

RADON: Repairable Atomic Data Object in Networks

Kishori M. Konwar, N. Prakash, Nancy Lynch, Muriel Médard

Department of EECS, Massachusetts Institute of Technology
 {kishori, lynch}@casil.mit.edu, {prakashn, medard}@mit.edu

Abstract. Erasure codes offer an efficient way to decrease storage and communication costs while implementing atomic memory service in asynchronous distributed storage systems. In this paper, we provide erasure-code-based algorithms having the additional ability to perform background repair of crashed nodes. A repair operation of a node in the crashed state is triggered externally, and is carried out by the concerned node via message exchanges with other active nodes in the system. Upon completion of repair, the node re-enters active state, and resumes participation in ongoing and future read, write, and repair operations. To guarantee liveness and atomicity simultaneously, existing works assume either the presence of nodes with stable storage, or presence of nodes that never crash during the execution. We demand neither of these; instead we consider a natural, yet practical network stability condition $N1$ that only restricts the number of nodes in the crashed/repair state during broadcast of any message.

We present an erasure-code based algorithm $RADON_C$ that is always live, and guarantees atomicity as long as condition $N1$ holds. In situations when the number of concurrent writes is limited, $RADON_C$ has significantly improved storage and communication cost over a replication-based algorithm $RADON_R$, which also works under $N1$. We further show how a slightly stronger network stability condition $N2$ can be used to construct algorithms that never violate atomicity. The guarantee of atomicity comes at the expense of having an additional phase during the read and write operations.

1 Introduction

We consider the problem of designing algorithms for distributed storage systems (DSSs) that offer consistent access to stored data. Large scale DSSs are widely used by several industries, and also widely studied by academia for a variety of applications ranging from e-commerce to sequencing genomic-data. The most desirable form of consistency is atomicity, which in simple terms, gives the users of the data service the impression that the various concurrent read and write operations take place sequentially. Implementations of atomicity on an asynchronous system under message passing framework, in the presence of failures, is often challenging. Traditional implementations [1], [2] use replication of data as the mechanism of fault-tolerance; however they suffer from the problem of having high storage cost, and communication costs for read and write operations.

Erasure codes provide an efficient way to decrease storage and communication cost in atomicity implementations. An $[n, k]$ erasure code splits the value v , say of size 1 unit into k elements, each of size $\frac{1}{k}$ units, creates n coded elements, and stores one coded element per server. The size of each coded element is also $\frac{1}{k}$ units. A class of erasure codes known as Maximum Distance Separable (MDS) codes have the property that value v can be reconstructed from any k out of these n coded elements. While it is known that usage of erasure codes in asynchronous decentralized storage systems do not offer all the advantages as in synchronous centralized systems [3], erasure code based algorithms like in [4], [5], [6], or [7] for implementing consistent memory service offer significant storage and communication cost savings over replication based algorithms, in many regimes of operation. For instance CASGC [6] improves the costs under the scenario when the number of writes concurrent with a read is known to be limited, whereas SODA [7] trades-off write cost in order to optimize storage cost, which is meaningful in systems with infrequent writes. Both CASGC and SODA are based on MDS codes.

In this work, we consider the additional important issue of repairing crashed nodes without disrupting the storage service. Failure of storage nodes is a norm rather than an exception in large scale DSSs today, primarily because of the usage of commodity hardware for affordability and scalability reasons. Replication based algorithms in [1], [2] and erasure-code based algorithms in [4], [6], or [7] do not consider repair of crashed nodes; instead assume that a crashed node remains so for the rest of the execution. Algorithms in [5], [8] consider background repair of crashed nodes; however they assume either the presence of nodes having stable storage, whose content is unaffected by crashes, or presence of a subset of nodes that never crash during the entire execution. We relax both these assumptions in this work. In our model, any one of the storage nodes can crash; further, we assume that a crashed node loses all its data, both volatile as well as stable storage. A repair operation of a node in the crashed state is triggered externally, and is carried out by the concerned node via message exchanges with other active nodes in the system. Upon completion of repair, the node (with the same id) re-enters active state, and resumes participation in ongoing and future read, write, and repair operations.

It is natural to expect a restriction on the number of crash and repair operations in relation to the read and write operations; the authors of [8] show an impossibility result in this direction, for guaranteeing liveness and atomicity, simultaneously. We formulate network stability conditions $N1$ and $N2$, which can be used to limit the number of crash and repairs operations overlapping with a client operation. These conditions are algorithm independent, and most likely to be satisfied in any practical storage network. At a high level, the condition $N1$ restricts the set of servers that can be in the crashed or repair state any time a process (client or server) *pings* all the n servers with corresponding messages. Condition $N2$ is slightly stronger than $N1$, and restricts the set of servers that can be in the crashed or repair state if the process wants to *ping-pong* a fraction of the servers. In a ping-pong, it is expected that the servers which receive a message also respond back to the sender of the message.

1.1 Summary of Our Contributions

We first present an impossibility result for an asynchronous DSS allowing background repair of crashed nodes, where there is no restriction on the number of crash and repair operations that occur during a client operation. We show that it is impossible to simultaneously achieve liveness and atomicity in such a system, even if all the crash and repair operations occur sequentially during the execution (i.e., at most one node remains in the crash or repair state at any point during the execution).

We then consider the problem of erasure-code based algorithm design under the network stability condition $N1$. We present the algorithm in two stages. First we present a replication-based algorithm $RADON_R$, which performs background-repair, and guarantees atomicity and liveness of operations under $N1$, if more than $3/4^{\text{th}}$ of all servers remain active during any ping operation. The write and read phases are almost identical to those of the ABD algorithm [1], except that during a write we expect responses from more than $3/4^{\text{th}}$ of all the servers, while in ABD responses are expected only from a majority of servers. A repair operation in $RADON_R$ is simply a read operation initiated by the concerned server. Thus the algorithm itself is simple; however, the proof of atomicity gets complicated because of the fact that a repair operation can potentially restore the contents of a node to a version that is older than what was present before the crash. We show how the network stability condition can be used to prove atomicity, and this proof is the key takeaway from $RADON_R$ towards constructing the erasure-code based algorithm.

Our erasure-code based algorithm $RADON_C$ uses $[n, k]$ MDS codes, and is a natural adaptation of $RADON_R$ for the usage of codes. A key challenge while using erasure codes is ensuring liveness of read operations, in the presence of concurrent write operations. Various techniques are known in literature to handle this challenge; for instance, [5] assumes synchronous write phases, [6] limits the number of writes concurrent with a read, while [7] uses an $O(n^2)$ write protocol to guarantee liveness of reads. In this work, like in [6], we make the assumption that the number of write operations concurrent with any read operation is limited by a parameter δ , which is known a priori. However, the usage of the concurrency bound differs from that of the CASGC algorithm in [6]; for instance, CASGC has three rounds for write operations, while $RADON_C$ uses only two rounds. In $RADON_C$, each server maintains a list of up to $\delta + 1$ coded elements, corresponding to the latest $\delta + 1$ versions received as a result of the various write operations. In comparison with $RADON_R$ where a writer expects responses from more than $3/4^{\text{th}}$ of all servers, a write operation in $RADON_C$ expects responses from more than $\frac{3n+k}{4}$ servers. During a read operation, the client reads the lists from more than $\frac{n+k}{2}$ nodes before decoding the value v . Like in $RADON_R$, a repair operation in $RADON_C$ is essentially a read operation by the concerned node; however this time the concerned node creates a list (instead of just one version) by decoding as many possible versions that it can from the $\lceil \frac{n+k}{2} \rceil$ responses. Liveness and atomicity of operations are proved under network stability condition N1, if more than $\frac{3n+k}{4}$ servers remain active during any ping operation. $RADON_C$ has substantially improved storage and communication costs than $RADON_R$, when the concurrency bound δ is limited; see Table 1 for a comparison.

In both $RADON_R$ and $RADON_C$, violation of the network stability condition N1 can result in executions that are not atomic, which might not be preferable in certain applications. The choice of consistency over liveness, or vice versa, is the subject matter of a wide range of discussions and perspectives among system designers and software engineers. For example, BigTable, a DSS by Google, prefers safety over liveness [9], whereas, Amazon’s Dynamo does not compromise liveness but settles for *eventual consistency* [10]. Our third algorithm $RADON_R^{(S)}$, which is replication-based, is designed to guarantee atomicity during every execution. Liveness is guaranteed under the slightly more stringent condition of N2, with more than $3/4^{\text{th}}$ of all servers remaining active during any ping-pong operation. The guarantee of atomicity of every execution also needs extra phases for read and write operations, when compared to $RADON_R$. The design of an erasure-coded version of $RADON_R^{(S)}$ that never violates atomicity, is an interesting direction that we leave out for future work.

Algorithm	Write Cost	Read Cost	Storage Cost	Safe under	Live under
$RADON_R$	n	$2n$	n	N1	N1
$RADON_C$	$\frac{n}{k}$	$(\delta + 2)\frac{n}{k}$	$(\delta + 1)\frac{n}{k}$	N1	N1
$RADON_R^{(S)}$	n	$2n$	n	<i>always</i>	N2

Table 1. Performance comparison of $RADON_R$, $RADON_C$ and $RADON_R^{(S)}$, where n is the number of servers, and δ is the maximum number of writes concurrent with a read or a repair operation. See Section 7 for a justification of the costs.

1.2 Other Related Work

Dynamic Reconfiguration: Our setting is closely related to the problem of implementing a consistent memory object in a dynamic setting, where nodes are allowed to voluntarily leave and join the

network. The problem involves dynamic reconfiguration of the set of nodes that take part in client operations, which is often implemented via a *reconfig* operation that is initiated by any of the participating processes, including the clients. Any node that wants to leave/join the network makes an announcement, via a *leave/join* operation, before doing so. The problem is extensively studied in the field of distributed algorithms [11], [12], [13], [14], [15]; review and tutorial articles appear in [16], [17], [18].

In our context, the problem of node repair could in fact be thought of as one of dynamic reconfiguration, wherein an involuntary crash is simulated by a voluntary leave operation without an explicit announcement. In this case, a new node joins as a replacement node via the *join* operation, which can be considered as the analogue of a *repair* operation. In the setting of dynamic reconfiguration, every node has a distinct identity; thus the replacement node joins the network with a new identity that is different from the identity of the crashed node [16]. This demands a reconfiguration of the set of participating nodes after every repair. Such reconfigurations get in the way of client operations, and add to the latency of read and write operations [18], in practical implementations. Clearly, a repair operation as considered in this work does not demand any reconfiguration, since a repaired node has the same identity as the crashed node. Also, the current work shows that modeling repair via a static system, permits design of algorithms where clients remain oblivious to the presence of repair operations. Furthermore, addressing storage and communication costs is not the focus of the works in dynamic reconfigurations; specifically, it is not known as to how erasure codes can be advantageously used in dynamic settings. Our *RADON_C* algorithm shows that when repair is carried out under a static model, it is indeed possible to advantageously use erasure to reduce costs, when the number of concurrent writes are limited.

We make additional comparisons between our model and results to those found in works on dynamic reconfiguration. Several impossibility results exist in the context of implementing a dynamic atomic register and simultaneously guaranteeing liveness; the authors in [13] argue impossibility if there are infinitely many reconfigurations during an execution, while the authors in [14] argue an impossibility when there is no upper bound on message delay. We see, not surprisingly, that even in the problem of repair, we need to suitably limit the number of crash and repair operations that occur in an execution, even if all crash and repairs are sequentially ordered. In [15], the authors implement a dynamic atomic register under a model that has an (unknown) upper bound D on any point-to-point message delay, and where the number of reconfigurations in any D units of time is limited. Our network condition $N1$ is similar, except that 1) we limit the number of crash and repairs during any broadcast messaging, instead of point-to-point messaging, and 2) we do not assume any bound on the message delay. In practice, limiting number of repairs during broadcast instead of every point-to-point messaging offers resiliency against *straggler* nodes, which refer to the nodes having the worst delays among all nodes. We would also like to note that the algorithm in [15] does not guarantee atomicity, if the number of reconfigurations in D units of time is higher than a set number. This appears similar to *RADON_R*, where atomicity is not guaranteed if we do not satisfy stability condition $N1$. While we show how the slightly tighter model $N2$ can be used to always guarantee atomicity, it is an interesting question as to whether the model $N2$ can be adopted in the work of [15] so as to always guarantee atomicity.

Repair-Efficient Erasure Codes for Distributed Storage: Recently, a large class of new erasure/network codes for storage have been proposed (see [19] for a survey), and also tested in networks [20], [21], [22], where the focus is efficient storage of immutable data, such as, archival data. These new codes are specifically designed to optimize performance metrics like repair-bandwidth

and repair-time (of failed servers), and offer significant performance gains when compared to the traditional Reed-Solomon MDS codes [23]. It needs to be explored if these codes can be used in conjunction with the *RADON_C* algorithm, to further improve the performance costs.

Other Works on using Erasure Codes: Applications of erasure codes to Byzantine fault tolerant DSSs are discussed in [24], [25], [26]. In [3], the authors consider algorithms that use erasure codes for emulating *regular* registers. Regularity [27], [28] is a weaker consistency notion than atomicity.

The rest of the document is organized as follows. Our system model appears in Section 2. The impossibility result, and the network stability conditions appear in Section 3. The three algorithms appear in Sections 4, 5 and 6, respectively. In Section 7, we discuss the storage and communication costs of the algorithms. Section 8 concludes the paper. Proofs of various claims appear in the Appendix.

2 Models and definitions

Processes and Asynchrony: We consider a distributed system consisting of *asynchronous* processes, each with a unique identifier (ID), of three types: a set of *readers*, $\mathcal{R}$; a set of *writers*, $\mathcal{W}$; and a set of n *servers*, $\mathcal{S}$. The readers and writers are together referred to as clients. The set $\mathcal{R} \cup \mathcal{W} \cup \mathcal{S}$ forms a totally ordered set under some defined relation ($>$). The reader and writer processes initiate *read* and *write* operations respectively, and communicate with the servers using messages. A reader or writer can invoke a new operation only after all previous operations invoked by it has completed. The property is referred to as the *well-formedness* property of an execution. We assume that every client/server is connected to every other server via a reliable communication link; thus as long as the destination process is non-faulty, any message sent on the link eventually reaches the destination process.

Crash and Recovery: A client may fail at any point during the execution. At any point during the execution, a server can be in one (and only one) of the following three states: *active*, *crashed* or *repair*. A crash event triggers a server to enter the *crashed* state from an *active* state. The server remains in the *crashed* state for an arbitrary amount of time, but eventually is triggered by a repair event to enter the *repair* state. Crash and repair events are assumed to be externally triggered. A server in the *repair* state can experience another crash event, and go back to the *crashed* state. A server in the *crashed* state does not perform any local computation. The server also does not send or receive messages in the *crashed* state, i.e., any message reaching the server in a *crashed* state is lost. A server which enters the *repair* state has all its local state variables set to default values, i.e., a crash event causes the server to lose all its state variables. A server in the *repair* state can perform computations like in the *active* state.

Atomicity and Liveness: We aim to implement only one atomic read/write memory object, say x , under the MWMR setting on a set of servers, because any shared atomic memory can be emulated by composing individual atomic objects. The object value v comes from some set V ; initially v is set to a distinguished value v_0 ($\in V$). Reader r requests a read operation on object x . Similarly, a write operation is requested by a writer w . Each operation at a non-faulty client begins with an *invocation step* and terminates with a *response step*. An operation is *incomplete* when its invocation step does not have the associated response step; otherwise it is *complete*.

By *liveness of a read or a write operation*, we mean that during any well-formed execution, any read or write operation respectively initiated by a non-faulty reader or writer completes, despite

the crash failure of any other client. By *liveness of repair* associated with a crashed server, we mean that the server which enters a crashed state eventually re-enters the active state, unless it experiences a crash event during every repair operation that the server attempts. The liveness of repair holds despite the crash failure of any other client.

Background on Erasure coding: In $RADON_C$, we use an $[n, k]$ linear MDS code [30] over a finite field $\mathbb{F}_q$ to encode and store the value v among the n servers. An $[n, k]$ MDS code has the property that any k out of the n coded elements can be used to recover (decode) the value v . For encoding, v is divided¹ into k elements $v_1, v_2, \dots, v_k$ with each element having size $\frac{1}{k}$ (assuming size of v is 1). The encoder takes the k elements as input and produces n coded elements $c_1, c_2, \dots, c_n$ as output, i.e., $[c_1, \dots, c_n] = \Phi([v_1, \dots, v_k])$, where Φ denotes the encoder. For ease of notation, we simply write $\Phi(v)$ to mean $[c_1, \dots, c_n]$. The vector $[c_1, \dots, c_n]$ is referred to as the codeword corresponding to the value v . Each coded element c_i also has size $\frac{1}{k}$. In our scheme we store one coded element per server. We use Φ_i to denote the projection of Φ on to the i^{th} output component, i.e., $c_i = \Phi_i(v)$. Without loss of generality, we associate the coded element c_i with server i , $1 \leq i \leq n$.

Storage and Communication Cost: We define the total storage cost as the size of the data stored across all servers, at any point during the execution of the algorithm. The communication cost associated with a read or write operation is the size of the total data that gets transmitted in the messages sent as part of the operation. We assume that metadata, such as version number, process ID, etc. used by various operations is of negligible size, and is hence ignored in the calculation of storage and communication cost. Further, we normalize both the costs with respect to the size of the value v ; in other words, we compute the costs under the assumption that v has size 1 unit.

3 Network Stability Conditions

3.1 An Impossibility Result

The crash and recovery model described in Section 2 does not impose any restriction on the *rate of crash events, and repair operations* that happen in the system. In other words, the model described above does not limit in any manner the number of crash events/repair operations, which can overlap with any a client operation. In [8], the authors showed that without such restrictions, it is impossible to implement a shared atomic memory service, which guarantees liveness of operations. Below, we state an impossibility result which holds even if there is at most one server in the crashed/repair state at any point during the execution. We then introduce network stability conditions that enable us impose restrictions on the number of crash/repair events that overlap with any operation.

Theorem 1. *It is impossible to implement an atomic memory service that guarantees liveness of reads and writes, under the system model described in Section 2, even if 1) there is at most one server in the crashed/repair state at any point during the execution, and 2) every repair operation completes, and takes the repaired server back to the active state.*

3.2 Network Stability Conditions $N1$ and $N2$

We begin with the notions of a group-send operation, and effective consumption of a message.

¹ In practice v is a file, which is divided into many stripes based on the choice of the code, various stripes are individually encoded and stacked against each other. We omit details of representability of v by a sequence of symbols of $\mathbb{F}_q$, and the mechanism of data striping, since these are fairly standard in the coding theory literature.

group-send operation: The group-send operation is used to abstract the operation of a process sending a list of n messages $\{m_1, \dots, m_n\}$ to the set of all n servers $\{s_1, \dots, s_n\} = \mathcal{S}$, where message m_i is send to server $s_i, 1 \leq i \leq n$. Note that this is a mere abstraction of the process sending out n point-to-point messages sequentially to n servers, without interleaving the “send” operations with any significant local computations or waiting for any external inputs. The operation is no more powerful then sending n consecutive messages. The operation is written as $group\text{-}send([m_1, m_2, \dots, m_n])$. In the event $m_i = m, \forall i$, we simply write $group\text{-}send(m)$. Our model allows the sender to fail while executing the $group\text{-}send$ operation, in which case only a subset of the n servers receive their corresponding messages.

Effective Consumption: We say a process effectively consumes a message m , if it receives m , and executes all steps of the algorithm that depend only on the local state of the process, and the message m ; in other words, the process executes all the steps that do not require any further external messages.

Definition 1 (Network Stability Conditions). Consider a process p executing a $group\text{-}send([m_1, m_2, \dots, m_n])$ operation, and consider the following statements:

(a) (i) There exists a subset $\mathcal{S}_\alpha \subseteq \mathcal{S}$ of $|\mathcal{S}_\alpha| = \lceil \alpha n \rceil$ servers, $0 < \alpha < 1$, all of which effectively consume their respective messages from the $group\text{-}send$ operation, and (ii) all the servers in $\mathcal{S}_\alpha$ remain in the active state during the interval $[T_1 T_2]$, where T_1 denotes the point of time of invocation of the $group\text{-}send$ operation, and T_2 denotes the earliest point of time in the execution at which all of the servers in $\mathcal{S}_\alpha$ complete the effective consumption of their respective messages.

(b) Further, if effective consumption of the message m_i by server s_i involves sending a response back to the process p , for all $s_i \in \mathcal{S}_\alpha$, then all servers in $\mathcal{S}_\alpha$ remain in the active state during the interval $[T_1 T_3]$, where T_3 denotes the earliest point of time in the execution at which the process p completes effective consumption of the responses from the all the servers in $\mathcal{S}_\alpha$.

If the network satisfies Statement (a) for every execution of a $group\text{-}send$ operation by any process, we say that it satisfies network stability condition $N1$ with parameter α . If the network satisfies Statements (a) and (b) for every execution of a $group\text{-}send$ operation by any process, we say that it satisfies network stability condition $N2$ with parameter α .

Clearly, $N2$ implies $N1$. Note that the set $\mathcal{S}_\alpha$ which needs to satisfy the conditions need not be the same for various invocations of $group\text{-}send$ operations by either the same or distinct processes. Also, note that in condition $N2$, the process p might crash before completing the effective consumption of the responses from the servers in $\mathcal{S}_\alpha$. In this case we only expect Statement (a) to be satisfied, and not Statement (b). Furthermore, in both $N1$ and $N2$, we do not expect any of these statements to be true, if process p crashes after partial execution of the $group\text{-}send$ operation.

4 The $RADON_R$ Algorithm

In this section, we present the $RADON_R$ algorithm, and prove its liveness and atomicity properties for networks that satisfy the network condition $N1$ with $\alpha > \frac{3}{4}$. We begin with some useful notation. Tags are used for version control of the object values. A tag t is defined as a pair (z, w) , where $z \in \mathbb{N}$ and $w \in \mathcal{W}$ denotes the ID of a writer. We use $\mathcal{T}$ to denote the set of all the possible tags. For any two tags $t_1, t_2 \in \mathcal{T}$, we say $t_2 > t_1$ if (i) $t_2.z > t_1.z$ or (ii) $t_2.z = t_1.z$ and $t_2.w > t_1.w$. Note that $(\mathcal{T}, >)$ is a totally ordered set.

The protocols for writer, reader, and servers are shown in Fig. 1. Each server stores two state variables (i) (t_{loc}, v_{loc}) - a tag and value pair, initially set to (t_0, v_0) , (ii) *status* - a variable that can be in either *active* or *repair* state.

Fig. 1 The protocols for writer, reader, and any server $s \in S$ in $RADON_R$.

write(v): <u>get-tag:</u> $group_send(QUERY_TAG)$ Await responses from majority Select the max tag t^* <u>put-data:</u> $t_w = (t^*.z + 1, w)$ $group_send((PUT_DATA, (t_w, v)))$ Terminate after $\lceil \frac{3n+1}{4} \rceil$ acks. read: <u>get-data:</u> $group_send(QUERY_TAG_DATA)$ Await responses from majority Select (t_r, v_r), with max tag. <u>put-data :</u> $group_send((PUT_DATA, (t_r, v_r)))$ Wait for $\lceil \frac{3n+1}{4} \rceil$ acks Return v_r Server $s \in S$: <u>State Variables:</u> $(t_{loc}, v_{loc}) \in \mathcal{T} \times \mathcal{V}$, initially (t_0, v_0) $status \in \{active, repair\}$, initially <i>active</i>	<u>get-tag-resp, recv QUERY-TAG from writer w:</u> if $status = active$ then Send t_{loc} to w <u>get-data-resp, recv QUERY-TAG-DATA from reader r:</u> if $status = active$ then Send (t_{loc}, v_{loc}) to r <u>put-data-resp, recv PUT-DATA, (t, v) from client c :</u> if $status = active$ then if $t > t_{loc}$ then $(t_{loc}, v_{loc}) \leftarrow (t, v)$ Send ack to c. <u>init-repair :</u> $status \leftarrow repair$ $(t_{loc}, v_{loc}) \leftarrow (t_0, v_0)$ $group_send(REPAIR_TAG_DATA)$ Await responses from majority Select (t_{rep}, v_{rep}), for max tag $(t_{loc}, v_{loc}) \leftarrow (t_{rep}, v_{rep})$ $status \leftarrow active$ <u>init-repair-resp, recv REPAIR-TAG-DATA from s':</u> if $status = active$ then Send (t_{loc}, v_{loc}) to s'
---	--

The write and read operations are very similar to those in the ABD algorithm [1], and each consists of two phases. In the first phase, *get-tag*, of a write operation π , the writer queries all servers for local tags, awaits responses from a majority of servers, and selects the maximum tag t^* from among the responses. Next, the writer executes the *put-data* phase, during which a new tag $t_w = tag(\pi)$ is created by incrementing the integer part of t^* , and by incorporating the writer's own ID. The writer then sends pair (t_w, v) to all servers, and awaits acknowledgments (acks) from $\lceil \frac{3n+1}{4} \rceil$ servers before completing the operation. The two phases are identical to those of the ABD algorithm [1], except for the fact that during the second phase, ABD expects acks from only a majority of servers, whereas here we need from $\lceil \frac{3n+1}{4} \rceil$ servers. During a read operation ρ , the reader in the *get-data* phase queries all the servers in S for the respective local tag and value pairs. Once it receives responses from a majority of servers in S , it picks the pair with the highest tag, which we designate as $t_r = tag(\pi)$. In the subsequent *put-data* phase, the reader writes back the tag t_r and the corresponding value v_r to all servers, and terminates after receiving acknowledgments from $\lceil \frac{3n+1}{4} \rceil$ servers. Once again, we remark that both phases in the read are identical to those of the ABD algorithm, except for the difference in the number of the servers from which acks are

expected in the second write-back phase. Note that, during both the write and operations, a server responds to an incoming message only if it is in the active state.

A repair operation is initiated via the action *init-repair*, by an external trigger, at a server which is in the crashed state. Note that we do not explicitly define a *crashed* state since a crash is not a part of the algorithm. We assume that as soon as the repair operation starts, the variable *status* is set to the *repair* state, and also the local (tag, value) pair is set to the default state (t_0, v_0) . The repair operation is essentially the first phase of the read operation, during which the server queries all the servers for the respective local tag and value pairs, and stores the tag and value pair corresponding to the highest tag after receiving responses from a majority of servers. Finally, the repair operation is terminated setting variable *status* to *active* state. A server in S responds to a request generated from *init-repair* phase only if it is in the active state.

4.1 Analysis of $RADON_R$

Liveness of read, write and repair operations in $RADON_R$ follows immediately if we assume condition $N1$ with $\alpha > \frac{3}{4}$. This is because liveness of any operation depends on sufficient number of responses from the servers during the various phases of the operation. From Fig. 1, we know that the maximum number of responses that is expected in any phase is $\lceil \frac{3n+1}{4} \rceil$, which is guaranteed under $N1$ with $\alpha > \frac{3}{4}$.

The tricky part is to prove atomicity of reads and writes. The proof is based on Lemma 13.16 of [31], a restatement of which can be found in [29]. Consider two completed write operations π_1 and π_2 , such that, π_2 starts after the completion of π_1 . For any completed write operation π , we define $tag(\pi) = t_w$, where t_w is the tag which the writer uses in the *put-data* phase. In this case, one of the requirements the algorithm needs to satisfy to ensure atomicity is $tag(\pi_2) > tag(\pi_1)$. While this fact is straightforward to prove for an algorithm like ABD, which does not have background repair, in $RADON_R$, we need to consider the effect of those repair operations that overlap with π_1 , and also those that occur in between π_1 and π_2 . The point to note is that such repair operations can potentially restore the contents of the repaired node such that the restored tag is less than $tag(\pi_1)$. We then need to show the absence of propagation of older tags (older than $tag(\pi_1)$) into a majority of nodes, due to a sequence of repairs which happen before π_2 decides its tag. We do this via the following two observations: 1) In Lemma 1, we show that any successful repair operation, which begins after a point of time T , always restores value to one, which corresponds to a tag which is at least as high as the minimum of the tags stored in any majority of active servers at time T . This fact is in turn used to prove a similar property for reads and writes, as well. 2) We next show (as part of proof of Theorem 3), under the assumption of $N1$ with $\alpha > 3/4$, the existence of a point of time T before the completion of π_1 such that a majority of nodes are active at T , and all of whose tags are at least as high as $tag(\pi_1)$. The two steps are together used to prove that $tag(\pi_2) > tag(\pi_1)$. A similar sequence of steps are used to show atomicity properties of read operations, as well.

For a completed read operation π , $tag(\pi) = t_r$, where t_r is the tag corresponding to the value v_r returned by the reader. For a completed repair π , $tag(\pi) = t_{rep}$, where t_{rep} is the tag corresponding to the value restored during the repair operation.

Lemma 1. *Let β denote a well-formed execution of $RADON_R$. Suppose T denotes a point of time in β such that there exists a majority of servers $\mathcal{S}_m$, $\mathcal{S}_m \subset \mathcal{S}$ all of which are in the active state at time T . Also, let t_s denote the value of the local tag at server $s \in \mathcal{S}_m$, at time T . Then, if π denotes any completed repair or read operation that is initiated after time T , we have $tag(\pi) \geq \min_{s \in \mathcal{S}_m} t_s$.*

Also, if π denotes any completed write operation that is initiated after time T , we have $\text{tag}(\pi) > \min_{s \in \mathcal{S}_m} t_s$.

Theorem 2 (Liveness). *Let γ denote a well-formed execution of RADON_R , under the condition $N1$ with $\alpha > \frac{3}{4}$. Then every operation initiated by a non-faulty client completes.*

Theorem 3 (Atomicity). *Every execution of the RADON_R algorithm operating under the $N1$ network stability condition with $\alpha > \frac{3}{4}$, is atomic.*

We note that, though Lemma 1 gives a result about completed operations, condition $N1$ is not a prerequisite for the result in Lemma 1. In other words, the result in Lemma 1 holds for any completed operation, even if condition $N1$ is violated. As we will see, this is an important fact that we will use to establish atomicity of $\text{RADON}_R^{(S)}$ for any execution.

5 Algorithm RADON_C

In this section, we present the erasure-code based RADON_C algorithm for implementing atomic memory service, and performing repair of crashed nodes. The algorithm uses $[n, k]$ MDS codes for storage. Liveness and atomicity are guaranteed under the following assumptions: 1) the $N1$ network stability condition with $\alpha \geq \frac{3n+k}{4n}$, 2) the number of write operations concurrent with a read or repair operation is at most δ . The precise definition of concurrency depends on the algorithm itself, and appears later in this section. The RADON_C algorithm has significantly reduced storage and communication cost requirements than RADON_R , when δ is limited.

The algorithm (see Fig. 2) is a natural generalization of the RADON_R algorithm accounting for the fact that we use MDS codes. The write operation has two phases, where the first phase finds the maximum tag in the system based on majority responses. During the second phase, the writer computes the coded elements for each of the n servers and uses the group-send operation to disperse them. The *group-send* operation here uses a vector of length n , where the i^{th} element denotes the message for the i^{th} server, $1 \leq i \leq n$. Each server keeps a *List* of up to $(\delta + 1)$ (tag, coded-element) pairs. Every time a (tag, coded-element) message arrives from a writer, the pair gets added to the *List*, which is then pruned to at most $(\delta + 1)$ pairs, corresponding to the highest tags. The writer terminates after getting acks from $\lceil \frac{3n+k}{4} \rceil$ servers.

During a read operation, the reader queries all servers for their entire local *Lists*, and awaits responses from $\lceil \frac{n+k}{2} \rceil$ servers. Once the reader receives *Lists* from $\lceil \frac{n+k}{2} \rceil$ servers, it selects the highest tag t_r whose corresponding value v_r can be decoded using the coded elements in the lists. The read operation completes following a write-back of (t_r, v_r) using the *put-data* phase.

The repair operation is very similar to the first phase of the read operation, during which a server collects lists from $\lceil \frac{n+k}{2} \rceil$ servers. But this time, the server figures out the set of all the possible tags that can be decoded from among the *Lists*, and prunes the set to the highest $(\delta + 1)$ tags. The repaired *List* then consists of (tag, coded-element) pairs corresponding these (at most) $(\delta + 1)$ tags. Assuming repair of server i , the creation of a coded-element corresponding to a value v involves first decoding the value v , and then computing $\Phi_i(v)$ (referred to as re-encoding in Fig. 2).

5.1 Analysis of RADON_C

Throughout this section, we assume network stability condition $N1$ with $\alpha \geq \frac{3n+k}{4n}$. Tags for completed read and write operations are defined in the same manner as we did for RADON_R ; we

Fig. 2 The protocols for write, reader, and any server $s_i \in \mathcal{S}$ in $RADON_C$.

write(v):	$List \subseteq \mathcal{T} \times \mathcal{C}_s$, initially $\{(t_0, \Phi_i(v_0))\}$
<u>get-tag:</u>	<u>get-tag-resp, recv QUERY-TAG from writer w:</u>
group-send(QUERY-TAG)	if $status = active$ then
Await responses from majority	$t^* = \max_{(t,c) \in List} t$
Select the max tag t^*	Send t^* to w
<u>put-data:</u>	<u>get-data-resp, recv QUERY-LIST from reader r:</u>
$t_w = (t^*.z + 1, w)$.	if $status = active$ then
$code\text{-}elems = [(t_w, c_1), \dots, (t_w, c_n)]$, $c_i = \Phi_i(v)$	Send $List$ to r
group-send(CODE-ELEMENTS, $code\text{-}elems$)	<u>put-data-resp, recv CODE-ELEMENTS, (t, c_i) from p:</u>
Terminate after $\lceil \frac{3n+k}{4} \rceil$ acks	if $status = active$ then
read:	$List \leftarrow List \cup \{(t, c_i)\}$
<u>get-data:</u>	if $ List > \delta + 1$ then
group-send(QUERY-LIST)	Retain the (tag, coded-element) pairs for the $\delta + 1$
Wait for $\lceil \frac{n+k}{2} \rceil$ $Lists$	highest tags in $List$, and delete the rest.
Select the max tag, t_r , whose corresponding value, v_r , is decodable using the $Lists$.	Send ack to p .
<u>put-data:</u>	<u>init-repair :</u>
$code\text{-}elems = [(t_r, c_1), \dots, (t_r, c_n)]$, $c_i = \Phi_i(v_r)$	$status \leftarrow repair$
group-send(CODE-ELEMENTS, $code\text{-}elems$)	group-send(REPAIR-LIST)
Wait for $\lceil \frac{3n+k}{4} \rceil$ acks	Wait for $\lceil \frac{n+k}{2} \rceil$ $Lists$
Return v_r	Find (tag, value) pairs decodable from $Lists$.
Server $s_i \in \mathcal{S}$:	Restore local $List$ via re-encoding and retaining the
<u>State Variables:</u>	(tag, coded-element) pairs corresponding to at most $\delta + 1$
$status \in \{active, repair\}$, initially $active$	highest tags, from the above pairs
	$status \leftarrow active$
	<u>init-repair-resp, recv REPAIR-LIST from server s':</u>
	if $status = active$ then
	Send $List$ to s'

avoid repeating them here. We first discuss liveness properties of $RADON_C$. Let us first consider liveness of repair operations. Towards this, note from the algorithm in Fig. 2 that a repair operation never gets stuck even if it does not find any set of k $Lists$ among the responses, all of which have a common tag. In such a case, the algorithm allows the possibility that the repaired $List$ is simply empty, at the point of execution when the server re-enters the active state. In other words, liveness of a repair operation is trivially proved, i.e., a server in a repair state always eventually reenters the active state, as long as it does not experience a crash during the repair operation. The triviality of liveness of repair operations, observed above, does not extend to read operations. For a read operation to complete the *get-data* phase, it must be able to find a set of k $Lists$ among the responses all of which contain coded-elements corresponding to a common tag; otherwise a read operation gets stuck. The discussion above motivates the following definitions of valid read and valid repair operations.

Definition 2 (Valid Read and Repair Operations). *A read operation will be called as a valid read if the associated reader remains alive at least until the reception of the $\lceil \frac{n+k}{2} \rceil$ responses during the *get-data* phase. Similarly, a repair operation will be called a valid repair if the associated server does not experience a further crash event during the repair operation.*

Definition 3 (Writes Concurrent with a Valid Read (Repair)). Consider a valid read (repair) operation π . Let T_1 denote the point of initiation of π . For a valid read, let T_2 denote the earliest point of time during the execution when the associated reader receives all the $\lceil \frac{n+k}{2} \rceil$ responses. For a valid repair, let T_2 denote the point of time during the execution when the repair completes, and takes the associated server back to the active state. Consider the set $\Sigma = \{\sigma : \sigma \text{ is any write operation that completes before } \pi \text{ is initiated}\}$, and let $\sigma^* = \arg \max_{\sigma \in \Sigma} \text{tag}(\sigma)$. Next, consider the set $\Lambda = \{\lambda : \lambda \text{ is any write operation that starts before } T_2 \text{ such that } \text{tag}(\lambda) > \text{tag}(\sigma^*)\}$. We define the number of writes concurrent with the valid read (repair) operation π to be the cardinality of the set Λ .

The above definition captures all the write operations that overlap with the read, until the time the reader has all data needed to attempt decoding a value. However, we ignore those write operations that might have started in the past, and never completed yet, if their tags are less than that of any write that completed before the read started. This allows us to ignore write operations due to failed writers, while counting concurrency, as long as the failed writes are followed by a successful write that completed before the read started.

The following lemma could be considered as the analogue of Lemma 1 for RADON_C . The first part of the lemma shows that under N1 with $\alpha \geq \frac{3n+k}{4n}$, the repaired *List* is never empty; there is always at least one (tag, coded-element) pair in the repaired *List*. Parts 2 and 3 are used to prove liveness and atomicity of client operations.

Lemma 2. Consider any well-formed execution β of RADON_C operating under the network stability condition N1 with $\alpha \geq \frac{3n+k}{4n}$. Further assume that the number of writes concurrent with any valid read or repair operation is at most δ . For any operation π , consider the set $\Sigma = \{\sigma : \sigma \text{ is a read or a write in } \beta \text{ that completes before } \pi \text{ begins}\}$, and also let $\sigma^* = \arg \max_{\sigma \in \Sigma} \text{tag}(\sigma)$. Then, the following statements hold:

- If π denotes a completed repair operation on a server $s \in S$, then the repaired *List* of server s due to π contains the pair $(\text{tag}(\sigma^*), c_s^*)$.
- If π denotes a read operation associated with a non-faulty reader r , and further, if S_1 denotes the set of $\lceil \frac{n+k}{2} \rceil$ servers whose responses, say $\{L_\pi(s), s \in S_1\}$, are used by r to attempt decoding of a value in the get-data phase, then there exists $S_2 \subseteq S_1$, $|S_2| = k$, such that $\forall s \in S_2, (\text{tag}(\sigma^*), c_s^*) \in L_\pi(s)$.
- If π denotes a write operation associated with a non-faulty writer w , and further if S_1 denotes the set of majority servers whose responses are used by w to compute max-tag in the get-tag phase, then there exists a server $s \in S_1$, whose response tag $t_s \geq \text{tag}(\sigma^*)$.

Here, c_s^* denotes the coded-element of server s for value v^* , associated with $\text{tag}(\sigma^*)$.

Theorem 4 (Liveness). Let β denote a well-formed execution of RADON_C , operating under the N1 network stability condition with $\alpha \geq \frac{3n+k}{4n}$ and δ be the maximum number of write operations concurrent with any valid read or repair operation. Then every operation initiated by a non-faulty client completes.

Theorem 5 (Atomicity). Any execution of RADON_C , operating under condition N1 with $\alpha \geq \frac{3n+k}{4n}$, is atomic, if the maximum number of write operations concurrent with a valid read or repair operation is δ .

6 The $RADON_R^{(S)}$ Algorithm

In this section, we present the $RADON_R^{(S)}$ algorithm having the property that every execution is atomic. Liveness is guaranteed under the slightly stronger network stability condition $N2$ with $\alpha > \frac{3}{4}$. In comparison with $RADON_R$, the algorithm has extra phases for both read and write operations, in order to guarantee safety of every execution.

Fig. 3 The protocols for writer, reader, and any server $s \in \mathcal{S}$ in $RADON_R^{(S)}$.

write(v): <u>get-tag:</u> $group\text{-}send(QUERY\text{-}TAG)$ Await responses from majority Select the max tag t^* <u>put-data:</u> $t_w = (t^*.z + 1, w)$. $group\text{-}send((PUT\text{-}DATA, (t_w, v)))$ Wait for $\lceil \frac{3n+1}{4} \rceil$ acks (say from $\mathcal{S}_\alpha$) <u>confirm-data:</u> $group\text{-}send((CONFIRM\text{-}DATA, t_w))$ Terminate after acks from majority from among servers in $\mathcal{S}_\alpha$ read: <u>get-data:</u> $group\text{-}send(QUERY\text{-}TAG\text{-}DATA)$ Await responses from majority Select (t_r, v_r), with max tag. <u>put-data :</u> $group\text{-}send((PUT\text{-}DATA, (t_r, v_r)))$ Wait for $\lceil \frac{3n+1}{4} \rceil$ acks (say from $\mathcal{S}_\alpha$) <u>confirm-data:</u> $group\text{-}send((CONFIRM\text{-}DATA, t_r))$ Await acks from a majority of servers in $\mathcal{S}_\alpha$ Return v_r Server $s \in \mathcal{S}$: <u>State Variables:</u> $(t_{loc}, v_{loc}) \in \mathcal{T} \times \mathcal{V}$, initially (t_0, v_0) $status \in \{active, repair\}$, initially <i>active</i> $Seen \subseteq \mathcal{T} \times \{\mathcal{W} \cup \mathcal{R}\}$, initially empty	<u>get-tag-resp, recv QUERY-TAG from writer w:</u> if $status = active$ then Send t_{loc} to w <u>get-data-resp, recv QUERY-TAG-DATA from reader r:</u> if $status = active$ then Send (t_{loc}, v_{loc}) to r <u>put-data-resp, recv (PUT-DATA, (t, v)) from c :</u> if $status = active$ then if $t > t_{loc}$ then $(t_{loc}, v_{loc}) \leftarrow (t, v)$ $Seen \leftarrow Seen \cup \{(t, c)\}$ Send ack to c. <u>confirm-data-resp, recv (CONFIRM-DATA, t) from c:</u> if $status = active$ then if $(t, c) \in Seen$ then Remove (t, c) from $Seen$ Send ack to client c. <u>init-repair :</u> $status \leftarrow repair$ $(t_{loc}, v_{loc}) \leftarrow (t_0, v_0)$ $Seen \leftarrow \emptyset$ $group\text{-}send(REPAIR\text{-}TAG\text{-}DATA)$ Await responses from majority. Select (t_{rep}, v_{rep}), with max tag $(t_{loc}, v_{loc}) \leftarrow (t_{rep}, v_{rep})$ $status \leftarrow active$ <u>init-repair-resp, recv REPAIR-TAG-DATA from s':</u> if $status = active$ then Send (t_{loc}, v_{loc}) to s'
--	--

The write operation has three phases (see Fig. 3). The first two phases are identical to those of $RADON_R$ during which the writer queries for the local tags, and then sends out the new (tag, value) pair, respectively. In the third phase, called *confirm-data*, the writer ensures the presence of at least a majority of servers, which the writer knows for sure that received its data during the second phase, *put-data*. In order to facilitate the *confirm-data* phase, the servers maintain a *Seen* variable. Any time the server receives a value from a writer, the server adds the corresponding (tag,

writer ID) pair to the *Seen* list. Next, during the *confirm-data-resp* phase, the server responds to the writer only if this (tag, writer ID) pair is part of the *Seen* variable. The idea is that if the server experiences a crash and a successful repair operation in between the *put-data* and *confirm-data* phases, the server no longer has the (tag, writer ID) pair in its *Seen* variable, and hence does not respond to the *confirm-data* phase. This is because, a crash removes all state variables, including *Seen*, and the repair algorithm (see Fig. 3) simply restores the *Seen* variable to its default value, the empty set. Further, by ensuring that the writer expects acks from among a majority of servers in *confirm-data*, from among the $\frac{3n+1}{4}$ servers whose acks were obtained during *put-data*, we can guarantee that any execution is atomic.

The read operation also has three phases, first two of which are identical to those of $RADON_R$, except for the use of the *Seen* variable in the server during the *put-data* phase. The third phase is the *confirm-data* phase as in the write operation. The repair operation has one phase, and is nearly exactly identical to that of $RADON_R$. Note that the *Seen* variable gets reset to its initial value during repair.

6.1 Analysis of $RADON_R^{(S)}$

We overview the proofs of liveness and atomicity before formal claims. For liveness of writes, we assume $N2$ with $\alpha > \frac{3}{4}$, and argue the existence of a majority $\mathcal{S}_m$ of servers all of which remain active from the point of time at which the *group-send* operation gets initiated in the *put-data* phase, till the point of time all the servers in $\mathcal{S}_m$ effectively consume requests for *confirm-data* from the writer. In this case, write operation completes after receiving acks from servers in $\mathcal{S}_m$ during the *confirm-data* phase. The set $\mathcal{S}_m$ exists because, under $N2$ with $\alpha > \frac{3}{4}$, a set $\mathcal{S}_\alpha$ of $\lceil \frac{3n+1}{4} \rceil$ servers remain alive from the start of the group-send, till the effective consumption of the acks by the writer in *put-data* phase. Also, a second set $\mathcal{S}'_\alpha$ of $\lceil \frac{3n+1}{4} \rceil$ servers remain active from the start of the group-send in the *confirm-data* phase, till all servers in $\mathcal{S}'_\alpha$ complete the respective effective consumption from this group-send. We note that $\mathcal{S}'_\alpha \cap \mathcal{S}_\alpha$ is at least a majority. We next use the observation that the *group-send* operation in the *confirm-data* phase forms part of the effective consumption of the last of the acks in the *put-data* phase. Using this, we argue that the servers in $\mathcal{S}'_\alpha \cap \mathcal{S}_\alpha$ remain active till they effectively consume message from *group-send* operation of the *confirm-data* phase, and thus $\mathcal{S}'_\alpha \cap \mathcal{S}_\alpha$ is a candidate for $\mathcal{S}_m$. The liveness of read is similar to that of write, while liveness of repair is straightforward under $N2$ with $\alpha > \frac{3}{4}$.

Towards proving atomicity of reads and writes, we first define tags for completed reads, writes and repair operations exactly in the same manner as we did in $RADON_R$. Consider two completed write operations π_1 and π_2 such that π_2 starts after the completion of π_1 , and we need to show that $tag(\pi_2) > tag(\pi_1)$. As in $RADON_R$, we do this in two parts: Lemma 1 holds as it is for $RADON_R^{(S)}$ as well. Recall that Lemma 1 essentially shows that if a majority of active nodes is locked-on to any particular tag, say t' , at a specific point of time T during the execution of the algorithm, then any repair operation which begins after the time T always restores the tag to one which is at least as high as t' . The challenge now is to show the existence of these favorable points of time instants T as needed in the assumption of the lemma. While in $RADON_R$, we used the $N1$ to argue this, in $RADON_R^{(S)}$, we do not use $N2$; instead we rely on the third *confirm-data* phase of the first write operation π_1 .

Theorem 6 (Liveness). *Let β denote a well-formed execution of $RADON_R^{(S)}$ under condition $N2$ with $\alpha > \frac{3}{4}$. Then every operation initiated by a non-faulty client completes.*

Theorem 7 (Atomcity). *Every execution of the $RADON_R^{(S)}$ algorithm is atomic.*

7 Storage and Communication Costs of Algorithms

We give a justification of storage and communication cost numbers of the three algorithms, appearing in Table 1. Recall that the size of value v is assumed to be 1 and also that we ignore the costs due to metadata. It is clear that both $RADON_R$ and $RADON_R^{(S)}$ have storage cost n , write cost n , and read cost $2n$ (due to write back). For $RADON_C$, each server stores at most $\delta + 1$ coded-elements, where each element has size $\frac{1}{k}$. Thus storage cost of $RADON_C$ is $(\delta + 1)\frac{n}{k}$. The write cost of $RADON_C$ is simply $\frac{n}{k}$, and the contribution comes from the writer sending one coded-element to each of the n servers. For a read, getting the entire *Lists* during the *get – data* phase incurs a cost of $(\delta + 1)\frac{n}{k}$. The write-back phase incurs an additional cost of $\frac{n}{k}$. Thus, the total read cost in $RADON_C$ is $(\delta + 2)\frac{n}{k}$.

8 Conclusions

In this paper, we provided an erasure-code-based algorithm for implementing atomic memory, having the ability to perform repair of crashed nodes in the background, without affecting client operations. We assumed a static model with a fixed, finite set of nodes, and also a practical network condition *N1* to facilitate repair. We showed how the usage of MDS codes significantly improve storage and communication costs over a replication based solution, when the number of writes concurrent with a read or repair is limited. Liveness and atomicity are guaranteed as long as *N1* is satisfied; however violation of *N1* can lead to non-atomic executions. We further showed how a slightly stringent network condition *N2* can be used to construct a replication based algorithm that always guarantees atomicity. Ongoing efforts include exploring possibility of using repair-efficient erasure codes [19] in $RADON_C$, and testbed evaluations on cloud based infrastructure.

9 Acknowledgments

The work is supported in part by AFOSR under grants FA9550-14-1-043, FA9550-14-1-0403, and in part by NSF under awards CCF-1217506, CCF-0939370.

References

1. H. Attiya, A. Bar-Noy, and D. Dolev, “Sharing memory robustly in message passing systems,” *Journal of the ACM*, vol. 42(1), pp. 124–142, 1996.
2. R. Fan and N. Lynch, “Efficient replication of large data objects,” in *Distributed algorithms*, ser. Lecture Notes in Computer Science, 2003, pp. 75–91.
3. A. Spiegelman, Y. Cassuto, G. Chockler, and I. Keidar, “Space Bounds for Reliable Storage: Fundamental Limits of Coding,” in *Proceedings of the International Conference on Principles of Distributed Systems (OPODIS2015)*, 2015.
4. M. K. Aguilera, R. Janakiraman, and L. Xu, “Using erasure codes efficiently for storage in a distributed system,” in *Proceedings of International Conference on Dependable Systems and Networks (DSN)*, 2005, pp. 336–345.
5. P. Dutta, R. Guerraoui, and R. R. Levy, “Optimistic erasure-coded distributed storage,” in *Proceedings of the 22nd international symposium on Distributed Computing (DISC)*, Berlin, Heidelberg, 2008, pp. 182–196.
6. V. R. Cadambe, N. A. Lynch, M. Médard, and P. M. Musial, “A coded shared atomic memory algorithm for message passing architectures,” in *Proceedings of 13th IEEE International Symposium on Network Computing and Applications (NCA)*, 2014, pp. 253–260.

7. K. M. Konwar, N. Prakash, E. Kantor, N. Lynch, M. Medard, and A. A. Schwarzmann, "Storage-optimized data-atomic algorithms for handling erasures 124 and errors in distributed storage systems," in *30th IEEE International Parallel & Distributed Processing Symposium (IPDPS)*, 2016.
8. R. Guerraoui, R. R. Levy, B. Pochon, and J. Pugh, "The collective memory of amnesic processes," *ACM Trans. Algorithms*, vol. 4, no. 1, pp. 1–31, 2008.
9. F. Chang, J. Dean, S. Ghemawat, W. C. Hsieh, D. A. Wallach, M. Burrows, T. Chandra, A. Fikes, and R. E. Gruber, "Bigtable: A distributed storage system for structured data," *ACM Trans. Comput. Syst.*, vol. 26, no. 2, pp. 4:1–4:26, jun 2008.
10. G. DeCandia, D. Hastorun, M. Jampani, G. Kakulapati, A. Lakshman, A. Pilchin, S. Sivasubramanian, P. Vosshall, and W. Vogels, "Dynamo: Amazon's highly available key-value store," in *Proceedings of Twenty-first ACM SIGOPS Symposium on Operating Systems Principles*, ser. SOSP '07. New York, NY, USA: ACM, 2007, pp. 205–220.
11. N. Lynch and A. A. Shvartsman, "RAMBO: A reconfigurable atomic memory service for dynamic networks," in *Proceedings of 16th International Symposium on Distributed Computing (DISC)*, 2002, pp. 173–190.
12. M. K. Aguilera, I. Keidar, D. Malkhi, and A. Shraer, "Dynamic atomic storage without consensus," *Journal of the ACM*, pp. 7:1–7:32, 2011.
13. A. Spiegelman and I. Keidar, "On liveness of dynamic storage," *CoRR*, vol. abs/1507.07086, 2015. [Online]. Available: <http://arxiv.org/abs/1507.07086>
14. R. Baldoni, S. Bonomi, A. M. Kermarrec, and M. Raynal, "Implementing a register in a dynamic distributed system," in *Distributed Computing Systems, 2009. ICDCS '09. 29th IEEE International Conference on*, June 2009, pp. 639–647.
15. H. Attiya, H. C. Chung, F. Ellen, S. Kumar, and J. L. Welch, "Simulating a shared register in an asynchronous system that never stops changing - (extended abstract)," in *Distributed Computing - 29th International Symposium, DISC 2015, Tokyo, Japan, October 7-9, 2015, Proceedings*, 2015, pp. 75–91.
16. M. K. Aguilera, I. Keidar, D. Malkhi, J. P. Martin, and A. Shraery, "Reconfiguring replicated atomic storage: A tutorial," *Bulletin of the EATCS*, vol. 102, pp. 84–081, 2010.
17. A. Spiegelman, I. Keidar, and D. Malkhi, "Dynamic reconfiguration: A tutorial," *OPODIS 2015*, 2015.
18. P. Musial, N. Nicolaou, and A. A. Shvartsman, "Implementing distributed shared memory for dynamic networks," *Communications of the ACM*, vol. 57, no. 6, pp. 88–98, 2014.
19. A. G. Dimakis, K. Ramchandran, Y. Wu, and C. Suh, "A survey on network codes for distributed storage," *Proceedings of the IEEE*, vol. 99, no. 3, pp. 476–489, 2011.
20. C. Huang, H. Simitci, Y. Xu, A. Ogus, B. Calder, P. Gopalan, J. Li, and S. Yekhanin, "Erasure coding in windows azure storage," in *Proc. USENIX Annual Technical Conference (ATC)*, 2012, pp. 15–26.
21. M. Sathiamoorthy, M. Asteris, D. Papailiopoulos, A. G. Dimakis, R. Vadali, S. Chen, and D. Borthakur, "XORing elephants: novel erasure codes for big data," in *Proceedings of the 39th international conference on Very Large Data Bases*, 2013, pp. 325–336.
22. K. V. Rashmi, P. Nakkiran, J. Wang, N. B. Shah, and K. Ramchandran, "Having your cake and eating it too: Jointly optimal erasure codes for i/o, storage, and network-bandwidth," in *13th USENIX Conference on File and Storage Technologies (FAST)*, 2015, pp. 81–94.
23. I. S. Reed and G. Solomon, "Polynomial codes over certain finite fields," *Journal of the society for industrial and applied mathematics*, vol. 8, no. 2, pp. 300–304, 1960.
24. C. Cachin and S. Tessaro, "Optimal resilience for erasure-coded byzantine distributed storage," in *Proceedings of International Conference on Dependable Systems and Networks (DSN)*, 2006, pp. 115–124.
25. D. Dobre, G. Karamé, W. Li, M. Majuntke, N. Suri, and M. Vukolić, "Powerstore: proofs of writing for efficient and robust storage," in *Proceedings of the 2013 ACM SIGSAC conference on Computer & communications security*, 2013, pp. 285–298.
26. J. Hendricks, G. R. Ganger, and M. K. Reiter, "Low-overhead byzantine fault-tolerant storage," *ACM SIGOPS Operating Systems Review*, vol. 41, no. 6, pp. 73–86, 2007.
27. L. Lamport, "On interprocess communication," *Distributed computing*, vol. 1, no. 2, pp. 86–101, 1986.
28. C. Shao, J. L. Welch, E. Pierce, and H. Lee, "Multiwriter consistency conditions for shared memory registers," *SIAM Journal on Computing*, vol. 40, no. 1, pp. 28–62, 2011.
29. K. M. Konwar, N. Prakash, M. Medard, and N. Lynch, "RADON: Repairable atomic data object in networks," *CoRR*, vol. abs/1605.05717, 2016.
30. W. C. Huffman and V. Pless, *Fundamentals of error-correcting codes*. Cambridge university press, 2003.
31. N. A. Lynch, *Distributed Algorithms*. Morgan Kaufmann Publishers, 1996.

Appendix

A Proof of Theorem 1

The theorem is restated for convenience.

Theorem 8. (Theorem 1) *It is impossible to implement an atomic memory service that guarantees liveness of reads and writes, under the system model described in Section 2, even if 1) there is at most one server in the crashed/repair state at any point during the execution, and 2) every repair operation completes, and takes the repaired server back to the active state.*

Proof. We prove this result by contradiction, by assuming an algorithm A_{alg} that guarantees liveness and atomicity, and also is such that every repair operation completes, and takes the repaired server back to the active state. Let the initial value stored in the system be $v_0 \in V$, where V is the domain of all values. Consider a non-faulty writer w , and suppose w initiates a write operation π^w with the value v_1 , such that $v_1 \neq v_0$. Let $\mathcal{S}_w \subseteq \mathcal{S}$ be the set of all servers that writer w sends messages before w expects any response from any of the servers in $\mathcal{S}$. Without loss of generality² let $\mathcal{S}_w = \{s_1, s_2, \dots, s_k\}$, for some $k \leq n$, and let m_i denote the message sent by w to server s_i , $1 \leq i \leq k$. Note that if w sends two or messages to a particular server, say s_1 , all these can be combined into m_1 , since all these messages are sent without expecting any response.

Consider an execution which starts with all the servers in the active state, the operation π^w begins, messages get sent out to servers in $\mathcal{S}_w$. Delay the messages such that message m_1 arrives at server s_1 before any other server in $\mathcal{S}_w$ receives the respective message. Assume that s_1 is in the crashed state when m_1 arrives, so s_1 does not receive m_1 . Further assume that all the other servers are in the active state at this point of execution. Let server s_1 undergo a successful repair operation, before any other server in $\mathcal{S}_w$ receives its respective message. Next, consider the case when server s_2 receives the message m_2 , and delay the messages to all other servers, and assume that s_2 is in the crashed state when m_2 arrives. The sequence of crash and repair can be carried out in this manner one-by-one for every server in $\mathcal{S}_w$, where all these servers end up losing the writer message, though they get repaired. Now if the algorithm is such that the writer expects a response from any of the servers in $\mathcal{S}$, clearly it will not happen, since no server in $\mathcal{S}$ has received any message from w while the server is in the active state. Thus liveness of write is compromised.

We next consider the case when the writer decides to terminate without expecting any response from any server in $\mathcal{S}$, and show that such a method of guaranteeing liveness results in violation of atomicity. Let us call this execution fragment (as discussed above) with such a write as $\beta^w(v_1)$. After the write π^w completes, a read π^r associated with a non-faulty reader, begins. By liveness of read, and atomicity, the read must return v_1 . Let the execution fragment associated with the read be denoted as β^r , so that the overall execution fragment under consideration is $\beta^w(v_1) \circ \beta^r$. Next, consider the execution fragment $\beta^w(v'_1)$ obtained by replacing v_1 with v'_1 such that $v'_1 \neq v_1$. Since a crash causes a server to lose its entire state, it is clear that to the reader r there is no distinction between the state of the system after $\beta^w(v_1)$, and the state of the system after $\beta^w(v'_1)$. In this case, if we consider the execution $\beta^w(v'_1) \circ \beta^r$, the read returns v_1 ($\neq v'_1$), since in the execution $\beta^w(v_1) \circ \beta^r$ also, r returned v_1 . However it violates atomicity of $\beta^w(v'_1) \circ \beta^r$, which completes the proof.

² Clearly, the writer must send a message to at least one server, so we ignore the trivial case when $\mathcal{S}_w$ is empty.

B Proof of Lemma 1

The lemma is restated for convenience.

Lemma 3 (Lemma 1). *Let β denote a well-formed execution of the RADON_R algorithm. Suppose T denotes a point of time in the execution β such that there exists a majority of servers $\mathcal{S}_m$, $\mathcal{S}_m \subset \mathcal{S}$ all of which are in the active state at the time T . Also, let t_s denote the value of the local tag at server s , at time T . Then, if π denotes any completed repair or read operation that is initiated after time T , we have $\text{tag}(\pi) \geq \min_{s \in \mathcal{S}_m} t_s$. Also, if π denotes any completed write operation that is initiated after time T , then we have $\text{tag}(\pi) > \min_{s \in \mathcal{S}_m} t_s$.*

Proof. We use ρ to denote $\min_{s \in \mathcal{S}_m} t_s$. Also, for any state variable $x(s)$ that is stored in server s , we write $x(s)|_T$ to denote the value x at time T . Below, we separately consider the cases when π denotes a successful repair, read and write operations, in this respective order.

(a) π is a successful repair operation: We prove the statement by contradiction, by starting with the assumption that $\text{tag}(\pi) < \rho$. Let T_π denote the point of time in the execution β at which the operation π completes. Let Π'_R denote the set of all successful repair operations which start after the time T , but start before T_π , and is such that $\forall \pi' \in \Pi'_R$, we have $\text{tag}(\pi') < \rho$. Clearly, $\pi \in \Pi'_R$. Let $\pi^* \in \Pi'_R$ denote the repair operation, which completes first. Note that π^* exists since the set Π'_R is finite. Now, let $\hat{\mathcal{S}}$ denote the set of majority servers based on whose responses the operation π^* completed. Clearly, $|\hat{\mathcal{S}} \cap \mathcal{S}_m| \geq 1$. For any server $s \in \hat{\mathcal{S}} \cap \mathcal{S}_m$, let T_s denote the point of time in the execution at which the server s responds to π^* with its local (tag, value) pair. Clearly, the server s must have remained in the active state during the entire interval $[T, T_s]$. This follows because s is active at time T , π^* is the first completed repair operation that started after T , and due to the fact that a server responds to a repair request only if it is in the active state. In this case, we know that³ $t_{loc}(s)|_{T_s} \geq t_{loc}(s)|_T \geq \rho$ for any s in $\hat{\mathcal{S}} \cap \mathcal{S}_m$. Therefore, we have $\text{tag}(\pi^*) = \max_{s \in \hat{\mathcal{S}}} t_{loc}(s)|_{T_s} \geq \max_{s \in \hat{\mathcal{S}} \cap \mathcal{S}_m} t_{loc}(s)|_{T_s} \geq \rho$, which contradicts the existence of $\pi^* \in \Pi'_R$. From, this we conclude that the set Π'_R must be empty to avoid contradictions, and hence $\text{tag}(\pi) \geq \rho$.

(b) π is a successful read operation: We prove this by contradiction by starting with the assumption that $\text{tag}(\pi) < \rho$. Let $\hat{\mathcal{S}}$ denote the set of majority servers based on whose responses during the *get-data* phase (see Fig. 1), the read operation completed. As in Part a), we know that $|\hat{\mathcal{S}} \cap \mathcal{S}_m| \geq 1$. In this case, let T_s denote the point of time during the execution at which the server $s \in \hat{\mathcal{S}} \cap \mathcal{S}_m$ responded to the reader. Next, note that in the *get-data* phase, the reader picks the response with the highest tag. Thus, since we assume that $\text{tag}(\pi) < \rho$, it must be true that $t_{loc}(s)|_{T_s} < \rho$, $s \in \hat{\mathcal{S}} \cap \mathcal{S}_m$. Since the server $s \in \hat{\mathcal{S}} \cap \mathcal{S}_m$ is active at time T such that $t_{loc}(s)|_T \geq \rho$, this would imply that server s experienced a crash event after time T , and came back to the *active* state before the time T_s via a successful repair operation ϕ such that $\text{tag}(\phi) < \rho$. But then, this contradicts Part a) of the theorem which we proved above, and hence we conclude that $\text{tag}(\pi) \geq \rho$.

(c) π is a write operation: Once again we prove via contradiction, by starting with the assumption that $\text{tag}(\pi) \leq \rho$. Let $\hat{\mathcal{S}}$ denote the set of majority servers based on whose responses during the *get-tag* phase, the writer determined $\text{tag}(\pi)$. We know from the algorithm that $\text{tag}(\pi)$ is strictly larger than all the tags among the responses from $\hat{\mathcal{S}}$. Since $|\hat{\mathcal{S}} \cap \mathcal{S}_m| \geq 1$, we argue like in Part b), and arrive at a contradiction to Part a).

³ Any read or write operation cannot decrease the local tag that is stored in an active server.

C Proof of Theorem 3

The theorem is restated here for convenience.

Theorem 9 (Theorem 3). *Every execution of the $RADON_R$ algorithm operating under the $N1$ network stability condition with $\alpha > \frac{3}{4}$, is atomic.*

C.1 Some Preliminaries

Partial Order on read and write operations Consider any well-formed execution β of $RADON_R$, all of whose invoked read or write operations complete. Let Π_{RW} denote the set of all completed read and write operations in β . We first define a partial order ($\prec$) on Π_{RW} . Towards this, recall that for any completed write operation π , we defined $tag(\pi)$ as the tag created by the writer during the *write-put* phase. Also, recall that for any completed read operation π , we define $tag(\pi)$ as the tag corresponding to the value returned by the read. The partial order ($\prec$) in Π_{RW} is defined as follows: For any $\pi, \phi \in \Pi_{RW}$, we say $\pi \prec \phi$ if one of the following holds: (i) $tag(\pi) < tag(\phi)$, or (ii) $tag(\pi) = tag(\phi)$, and π and ϕ are write and read operations, respectively. The proof of atomicity is based on the following lemma, which is simply a restatement of the sufficiency condition for atomicity presented in [31].

Lemma 4. *Consider any well-formed execution β of the algorithm, such that all the invoked read and the write operations complete. Now, suppose that all the invoked read and write operations in β can be partially ordered by an ordering $\prec$, so that the following properties are satisfied:*

- P1. *The partial order ($\prec$) is consistent with the external order of invocation and responses, i.e., there are no operations π_1 and π_2 , such that π_1 completes before π_2 starts, yet $\pi_2 \prec \pi_1$.*
- P2. *All operations are totally ordered with respect to the write operations, i.e., if π_1 is a write operation and π_2 is any other operation then either $\pi_1 \prec \pi_2$ or $\pi_2 \prec \pi_1$.*
- P3. *Every read operation returns the value of the last write preceding it (with respect to $\prec$), and if no preceding write is ordered before it, then the read returns the initial value of the object.*

Then, the execution β is atomic.

C.2 Proof of Atomicity under $N1$ with $\alpha > 3/4$

We need to prove the properties $P1$, $P2$ and $P3$ of Lemma 4. We do this under $N1$ with $\alpha > 3/4$, using Lemma 1. Let ϕ and π denote two operations in Π_{RW} such that ϕ completed before π started. Also, let c_ϕ and c_π denote the clients that initiated the operations ϕ and π , respectively.

Property P1 We want to show that $\pi \not\prec \phi$. We show this in detail only for the case when ϕ and π are both write operations. The proofs of other three cases⁴ are similar, and hence omitted. By virtue of the definition of the partial order $\prec$, it is enough to prove that $tag(\pi) > tag(\phi)$. Consider the *put-data* phase of ϕ , where the writer sends the pair (t_w, v) to all servers via the *group-send* operation. Under the condition $N1$ with $\alpha > 3/4$, we know that there exists a set $\mathcal{S}_\alpha \subseteq \mathcal{S}$ of $|\mathcal{S}_\alpha| \geq \lceil \frac{3n+1}{4} \rceil$ servers all of which remain in the active state during the interval $[T_1, T_2]$ where T_1

⁴ These correspond to the case when ϕ and π are both read operations, and the cases where one of them is a write and the other is a read.

denotes the point of time of invocation of the group-send operation, and T_2 denotes the earliest point of time during the execution where all of the servers in $\mathcal{S}_\alpha$ complete effective consumption (including sending ack to the writer c_ϕ) of the message (t_w, v) . Also, let $\mathcal{S}' \subseteq \mathcal{S}$ denote the set of $\lceil \frac{3n+1}{4} \rceil$ servers whose acks are used by the writer to decide the completion of the write operation. Clearly, $|\mathcal{S}' \cap \mathcal{S}_\alpha| > \frac{n}{2}$. Let T denote the earliest point of time during the execution when all servers in $\mathcal{S}' \cap \mathcal{S}_\alpha$ complete their respective effective consumption of the message (t_w, v) . In this case note that a) T occurs before the point of completion of the write operation, b) all servers in $\mathcal{S}' \cap \mathcal{S}_\alpha$ are in the active state at T , and c) $t_{loc}(s)|_T \geq \text{tag}(\phi), \forall s \in \mathcal{S}' \cap \mathcal{S}_\alpha$. We now apply Lemma 1 to conclude that $\text{tag}(\pi) > \text{tag}(\phi)$.

Property P2 This follows from the construction of tags, and the definition of the partial order ($\prec$).

Property P3 This follows from the definition of partial order ($\prec$), and by noting that value returned by a read operation π is simply the value associated with $\text{tag}(\pi)$.

D Proof of Lemma 2

The lemma is restated below for easy reference.

Lemma 5 (Lemma 2). *Consider any well-formed execution β of RADON_C operating under the network stability condition N1 with $\alpha \geq \frac{3n+k}{4n}$. Further assume that the number of writes concurrent with any valid read or repair operation is at most δ . For any operation π , consider the set $\Sigma = \{\sigma : \sigma \text{ is a read or a write in } \beta \text{ that completes before } \pi \text{ begins}\}$, and also let $\sigma^* = \arg \max_{\sigma \in \Sigma} \text{tag}(\sigma)$. Then, the following statements hold:*

- If π denotes a completed repair operation on a server $s \in \mathcal{S}$, then the repaired List of server s due to π contains the pair $(\text{tag}(\sigma^*), c_s^*)$.
- If π denotes a read operation associated with a non-faulty reader r , and further, if $\mathcal{S}_1$ denotes the set of $\lceil \frac{n+k}{2} \rceil$ servers whose responses, say $\{L_\pi(s), s \in \mathcal{S}_1\}$, are used by r to attempt decoding of a value in the get-data phase, then there exists $\mathcal{S}_2 \subseteq \mathcal{S}_1$, $|\mathcal{S}_2| = k$, such that $\forall s \in \mathcal{S}_2, (\text{tag}(\sigma^*), c_s^*) \in L_\pi(s)$.
- If π denotes a write operation associated with a non-faulty writer w , and further if $\mathcal{S}_1$ denotes the set of majority servers whose responses are used by w to compute max-tag in the get-tag phase, then there exists a server $s \in \mathcal{S}_1$, whose response tag $t_s \geq \text{tag}(\sigma^*)$.

Here, c_s^* denotes the coded-element of server s corresponding to the value v^* , associated with $\text{tag}(\sigma^*)$.

We prove the lemma separately for the cases of repair and read (Parts 1 and 2). The proof for the third part for the case of write operations is similar to that of Part 2, and hence omitted.

Proof of Part 1 of Lemma 5 Consider the set Σ and the operation σ^* as defined in the statement of Lemma 5. Without loss of generality, let us assume that σ^* is a write operation. Since we assume condition N1 with $\alpha \geq \frac{3n+k}{4n}$, there exists a set $\mathcal{S}_\alpha$ of $\lceil \frac{3n+k}{n} \rceil$ servers that respects N1 for the group-send operation (say gp^*) in the put-data phase of σ^* . If $\mathcal{S}_1$ denotes the set of $\lceil \frac{3n+k}{n} \rceil$ servers, whose responses are used by the writer to decide termination, we then know that 1) $|\mathcal{S}_\alpha \cap \mathcal{S}_1| \geq \lceil \frac{n+k}{2} \rceil$, and 2) if T_{prop} denotes the earliest point of time during the execution when all the servers in

$\mathcal{S}_{prop} = \mathcal{S}_\alpha \cap \mathcal{S}_1$ complete effective consumption of their respective messages from the group-send operation gp^* , then every server in $\mathcal{S}_{prop}$ remains active at T_{prop} , and has not experienced a crash after its effective consumption, until T_{prop} . Our goal is to show that the repair operation π always receives at least k responses from among the servers in $\mathcal{S}_{prop}$, and must be able to decode (and then re-encode) the value corresponding to $\text{tag}(\sigma^*)$. Below we consider the effects of concurrent writes having higher tags, and repairs before π starts, both of which can potentially remove coded elements corresponding to $\text{tag}(\sigma^*)$, from lists of various servers. We show under the assumptions of the lemma, that neither of these cause a problem.

Let us first consider the effect of concurrent writes. Towards this, consider the set Λ of writes concurrent with the valid repair operation π (see Definition 3). Recall that $\Lambda = \{\lambda : \lambda \text{ is a write operation that starts before } \pi, \text{ such that } \text{tag}(\lambda) > \text{tag}(\sigma^*)\}$. By assumption on the lemma, we know that $|\Lambda| \leq \delta$. In this case, it is clear that if a server $s \in \mathcal{S}_{prop}$ does not crash in the interval $[T_{prop}, T]$, the $\text{List}(s)|_T$ contains the pair corresponding to $\text{tag}(\sigma^*)$, for any T such that $T_{prop} \leq T \leq T_{end}(\pi)$. Here $T_{end}(\pi)$ denotes the point of completion of π .

Let us next consider the effect of repairs, let $\tilde{\Pi} = \{\tilde{\pi} : \text{a repair which start after } T_{prop}, \text{ but also start before the completion of } \pi\}$. Clearly, $\pi \in \tilde{\Pi}$. Let $\tilde{\pi}^* \in \tilde{\Pi}$ denote the repair operation that completes first. Clearly, it must be true that $T_{prop} < T_{end}(\tilde{\pi}^*) \leq T_{end}(\pi)$. We prove Part 1 of the Lemma 5 for $\tilde{\pi}^*$ first. Using this result, we prove the lemma for the repair operation in $\tilde{\Pi}$ which completes second. We continue in an inductive manner (on the finite set $\tilde{\Pi}$), until we hit π . Towards proving the lemma for $\tilde{\pi}^*$, consider the group-send operation, where $\tilde{\pi}^*$ requests for local *Lists* from all servers. Let $\mathcal{S}_\theta \subset \mathcal{S}_{prop}$ denote the servers among $\mathcal{S}_{prop}$ which are not in the active state when the repair request arrives. Also, let $\mathcal{S}_a \subset \mathcal{S}$ denote the set of all servers which are in the active state when the repair request arrives. Clearly, $|\mathcal{S}_a| \leq n - |\mathcal{S}_\theta|$. Next, let $\mathcal{S}_{ack} \subset \mathcal{S}_a$ denote the set of $\lceil \frac{n+k}{2} \rceil$ servers based on whose responses the repair operation $\tilde{\pi}^*$ completes. Now, since $\mathcal{S}_{prop} \setminus \mathcal{S}_\theta \subset \mathcal{S}_a$, we have

$$(\mathcal{S}_{prop} \setminus \mathcal{S}_\theta) \cup \mathcal{S}_{ack} \subset \mathcal{S}_a \quad (1)$$

$$\implies |\mathcal{S}_{prop} \setminus \mathcal{S}_\theta| + |\mathcal{S}_{ack}| - |(\mathcal{S}_{prop} \setminus \mathcal{S}_\theta) \cap \mathcal{S}_{ack}| \leq |\mathcal{S}_a| \quad (2)$$

$$\implies |\mathcal{S}_{prop}| - |\mathcal{S}_\theta| + \lceil \frac{n+k}{2} \rceil - |(\mathcal{S}_{prop} \setminus \mathcal{S}_\theta) \cap \mathcal{S}_{ack}| \leq n - |\mathcal{S}_\theta| \quad (3)$$

$$\implies |(\mathcal{S}_{prop} \setminus \mathcal{S}_\theta) \cap \mathcal{S}_{ack}| \geq k, \quad (4)$$

where the last inequality follows from our earlier observation that $|\mathcal{S}_{prop}| \geq \lceil \frac{n+k}{2} \rceil$. Next, note that any server s in $(\mathcal{S}_{prop} \setminus \mathcal{S}_\theta) \cap \mathcal{S}_{ack}$ remains active from T_{prop} until the point when s responds to the repair request from $\tilde{\pi}^*$. This follows because of the facts that 1) s is active at T_{prop} , 2) a server responds to a repair request only if it is in the active state, and 3) since $\tilde{\pi}^*$ is the first repair operation that completes after T_{prop} . Also, recall our earlier observations that 1) if a server $s \in \mathcal{S}_{prop}$ does not crash in the interval $[T_{prop}, T]$, then $\text{List}(s)|_T$ contains the pair corresponding to $\text{tag}(\sigma^*)$, for any T such that $T_{prop} \leq T \leq T_{end}(\pi)$, and 2) $T_{prop} < T_{start}(\tilde{\pi}^*) < T_{end}(\tilde{\pi}^*) \leq T_{end}(\pi)$. In this case, we know that the responses of all the servers in $(\mathcal{S}_{prop} \setminus \mathcal{S}_\theta) \cap \mathcal{S}_{ack}$ to $\tilde{\pi}^*$, contain the pair corresponding to $\text{tag}(\sigma^*)$. From (4), it follows that the repaired list for $\tilde{\pi}^*$, before pruning to $(\delta + 1)$ entries, contains the pair corresponding to $\text{tag}(\sigma^*)$. Finally the fact that $\text{tag}(\sigma^*)$ is among the highest $\delta + 1$ tags, and hence part of the pruned list, follows from our earlier observations⁵ that

⁵ Note that Λ need not be the set of writes concurrent with $\tilde{\pi}^*$. The above argument where we say that the pruned list, after the repair $\tilde{\pi}^*$, is of size at most $\delta + 1$ can be argued entirely based on Λ itself.

1) $|A| \leq \delta$, and 2) $T_{end}(\tilde{\pi}^*) \leq T_{end}(\pi)$. This completes our proof of Part 1 of Lemma 5 for the repair operation $\tilde{\pi}^*$.

We next prove the lemma for the repair operation $\pi_2 \in \tilde{\Pi}$, which completes second. The proof is mostly identical, and we will only highlight the place where we use the result on $\tilde{\pi}^*$. Clearly, since we carry out the induction only until we hit π , it must be true that $T_{prop} < T_{end}(\pi_2) \leq T_{end}(\pi)$. Consider the group-send operation, where π_2 requests for local *Lists* from all servers. Let $\mathcal{S}_\theta^{(2)} \subset \mathcal{S}_{prop}$ denote the servers among $\mathcal{S}_{prop}$ which are not in the active state when the repair request arrives. Also, let $\mathcal{S}_a^{(2)} \subset \mathcal{S}$ denote the set of all servers which are in the active state when the repair request arrives. As before, $|\mathcal{S}_a^{(2)}| \leq n - |\mathcal{S}_\theta^{(2)}|$. Next, let $\mathcal{S}_{ack}^{(2)} \subset \mathcal{S}_a^{(2)}$ denote the set of $\lceil \frac{n+k}{2} \rceil$ servers based on whose responses the repair operation π_2 completes. Along the lines of (1)-(4), one can show that $|(\mathcal{S}_{prop} \setminus \mathcal{S}_\theta^{(2)}) \cap \mathcal{S}_{ack}^{(2)}| \geq k$. Next, if we consider the set $(\mathcal{S}_{prop} \setminus \mathcal{S}_\theta^{(2)}) \cap \mathcal{S}_{ack}^{(2)}$, at most one of the servers in this set would have undergone a crash after the time T_{prop} , and got repaired before the time the server responded to π_2 . Note that more than one repair operation on $(\mathcal{S}_{prop} \setminus \mathcal{S}_\theta^{(2)}) \cap \mathcal{S}_{ack}^{(2)}$ cannot happen, since this will contradict the assumption that π_2 is the second repair operation to complete after T_{prop} . Further, if one repair operation among a server in $(\mathcal{S}_{prop} \setminus \mathcal{S}_\theta^{(2)}) \cap \mathcal{S}_{ack}^{(2)}$ has indeed occurred, this must be the operation $\tilde{\pi}^*$ which we considered above. Further, we know that the repaired *List* due to $\tilde{\pi}^*$ contains the pair corresponding to $tag(\sigma^*)$. In other words, irrespective of whether one repair operation occurred among the servers in $(\mathcal{S}_{prop} \setminus \mathcal{S}_\theta^{(2)}) \cap \mathcal{S}_{ack}^{(2)}$, or not, the responses of all the servers in $(\mathcal{S}_{prop} \setminus \mathcal{S}_\theta^{(2)}) \cap \mathcal{S}_{ack}^{(2)}$ contain the pair corresponding to $tag(\sigma^*)$. The rest of the proof is similar to that of $\tilde{\pi}^*$, where we argue that the pruned list after repair contains the pair corresponding to $tag(\sigma^*)$. The rest of the induction is similar, and this completes the proof of Part 1 of Lemma 5.

Proof of Part 2 of Lemma 5 The proof follows mostly along the lines of proof of Part 1 of the lemma. We will only highlight the main steps here. Consider the set Σ and the operation σ^* as defined in the statement of the lemma. Without loss of generality, let us assume that σ^* is a write operation. Also, define the time T_{prop} and the set $\mathcal{S}_{prop}$ exactly in the same way as what we defined in the proof of Part 1 of the lemma. Let T_1 denote the earliest point of time during the execution when the reader receives responses from all the servers in $\mathcal{S}_1$, where $\mathcal{S}_1$ is as defined in the statement of this lemma. Consider the set of writes concurrent with the valid read operation π . Recall from Definition 3 that $A = \{\lambda : \lambda \text{ is a write operation which starts before time } T_1 \text{ such that } tag(\lambda) > tag(\sigma^*)\}$. From the assumption on concurrency in the lemma statement, we know that $|A| \leq \delta$. In this case, it is clear (like in the proof of Part 1 above) that if a server $s \in \mathcal{S}_{prop}$ does not crash in the interval $[T_{prop}, T_1]$, the $List(s)|_T$ contains the pair corresponding to $tag(\sigma^*)$, for any T such that $T_{prop} \leq T \leq T_1$. Now, if server $s \in \mathcal{S}_{prop}$ undergoes a crash and repair operation (say ρ) during the interval $[T_{prop}, T]$ (so that it is active again at T), we can argue exactly like in the proof of Part 1 above, and show that the repaired *List* due to ρ contains the pair corresponding to $tag(\sigma^*)$. This can be done by considering the set $\tilde{\Pi} = \{\tilde{\pi} : \text{a repair which start after } T_{prop}, \text{ but also start before } T_1\}$, and applying induction on $\tilde{\Pi}$ based on the order of completion times of the repair operations. This completes the proof of our claim about $List(s)|_T$.

The rest of the proof follows simply by noting $|\mathcal{S}_1 \cap \mathcal{S}_{prop}| \geq k$, and thus the value corresponding to $tag(\sigma^*)$ is surely a candidate for decoding, since we know that an $[n, k]$ linear MDS code can be uniquely decoded given any k out of the n coded-elements.

E Liveness: Proof of Theorem 4

The theorem is restated below for easy reference:

Theorem 10. (Theorem 4) *Let β denote a well-formed execution of $RADON_C$, operating under the N1 network stability condition with $\alpha \geq \frac{3n+k}{4n}$ and δ be the maximum number of write operations concurrent with any valid read or repair operation. Then every operation initiated by a non-faulty client completes.*

Proof. Liveness of writes depends only on sufficient number of responses in the two phases. The maximum number of responses expected in either of the two phases is $\frac{3n+k}{4n}$, which we know is guaranteed under N1 with $\alpha \geq \frac{3n+k}{4n}$. Liveness of reads follows by combining Lemma 2 (for decodability of a value), and the liveness of write operations (for the write-back phase).

F Atomicity: Proof of Theorem 5

The theorem is restated first:

Theorem 11. (Theorem 5) *Any execution of $RADON_C$, operating under condition N1 with $\alpha \geq \frac{3n+k}{4n}$, is atomic, if the maximum number of write operations concurrent with a valid read or repair operation is δ .*

Proof. The proof is based on Lemmas 4 and 2. In order to apply Lemma 4, consider any well-formed execution β of $RADON_C$, all of whose invoked read and write operations, denoted by the set Π_{RW} , complete. We define a partial order ($\prec$) on Π_{RW} like in the proof of Theorem 3 for case of $RADON_R$. To prove Property P1 of Lemma 4, consider two successful operations ϕ and π such that ϕ completes before π begins. Firstly, consider the case π is a write, and ϕ is either a read or write. We need to show that $tag(\pi) > tag(\phi)$, which we note follows directly from Part 3 of Lemma 2. Next, consider the case when consider the case π is a read, and ϕ is either a read or write. We need to show that $tag(\pi) \geq tag(\phi)$, which we note follows directly from Part 2 of Lemma 2. This completes the proof of Property P1. Proofs of Properties P2 and P3 are similar to those of the corresponding properties in Theorem 3, where we proved atomicity of $RADON_R$.

G Proof of Theorem 6

The theorem is restated for convenience.

Theorem 12. (Theorem 6) *Let β denote a well-formed execution of $RADON_R^{(S)}$ operating under condition N2 with $\alpha > \frac{3}{4}$. Then every operation initiated by a non-faulty client completes.*

We will prove that a write operation associated with a non-faulty client always completes, the proof for a read is similar and hence is omitted. The main step is to show the completion of the *confirm-data* phase. Consider the *put-data* phase, and note that under N2 with $\alpha > \frac{3}{4}$, we are guaranteed that there exists of set of $\mathcal{S}_\alpha \subset \mathcal{S}$ servers, such that 1) $|\mathcal{S}_\alpha| \geq \lceil \frac{3n+1}{4} \rceil$, and 2) every server in $\mathcal{S}_\alpha$ remains active from the point of time T_1 of initiation of the group-send operation of *put-data* phase till the point of time T'_1 , when the writer effectively consumes all responses (acks) from the servers in $\mathcal{S}_\alpha$. Next, let $\mathcal{S}_1 \subset \mathcal{S}$ denote the set of $\lceil \frac{3n+1}{4} \rceil$ whose acks are received by the

writer before moving on to the *confirm-data* phase. First of all note that the existence of the set $\mathcal{S}_1$ is clearly guaranteed under $N2$ with $\alpha > \frac{3}{4}$ (since the set $\mathcal{S}_\alpha$ is a candidate for $\mathcal{S}_1$). Secondly, we note that the group-send operation in the *confirm-data* phase forms part of the effective consumption of the last ack that is received from the servers in $\mathcal{S}_1$. This follows from the definition of effective-consumption, and by noting the execution of the group-send operation in the *confirm-data* phase does not depend on any more input after all the acks in the *put-data* phase are received. Let T_2 denote the point of time at which the group-send operation in the *confirm-data* phase gets initiated. Note that $T'_1 \geq T_2$, in fact if $\mathcal{S}_1 \neq \mathcal{S}_\alpha$, we have⁶ $T'_1 > T_2$. Next we apply the network condition to the group-send operation in the *confirm-data* phase. From the $N1$ part of $N2$, we know that there exists a $\mathcal{S}'_\alpha$ of $\lceil \frac{3n+1}{4} \rceil$ servers, all of which receive and effectively consume the message from the group-send operation, and remain active from T_2 till the point of time T'_2 when the last of the servers in $\mathcal{S}'_\alpha$ completes effective consumption. Now if we let $\mathcal{S}_\gamma = \mathcal{S}_\alpha \cap \mathcal{S}'_\alpha$, observe that 1) $|\mathcal{S}_\gamma| > \frac{n}{2}$, and 2) all the servers in $\mathcal{S}_\gamma$ remain active from T_1 till T'_2 . The second part follows from our earlier observation that $T'_1 \geq T_2$. In this case, we infer that all the servers in $\mathcal{S}_\gamma$ does indeed acknowledge back to writer as part of their effective consumption of the *confirm-data* message, and since $\mathcal{S}_\gamma \subset \mathcal{S}_\alpha$ is at least a majority, we conclude that the write operation associated with the non faulty writer eventually completes.

H Proof of Theorem 7

Theorem 13. (Theorem 7) *Every execution of the $RADON_R^{(S)}$ algorithm is atomic.*

H.1 Some Preliminaries

The proof is based on Lemma 4, and the equivalent of Lemma 1 for $RADON_R^{(S)}$, which we state below for the sake of completion:

Lemma 6. *Let β denote a well-formed execution of $RADON_R^{(S)}$. Suppose T denotes a point of time in β such that there exists a majority of servers $\mathcal{S}_m$, $\mathcal{S}_m \subset \mathcal{S}$ all of which are in the active state at time T . Also, let t_s denote the value of the local tag at server s , at time T . Then, if π denotes any completed repair or read operation that is initiated after time T , we have $tag(\pi) \geq \min_{s \in \mathcal{S}_m} t_s$. Also, if π denotes any completed write operation that is initiated after time T , we have $tag(\pi) > \min_{s \in \mathcal{S}_m} t_s$.*

Proof. Similar to the proof of Lemma 1.

Next, in order to apply Lemma 4, consider any well-formed execution β of $RADON_R^{(S)}$, all of whose invoked read and write operations, denoted by the set Π_{RW} , complete. Recall the discussion in Section 6, where we noted that tags for completed operations in $RADON_R^{(S)}$ are defined exactly

⁶ In this case, some of the acks from the servers in $\mathcal{S}_\alpha$ get effectively consumed only after the required number $\lceil \frac{3n+1}{4} \rceil$ have already been consumed, the last of which includes execution of the group-send operation of the *confirm-data* phase. We note that the effective consumption of these additional acks from servers in $\mathcal{S}_\alpha$ is the operation where server simply ignores these, which is not explicitly mentioned in the algorithm. We also note that the notion of atomicity of any sequences of effective consumptions that are local to a server, is implicitly used when we argue that $T'_1 > T_2$. By this we mean that if a server receives a message m_1 before m_2 , the effective consumption of message m_1 is assumed to be entirely completed before the effective consumption of the message m_2 starts.

as we had done for $RADON_R$. Thus, for any completed write operation π , we define $tag(\pi)$ as the tag created by the writer during the *write-put* phase. For any completed read operation π , we define $tag(\pi)$ as the tag corresponding to the value returned by the read. Further, we define a partial order ($\prec$) on Π_{RW} like in the proof of Theorem 3 for case of $RADON_R$. These are restated for the sake of completion: For any $\pi, \phi \in \Pi_{RW}$, we say $\pi \prec \phi$ if one of the following holds: (i) $tag(\pi) < tag(\phi)$, or (ii) $tag(\pi) = tag(\phi)$, and π and ϕ are write and read operations, respectively.

H.2 Proof of Atomicity

Property P1 Consider two successful operations ϕ and π such that ϕ completes before π begins. We want to prove that $\pi \not\prec \phi$. Consider the case when both ϕ and π are write operations (the other cases are similar, so only one case is discussed). By virtue of the definition of the partial order ($\prec$), it is enough to prove that $tag(\pi) > tag(\phi)$. Let $\mathcal{S}_\alpha$ and $\mathcal{S}_1$ respectively denote the set of servers whose responses were used by the writer during the *put-data* and *confirm-data* phases of ϕ . Let T denote the time of initiation of the *confirm-data* phase of ϕ . From the algorithm (see Fig. 3), we know that $\mathcal{S}_1 \subset \mathcal{S}_\alpha$. Further, based on the algorithm, it is clear that all servers in $\mathcal{S}_1$ (which is a majority) are active at time T , such that $t_{loc}(s)|_T \geq tag(\phi)$. In this case, we apply Lemma 6 to conclude that $tag(\pi) > tag(\phi)$.

Property P2 This follows from the construction of tags, and the definition of the partial order ($\prec$).

Property P3 This follows from the definition of partial order ($\prec$), and by noting that value returned by a read operation π is simply the value associated with $tag(\pi)$.