

Fides: Managing Data on Untrusted Infrastructure

Sujaya Maiyya

Danny Hyun Bum Cho

Divyakant Agrawal

Amr El Abbadi

University of California, Santa Barbara

{sujaya_maiyya, hyunbumcho, divyagrwal, elabbadi}@ucsb.edu

Abstract—Significant amounts of data are currently being stored and managed on third-party servers. It is impractical for many small scale enterprises to own their private datacenters, hence renting third-party servers is a viable solution for such businesses. But the increasing number of malicious attacks, both internal and external, as well as buggy software on third-party servers is causing clients to loose their trust in these external infrastructures. While small enterprises cannot avoid using external infrastructures, they need the right set of protocols to manage their data on untrusted infrastructures. In this paper, we propose *TFCCommit*, a novel atomic commitment protocol that executes transactions on data stored across multiple untrusted servers. To our knowledge, *TFCCommit* is the first atomic commitment protocol to execute transactions in an untrusted environment without using expensive Byzantine replication. Using *TFCCommit*, we propose an *auditable* data management system, *Fides*, residing completely on untrustworthy infrastructure. As an auditable system, *Fides* guarantees the detection of potentially malicious failures occurring on untrusted servers using tamper-resistant logs with the support of cryptographic techniques. The experimental evaluation demonstrates the scalability of our approach and the relatively low overhead of executing transactions on untrusted infrastructure.

Index Terms—Databases, Security, Malicious failures, Audits, Fault detection, ACID guarantees

Acknowledgements: This work is funded by NSF grants CNS-1703560 and CNS1815733.

I. INTRODUCTION

A fundamental problem in distributed data management is to ensure the atomic and correct execution of transactions. Any transaction that updates data stored across multiple servers needs to be executed atomically, i.e., either all the operations of the transaction are executed or none of them are executed. This problem has been solved using commitment protocols, such as Two Phase Commit (2PC) [14]. Traditionally, the infrastructure, and hence the servers storing the data, were considered trustworthy. A standard assumption was that if a server failed, it would simply crash; and unless a server failed, it executed the designated protocol correctly.

The recent advent of cloud computing and the rise of blockchain systems are dramatically changing the trust assumptions about the underlying infrastructure. In a cloud environment, clients store their data on third-party servers, located on one or more data centers, and they execute transactions on the data. The servers hosted in the data centers are vulnerable to external attacks or software bugs that can potentially expose a client's critical data to a malign agent (e.g., credit details exposed in Equifax data breach [2]). Further, a server may intentionally decide not to follow the protocol execution, either to improve its performance or for any other self-interest (e.g.,

next big cyber threat is speculated to be intentional data manipulation [1]).

The increasing popularity of blockchain is also exposing the challenges of storing data on non-trustworthy infrastructures. Applications such as supply chain management [19] execute transactions on data repositories maintained by multiple administrative domains that mutually distrust each other. Blockchains resort to expensive protocols (mining or byzantine replication) to tolerate malicious failures since for many applications, both the underlying infrastructure and the participating entities are untrusted.

In most existing databases, the prevalent approach to tolerate malicious failures is by replicating either the whole database or the transaction manager [4], [12], [13], [31], [35]. Practical Byzantine Fault Tolerance (PBFT) [7] by Castro and Liskov has become the predominant replication protocol used in designing data management systems residing on untrusted or byzantine infrastructure. These systems provide fault-tolerance in that the system makes progress in spite of byzantine failures; the replication masks these failures and ensures that non-faulty processes always observe correct and reliable information. *Fault tolerance is guaranteed only if at most one third of the replicas are faulty* [6].

In a relatively open and heterogeneous environment, knowing the number of faulty servers – let alone placing a bound on them – is unrealistic. In such settings, an alternate approach to tolerate malicious failures is *fault-detection* which can be achieved using *auditability*. Fault detection imposes no bound on the number of faulty servers – any server can fail maliciously but the failures are always detected as they are not masked from the correct servers; detection requires only one server to be correct at any given time. To guarantee fault detection through audits, tamper-proof logs have been proposed and widely used in systems such as PeerReview [15] and CATS [33].

Motivated by the need to develop a fault-detection based data management system, we make two major propositions in this paper. First, we develop a data management system, *Fides*¹, consisting of untrusted servers that may suffer arbitrary failures in all the layers of a typical database, i.e., the transaction execution layer, the distributed atomic commitment layer, and the datastore layer. Second, we propose a novel atomic commit protocol – *TrustFree Commitment* (TFCCommit) – an integral component of *Fides* that commits distributed transactions across untrusted servers while providing auditable

¹*Fides* is the Roman Goddess of trust and good faith.

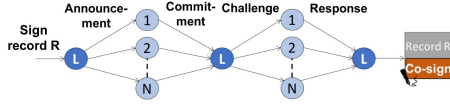


Fig. 1: Collective Signing.

guarantees. To our knowledge, TFCCommit is the first to solve the distributed atomic commitment problem in an untrusted infrastructure without using expensive byzantine replication protocols. Although we present Fides with TFCCommit as an integral component, TFCCommit can be disintegrated from Fides and used in any other design of a trust-free data management.

With detection being the focus rather than tolerance of malicious failures, Fides precisely identifies the point in the execution history at which a fault occurred, as well as the servers that acted maliciously. These guarantees provide two fold benefits: i) A malicious fault by a database server is eventually detected and undeniably linked to the malicious server, and ii) A benign server can always defend itself against falsified accusations. By providing auditability, Fides incentivises a server not to act maliciously. Furthermore, by designing a stand-alone commit protocol, TFCCommit, that leverages cryptography, we take the first step towards developing a full-fledged data management system that fully resides in untrusted infrastructures. We believe it is critical to start with a strong and solid atomic commitment building block that can be expanded to include fault tolerance and other components of a transaction management hierarchy.

Section II provides the necessary background used in developing a trust-free data management system. Section III discusses the architecture, system, and failure models of Fides. Section IV describes the auditable transaction model in Fides and also introduces TFCCommit. Experimental evaluation of TFCCommit is presented in Section V, followed by related work in Section VI. Section VII concludes the paper.

II. CRYPTOGRAPHIC PRELIMINARIES

Developing a data management system built on untrusted infrastructure relies heavily on many cryptographic tools. In this section, we provide the necessary cryptographic techniques used in developing Fides.

A. Collective Signing

Multisignature (multisig) is a form of digital signature that allows more than one user to sign a single record. Multisigs, such as Schnorr Multisignature [29], provide additional authenticity and security compared with single user's signature. Collective Signing (CoSi) [30], an optimization of Schnorr Multisigs, allows a *leader* to produce a record which then can be publicly validated and signed by a group of *witnesses*. CoSi requires two rounds of communication to produce a *collective signature* (co-sign) with the size and verification cost of a single signature. Figure 1 represents the four phases of CoSi where L is the leader and $1, 2, \dots, N$ are the witnesses:

Announcement: The leader *announces* the beginning of a new round to all the witnesses and sends the record R to be collectively signed.

Commitment: Each witness, in response, picks a random secret, which is used to compute the Schnorr commitment, x_{sch} . The witness then sends the commitment to the leader.

Challenge: The leader aggregates all the commits, $X = \sum x_{sch}$ and computes a Schnorr challenge, $ch = hash(X|R)$. The leader then broadcasts the challenge to all the witnesses.

Response: Each witness validates the record before computing a Schnorr-response, r_{sch} , using the challenge and its secret key. The leader collects and aggregates all the responses to finally produce a Schnorr multisignature.

The collective signature provides a **proof** that the record is produced by the leader and that all the witnesses signed it only after a successful validation. Anyone with the public keys of all the involved servers can verify the co-sign and the verification cost is the same as verifying a single signature. An invalid record will not produce enough responses to prove the authenticity of the record. We refer to the original work [30] for a detailed discussion of the protocol.

B. Merkle Hash Tree

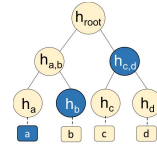


Fig. 2: Merkle Hash Tree example.

A merkle hash tree (MHT) [24] is a binary tree with each leaf node labeled with the hash of a data item and each internal node labeled with the hash of the concatenated labels of its children. Figure 2 shows an example of a MHT. The hash functions, h , used in MHTs are *one way hash functions* i.e., for a given input x , $h(x) = y$, such that, given y and h , it is computationally infeasible to obtain x . The hash function h must also be collision-free, i.e., it is highly unlikely to have two distinct inputs x and z that satisfies $h(x) = h(z)$. Any such hash function can be used to construct a MHT.

Data Authentication Using MHTs: MHTs are used to authenticate a set of data values [24] by requiring the prover, say Alice, to publicly share the root of the MHT, h_{root} , whose leaves form the data set. To authenticate a single data value, all that a verifier, say Bob, needs from Alice is a *Verification Object* ($\mathcal{VO}$) consisting of all the sibling nodes along the path from the data value to the root. The highlighted nodes in Figure 2 form the verification object for data item a , $\mathcal{VO}(a)$, which is of size $\log_2 n$. To authenticate data item a , Alice generates the $\mathcal{VO}(a)$, and provides the value of a and $\mathcal{VO}(a)$ to Bob. Given the value of a , Bob computes $h(a)$ and uses h_b from $\mathcal{VO}(a)$ to compute $h_{a,b} = h(h(a)|h(b))$ i.e., the hash of $h(a)$ concatenated with $h(b)$. Finally, using $h_{a,b}$ and $h_{c,d}$ sent in the $\mathcal{VO}(a)$, Bob computes the root, $h_{a,b,c,d} = (h_{a,b}|h_{c,d})$. Bob then compares the computed root, $h_{a,b,c,d}$, with the root publicly shared by Alice h_{root} . Assuming the use of a collision free hash function ($h(a_1) \neq h(a_2)$ where $a_1 \neq a_2$), it would be computationally infeasible for Alice to tamper with a 's

value such that the h_{root} published by Alice matches the root computed by Bob using the verification object.

III. FIDES ARCHITECTURE

Fides is a data management system built on untrusted infrastructure. This section lays the premise for Fides by presenting the system and failure models, and the audit mechanisms.

A. System Model

Fides is a distributed database of multiple servers; the data is partitioned into multiple shards and distributed on these servers (perhaps provisioned by different providers). Shards consist of a set of data items, each with a unique identifier. The system assumes neither the servers nor the clients to be trustworthy and hence, can behave arbitrarily. Servers and clients are uniquely identifiable using their public keys and are aware of all the other servers in the system. All message exchanges (client-server or server-server) are digitally signed by the sender and verified by the receiver [28].

The clients interact with the data via transactions consisting of read and write operations. The data can be either single-versioned or multi-versioned with each committed transaction generating a new version. Every data item has an associated read timestamp r_{ts} and a write timestamp w_{ts} , indicating the timestamp of the last transaction that read and wrote the item, respectively. When a transaction commits, it updates the timestamps of the accessed data items.

We choose a simplified design for a database server to minimize the potential for failure. As indicated in Figure 3a, each database server is composed of four components: an *execution layer* to perform transactional reads and writes; a *commitment layer* to atomically (i.e., all servers either **commit** or **abort** a transaction) terminate transactions; a *datastore* where the data shards are stored; and a *tamper-proof log*.

As individual servers are not trusted, we replace the local transaction logs used in traditional protocols with a *globally replicated tamper-proof log* (this approach is inspired by blockchain). The log – a linked-list of transaction *blocks* linked using cryptographic hash pointers – guarantees immutability. Global replication of the log guarantees that even if a subset (but not all) of the servers collude to tamper the log, the transaction history is persistent.

B. Failure model

In Fides, a server that fails maliciously can behave arbitrarily, i.e., send arbitrary messages, drop messages, or corrupt the data it stores. Fides assumes that each server and client is computationally bounded and is incapable of violating any cryptographic primitives such as forging digital signatures or breaking one-way hash functions – the operations that typically require brute force techniques.

Let n be the total number of servers and f the maximum number of faulty servers. Fides tolerates up to $n - 1$ faulty servers, i.e., $n > f$. To detect failures, Fides requires at least one server to be correct and failure-free (free of malicious, crash, or network partition failures) at a given time. This implies that the correct set of servers are not static and can

vary over time. This failure model is motivated by Dolev and Strong’s [9] protocol.

An individual server, comprising of four components as shown in Figure 3a, can fail at one or more of the components. A fault in the *execution layer* can return incorrect values; in the *commit layer* can violate transaction atomicity; in the *datastore* can corrupt the stored data values; and in the *log* can change or reorder the transaction history. These failures can be intentional (to gain application level benefits) or unintentional (due to software bugs or external attacks); Fides does not distinguish between the two.

A malicious client in Fides can send arbitrary messages or semantically incorrect transactions to a database server but later blame the server for updating the database inconsistently. To circumvent this, the servers store all digitally signed, unforgeable messages exchanged with the client. This message log serves as a proof against a falsified blame when a client’s transaction sends the database to a semantically inconsistent state.

C. Auditing Fides

Auditability has played a key role in building dependable distributed systems [15], [32], [33]. Fides provides auditability: the application layer or an external auditor can audit individual servers with an intent to either detect failures or verify correct behavior.

Fides guarantees that any failure, as discussed in Section III-B, will be detected in an offline audit. Fides focuses on failure detection rather than prevention; detection includes identifying (i) the **precise point** in transaction history where an anomaly occurred, and (ii) the exact misbehaving server(s) that is irrefutably linked to a failure. The auditor is considered to be a powerful external entity and during each audit:

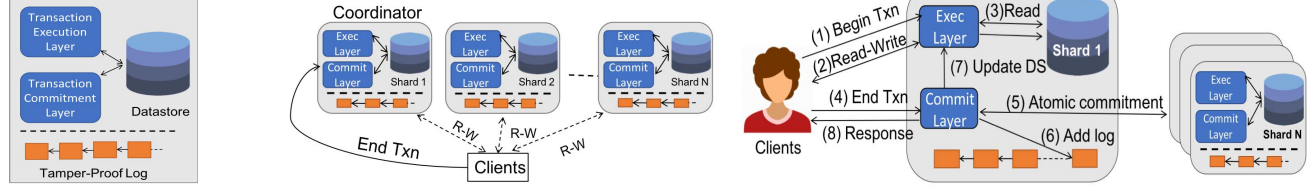
(i) The auditor gathers the tamper-proof logs from all the servers before the auditing process.

(ii) Given that at least one server is correct, from the set of logs collected from all servers, the auditor identifies the *correct* and *complete* log (how is explained in detail in Section IV-D). The auditor uses this log to audit the servers.

Optimizations such as checkpointing can be used to minimize the log storage space at each server; these optimizations are orthogonal and hence not discussed further. If the audit uncovers any malicious activity, a practical solution can be to penalize the misbehaving server in legal, monetary, or other forms specific to the application. This discourages a server from acting maliciously.

IV. FIDES

In this section we present Fides: an *auditable* data management system built on untrusted infrastructure. The basic idea is to integrate cryptographic techniques such as digital signatures (public and private key encryption), collective signing, and Merkle Hash Trees (MHT) with the basic transaction execution in database systems. This integration results in *verifiable* transaction executions in an environment where the database servers cannot be trusted.



(a) Components of a db server.

(b) Client interactions in Fides

(c) Transaction life-cycle in Fides

Fig. 3: The database server components and the client-server interactions in Fides

A. Overview

Figure 3b illustrates the overall design of Fides. The clients read and write relevant data by directly interacting with the appropriate database partition server (this can be accomplished by linking the client application with a run-time library that provides a lookup and directory service for the database partitions). The architecture intentionally avoids the layer of front-end database servers (e.g., Transaction Managers) to coordinate executing transactional reads and writes as these front-end servers may themselves be vulnerable and exhibit malicious behavior by relaying incorrect reads/writes. Hence, in Fides all data-accesses are managed directly between the client and the relevant database server.

Since data-accesses are handled with minimal synchronization among concurrent activities, the burden of ensuring the correct execution of transactions occurs when a transaction is *terminated*. We use a simplified setup where one *designated* server acts as the transaction *coordinator* responsible for terminating all transactions. The coordinator is also an untrusted database server that has additional responsibilities only during the termination phase.

When a client application decides to terminate its transaction, it sends the termination request to the designated coordinator; all other database servers act as cohorts during the termination phase. For ease of exposition, we first present a termination protocol executed globally involving *all* database servers, irrespective of the shards accessed in that transaction. The global execution implies transactions are terminated and ordered *sequentially*. We relax this requirement and allow different coordinators for concurrent transactions in [3].

The following is an overview of the client-server interaction: a typical life-cycle of a transaction as depicted in Figure 3c.

- 1. Begin transaction:** A client starts accessing the data by first sending a *Begin Transaction* request to all the database servers storing items read or written by the transaction.
- 2. Read-write request:** The client then sends requests to each server indicating the data items to be read and written.
- 3. Read-write response:** The servers respond to a read requests by fetching the data from the datastore and relaying it to the client. The write requests are buffered.
- 4. End Transaction:** After completing data access, the client sends *End Transaction* to the coordinator which coordinates the commitment to ensure transaction correctness (i.e., serializability) and transaction atomicity (i.e., all-or-nothing property).

key	description
TxnId	commit timestamp of txn
R set	list of $\langle id : value, r_{ts}, w_{ts} \rangle$
W set	list of $\langle id : new_val, old_val, r_{ts}, w_{ts} \rangle$
$\sum roots$	MHT roots of shards
decision	commit or abort
h	hash of previous block
co-sign	a collective signature of participants

TABLE I: Details stored in each block

5. Atomic commitment: The coordinator and the cohorts collectively execute the atomic commit protocol – TFCCommit – and decide either to **commit** or **abort** the transaction. The commitment produces a block (i.e., an entry in the log) containing the transaction details. If the decision is **commit**, then the next two steps are performed.

6. Add log: All servers append the block in a consistent order to their logs, thus creating a globally replicated log.

7. Update datastore: The datastore is updated based on the buffered writes, if any, along with updating the timestamps r_{ts} and w_{ts} of the data items accessed in the transaction.

8. Response: The coordinator responds to the client informing whether the transaction was committed or aborted.

The log, stored as a linked-list of blocks, encompasses the transaction details essential for auditing. It is vital to understand the structure of each block before delving deeper into transaction execution details. Every block stores the information shown in Table I. Although a block can store multiple transactions, for ease of explanation, *we assume that only one transaction is stored per block*.

As indicated in Table I, each transaction is identified by its commit timestamp, assigned by the client that executed this transaction. Any timestamp that supports total ordering can be used by the client – e.g., a Lamport clock with $\langle client_id : client_time \rangle$ – as long as all clients use the same timestamp generating mechanism.

A block contains the transaction read and write sets consisting of three vital pieces of information: 1) the data-item identifiers that are read/written, 2) the values of items read and the new values written; the *old_val* in the write set is populated only for blind writes, and 3) the latest read r_{ts} and write w_{ts} timestamps of those data items at the time of access.

The blocks also contain: the Merkle Hash Tree roots of the shards involved in the transaction (explained more in Section IV-B); the **commit** or **abort** transaction decision; the hash of the previous block forming a chain of blocks linked

by their hashes; and finally, a collective signature of all the servers (how and why are explained in Section IV-C).

B. Transaction Execution

This section describes the correct mechanism of executing transactions (reads and writes) and discusses techniques to detect deviations from the expected behavior.

1) *Correct Behavior*: With regard to transaction execution, a correct database server is responsible for the following actions: (i) return the values and timestamps of data-items specified in the read requests, and (ii) buffer the values of data-items updated in the transaction and if the transaction successfully commits, update the datastore based on the buffered writes. We explain how a correct server achieves these actions.

Reads and Writes: A client sends a *begin transaction* message to all the database servers storing the items read or written by the transaction. The client then sends a *Read* request consisting of the data-item ids to the respective servers. For example, if a transaction reads data item x from server $S1$ and item y from server $S2$, the client sends $Read(x)$ to $S1$ and $Read(y)$ to $S2$. The servers respond with the data values **along with** the associated read r_{ts} and write w_{ts} timestamps.

The client then sends the *Write* message with the data-item ids and their updated values to the respective servers. For example, if a transaction writes data item x in server $S1$ with value 5 and item y in server $S2$ with value 10, the client sends $Write(x,5)$ to $S1$ and $Write(y,10)$ to $S2$. The servers buffer these updates and respond with an acknowledgement. To support blind writes, the acknowledgement includes the old values and associated timestamps of the data-items that are being written but were not read before.

After completing the data accesses, the client sends the *end transaction* request – sent only to the designated coordinator – consisting of the read and the write set: a list of data item ids, respective r_{ts} and w_{ts} returned by the servers, and the values read and the new values written. The coordinator then executes TFCommit among all the servers to terminate (commit or abort) the transaction (explained in detail in Section IV-C). If all the involved servers decide to commit the transaction, each involved server constructs a Merkle Hash Tree (MHT) (Section II-B) of its data shard with all the data items – with updated values – as the leaves of the tree and with the root node $root_{mht}$. The read and write sets and MHT roots become part of the block in the log once the transaction is committed.

Updating the datastore: If the transaction commits, the servers involved in the transaction update the data values in their datastores based on the buffered writes. The servers also update the read and write timestamps of the data items accessed in the transaction to its commit timestamp.

The data can be single-versioned or multi-versioned. For multi-versioned data, when a transaction commits, a correct server additionally creates a new version of the data items accessed in the transaction *while maintaining the older versions*. Although an application using Fides can choose between single-versioned or multi-versioned data, multi-versioned data can provide **recoverability**. If a failure occurs, the data can

be reset to the last sanitized version and the application can resume execution from there.

2) *Detecting Malicious Behavior*: With regard to transaction execution, a server may misbehave by: (i) returning inconsistent values of data-items specified in the read requests; and (ii) buffering incorrect values of data-items updated in the transaction or updating the datastore incorrectly.

(i) **Incorrect Reads**: All faults in Fides are detected by an auditor during an audit. As mentioned in Section III-C, during an audit, the auditor collects the log from all servers and constructs the correct and complete log.

To detect an incorrect read value returned by a malicious server, the auditor must know the expected value of the data-item. The read and write sets in each log entry contains the information on the updated value of a written item and the read value of a read item. Note that in our simplifying assumption (which will be relaxed later), each block contains *only one* transaction and the transactions are committed sequentially with the log reflecting this sequential order. By traversing the log, at each entry, the auditor knows the most recent values of a given data item. We leverage this to identify incorrectly returned values.

Lemma 1: The auditor detects an incorrect value returned for a data item by a malicious server.

Proof: Consider a transaction T_i that committed at timestamp ts_i and stored in the log at block b_i . Assume transaction T_i read an item x and updated it. Let b_j be the first block after b_i to access the same data item x – where $j > i$, indicating that transaction T_j in b_j committed **after** the transaction T_i in b_i . The read value of x in b_j must reflect the value written in b_i ; if the values differ, an anomaly is detected. $\square$

(ii) **Incorrect Writes**: The effect of incorrectly buffering a write or incorrectly updating the datastore is the same: the datastore ends up in an inconsistent state. The definition of incorrect datastore depends on the type of data: for single versioned data, the latest state of data (data values and timestamps) in the datastore is incorrect; for multi-versioned data, one or more versions of the data are incorrect. We discuss techniques to detect incorrect datastore for both types of data.

To detect an inconsistent datastore, we use the data authentication technique proposed by Merkle [24] discussed in Section II-B. To use this technique, the auditor requires the read and written values in each transaction and the resultant Merkle Hash Tree (MHT) root – all pieces of information stored within each block.

Multi-versioned data: For multi-versioned data, the audit policy can involve auditing a single version chosen arbitrarily or exhaustively auditing all versions starting from either the first version (block 0) or the latest version. We explain auditing a single version, which can easily be extended to exhaustively auditing all versions.

Let T_i be a transaction committed at timestamp ts that read and wrote data item x stored in server S_k . Assume the auditor audits server S_k at version ts . Once the auditor notifies the server about the audit, the server constructs the Merkle Hash

Tree with the data at version ts as the leaves; S_k then shares the *Verification Object* $\mathcal{VO}$ —consisting of all the sibling nodes along the path from the data x to the root—with the auditor.

The log entry corresponding to transaction T_i stores the value read for item x and the new value written. The auditor uses (i) the $\mathcal{VO}$ sent by S_k , and (ii) the hash of x 's value stored in the write set of the log, to compute the expected MHT root for the data in S_k . The auditor then compares the computed root with the one stored in the log. A mismatch indicates that the data at version ts is incorrect.

Single-versioned data: For single versioned data, the correctness is only with respect to the latest state of the data. Hence, rather than using an arbitrary block to obtain the MHT root of server S_k , the auditor uses the latest block in the log that accessed the data in S_k to obtain the latest MHT root. The other steps are similar to multi-versioned data.

Lemma 2: The auditor detects an inconsistent datastore. For multi-versioned data, the auditor detects the precise version at which the datastore became inconsistent.

Proof: Detection is guaranteed since Merkle Hash Trees (MHT) use collision-free hash functions (i.e., $h(x) \neq h(y)$ where $x \neq y$), and a malicious server cannot update a data value such that the MHT root stored in the block matches the root computed by the auditor using the verification object sent by the server. For multi-versioned data, the auditor identifies the precise version at which data corruption occurred by systematically authenticating all blocks in the log until a version with mismatching MHT roots is detected. $\square$

C. Transaction Commitment

This section describes how transactions are terminated in Fides and presents a novel distributed atomic commitment protocol—**TrustFree Commit** (TFCommit)—that handles malicious failures. This section also discusses techniques to detect failures if a server deviates from the expected behavior. With regard to transaction commitment, a correct database server is responsible for the following: (i) Ensure transaction isolation (i.e., strict serializability); (ii) Ensure atomicity—either all servers commit the transaction or no servers commit the transaction; and (iii) Ensure *verifiable* atomicity.

1) **Correct Behavior: Transaction Isolation:** Transaction isolation determines how the impact of one transaction is perceived by the other transactions. In Fides, even though multiple transactions can execute, i.e., read and write data, concurrently, Fides provides serializable executions in which concurrent transactions seem to execute in sequence. To do so, servers in Fides abort a transactions if it cannot be serialized with already committed transactions in the log. The read r_{ts} and write w_{ts} timestamps associated with each data item is used to detect non-serializable transactions. The latest timestamps can be obtained from either the datastore or the transaction log. Similar to timestamp based optimistic concurrency control mechanism, at commit time, a server checks if the data accessed in the terminating transaction has been updated since they were read. If yes, the server chooses to abort the transaction.

Atomicity and Verifiability: Consider a traditional atomic commit protocol that provides atomicity: Two Phase Commit (2PC) [14]. 2PC guarantees atomicity provided servers are benign and trustworthy. It is a centralized protocol where one server acts as a coordinator and the others act as cohorts. To terminate a transaction, the coordinator collects **commit** or **abort** votes from all cohorts, and decides to **commit** the transaction *only if* all the cohorts choose to commit, and otherwise decides to **abort**. The decision is then asynchronously sent to the client and the cohorts. 2PC is sufficient to ensure atomicity if servers are trustworthy; but in untrusted environments, 2PC is inadequate as a cohort or the coordinator may maliciously lie about the decision.

To make 2PC trust-free, we combine 2PC with a multi-signature scheme, Collective Signing or CoSi (Section II-A): a two-round protocol where a set of processes collectively sign a given record using their private keys and random secrets. CoSi guarantees that a record (or in our case block) produced by a leader (or coordinator) is validated and signed by all the witnesses (or cohorts) and that *if any of the involved processes lied in any of the phases, the resulting signature will be incorrect*. A signature is bound to a single record; any process with the public keys of all the processes can verify whether the signature is valid and *corresponds* to that record.

We propose a novel approach of integrating 2PC with CoSi to achieve the atomicity properties of 2PC *and* the verifiable properties of CoSi. The basic idea is that the coordinator, similar to 2PC, collects **commit** or **abort** votes from the cohorts, forms a decision, and encapsulates the transaction details including the decision in a block. The coordinator then sends the block to be verified and collectively signed by the cohorts. An incorrect block (either with inaccurate transaction details or wrong decision) produced by a malicious coordinator will not be accepted by correct servers, thus resulting in an invalid signature that can be easily verified by an auditor.

A successful round of TFCommit produces a block to be appended to the log *in a consistent order* by all servers. For ease of exposition, this section presents TFCommit with two main assumptions: (i) the transactions are committed sequentially to avoid forks in the log; and (ii) all servers participate in transaction termination—even the servers that did not partake in transaction execution—to have identical block order in their logs. We relax these assumptions and discuss various techniques to scale TFCommit, which is published in the extended technical report [3] due to lack of space.

Recall from Table I all the details stored in each block. Once a block is cosigned and logged by all servers, it is immutable; hence, all the details must be filled in during different phases of TFCommit. However, to ensure atomicity and verifiability of TFCommit, we only need the transaction id, its decision, and the co-sign. Other details such as the *Read* and *Write* sets, Merkle Tree roots, and hashes are necessary to detect other failures including isolation violation and data corruption.

The protocol:

A client, $\mathcal{A}$, upon finishing transaction execution, sends a signed $\mu = \langle \text{end_transaction}(ts_i, R \text{ set-} W \text{ set}) \rangle_{\sigma_{\mathcal{A}}}$ request

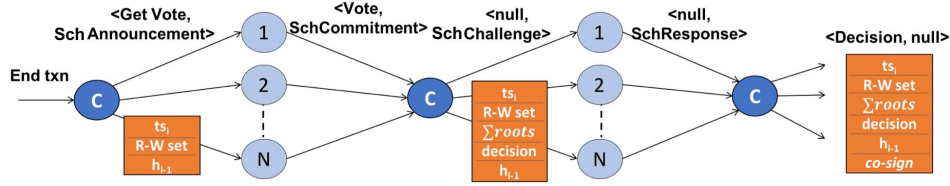


Fig. 4: Different phases and block generation progress made in each phase of TFCCommit

to the coordinator, where ts_i is a client-assigned commit timestamp of the transaction and $R\ set-W\ set$ is the read and write sets consisting of data item ids, values read and new values written, and the timestamps r_{ts} , and w_{ts} . The servers ignore any *end transaction* request with a timestamp lower than the latest committed timestamp. Since the datastore is updated only after a successful commit of a transaction, the servers are unaffected if a malicious client does not send the *end_transaction* request.

TFCCommit is a 3-round protocol involving 5 phases of communication as shown in Figure 4. Since TFCCommit merges 2PC with CoSi, we indicate each phase by a mapping of $\langle 2PC\ phase, CoSi\ phase \rangle$. Figure 4 shows the phases as well as the progress made in constructing the block at each phase. The phases of TFCCommit are:

1) **<GetVote, SchAnnouncement>**: Upon receiving the $\mu = \langle end_transaction(ts_i, R\ set-W\ set) \rangle_{\sigma_A}$ request from the client, to commit a transaction T , the coordinator C prepares a partially filled block, $b_i = [ts_i, Rset - Wset, h_{i-1}]$, containing the commit timestamp, read and write sets, and hash of the previous block. C then encapsulates the signed client request μ and sends the $\langle get_vote(b_i, \mu) \rangle_{\sigma_C}$ message to all the cohorts.

2) **<Vote, SchCommitment>**: Every cohort $\mathcal{H}$ verifies both the *get_vote* message and the encapsulated client request, and computes the Schnorr-commitment (x_{sch}) for CoSi. Then, *only the cohorts that are part of the transaction*, perform the following actions. A cohort involved in the transaction locally decides whether to **commit** or **abort** the transaction. If the cohort locally decides to **commit**, then it constructs a Merkle Hash Tree (MHT) (Section II-B) of its shard with all the data items as leaves of the MHT and with the root node $root_{mht}$. The MHT reflects all the updates in T_i assuming that T_i be committed; since MHT computation is done in memory, the datastore is unaffected if T_i eventually aborts. (The MHT root is required for datastore authentication, as explained in Section IV-B2.) The involved cohorts then send $\langle vote(decision, root_{mht}, x_{sch}) \rangle_{\sigma_H}$ whereas the cohorts not part of the transaction send $\langle vote(x_{sch}) \rangle_{\sigma_H}$ to the coordinator. As the coordinator is also involved in co-signing, it produces the appropriate vote message.

3) **<null, SchChallenge>**: In this phase, the coordinator C collects all the cohort responses and checks if any cohort (or itself) involved in the transaction decided to abort. If none, it chooses **commit**, otherwise **abort**. It then aggregates all the MHT roots of the involved cohorts ($roots = \sum root_{mht}$), and fills the roots field in the block b_i along with the decision field.

If any involved cohorts chose **abort**, the respective roots will be missing in the block. Finally, the coordinator aggregates the Schnorr-commitments $X_{sch} = \sum x_{sch}$ from all the servers and computes the Schnorr-challenge by concatenating and hashing X_{sch} with b_i i.e., $ch = h(X_{sch} || b_i)$. The coordinator then sends $\langle challenge(ch, X_{sch}, b_i) \rangle_{\sigma_C}$ to all cohorts.

4) **<null, SchResponse>**: In this phase, every cohort $\mathcal{H}$, checks if the decision within the block b_i is **abort**, and if so, b_i should have some missing roots; if the decision is **commit**, b_i should have all the roots from the involved servers. Every involved cohort that sent the MHT root in the *vote* phase verifies if its corresponding root in the block is the same as the one it sent. Cohorts also verify whether a potentially malicious coordinator computed the challenge, ch , correctly by hashing the concatenated X_{sch} and b_i , both of which were sent in the challenge message. A cohort then computes the Schnorr-response r_i using its secret key and the challenge ch , and sends $\langle response(r_i) \rangle_{\sigma_H}$ to the coordinator.

5) **<Decision, null>**: The coordinator collects all the Schnorr-responses and aggregates them, $R_{sch} = \sum r_{sch}$, to form the collective signature represented by $\langle ch, R_{sch} \rangle$. Intuitively, the challenge ch is computed using the block; and the Schnorr-response R_{sch} requires the private keys of the servers, thus the signature binds the block with the public keys of the servers. The coordinator then updates the *co-sign* field in the block and sends the finalized block to the client and the cohorts. If the decision is **commit**, all servers append block b_i to their log and update their respective datastores.

The client, with the public keys of all the servers, verifies the *co-sign* before accepting the decision – even an aborted transaction must be signed by all the servers. If the verification fails, the client detects an anomaly and may trigger an audit, which may halt the progress in the system.

TFCCommit, similar to 2PC, can be blocking if either the coordinator or any cohort fails (crash or malicious). TFCCommit can be made non-blocking by adding another phase that makes the chosen value available, as in the case of Three Phase Commit, Paxos-Commit, or Paxos Atomic Commit, all of them discussed in [22]; we leave this extension for future work.

2) **Detecting Malicious Behavior**: A correct execution of TFCCommit ensures serializable transaction isolation, atomicity, and verifiable commitment. However, a malicious server can (i) violate the isolation guarantees by committing non-serializable transactions; (ii) a malicious coordinator can break atomicity by convincing some servers to commit a transactions and others to abort; or (iii) a server can send wrong cryptographic values during co-signing to violate verifiability.

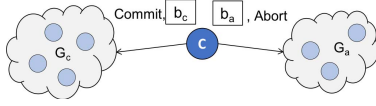


Fig. 5: Atomicity violation of TFCCommit

Lemma 3: The auditor detects serializability violation.

Proof: Transaction execution is based on executing read and write operations in the timestamp order. The transactions are ordered based on the timestamps, which are monotonically increasing. If a transaction has done a conflicting access inconsistent with the timestamp order, it leads to one of the following conflicts: 1) RW-conflict: a transaction with a smaller timestamp read a data-item with a larger timestamp; 2) WW-conflict: a transaction with a smaller timestamp wrote a data-item that was already updated with a larger timestamp; 3) WR-conflict: a transaction with a smaller timestamp wrote a data-item after it was read by a transaction with a larger timestamp. For each transaction audited, the auditor verifies if any of the above violations exist, and if so, the auditor detects the server responsible for the violation to be misbehaving. This is equivalent to verifying that no cycle exists in the Serialization Graph of the transactions being audited. $\square$

Lemma 4: The auditor or a correct server detects incorrect cryptographic values for CoSi sent by a malicious server – which hampers verifiability of TFCCommit.

Proof: Refer to the extended technical report [3] for the proof.

Lemma 5: The auditor or a correct server detect atomicity violation of TFCCommit.

Proof: Recall that the coordinator C collects votes in phase two of TFCCommit, forms the decision, and sends the partial block containing the decision in the *challenge* message. Consider Figure 5 where a malicious coordinator sends block b_c with commit decision to group G_c and block b_a with abort decision to group G_a . More precisely, the coordinator sends $\langle \text{challenge}(ch, X_{sch}, b_c) \rangle_{\sigma_C}$ to G_c (X_{sch} is the aggregated Schnorr-commits) and $\langle \text{challenge}(ch, X_{sch}, b_a) \rangle_{\sigma_C}$ to G_a . Since the decision is part of the block, the two blocks b_c and b_a have to be different if the coordinator violates atomicity. But with respect to the challenge ch , there are two possibilities, both producing invalid signatures:

- **Case 1:** Coordinator sends the same challenge ch computed using block b_c (or b_a) to both groups.

Any correct server in the group G_a will recompute the challenge using the block it received, b_a , and immediately recognize that the challenge sent by the coordinator does not correspond to the block b_a . (Alternatively, if the coordinator used b_a to compute the challenge ch , then servers in G_c will detect the anomaly.) Even if the servers in one group, say G_a , collude with the coordinator and do not expose the anomaly, the challenge ch corresponds only to block b_c . The auditor, while auditing a server in group G_a , detects that the co-sign in block b_a is invalid as it does not correspond to that block.

- **Case 2:** Coordinator sends the challenge ch computed using block b_c to group G_c and the challenge ch' computed using block b_a to group G_a .

In the final step of TFCCommit, the servers in group G_c will use ch to compute the Schnorr-response, whereas the servers in group G_a will use ch' to compute the Schnorr-response. Given that the final collective signature can be tied only to a single block, the co-sign does not correspond to either b_c or b_a , hence producing a wrong signature. $\square$

D. Transaction Logging

The transaction log in Fides is a tamper-proof, globally replicated log. When a transaction commits after a successful round of TFCCommit, all servers append the newly produced block to their logs.

Detecting Malicious Behavior: One or more faulty servers can collude (but not all at once) to (i) tamper an arbitrary block, (ii) reorder the blocks, or (iii) omit the tail of the log (last few blocks). The auditor collects logs from all the servers and uses the collective signature stored in each block to detect an incorrect log.

Lemma 6: Given a set of logs collected from all servers, the auditor detects all incorrect logs – logs with arbitrary blocks that are modified or logs with reordered blocks.

Proof: The collective signature in each block prevents a malicious server from manipulating that block once it is appended to the log. The signature is tied specifically to one block and if the contents of the block are manipulated, the signature verification will fail. One or more malicious servers cannot tamper with an arbitrary block successfully without the cooperation of *all* the servers. And since the hash of the previous block is part of a log entry, unless *all* the servers collude, the blocks cannot be successfully re-ordered. $\square$

Lemma 7: From the logs collected from all servers, the auditor detects all incomplete logs – logs with missing tail entries.

Proof: A subset of servers cannot successfully modify arbitrary blocks in the log (proof in Lemma 6) but they can omit the tail of the log. During an audit, the auditor gathers the logs from all the servers. At least one correct server exists with the complete log – which can easily be verified for correctness by validating the collective signature and hash pointer in each block. The auditor uses this complete and verified log to detect that one or more servers store an incomplete log. $\square$

E. Correctness of Fides

Definition 1: Verifiable ACID properties

In transaction processing, ACID refers to the four key components: Atomicity, Consistency, Isolation, and Durability.

We define *v-ACID* as the ACID properties that can be verified. *v-ACID* indicates that a database system provides verifiable evidence that the ACID guarantees are upheld. This definition is useful when individual database servers are untrusted and may violate ACID – in which case the system must allow verifying and detecting the violations.

Theorem 1: Fides provides Verifiable ACID guarantees.

Proof: Fides guarantees that an external auditor can verify if the database servers provide ACID guarantees or not. The first

step in the verification is for the auditor to obtain a *correct* and *complete* log. Given the assumption that at least one server is correct at a given time, Lemmas 6 and 7 prove that during an audit, the auditor always identifies correct and complete log.

Lemma 5 proves that *Atomicity* violation is verifiable; Lemma 2 proves that the auditor verifies if the effect of a transaction resulted in an inconsistent database when a server buffers inconsistent writes, i.e., verifiable *Consistency*; Lemma 3 proves that the *Isolation* guarantee which ensures serializable transaction execution is verifiable; and finally, Lemmas 1 and 2 verify if the effects of committed transactions are *Durable*. Hence, an auditor verifies whether the servers in Fides uphold ACID properties.

Note that multiple ACID violations can exist in the transaction execution. Since the log is sequential, the auditor identifies the first occurrence of any of these violations and the blocks after that need not be audited as everything past that violation can be incorrect and hence irrelevant to a correct execution. $\square$

V. EVALUATION

In this section, we discuss the experimental evaluation of TFCCommit. Our goal is to measure the overhead incurred in executing an atomic commit protocol on untrusted infrastructure. The focus of Fides and TFCCommit is *fault detection* in a non-replicated system, hence solutions based on replication that typically use PBFT [7] are orthogonal to TFCCommit.

In evaluating TFCCommit, we measure the performance using two aspects: *commit latency* - time taken to terminate a transaction once the client sends *end transaction* request, and *throughput* - the number of transactions committed per second; TFCCommit was implemented in Python. We deployed multiple database servers on a single Amazon AWS datacenter (US-West-2 region) where each server was an EC2 m5.xlarge vm consisting of 4 vCPUs, 16 GiB RAM and upto 10 Gbps network bandwidth. Unless otherwise specified in the experiment, each database server stores a single shard (or partition) of data consisting of 10000 data items.

To evaluate the protocol, we used Transactional-YCSB-like benchmark [8] consisting of transactions with read-write operations. Each transaction consisted of 5 operations on different data items thus generating a multi-record workload. The data items were picked at random from a pool of all the data partitions combined, resulting in distributed transactions. Although we presented TFCCommit and Fides with the simplifying assumption of one transaction per block, in the experiments, we typically stored 100 non-conflicting transactions in each block. Every experimental run consisted of 1000 client requests and each data point plotted in this section is an average of 3 runs.

A. TFCCommit vs. 2PC

As a first step, we compare the trust-free protocol TFCCommit with its trusted counterpart Two Phase Commit [14]. TFCCommit is essentially 2PC combined with the cryptographic primitives (Co-Signing and Merkle Hash Trees) which results

in an additional phase due to the trust-free nature. Thus, comparing TFCCommit with 2PC highlights the overhead incurred by TFCCommit to operate in an untrusted setting. Both 2PC and TFCCommit are implemented such that transactions are terminated and blocks are produced *sequentially* so that the log does not have forks.

Figure 6a contrasts the performance of 2PC vs. TFCCommit. We increase the number of servers and measure commit latency and throughput. In this experiment, each block stores a *single* transaction so that we can measure the overhead induced by TFCCommit *per transaction*.

As indicated in the figure, the average latency to commit a single transaction in an untrusted setting is approximately 1.8x more than a trusted environment. The throughput for 2PC is approximately 2.1x higher than TFCCommit. TFCCommit performs additional computations compared with 2PC: Merkle Hash Tree (MHT) updates to compute new roots after each transaction, collective signature on each block, and an additional phase. In spite of the additional computing and achieving trust-free atomic commitment, TFCCommit is only 1.8x slower than 2PC. Having shown the overhead of TFCCommit as compared to 2PC, the following experiments measure the performance of TFCCommit by varying different parameters.

B. Number of transactions per block

In this experiment, we fix the number of servers to 5 and increase the load on the system by increasing the number of transactions stored within each block. Each database server consisted of 10000 data items. Figure 6b indicates the average latency to commit a single transaction and the throughput while increasing number of transactions stored within each block from 2 to 120. The latency to commit a single transaction reduces by 2.8x (from 2.7ms to 0.7ms) and the throughput increases by 2.5x when 80 or more transactions are batched in a single block. This experiment highlights that even though the blocks are produced sequentially, the performance of TFCCommit can be significantly enhanced by processing multiple transactions in one block.

C. Number of shards

In this experiment, we measure the scalability of TFCCommit by increasing the number of database servers (each storing a shard of 10000 data items) from 3 to 9, while keeping the number of transaction per block constant (100 per block). Figure 6c depicts the experimental results. The throughput of TFCCommit increases by 47% and the commit latency reduces by 33% when the number of servers are increased from 3 to 9. Figure 6c also shows the most expensive operation in committing transactions i.e., Merkle Hash Tree (MHT) updates. Recall from Section IV-C that in TFCCommit, termination of each transaction requires computing the updated MHT root. Given that each block has 100 transactions, which in turn consists of 5 operations each, there are 500 operations in each block. With only 3 servers, all the operations access the three shards whereas with 9 servers, the 500 operations are spread across nine shards. Thus, the load per server reduces when

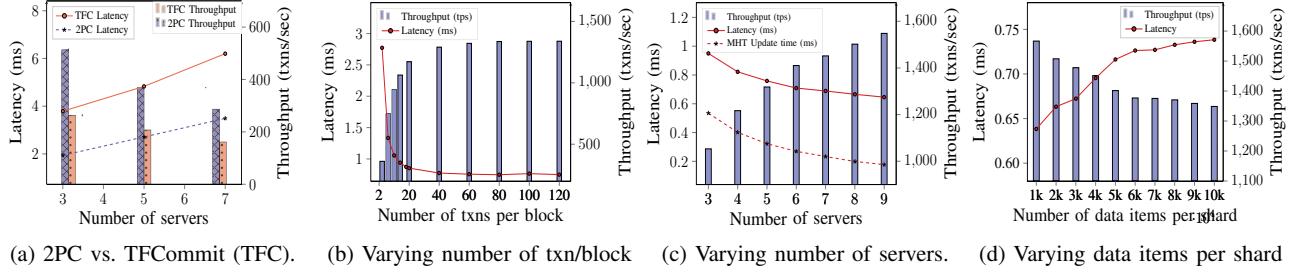


Fig. 6: Various performance evaluations of TFCCommit

there are more servers, resulting in the reduction of MHT update latencies. This experiment highlights that TFCCommit is scalable and performs well with increasing number of servers.

D. Number of data items

In the final set of experiments, we measure the performance of TFCCommit by varying the number of data items stored in each database server, while keeping a constant of 100 transactions per block and using 5 database servers. The number of items stored in each server increased from 1000 to 10000 to measure the commit latency and throughput of TFCCommit, as shown in Figure 6d. The commit latency increases by 15% and the throughput reduces by 14% with the increase in number of data items per shard. The performance fluctuation is due to the Merkle Hash Tree updates that varies with the number of data items. Updating a single leaf node in a binary hash tree with 1000 leaf nodes (data items) updates 10 nodes (from leaf to the root) and a tree with 10000 leaf nodes updates roughly 14 nodes. Thus, the performance of TFCCommit decreases with increasing number of data items stored within each server.

VI. RELATED WORK

The literature on databases that tolerate malicious failures is extensive [11]–[13], [31]. All of these solutions differ from *Fides* as they: assume a single non-partitioned database, rely on *replicating* the database to tolerate byzantine failures, and some also require a trusted component for correctness. Garcia-Molina et al. [12] were the earliest to propose a set of database schemes that tolerate malicious faults. Gashi et al. [13] discuss fault-tolerance other than just crash failures and provide a report composed of database failures caused by software bugs. HRDB by Vandiver et al. [31] propose a replication scheme to handle byzantine faults wherein a trusted coordinator delegates transactions to the replicas. The coordinator also orders the transactions and decides when to safely commit a transaction. Byzantium by Garcia et al. [11] provides an efficient replicated middleware between the client and the database to tolerate byzantine faults.

The advent of blockchains brought with it a set of technologies and protocols that manage data in untrusted environments. But these protocols and their applications are mostly limited to crypto-currencies and cannot be easily extended for large scale distributed data management. Although permissionless

blockchains such as Elastico [21] Omniledger [18], and RapidChain [34] discuss sharding, it is with respect to transactions, i.e., different servers execute different transactions to enhance performance but all of them maintain copies of same data, essentially acting as replicas of a single database. These solutions differ from *Fides* as they focus of replicated data rather than distributed data.

In the space of transaction commitment, proposals such as [4], [25], [35] tolerate malicious faults. Mohan et al. [25] integrated 2PC with Byzantine Fault-Tolerance (BFT) to make 2PC non-blocking and to prevent the coordinator from sending conflicting decisions. Zhao et al. [35] propose a commit protocol that tolerates byzantine faults at the coordinator by replicating the coordinator and executing BFT. Chainspace [4] proposes a commit protocol in a blockchain setting by replicating each shard and executing BFT per shard to agree on the transaction decision. All these solutions require replication and execute BFT on the replicas, and hence differ from TFCCommit. TFCCommit uses CoSi [30], to allow verifiability. CoSi has been adapted to make consensus more efficient in blockchains, e.g., ByzCoin [17]. To our knowledge, TFCCommit is the first to merge CoSi with atomic commitment.

Fides uses a tamper-proof log to audit the system and detect any failures across database servers; this technique has been studied for decades in distributed systems [15], [32], [33]. In [32], Yumerefendi et al. highlight the use of accountability – a mechanism to detect and expose misbehaving servers – as a general distributed systems design. They implement CATS [33] an accountable network storage system that uses secure message logs to detect and expose misbehaving nodes. PeerReview [15] generalizes this idea by building a practical accountable system that uses tamper-evident logs to detect and irrefutably identify the faulty nodes. More recent solutions such as BlockchainDB [10], BigchainDB [23], Veritas [5] use blockchain as a tamper-proof log to store transactions across fully or partially replicated databases. CloudBFT [27], on the other hand, tolerates malicious faults in the cloud by relying on tamper-proof hardware to order the requests.

The datastore authentication technique that uses Merkle Hash Trees (MHT) and Verification Objects was first proposed by Merkle [24]. The technique employed in *Fides* that enables verifying the datastore per transaction is inspired by the work of Jain et al. [16]. Their solution assumes a single outsourced database, and requires a central trusted site to store the MHT roots of the outsourced data and the transaction history. Many

works have looked at query correctness, freshness, and data provenance for static data but only few solutions such as [20] and [26] (apart from [16] discussed above) consider data updates, but they also assume a single outsourced database.

VII. CONCLUSION

Traditional data management systems typically consider crash failures only. With the increasing usage of the cloud, crowdsourcing, and the rise of blockchain, the need to store data on untrusted servers has risen. The typical approach for achieving fault-tolerance, in general, uses replication. However, given the strict bounds on consensus in malicious settings, alternative approaches need to be explored. In this paper, we propose Fides, an auditable data management system designed for infrastructures that are *not* trusted. Instead of using replication for fault-tolerance, Fides uses fault-detection to discourage malicious behavior. An integral component of any distributed data management system is the commit protocol. We propose TFCOMMIT, a novel distributed atomic commitment protocol that executes transactions on untrusted servers. Since every server in Fides is untrusted, Fides replaces traditional transaction logs with a tamper-proof log similar to blockchain. The tamper-proof log stores all the necessary information required to audit the system and detect any failures. In conclusion, Fides provides a correct and solid data management system built on untrusted infrastructure that provides a solid foundation for future expansions to include other features of data management.

REFERENCES

- [1] Cyber Threat Data Manipulation. <https://www.theguardian.com/technology/2015/sep/10/cyber-threat-data-manipulation-us-intelligence-chief>. Accessed: 2019-07-10.
- [2] Equifax Data Breach. <https://investor.equifax.com/news-and-events/news/2017/09-15-2017-224018832>. Accessed: 2017-09-15.
- [3] Fides Technical Report. https://sites.cs.ucsb.edu/~sujaya_maiyya/assets/papers/Fides.pdf. Accessed: 2020-01-11.
- [4] M. Al-Bassam, A. Sonnino, S. Bano, D. Hryczyszyn, and G. Danezis. Chainspace: A shared smart contracts platform. *arXiv preprint arXiv:1708.03778*, 2017.
- [5] L. Allen, P. Antonopoulos, A. Arasu, J. Gehrke, J. Hammer, J. Hunter, R. Kaushik, D. Kossmann, J. Lee, R. Ramamurthy, et al. Veritas: Shared verifiable databases and tables in the cloud. *CIDR*, 2019.
- [6] G. Bracha and S. Toueg. Asynchronous consensus and broadcast protocols. *Journal of the ACM (JACM)*, 32(4):824–840, 1985.
- [7] M. Castro, B. Liskov, et al. Practical byzantine fault tolerance. In *OSDI*, volume 99, pages 173–186, 1999.
- [8] B. F. Cooper, A. Silberstein, E. Tam, R. Ramakrishnan, and R. Sears. Benchmarking cloud serving systems with ycsb. In *Proceedings of the 1st ACM symposium on Cloud computing*, pages 143–154. ACM, 2010.
- [9] D. Dolev and H. R. Strong. Authenticated algorithms for byzantine agreement. *SIAM Journal on Computing*, 12(4):656–666, 1983.
- [10] M. El-Hindi, C. Binnig, A. Arasu, D. Kossmann, and R. Ramamurthy. Blockchaindb: a shared database on blockchains. *Proceedings of the VLDB Endowment*, 12(11):1597–1609, 2019.
- [11] R. Garcia, R. Rodrigues, and N. Preguiça. Efficient middleware for byzantine fault tolerant database replication. In *European Conference on Computer Systems (EuroSys)*, pages 107–122. ACM, 2011.
- [12] H. Garcia Molina, F. Pittelli, and S. Davidson. Applications of byzantine agreement in database systems. *ACM Transactions on Database Systems (TODS)*, 11(1):27–47, 1986.
- [13] I. Gashi, P. Popov, V. Stankovic, and L. Strigini. On designing dependable services with diverse off-the-shelf sql servers. In *Architecting Dependable Systems II*, pages 191–214. Springer, 2004.
- [14] J. N. Gray. Notes on data base operating systems. In *Operating Systems*, pages 393–481. Springer, 1978.
- [15] A. Haeberlen, P. Kouznetsov, and P. Druschel. Peerreview: Practical accountability for distributed systems. *ACM SIGOPS operating systems review*, 41(6):175–188, 2007.
- [16] R. Jain and S. Prabhakar. Trustworthy data from untrusted databases. In *2013 IEEE 29th International Conference on Data Engineering (ICDE)*, pages 529–540. IEEE, 2013.
- [17] E. K. Kogias, P. Jovanovic, N. Gailly, I. Khoffi, L. Gasser, and B. Ford. Enhancing bitcoin security and performance with strong consistency via collective signing. In *25th {USENIX} Security Symposium ({USENIX} Security 16)*, pages 279–296, 2016.
- [18] E. Kokoris-Kogias, P. Jovanovic, L. Gasser, N. Gailly, E. Syta, and B. Ford. Omniledger: A secure, scale-out, decentralized ledger via sharding. In *2018 IEEE Symposium on Security and Privacy (SP)*, pages 583–598. IEEE, 2018.
- [19] K. Korpela, J. Hallikas, and T. Dahlberg. Digital supply chain transformation toward blockchain integration. In *proceedings of the 50th Hawaii international conference on system sciences*, 2017.
- [20] F. Li, M. Hadjieleftheriou, G. Kollios, and L. Reyzin. Dynamic authenticated index structures for outsourced databases. In *Proceedings of the 2006 ACM SIGMOD international conference on Management of data*, pages 121–132. ACM, 2006.
- [21] L. Luu, V. Narayanan, C. Zheng, K. Baweja, S. Gilbert, and P. Saxena. A secure sharding protocol for open blockchains. In *Proceedings of the 2016 ACM SIGSAC Conference on Computer and Communications Security*, pages 17–30. ACM, 2016.
- [22] S. Maiyya, F. Nawab, D. Agrawal, and A. E. Abbadi. Unifying consensus and atomic commitment for effective cloud data management. *Proceedings of the VLDB Endowment*, 12(5):611–623, 2019.
- [23] T. McConaghy, R. Marques, A. Müller, D. De Jonghe, T. McConaghy, G. McMullen, R. Henderson, S. Bellemare, and A. Granzotto. Bigchaindb: a scalable blockchain database. *white paper, BigChainDB*, 2016.
- [24] R. C. Merkle. A certified digital signature. In *Conference on the Theory and Application of Cryptology*, pages 218–238. Springer, 1989.
- [25] C. Mohan, R. Strong, and S. Finkelstein. Method for distributed transaction commit and recovery using byzantine agreement within clusters of processors. In *Proceedings of the second annual ACM symposium on Principles of distributed computing*, pages 89–103. ACM, 1983.
- [26] M. Narasimha and G. Tsudik. Authentication of outsourced databases using signature aggregation and chaining. In *International conference on database systems for advanced applications*, pages 420–436. Springer, 2006.
- [27] R. Nogueira, F. Araújo, and R. Barbosa. Cloudbft: elastic byzantine fault tolerance. In *2014 IEEE 20th Pacific Rim International Symposium on Dependable Computing*, pages 180–189. IEEE, 2014.
- [28] R. L. Rivest, A. Shamir, and L. Adleman. A method for obtaining digital signatures and public-key cryptosystems. *Communications of the ACM*, 21(2):120–126, 1978.
- [29] C.-P. Schnorr. Efficient signature generation by smart cards. *Journal of cryptography*, 4(3):161–174, 1991.
- [30] E. Syta, I. Tamas, D. Visher, D. I. Wolinsky, P. Jovanovic, L. Gasser, N. Gailly, I. Khoffi, and B. Ford. Keeping authorities’ honest or bust? with decentralized witness cosigning. In *2016 IEEE Symposium on Security and Privacy (SP)*, pages 526–545. Ieee, 2016.
- [31] B. Vandiver, H. Balakrishnan, B. Liskov, and S. Madden. Tolerating byzantine faults in transaction processing systems using commit barrier scheduling. In *ACM SIGOPS Operating Systems Review*, volume 41, pages 59–72. ACM, 2007.
- [32] A. R. Yumerefendi and J. S. Chase. The role of accountability in dependable distributed systems. In *Proceedings of HotDep*, volume 5, pages 3–3. Citeseer, 2005.
- [33] A. R. Yumerefendi and J. S. Chase. Strong accountability for network storage. *ACM Transactions on Storage (TOS)*, 3(3):11, 2007.
- [34] M. Zamani, M. Movahedi, and M. Raykova. Rapidchain: Scaling blockchain via full sharding. In *Proceedings of the 2018 ACM SIGSAC Conference on Computer and Communications Security*, pages 931–948. ACM, 2018.
- [35] W. Zhao. A byzantine fault tolerant distributed commit protocol. In *Third IEEE International Symposium on Dependable, Autonomic and Secure Computing (DASC 2007)*, pages 37–46. IEEE, 2007.