Now, it holds a list of triples of the form $\langle F_s(k_i), v_i, E_{PRX}(k_i) \rangle$. The DB inserts these values into a table which is indexed by the (blinded) key $F_s(k_i)$. At the end, the DB has a table of entries of the form $\langle F_s(k_i), T[k_i], E_{PRX}(k_i) \rangle$, where $T[k_i]$ is (in general) a list of all the v_i 's that appeared with this key (or simply the number of times a client inputted k_i in the case of threshold set intersection).
 3. The DB uses some predefined function f to partition the table into two parts: R , which consists of the rows whose keys should be revealed, and H , which consists of the rows whose keys should remain hidden. Then, it publishes the two-tuples $\langle T[k_i], E_{PRX}(k_i) \rangle$ in R and simply the values $T[k_i]$ in H .
 4. The proxy goes over the table R and replaces all the encrypted $E_{PRX}(k_i)$ entries with their decrypted key k_i .

An alternative realization of Steps 3 and 4 is for the database to query the proxy directly and have it decrypt the $E_{PRX}(k_i)$'s in R . And then the database can publish all tuples $\langle k_i, T[k_i] \rangle \in R$ and $T[k_i] \in H$ itself.

Remark. In some cases, it might be useful to hide the $T[k_i]$ value of R and H from the participants (except of the DB) and instead reveal only the keys of R . This can be easily done by modifying Step 3 and letting the DB send only $E_{PRX}(k_i)$ for all k in R . Then, the proxy (or database) can publish only the decryption of these values.

Also, note that in the case where the values are always one, *i.e.*, the participants only want to increment a counter for some key, the table R simply consists of keys and their frequencies, and H

²In fact, given participant asynchrony, the proxy can continuously accept participant inputs, buffering them only until some queue is filled or some time period passes, rather than waiting for *all* participants to submit inputs.

is simply a frequency table of all the unrevealed keys. So at least in this case, it seems that this additional information is not typically harmful.

Eliminating subliminal channels. Public-key encryption schemes use randomness (in addition to the public key) to encrypt a message. This randomness can be used as a subliminal channel between the participants to the database or to the proxy. To solve this problem, we use an encryption scheme that supports re-randomization of ciphertexts; that is, given an encryption of x with randomness b , it should be possible to recompute an encryption of y under fresh randomness b' (without knowing the private key). Now we can eliminate the subliminal channel by asking the proxy to re-randomize the ciphertexts which are encrypted under the DB's public key (at Step 1), and similarly ask the DB to re-randomize the ciphertexts which are encrypted under the Proxy's public key in Step 3.

3.3 Concrete Instantiation of the Cryptographic Primitives

In the following section, we assume that the input keys are represented by m -bit strings. We assume that m is not very large (*e.g.*, less than 192–256); otherwise, one can hash the input keys and apply the protocol to resulting hashed values.

Public Parameters. Our implementation mostly employs Discrete-Log based schemes. In the following, g is a generator of a multiplicative group $\mathbb{G}$ of prime order p for which the decisional Diffie-Hellman (DDH) assumption holds. We publish (g, p) during initialization and assume the existence of algorithms for multiplication (and thus also for exponentiation) in $\mathbb{G}$. We let Z_p^* denote the multiplicative group modulo p , that is, the set $\{1, \dots, p-1\}$ with the group operation multiplication modulo p .

ElGamal Encryption. We will use ElGamal encryption over the group $\mathbb{G}$. The private key is a random element a from Z_p^* , and the public key is the pair $(g, h = g^a)$. To encrypt a message $x \in \mathbb{G}$, we choose a random b from Z_p^* and compute $g^b, x \cdot h^b$. To decrypt the ciphertext (A, B) , compute $B/A^a = B \cdot A^{-a}$ (where $-a$ is the multiplicative inverse of a in Z_p^*). We will also need the ability to re-randomize the ciphertext; that is, given an encryption of x with randomness b , it should be possible to recompute an encryption of y under fresh randomness b' . This is done as follows: Given (A, B) , choose a random $b' \in Z_p^*$ and compute $(A \cdot g^{b'}, B \cdot h^{b'})$.

Naor-Reingold PRF [27]. The key s of the function $F_s : \{0, 1\}^m \rightarrow \mathbb{G}$ contains m values $(s_1, \dots, s_m)$ chosen randomly from Z_p^* . Given m -bit string $k = x_1 \dots x_m$, the value of $F_s(k)$ is $g^{\prod_{i=1}^m s_i^{x_i}}$, where the exponentiation is computed in the group $\mathbb{G}$.

Oblivious-Transfer [25, 31]. To implement the sub protocol of Step 1, we will need an additional cryptographic tool called Oblivious Transfer (OT). In an OT protocol, we have two parties: sender and receiver. The sender holds two strings (α, β) , and the receiver has a selection bit c . At the end of the protocol, the receiver learns a *single* string: α if $c = 0$, and β if $c = 1$. In addition, the sender learns nothing (in particular, it does not know the value of the selector c).

3.3.1 Client-proxy protocol to blind input keys

Our construction is inspired by a protocol for oblivious evaluation of the PRF F , which is explicit in [11] and implicit in [25, 26]. We believe that this construction might have further applications.

- **Parties:** Participant, Proxy.

- **Inputs.** *Participant:* a single m -bit string $k = (x_1 \dots x_m)$; *Proxy:* a secret key $s = (s_1, \dots, s_m)$ of Naor-Reingold PRF F .

1. Proxy chooses m random values $u_1, \dots, u_m$ from Z_p^* and an additional random $r \in Z_p^*$. Then for each $1 \leq i \leq m$, the proxy and the participant invoke the OT protocol where proxy is the sender with inputs $(u_i, s_i \cdot u_i)$ and receiver uses x_i as his selector bit. That is, if $x_i = 0$, the participant learns u_i and otherwise it learns $s_i \cdot u_i$. The proxy also sends the value $\hat{g} = g^{r/\Pi u_i}$. (All these steps can be applied in parallel.)
2. The participant multiplies together the values received in the OT stage. Let M denote this value. Then, it computes $\hat{g}^M = (g^{\Pi_{i=1} s_i})^r = F_s(k)^r$. Finally, the participant chooses a random element a from Z_p^* and encrypts $F_s(k)^r$ under the public key $\text{DB} = (g, h)$ of the database. The participant sends the result $(g^a, F_s(k)^r \cdot h^a)$ to the proxy.
3. The proxy raises the received pair to the power of r' , where r' is the modular inverse of r modulo p . It also re-randomizes the resulting ciphertext.

Correctness. Recall that $\mathbb{G}$ has a prime order p . Hence, when the pair $(g^a, F_s(k)^r \cdot h^a)$ is raised to the power of $r' = r^{-1}$, the result is $(g^{ar'}, F_s(k) \cdot h^{ar'})$, which is exactly $E_{\text{DB}}(F_s(k))$. Thus, the protocol is correct.

Privacy. All the proxy sees is the random tuple $(u_1, \dots, u_m, r)$ and $E_{\text{DB}}(F_s(k)^r)$. This view gives no additional information except of $E_{\text{DB}}(F_s(k))$. (Formally, the view can be perfectly simulated given $E_{\text{DB}}(F_s(k))$.) On the other hand, we claim that all the participant sees is a sequence of random values and therefore it also learns nothing. Indeed, the participant sees the vector $(s_1^{x_1} \cdot u_1, \dots, s_m^{x_m} \cdot u_m)$, whose entries are randomly distributed over $\mathbb{G}$, as well as the value $\hat{g} = (g^{1/\Pi u_i})^r$. Since r is randomly and independently chosen from Z_p^* , and since $\mathbb{G}$ has a prime order p , the element $\hat{g}$ is also uniformly and independently distributed over $\mathbb{G}$.³

3.3.2 Implementing Oblivious Transfer

In general, oblivious transfer is an expensive public-key operation (*e.g.*, it may take two exponentiations per single invocation). In the above protocol, then, we execute an OT protocol for each *bit* of the participants input k (which would result, for example, in 64 exponentiations just to input a single IP address). However, Ishai *et al.* [17] show how to reduce the amortized cost of OT to be as fast as matrix multiplication. This “batch OT” protocol uses a standard OT protocol as building block. We implemented this batch OT protocol on top of the basic OT protocol of [25].

3.4 Efficiency of our Protocol

The round complexity of the protocol is constant, and the communication complexity is linear in the number of items. The computational complexity of the protocol is dominated by cryptographic operations. For each m -bit input key, we have the following amortized complexity: (1) The participant who holds the input key computes 4 exponentiations, as well as m modular multiplication / symmetric-key operations. (2) The proxy computes 5 exponentiations, as well as m modular multiplication / symmetric-key operations. (3) The database computes 4 exponentiations.

³This is true as long as $g^{1/\Pi u_i} \not\equiv 1 \pmod{p}$, which happens with all but negligible probability.

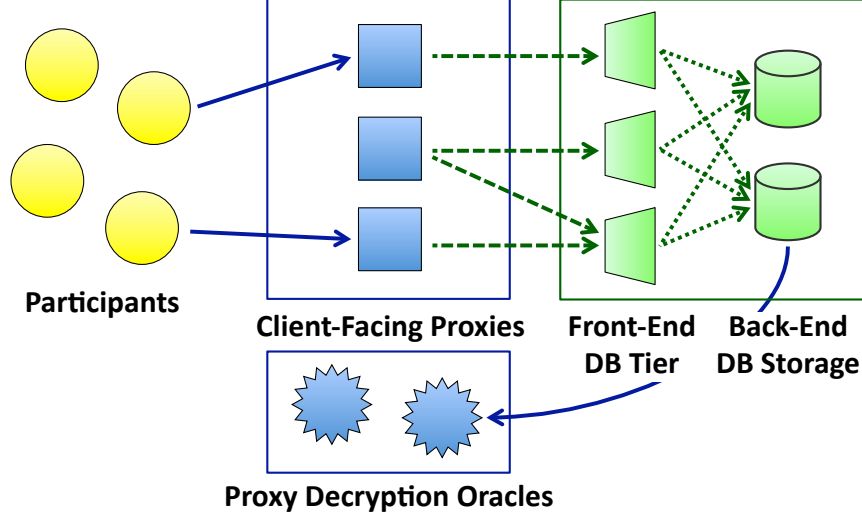


Figure 2: Distributed proxy and database architecture

4 Distributed Architecture and Implementation

In our system, both the proxy and database logical components can be physically replicated in a relatively straightforward manner. In particular, our design can scale out horizontally to handle higher loads, by increasing the number of proxy and/or database replicas, and then distributing requests across these replicas. Our distributed architecture is shown in Figure 2.

4.1 Proxy: Client-Facing Proxies and Decryption Oracles

One administrative domain can operate any number of proxies. Each proxy’s functionality may be logically divided into two components: handling client requests, and serving as decryption oracles for the database when a particular key should be exposed. None of these proxies need to interact, other than having all client-facing proxies use the same secret s to key the pseudorandom function F and all decryption-oracle proxies use the same public/private key PRX . In fact, these two proxies play different logical roles in our system and could even be operated by two different administrative domains. In our current implementation, all proxies register with a single group membership server, although a distributed group membership service could be implemented for additional fault tolerance [5, 38].

To discover a client-facing proxy, a client contacts this group membership service, which returns a proxy IP address in round-robin order (this could be replaced by any technique for server selection, including DNS, HTTP redirection, or a local load balancer). To submit its inputs, a client connects with this proxy and then executes an amortized Oblivious Transfer (OT) protocol on its input batch. This results in the proxy learning $\langle E_{\text{DB}}(F_s(k_i)), E_{\text{DB}}(v_i), E_{\text{DB}}(E_{\text{PRX}}(k_i)) \rangle$ for each input tuple, which it pushes onto an internal queue. (While the Section 3.3 only described the use of ElGamal encryption, its special properties are only needed for $E_{\text{DB}}(F_s(k_i))$; the other public-key operations can be RSA, which we use in our implementation.) When this queue reaches a certain length (10000 in our implementation), the proxy randomly permutes (shuffles) the items in the queue, and sends them to a database server.

The database, upon determining that a key k_i ’s value satisfies f , sends $E_{\text{PRX}}(k_i)$ to a proxy-decryption oracle. The proxy-decryption oracle decrypts $E_{\text{PRX}}(k_i)$ and returns k_i to the database

for storage and subsequent release to other participants in the system.

4.2 Database: Front-end Decryption and Back-end Storage

The database component can also be replicated. Similar to the proxy, we again separate database functionality into two parts: the *front-end* module that handles proxy submissions and decrypts inputs, and a *back-end* module that serves as a storage layer. Each logical module can in turn be replicated in a similar fashion as the proxy.

The servers comprising the front-end database tier do not need to interact, other than being configured with the same public/private keypair DB. Thus, any front-end database can decrypt the $E_{\text{DB}}(F_s(k_i))$ input supplied by a proxy, and the proxies can load balance input batches across these database servers.

The back-end database storage, on the other hand, needs to be more tightly coordinated, as we ultimately need to aggregate all $F_s(k_i)$'s together, no matter which proxy or front-end database processed them. Thus, the back-end storage tier partitions the keyspace of all 1024-bit strings over all storage nodes (using consistent hashing [19]). All such front-end and back-end database instances also register with a group membership server, which the front-end servers contact to determine the list of back-end storage nodes. Upon decrypting an input, the front-end node determines which back-end storage node is assigned the resulting key $F_s(k_i)$, and sends the tuple $\langle F_s(k_i), v_i, E_{\text{PRX}}(k_i) \rangle$ to this storage node.

As these storage nodes each accumulate a horizontal portion of the entire table T , they test the value column for their local table to see if any keys satisfy F . For each such row, the storage node sends the tuple $\langle F_s(k_i), T[k_i], E_{\text{PRX}}(k_i) \rangle$ to a proxy-decryption oracle.

4.3 Prototype Implementation

Our design is implemented in roughly 5,000 lines of C++. All communication between system components—client, front-end proxy, front-end database, back-end database storage, and proxy-decryption oracle—is over TCP using BSD sockets. We use the GnuPG library for large numbers (bignums) and cryptographic primitives (*e.g.*, RSA, ElGamal, and AES). The Oblivious Transfer protocol (and its amortized variant) were implemented from scratch, and comprised a total of 625 lines of code. All RSA encryption used a 1024-bit key, and ElGamal used a corresponding 1024-bit group size. AES-256 was used in the batch OT and its underlying OT primitive. The back-end database simply stores table rows in memory, although we plan to replace this with a durable key-value store (*e.g.*, BerkeleyDB [28]).

5 Performance Evaluation

We wish to evaluate our system along three primary dimensions: (a) given fixed computing resources, what is the throughput of our system as a function of the size of the input set? (b) what are the primary factors limiting throughput? and (c) how does the throughput scale with increasing computing resources? In each case, we are concerned with both (1) how long it takes for clients to send key-value pairs to the proxy during the OT phase (*proxy throughput*) and (2) how long it takes for the DB to decrypt and identify keys with values that satisfy the function f (*DB throughput*). We have instrumented our code to measure both. For a given experiment requiring the proxy to process n keys, proxy throughput is defined as n divided by the time it takes between when the first client contacts any client-facing proxy and when the last key is processed by some client-facing proxy. Similarly, database throughput is defined as the number of keys processed between when the

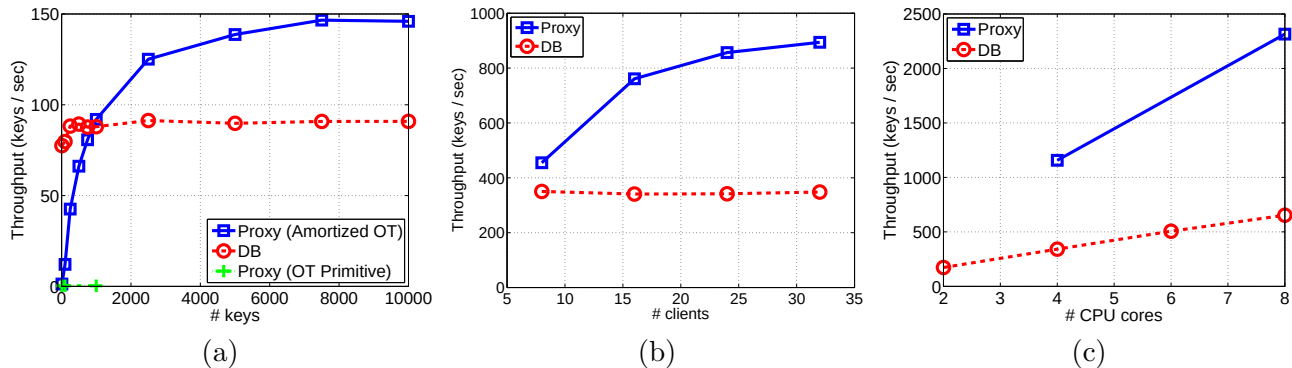


Figure 3: *Scaling*: Effect of (a) number of keys, (b) number of participants, and (c) number of proxy/database replicas.

first client-facing proxy forwards keys to some database front-end and when the database back-end storage processes the last submitted keys.

Our experiments were run on multiple machines. The servers (proxy and DB) were operated on HP DL160 servers (quad-core Intel Xeon 2.00 Ghz machines with 4 GB RAM running CentOS Linux). These machines can perform 1024-bit ElGamal encryption in 2.2 ms, ElGamal decryption in 2.5 ms, RSA encryption in 0.5 ms, and RSA decryption in 2.8 ms. Due to resource limitations, the clients were run on different machines depending on the experiment. The machines used for the clients were either of the same configuration as the servers, or on either (A) Sun SunFire X4100 servers with two dual-core 2.2 GHz Opteron 275 processors (four 64-bit cores) with 16GB RAM running CentOS, or (B) Dell PowerEdge 2650 servers with two 2.2 GHz Intel Xeon processors and 5 GB of memory, also running Linux.

5.1 Scaling and Bottleneck Analysis

Effect of number of keys (Figure 3a). The input trace to our system is parameterized by the number of clients and by the number of keys they each submit. In Figure 3a we measure the throughput of our system as a function of the number of keys. More precisely, we run a single client, a single proxy, and a single DB in order to measure single-CPU-core proxy throughput and single-CPU-core DB throughput. The top curve shows proxy throughput when the proxy and client utilize the amortized OT protocol, the middle curve shows DB throughput, and the bottom partial curve shows proxy throughput when the proxy and client utilize only the standard OT primitive, which does not include our amortization-based extensions. The throughput of the OT primitive is exceedingly low (less than one key per second), which is why it was not evaluated on the full range of x-values.

Proxy throughput scales well with the number of incoming keys when the client and proxy utilize the amortized OT protocol. Throughput increases with increasing numbers of keys because the amortized OT calls the primitive OT a fixed number of k times regardless of the number of input addresses n . With small n (e.g., up to 1000), the cost of these calls to the primitive OT dominate overall execution time and leave the proxy underutilized. However, as the size of the input set increases, the cost of encrypting keys on the client becomes the primary bottleneck, which is why there is a minimal increase in throughput above $n = 8000$.

DB throughput, on the other hand, does not scale with the number of keys. The reason for this is that the intensive work on the DB is decryption, which is performed in batch, and it is therefore entirely CPU limited. The DB becomes CPU limited at 10 keys and hence remains CPU limited

Global		Within amortized OT						
wait	encrypt	wait	pow	AES	arith	rand	other	OT primitive
60%	1%	0%	16%	4%	4%	4%	2%	7%

Table 2: Breakdown of proxy resource usage

at 10,000 keys (*i.e.*, latency goes up and throughput remains constant). We noted earlier that the machines the DB and proxy run on require 2.5 milliseconds per decryption. Since the DB has to perform 3 decryptions per key, the DB therefore has a maximum throughput of 135 keys per second on a single CPU core (assuming no queuing delays and no other computation). Figure 3a shows that our DB implementation achieves throughput of roughly 90 keys per second.

The amortized OT introduces a trade-off between the proxy-client message overhead and memory consumption. In fact, the amortized OT protocol is the only component of our system whose memory footprint can be nontrivial. In particular, the memory footprint of the OT protocol is essentially determined by the size of the y matrix in [17], which requires precisely $n \times 32 \times 2 \times 1024 / 8 = 8196n$ bytes (*i.e.*, we assume 32 bits per key, the 2 values for the OT primitive, 1024-bit keys, and 8 bits per byte). For $n = 10,000$ keys, for example, this requires 82 MB. The proxy computes this matrix and then sends it to the client. It can therefore either compute it in one batch or split it up and send the rows individually. Computing it in batch is faster but more memory intensive. For subsequent experiments we chose to send up to 5,000 keys at a time because Figure 3a shows that there is little to gain from larger batch sizes (and memory is very affordable), but this parameter is trivially configurable if another setting is desired.

Effect of number of participants (Figure 3b). Here we evaluate the throughput of our system as a function of the number of clients sending keys. In this experiment we limit the proxy and DB to one server machine each. Four client-facing proxy processes are launched on one machine and four front-end DB processes are launched on the other. They can therefore potentially utilize all eight cores on these two machines. Figure 3b shows that the proxy scales well with the number of clients. Proxy throughput increases by nearly a factor of two between 8 and 32 clients. This tells us that when communicating with a single client, a proxy spends a substantial fraction of its time idling. The four proxies in this experiment are not CPU limited until they are handling 32 clients, at which the throughput approaches 900 keys per second. The DB, on the other hand, is CPU-bound throughout. It has a throughput of approximately 350 keys per second independent of the number of clients.

Effect of number of replicas (Figure 3c). Finally, we wish to analyze how our distributed architecture scales with the available computing resources. In this experiment we provide up to 8 cores across two machines to each of the proxy and DB front-ends. While the proxy is evaluated on 64 clients, computing resource constraints meant that the DB is evaluated on 32 clients.

Both our proxy and DB scale linearly with the number of CPU cores allocated to them. Throughput for the DB with 2 cores when handling 32 clients was over 173 keys per second whereas at 8 cores the throughput was 651 keys per second—a factor of 3.75 increase in throughput for a factor of 4 increase in computing resources. The proxy has throughput of 1159 keys per second when utilizing 4 cores and 2319 when utilizing 8 cores—an exact factor of 2 increase in throughput for an equal increase in computing resources. This clearly demonstrates that our protocol, architecture, and implementation can scale up to handle very large data sets. In particular, our entire system could handle input sizes on the order of millions of keys in hours.

Micro-benchmarks. To gain a deeper understanding of the factors limiting the scalability of

Global		Within amortized OT						
wait	encrypt	wait	pow	AES	arith	rand	other	OT primitive
0%	40%	31%	16%	2%	1%	0%	3%	7%

Table 3: Breakdown of client resource usage

our design, we instrumented the code to account for how the client and proxy were spending their CPU cycles. While the DB is entirely CPU bound due only to decryptions (*i.e.*, its limitations are known), the proxy and client engage in the oblivious OT protocol whose bottlenecks are less clear. In Tables 2 and 3, we therefore show the fraction of time the client and proxy, respectively, spend performing various tasks needed for their exchange. In this experiment we have a single client send keys to a single proxy at the maximum achievable rate.

At the highest level, we split the tasks performed into (a) waiting (called “wait”), (b) encrypting or decrypting values (“encrypt”), or (c) engaging in the amortized OT protocol. We further split work within the amortized OT protocol into time spent waiting, performing power modulo (“pow” or powmod), calling AES256, performing basic arithmetic such as multiplication, division, or finding multiplicative inverses (“arith”), generating random numbers (“rand”), calling the OT primitive, and any other necessary tasks (“other”) such as XOR’ing numbers, allocating or de-allocating memory, etc.

Table 2 shows that when communicating with a single client, the client-facing proxy spends more than 60% of its time idling while waiting for the client—it is *more than* 60% because some part of the 7% of time spent within the OT primitive is also idle time. The 60% idle time is primarily due to waiting for the client to encrypt k_i and $F_s(k_i)$. The single largest computational expense for the proxy is performing powmods at 16%; the remaining non-OT tasks add up to 15%. [60+16+15+7=98% is not equal to 100% probably due to time spent in automatic destructors (in particular, the matrices used for the amortized OT) and rounding issues.] In order to make the proxy more efficient, therefore, utilizing a bignum library with faster powmod and basic arithmetic would be advantageous.

The client also spends a non-trivial amount of time waiting—31% of total execution time—but substantially less than the proxy. It spends 40% of its time encrypting values. The reason this 40% does not match up with the 60% idle time of the proxy is because the proxy finishes its portion of the amortized OT before the client does its. That is, 20 out of the proxy’s 60% idle time is due to the client parsing the proxy’s y matrix (notation consistent with [17]) and 40 is due to the client encrypting its values. As with the proxy, the client would benefit from faster powmods, but encryption is clearly the major bottleneck. We noted before that the cryptographic library we use (GnuPG) performed public-key operations in approximately 2.5–2.8 ms. On the same servers, we benchmarked the Crypto++ library to provide RSA decryption in only 1.2 ms, fully 230% faster. Thus, we are currently changing our implementation to use these libraries, and we fully expect a similar gain in overall system throughput.

5.2 Feasibility of Supporting Applications

In this section, we revisit several potential applications of our system, and consider our results in light of their potential demands on request rate—the number of requests per unit time that must be satisfied, the number of keys which must be stored in the system, and the number of participants.

Anomaly detection. Network operators commonly run systems to detect and localize anomalous behavior within their networks. These systems dynamically track the traffic mix, for example the

volume of traffic over various links, or the degree of fanout from a particular host, and detect behavior that differs substantially from the statistical norm. For example, Mao *et al.* [24] found that most DDoS attacks observed within a large ISP were sourced by fewer than 10,000 source IPs, and generated 31,612 alarms over a four-week period (0.8 events per hour). In addition, Soule *et al.* [35] found that volume anomalies occurred at a rate of four per day on average, most of which involved fewer than several hundred source IPs. Finally, Ramachandran *et al.* [32] found were able to localize 4,963 Bobax-infected host IPs sending spam from a single vantage point. We envision our system could be used to improve accuracy of these techniques by correlating anomalies across ISP boundaries. We found our system could handle 10,000 IP addresses as keys, with a request rate of several hundred keys per second, even with several hundred participants. Given our system exceeds the requirements of anomaly detection, our system may enable the participants to “tune” their anomaly detectors to be more sensitive, and reduce false positive rates by leveraging other ISPs’ observations.

Cross-checking certificates. Multiple vantage points may be used to validate authenticity of information (such as a DNS reply or ssh certificate [29, 37]) in the presence of “man-in-the-middle” attacks. Such environments present potentially larger scaling challenges due to the potentially large number of keys that could be inserted. According to [18], most hosts execute fewer than 15 DNS lookups per hour, and according to [34], ssh hosts rarely authenticate with more than 30 remote hosts over long periods of time. Here, we envision our system could simplify the deployment of such schemes, by reducing the amount of information revealed about clients’ request streams. Under this workload (15 key updates per hour, with 30 keys per participating host), our system scales to support several hundred hosts with only a single proxy. Extrapolating out to larger workloads, our system can handle tens of thousands of clients storing tens of thousands of keys with under fifty proxy/database pairs.

Distributed ranking. Search tools such as Alexa and Google Toolbar collect information about user behavior to refine search results returned to users. However, such tools are occasionally labeled as *spyware* as they reveal information about the contents of queries performed by users. Our tool may be used to improve privacy of user submissions to these databases. It is estimated that Alexa Toolbar has 180,000 active users, and it is known that average web users browse 120 pages per day. Here, the number of participants is large, but the number of keys they individually store in the system is smaller. Extrapolating our results to 180,000 participants, and assuming several thousands of keys, our system can still process several hundred requests per second (corresponding to several hundred thousand clients) per proxy/database pair.

6 Conclusions

In this paper, we presented the design, implementation, and evaluation of a collaborative data-analysis system that is both scalable and privacy preserving. Since a fully-distributed solution would be complex and inefficient, our design divides responsibility between two independent parties—a proxy that obviously blinds the client inputs and a database that identifies the (blinded) keys that have values satisfying an evaluation function. The functionality of both the proxy and the database can be easily distributed for greater scalability and reliability. Experiments with our prototype implementation show that our system performs well under increasing numbers of keys, participants, and proxy/database replicas. The performance is well within the requirements of our motivating applications, such as collaborating to detect the malicious hosts responsible for DoS attacks or to validate the authenticity of information in the presence of man-in-the-middle attacks.

As part of our ongoing work, we plan to evaluate our system in the context of several real applications—first through a trace-driven evaluation and later by extending our prototype to run these applications. In addition, we plan to explore opportunities to deploy our system in practice. A promising avenue is distributed Internet monitoring infrastructures such as NetDimes [36] and the new M-Lab (Measurement Lab) initiative [22]. We believe our system could lower the barriers to collaborative data analysis over the Internet, enabling a wide range of new applications that could improve Internet security, performance, and reliability.

References

- [1] Alexa the Web Information Company, 2009. <http://www.alexa.com/>.
- [2] M. Allman, E. Blanton, V. Paxson, and S. Shenker. Fighting coordinated attackers with cross-organizational information sharing. In *HotNets*, November 2006.
- [3] A. Ben-David, N. Nisan, and B. Pinkas. FairplayMP: A system for secure multi-party computation. In *Proc. ACM Computer and Communications Security Conference*, October 2008.
- [4] P. Bogetoft, D. L. Christensen, I. Damgard, M. Geisler, T. Jakobsen, M. Krøigaard, J. D. Nielsen, J. B. Nielsen, K. Nielsen, J. Pagter, M. Schwartzbach, and T. Toft. Multiparty computation goes live. Cryptology ePrint Archive, Report 2008/068, 2008. <http://eprint.iacr.org/>.
- [5] M. Burrows. The Chubby lock service for loosely-coupled distributed systems. In *Proc. OSDI*, November 2006.
- [6] B. Chor, O. Goldreich, E. Kushilevitz, and M. Sudan. Private information retrieval. *Journal of the ACM*, 45(6), November 1998.
- [7] R. Dingledine, N. Mathewson, and P. Syverson. Tor: The second-generation onion router. In *Proc. 13th USENIX Security Symposium*, August 2004.
- [8] J. R. Douceur. The Sybil attack. In *Proc. Intl. Workshop on Peer-to-Peer Systems*, March 2002.
- [9] R. Fagin, M. Naor, and P. Winkler. Comparing information without leaking it. *Communications of the ACM*, 39(5):77–85, 1996.
- [10] M. J. Freedman, K. Nissim, and B. Pinkas. Efficient private matching and set intersection. In *Advances in Cryptology — EUROCRYPT*, May 2004.
- [11] M. J. Freedman, Y. Ishai, B. Pinkas, and O. Reingold. Keyword search and oblivious pseudorandom functions. In *Proc. Theory of Cryptography Conference*, February 2005.
- [12] Friend-of-a-Friend Project, 2009. <http://www.foaf-project.org/>.
- [13] S. Garriss, M. Kaminsky, M. J. Freedman, B. Karp, D. Mazières, and H. Yu. RE: Reliable email. In *NSDI*, May 2006.
- [14] O. Goldreich. *Foundations of Cryptography: Basic Applications*. Cambridge University Press, 2004.
- [15] Google Safe Browsing for Firefox, 2009. <http://www.google.com/tools/firefox/safebrowsing/>.
- [16] C. Hazay and Y. Lindell. Efficient protocols for set intersection and pattern matching with security against malicious and covert adversaries. In *Proc. Theory of Cryptography Conference*, March 2008.
- [17] Y. Ishai, J. Kilian, K. Nissim, and E. Petrank. Extending oblivious transfers efficiently. In *Advances in Cryptology — CRYPTO*, August 2003.
- [18] J. Jung, E. Sit, H. Balakrishnan, and R. Morris. DNS performance and the effectiveness of caching. *IEEE/ACM Trans. Networking*, 10(5), October 2002.

- [19] D. Karger, E. Lehman, T. Leighton, R. Panigrahy, M. Levine, and D. Lewin. Consistent hashing and random trees: Distributed caching protocols for relieving hot spots on the world wide web. In *ACM Symposium on Theory of Computing*, 1997.
- [20] L. Kissner and D. Song. Privacy preserving set operations. In *Advances in Cryptology — CRYPTO*, August 2005.
- [21] Y. Lindell and B. Pinkas. Privacy preserving data mining. In *Advances in Cryptology — CRYPTO*, August 2000.
- [22] M-Lab: Welcome to Measurement Lab, 2009. <http://www.measurementlab.net/>.
- [23] D. Malkhi, N. Nisan, B. Pinkas, and Y. Sella. Fairplay: A secure two-party computation system. In *Proc. USENIX Security*, August 2004.
- [24] Z. Mao, V. Sekar, O. Spatscheck, J. van der Merwe, and R. Vasudevan. Analyzing large DDoS attacks using multiple data sources. In *SIGCOMM Workshop on Large Scale Attack Defense*, September 2006.
- [25] M. Naor and B. Pinkas. Oblivious transfer with adaptive queries. In *Advances in Cryptology — CRYPTO*, August 1999.
- [26] M. Naor and B. Pinkas. Oblivious transfer and polynomial evaluation. In *Proc. Symposium on Theory of Computing*, May 1999.
- [27] M. Naor and O. Reingold. Number-theoretic constructions of efficient pseudorandom functions. In *Proc. Symposium on Foundations of Computer Science*, October 1997.
- [28] Oracle. Berkeley DB, 2009. <http://www.oracle.com/technology/products/berkeley-db/>.
- [29] L. Poole and V. S. Pai. ConfDNS: Leveraging scale and history to improve DNS security. In *Proc. Workshop on Real, Large Distributed Systems*, November 2006.
- [30] Privacy Rights Clearinghouse. A chronology of data breaches, January 2009. <http://www.privacyrights.org/ar/ChronDataBreaches.htm>.
- [31] M. Rabin. How to exchange secrets by oblivious transfer. Technical Report TR-81, Harvard Aiken Computation Laboratory, 1981.
- [32] A. Ramachandran and N. Feamster. Understanding the network-level behavior of spammers. In *Proc. ACM SIGCOMM*, September 2006.
- [33] H. Ringberg, A. Soule, and M. Caesar. Evaluating the potential of collaborative anomaly detection. Unpublished report, 2008.
- [34] S. Schechter, J. Jung, W. Stockwell, and C. McLain. Inoculating SSH against address harvesting. In *Proc. Network and Distributed System Security Symposium*, February 2006.
- [35] A. Soule, H. Ringberg, F. Silveira, J. Rexford, and C. Diot. Detectability of traffic anomalies in two adjacent networks. In *Passive and Active Measurement*, April 2007.
- [36] The DIMES Project, 2009. <http://www.netdimes.org/new/>.
- [37] D. Wendlandt, D. G. Andersen, and A. Perrig. Perspectives: Improving SSH-style host authentication with multi-path probing. In *Proc. USENIX Annual Technical Conference*, 2008.
- [38] Yahoo! Hadoop Team. Zookeeper. <http://hadoop.apache.org/zookeeper/>, 2009.
- [39] A. C. Yao. Protocols for secure computations. In *Proc. Symposium on Foundations of Computer Science*, November 1982.