

Coded Trace Reconstruction

Mahdi Cheraghchi, *Senior Member, IEEE*, Ryan Gabrys, *Member, IEEE*, Olga Milenkovic, *Fellow, IEEE*, and João Ribeiro, *Student Member, IEEE*

Abstract—Motivated by average-case trace reconstruction and coding for portable DNA-based storage systems, we initiate the study of *coded trace reconstruction*, the design and analysis of high-rate efficiently encodable codes that can be efficiently decoded with high probability from few reads (also called *traces*) corrupted by edit errors. Codes used in current portable DNA-based storage systems with nanopore sequencers are largely based on heuristics, and have no provable robustness or performance guarantees even for an error model with i.i.d. deletions and constant deletion probability. Our work is the first step towards the design of efficient codes with provable guarantees for such systems. We consider a constant rate of i.i.d. deletions, and perform an analysis of marker-based code-constructions. This gives rise to codes with redundancy $O(n/\log n)$ (resp. $O(n/\log \log n)$) that can be efficiently reconstructed from $\exp(O(\log^{2/3} n))$ (resp. $\exp(O(\log \log n)^{2/3})$) traces, where n is the message length. Then, we give a construction of a code with $O(\log n)$ bits of redundancy that can be efficiently reconstructed from $\text{poly}(n)$ traces if the deletion probability is small enough. Finally, we show how to combine both approaches, giving rise to an efficient code with $O(n/\log n)$ bits of redundancy which can be reconstructed from $\text{poly}(\log n)$ traces for a small constant deletion probability.

Index Terms—Coding, Deletion channel, DNA storage, Trace reconstruction.

I. INTRODUCTION

TRACE reconstruction was originally introduced by Batu, Kannan, Khanna, and McGregor [1], motivated by problems in sequence alignment, phylogeny, and computational biology. The setting for the problem is as follows: There is an unknown string $x \in \{0, 1\}^n$, and our goal is to reconstruct it. Towards this goal, we are allowed to ask for *traces* of x , which are obtained by sending x through a deletion channel. This channel independently deletes bits of x with a given deletion probability d . Traces are subsequences of x , and they are i.i.d. according to the output distribution of the deletion channel on input x . We wish to minimize the number of traces (or, in

other words, independent samples from the output distribution of the deletion channel on input x) required for reconstructing x with high probability.

Since its introduction, the problem of trace reconstruction has been studied from several different perspectives. Two of the main perspectives correspond to *worst-case* trace reconstruction [1]–[5], where the reconstruction algorithm must work simultaneously for all strings in $\{0, 1\}^n$, and *average-case* trace reconstruction [1]–[3], [6]–[8], where the reconstruction algorithm is only required to work with high probability, taken over the choice of string and the randomness of the reconstruction algorithm for a uniformly random string. More precisely, the goal of worst-case trace reconstruction is to design (or prove the existence of) a reconstruction algorithm Rec such that for every $x \in \{0, 1\}^n$ it holds that

$$\Pr_{T_1, \dots, T_t} [\text{Rec}(T_1, \dots, T_t) = x] \geq 1 - 1/n,$$

where the T_i are i.i.d. according to the output distribution of the deletion channel with deletion probability d on input x , for $t = t(n)$ as small as possible. On the other hand, for average-case reconstruction we only require that the reconstruction algorithm Rec has small *average* error probability, i.e.,

$$2^{-n} \sum_{x \in \{0, 1\}^n} \Pr_{T_1, \dots, T_t} [\text{Rec}(T_1, \dots, T_t) = x] \geq 1 - 1/n \quad (1)$$

with $t = t(n)$ as small as possible, where for each x the T_i are i.i.d. according to the output distribution of the deletion channel on input x .

The number of traces required for average-case trace reconstruction is, as expected, much smaller than that required for worst-case trace reconstruction. The problem in question has also been studied from a combinatorial coding perspective [9]–[12].

The above described results on average-case trace reconstruction can be interpreted from a coding-theoretic perspective: They state that there exist very large codebooks which can be reconstructed efficiently from relatively few traces.¹ However, no efficient encoders are known for such codes, and it may also be possible to further reduce the number of traces required for reconstruction by relaxing the size of the code.

This point of view naturally leads to the problem of *coded trace reconstruction*: The goal is to design high rate, efficiently encodable codes whose codewords can be efficiently reconstructed with high probability from very few traces with constant deletion probability. Here, “high rate” refers to a rate

¹If there exists a reconstruction algorithm Rec satisfying (1) for some t , then a simple averaging argument shows that there exists a set $\mathcal{G}$ of size at least 2^{n-1} such that, for all $x \in \mathcal{G}$, Rec reconstructs x with probability at least $1 - 2/n$ from t traces. Therefore, $\mathcal{G}$ is a large code that can be reconstructed from t traces.

Manuscript received April 8, 2019; revised March 6, 2020; accepted May 11, 2020. The material in this paper was presented in part at the 2019 IEEE Information Theory Workshop. The research was supported by the DARPA Molecular Informatics Program, the NSF and SRC SemiSynBio Award, and NSF grant #1618366. The work was also supported in part by the Center for Science of Information (CSOI), an NSF Science and Technology Center, under grant agreement CCF-0939370.

M. Cheraghchi was with the Department of Computing, Imperial College London, London SW7 2AZ, United Kingdom. He is now with the CSE Department, University of Michigan, Ann Arbor, MI 48109, USA (email: mahdich@umich.edu).

R. Gabrys and O. Milenkovic are with the ECE Department, University of Illinois at Urbana-Champaign, Urbana, IL 61801, USA (emails: ryan.gabrys@gmail.com; milenkov@illinois.edu).

J. Ribeiro is with the Department of Computing, Imperial College London, London SW7 2AZ, United Kingdom (email: j.lourenco-ribeiro17@imperial.ac.uk).

Copyright (c) 2017 IEEE. Personal use of this material is permitted. However, permission to use this material for any other purposes must be obtained from the IEEE by sending a request to pubs-permissions@ieee.org.

approaching 1 as the block length increases. We remark that in such a case, the number of traces must grow with the block length of the code. Coded trace reconstruction is also closely related to and motivated by the read process in portable DNA-based data storage systems, which we discuss below.

a) Motivation: A practical motivation for coded trace reconstruction comes from portable DNA-based data storage systems using DNA nanopores, first introduced in [13]. In DNA-based storage, a block of user-defined data is first encoded over the nucleotide alphabet $\{A, C, G, T\}$, and then transformed into moderately long strands of DNA through a DNA synthesis process. For ease of synthesis, the DNA strands are usually encoded to have balanced GC -content, so that the fraction of $\{A, T\}$ and $\{G, C\}$ bases is roughly the same. To recover the block of data, the associated strand of DNA is sequenced with nanopores, resulting in multiple corrupted reads of its encoding. Although the errors encountered during nanopore sequencing include both deletions/insertions as well as substitution errors, careful read preprocessing alignment [13] allows the processed reads to be viewed as traces of the data block's encoding. As a result, recovering the data block in question can be cast in the setting of trace reconstruction. Due to sequencing delay constraints², it is of great interest to minimize the number of reads required to reconstruct the data block.

The trace reconstruction procedures associated to the codes used by practical portable DNA-based storage systems [13], [14] are largely based on heuristics. The trace reconstruction algorithm proposed in [13] operates on carefully designed coded strings, but makes use of multiple sequence alignment algorithms which are notoriously difficult to analyze rigorously. The trace reconstruction algorithm from [14] does not make use of specific read-error correction codes and is a variation of the Bitwise Majority Alignment (BMA) algorithm originally introduced in [1]. However, the BMA algorithm is only known to be robust when the errors correspond to i.i.d. deletions and the fraction of errors is at most $O(1/\log n)$, where n denotes the blocklength of the code. Moreover, the proposed codes have been designed only for a fixed blocklength. As a result, the codes from [13], [14] have no robustness or performance guarantees for trace reconstruction even under i.i.d. deletions with constant deletion probability. Consequently, our work on coded trace reconstruction is the first step towards the development of codes with provable robustness guarantees and good performance for trace reconstruction in portable DNA-based data storage systems.

We note that, in practice, multiple strings are sequenced at the same time, and so one must have a way of telling which traces come from each string. In the original practical implementation of the DNA-based storage system that uses nanopore readouts [13], the codewords were augmented by prefix and suffix address strings, which are special short strings at large Hamming distance from each other. These strings are needed for indicating the position of the codewords in the long information-bearing string as well as for random access.

Given that address strings are the first to pass through the nanopores and that they are known to the receiver, a mixture of reads from different codewords can be automatically partitioned into clusters without using specialized algorithms. On the other hand, the follow-up work on portable DNA-based storage [14] relies on clustering algorithms for read reconstruction (see also [15]). With this in mind, it would also be interesting to study codes that simultaneously have good clustering and trace reconstruction properties, and we leave this as an open problem. We note that the study of codes with good clustering properties in relation to DNA-based data storage was initiated recently in [16], but trace reconstruction properties are not considered there.

A. Related work

Recently, there has been significant interest both in trace reconstruction and coding for settings connected to DNA-based data storage.

Regarding trace reconstruction, a chain of works [1]–[3], [6]–[8] has succeeded in substantially reducing the number of traces required for average-case trace reconstruction. The state-of-the-art result, proved by Holden, Pemantle, and Peres [8], states that $\exp(O(\log^{1/3} n))$ traces suffice to reconstruct a random n -bit string under arbitrary constant deletion probability. Much less is known about the worst-case trace reconstruction problem. The current best upper bound of $\exp(O(n^{1/3}))$ traces was proved concurrently by De, O'Donnell, and Servedio [4] and Nazarov and Peres [5] through algorithms that rely only on single-bit statistics from the trace. They also showed that this upper bound is tight for this restricted type of algorithms. The gap between upper and lower bounds for trace reconstruction for both the average- and worst-case settings is still almost exponential. The state-of-the-art lower bounds, both derived by Chase [17] improving upon an approach of Holden and Lyons [18], are close to $\log^{5/2} n$ traces for average-case trace reconstruction and $n^{3/2}$ traces in the worst-case setting. Trace reconstruction has also been studied over a large class of sticky channels [19], motivated by nanopore sequencers. A channel is said to be *sticky* if it preserves the block structure of the input, i.e., no input runs are completely deleted, and no new runs are added to it. In particular, the deletion channel is not a sticky channel.

Another line of work has focused on combinatorial settings related to coded trace reconstruction [9]–[12]. The reported works study the number of traces required for exact reconstruction when each trace is subjected to a bounded number of adversarial edit errors, and the string may be assumed to belong to a code from some general class. We note that Levenshtein [9] also studies a version of probabilistic trace reconstruction where the string is sent through a memoryless channel, but the deletion channel is not included in this family of channels. Other closely related combinatorial reconstruction problems consist in recovering a string from a subset of its substrings [20]–[22].

Coded trace reconstruction is also related to the multi-use deletion channel. In fact, the decoding problem for this channel can be interpreted as coded trace reconstruction with a fixed

²Current Oxford nanopore sequencers can process roughly 450 nucleotides per second.

number of traces. Some results are known about the capacity of this channel for small deletion probability [23], and about maximum likelihood decoding [24].

Recent works have also examined coding for other channel models inspired by DNA-based storage. The model studied in [25]–[28] views codewords as sets comprised of several sequences over some alphabet, and the (adversarial) errors consist of erasure and insertion of sequences in the set, as well as deletions, insertions, or substitutions within some of the sequences. Each sequence in the set corresponds to a different DNA strand whose contents are reconstructed via high throughput sequencing-based procedures that introduce errors. The goal is to recover the whole set from the erroneous sequences. This is fundamentally different from our model, as we focus on the correct reconstruction of a long single strand of DNA from multiple sequencing attempts and high probabilistic deletion error rates, which is especially relevant for portable nanopore DNA-based data storage systems. Another model similar to the one described above, motivated by the permutation channel, is studied in [29]. However, codewords there consist of multisets of symbols in some alphabet, and errors are comprised of deletions, insertions, substitutions, and erasures of symbols in the multiset. An information-theoretic treatment of related but abstracted models of DNA-based data storage may be found in [30], [31]. More recently, a model for clustering sequencing outputs according to the relevant DNA strand and codes that allow for correct clustering have been studied in [16].

Finally, although we focus on portable DNA-based storage systems with nanopore sequencers [13], [14], there has also been significant activity on practical aspects of other types of DNA-based storage, e.g., [32]–[34]. General overviews of the field can be found in [35], [36].

a) Concurrent and subsequent work on trace reconstruction: Several works on many different aspects of trace reconstruction have appeared concurrently and subsequently to the preparation and submission of this paper. Chase [17] improved upon the trace reconstruction lower bounds derived by Holden and Lyons [18]. Srinivasavaradhan et al. [37] studied symbolwise MAP decoding from a small number of traces. Davies, Racz, and Rashtchian [38] studied trace reconstruction over trees (in the graph-theoretic sense), where standard trace reconstruction corresponds to recovering a path. Ban et al. [39] considered the problem of trace reconstruction of distributions over small sets of strings (a setting commonly known as population recovery). Efficient population recovery algorithms have also been studied [40]. Krishnamurthy et al. [41] explored alternative and generalized forms of trace reconstruction (e.g., trace reconstruction over matrices and sparse strings). Sini and Yaakobi [42] made some progress on combinatorial trace reconstruction from a small number of deletions, insertions, and substitutions.

The problem of coded trace reconstruction has also been considered concurrently and independently by Abroshan et al. [43]. They focus on the setting with a constant *number* of deletions (as opposed to constant *rate*). The code they consider is obtained by concatenating several blocks, each block being a codeword of a Varshamov-Tenengolts code. Subsequently

to our work, Brakensiek, Li, and Spang [44] built upon our techniques and results. In terms of code constructions, they improve upon the number of traces required as a function of the rate. However, their construction for rate $1 - \epsilon$ requires superpolynomial preprocessing whenever $\epsilon = o\left(\frac{\log \log N}{\log N}\right)$, where N denotes the code block length. In contrast, all our code constructions in that regime are efficiently encodable and reconstructible.

B. Channel model

The channel model used for representing nanopore systems may be summarized as follows. For a given input string $x \in \{0, 1\}^n$, a deletion probability d , and an integer $t(n)$, the channel returns $t(n)$ traces of x . Each trace of x is obtained by sending x through a deletion channel with deletion probability d , i.e., the deletion channel deletes each bit of x independently with probability d , and outputs a subsequence of x containing all bits of x that were not deleted in order. The $t(n)$ traces are i.i.d. according to the output distribution of the deletion channel with deletion probability d on input x .

Given a code $\mathcal{C} \subseteq \{0, 1\}^n$, we say that $\mathcal{C}$ can be *efficiently reconstructed from $t(n)$ traces* if there exists a polynomial $p(n) = \Omega(n)$ and a polynomial-time algorithm Rec such that for every $c \in \mathcal{C}$ it holds that

$$\Pr_{T_1, \dots, T_t} [\text{Rec}(T_1, \dots, T_t) = c] \geq 1 - 1/p(n),$$

where the T_i are i.i.d. according to the output distribution of the deletion channel with deletion probability d on input c for $i = 1, \dots, t$. In words, the reconstruction algorithm Rec recovers every codeword $c \in \mathcal{C}$ from $t(n)$ i.i.d. traces of c with probability at least $1 - 1/p(n)$ over the probability distribution of the traces. We remark that this definition corresponds to worst-case trace reconstruction restricted to codewords of $\mathcal{C}$. The goal of coded trace reconstruction is to design *efficiently encodable* codes $\mathcal{C}$ that can be efficiently reconstructed from $t(n)$ traces for $t(n)$ as small as possible.

Remark 1: We have decided to require reconstruction success probability $1 - 1/p(n)$, where p is some polynomial, both for concreteness and for consistency when comparing our results to previous works on trace reconstruction (which generally also require success probability $1 - 1/p(n)$). The general tradeoff between code rate, number of traces, and success probability is not a focus of this work. However, it is natural to consider other settings (say, requiring only success probability $1 - o(1)$), and we leave this as an interesting direction for future research.

In view of the above, we note that the difference (in terms of number of traces required) by algorithms with some constant success probability $s > 1/2$ and algorithms with success probability $1 - 1/p(n)$ is relatively small. Indeed, given any algorithm Rec with success probability $s > 1/2$ using t traces, we can obtain an algorithm Rec' with success probability $1 - 1/n$ using $O(t \log n)$ traces. This is achieved by repeating Rec $O(\log n)$ times, each on a new batch of t traces, and choosing the most common output of Rec . By a straightforward application of the Chernoff bound, we obtain the desired result. Of course, there is still much room for improvement if

one aims for trace reconstruction with sublogarithmic number of traces.

C. Our contributions

We initiate the study of coded trace reconstruction for efficiently encodable, high-rate codes against a constant rate of deletions. The codes we construct have redundancy $o(n)$, where n is the message length, are efficiently encodable, and can either be efficiently reconstructed from significantly fewer traces than what is guaranteed (only *existentially*) by state-of-the-art average-case trace reconstruction results, or satisfy additional desirable properties for their use in DNA storage systems. More specifically, the following are our main contributions:

a) *Combining markers and worst-case trace reconstruction.*: We analyze the performance of marker-based constructions coupled with worst-case trace reconstruction algorithms. These constructions have the advantage that they can be easily adapted to work with a large range of inner codes.

At a high level, the construction operates by splitting an n -bit message into short blocks of length $O(\log^2 n)$, encoding each block with an inner code satisfying a certain constraint, and adding markers of length $O(\log n)$ between the blocks. The structure of the markers and the property of the inner code imply that, with high probability, we can split the traces into many shorter sub-traces associated with substrings of length $O(\log^2 n)$, and then apply the worst-case trace reconstruction algorithm on the sub-traces. Our main result in this context is Theorem 2.

Theorem 2: For every constant deletion probability $d < 1$, there exists an efficiently encodable code $\mathcal{C} \subseteq \{0, 1\}^{n+r}$ with redundancy $r = O(n/\log n)$ that can be efficiently reconstructed from $\exp(O(\log^{2/3} n))$ traces.

Note that, although the number of traces required by Theorem 2 does not improve upon the best result for average-case trace reconstruction [8], this result is interesting for several reasons: First, the codes provided by Theorem 2 are efficiently encodable. As mentioned before, average-case trace reconstruction results only imply the existence of codes with good trace reconstruction properties, and give no efficient encoders. Second, as we discuss below, the construction that leads to Theorem 2 can be easily adapted (by considering different instantiations of the inner code and reconstruction algorithms) to yield codes that either require much fewer traces to be reconstructed (with larger redundancy), or satisfy additional desirable properties for DNA storage. Finally, we remark that, with significantly more effort, we are able to extend these initial ideas to construct efficiently encodable codes with the same redundancy as Theorem 2 that only require $\text{poly}(\log n)$ traces for efficient reconstruction for some small enough constant deletion probability, which is a significant improvement upon the number of traces required by average-case trace reconstruction. We discuss this below.

b) *Adding a second level of markers.*: The inner code in the construction of Theorem 2 only needs to satisfy a simple constraint. This gives us a lot of freedom in instantiating the inner code, and so we are able to extend the construction above

in several ways. As a particularly interesting instantiation of the inner code, we can iterate the marker-based construction by further dividing each block of length $\log^2 n$ into blocks of length $(\log \log n)^2$ and adding markers of length $O(\log \log n)$ between them. The reconstruction procedure is then almost the same, except that, in this setting, it is almost guaranteed that reconstruction of a small fraction of blocks will fail. Nevertheless, this problem can be easily resolved by adding error-correction redundancy to the string to be encoded. This leads to the following result, which can be extended beyond two marker levels.

Theorem 3: For every constant deletion probability $d < 1$, there exists an efficiently encodable code $\mathcal{C} \subseteq \{0, 1\}^{n+r}$ with redundancy $r = O(n/\log \log n)$ that can be efficiently reconstructed from $\exp(O(\log \log n)^{2/3})$ traces.

c) *Codes with balanced GC-content.*: We take advantage of the fact that we can instantiate the marker-based constructions with a large range of inner codes to construct high-rate marker-based codes over the $\{A, C, G, T\}$ alphabet with three important properties: The codes are efficiently encodable, have *balanced GC-content* (which is important both for accurate DNA synthesis and sequencing), and provably require few traces to be efficiently reconstructed. We follow the same ideas as the marker-based constructions above, but with different markers and an inner code over a larger alphabet and with stronger constraints. In this context, we obtain the following results.

Theorem 4: For every constant deletion probability $d < 1$, there exists an efficiently encodable code $\mathcal{C} \subseteq \{A, C, G, T\}^{n+r}$ with redundancy $r = O(n/\log n)$ and balanced GC-content that can be efficiently reconstructed from $\exp(O(\log^{2/3} n))$ traces.

Theorem 5: For every constant deletion probability $d < 1$, there exists an efficiently encodable code $\mathcal{C} \subseteq \{A, C, G, T\}^{n+r}$ with redundancy $r = O(n/\log \log n)$ and balanced GC-content that can be efficiently reconstructed from $\exp(O(\log \log n)^{2/3})$ traces.

d) *Leveraging average-case trace reconstruction algorithms.*: The result of Theorem 2 may be further improved by considering a more careful design of the high-rate inner code to be used in the marker-based constructions, provided that the deletion probability is a small enough constant. This allows for using a modified version of an algorithm for average-case trace reconstruction described in [2] which leads to a substantial reduction in the number of traces required for reconstruction and barely any rate changes. Towards achieving this goal, we first design a low-redundancy code that can be efficiently reconstructed from polynomially many traces. The proposed coding scheme relies on the fact that we can efficiently encode n -bit messages into strings that are almost *subsequence-unique* (see Definition 12) via explicit constructions of *almost k -wise independent random variables* (see Section II-B). The average-case trace reconstruction algorithm from [2] operates on subsequence-unique strings, and a simple adaptation of the algorithm suffices for our approach. In summary, we have the following result.

Theorem 6: There exists an absolute constant $d^* > 0$ such that for all $d \leq d^*$ there exists an efficiently encodable code

$\mathcal{C} \subseteq \{0, 1\}^{n+r}$ with redundancy $r = O(\log n)$ that can be efficiently reconstructed from $\text{poly}(n)$ traces with deletion probability d .

An important step in our analysis is to show how to adapt this code for use as an inner code in the marker-based construction. Some care is needed, since the global structure of the strings we deal with changes significantly due to the presence of the markers. In particular, the bootstrapping method in the trace reconstruction algorithm from [2] no longer works, and we must find a way to circumvent this issue. Our findings for this scenario are described in the next theorem.

Theorem 7: There exists an absolute constant $d^* > 0$ such that for all $d \leq d^*$ there exists an efficiently encodable code $\mathcal{C} \subseteq \{0, 1\}^{n+r}$ with redundancy $r = O(n/\log n)$ that can be efficiently reconstructed from $\text{poly}(\log n)$ traces with deletion probability d .

A summary of our main results can be found in Table I. For simplicity, our exposition mostly focuses on constructions of *binary* codes, although it provides some guidelines and simple coding procedures for quaternary codes. It is also important to note that designing a code with a given rate R for coded trace reconstruction is inherently harder for smaller alphabets. Indeed, the existence of a binary code $\mathcal{C} \subseteq \{0, 1\}^n$ with rate R that can be efficiently encoded and reconstructed from t traces with error probability ϵ implies the existence of a Q -ary code $\mathcal{C}'$ (for $Q = 2^q$) with the same rate R that can be efficiently encoded and reconstructed from t traces with error probability at most $q\epsilon$. To see this, consider the code

$$\mathcal{C}' = \{(c^1, c^2, \dots, c^q) : c^i \in \mathcal{C}, i = 1, \dots, q\} \subseteq \{0, 1\}^{q \cdot n}.$$

Clearly, the code $\mathcal{C}'$ is Q -ary code of length n and rate R . Moreover, let T denote a trace of some $c' = (c^1, c^2, \dots, c^q) \in \mathcal{C}'$. Observe that the trace T^i obtained by replacing each Q -ary symbol in T by the i -th bit of its binary expansion has the same distribution as a trace of c^i . As a result, applying the transformation $T \mapsto T^i$ to each of the t traces of c' and running the reconstruction algorithm associated with $\mathcal{C}$ allows us to recover c^i with error probability at most ϵ . Since this holds for every $i = 1, \dots, q$, a union bound over all i shows that we can simultaneously recover $c^1, c^2, \dots, c^q$ from t traces of c' with error probability at most $q\epsilon$.

D. Organization

The paper is organized as follows: In Section II, we define relevant notation and discuss known results that we find useful in our subsequent derivations. We describe and analyze general marker-based constructions in Section III. Then, we show how to reduce the number of traces required for a small deletion probability in Section IV. We discuss some open problems in Section V.

II. NOTATION AND PRELIMINARIES

A. Notation

We denote the length of a string x by $|x|$, and its Hamming weight by $w(x) = |\{i : x_i \neq 0\}|$. Given two strings x and y over the same alphabet, we denote their concatenation by $x||y$. For a string x , we define $x[a, b] = (x_a, x_{a+1}, \dots, x_b)$

and $x[a, b] = (x_a, x_{a+1}, \dots, x_b)$. If $|x| = n$, we define $x[a, \cdot] = (x_a, x_{a+1}, \dots, x_n)$. We say that y is a subsequence of x if there exist indices $i_1 < i_2 < \dots < i_{|y|}$ such $x_{i_j} = y_j$. Moreover, y is said to be a substring of x if $y = x[a, a+|y|)$ for some $1 \leq a \leq |x| - |y| + 1$. Given two strings $x, y \in \{0, 1\}^n$, we write $x + y$ for the bitwise XOR of x and y . A *run of length ℓ* in a string x is a substring of x comprising ℓ identical symbols. Sets are denoted by calligraphic letters such as $\mathcal{S}, \mathcal{T}$. Random variables are denoted by uppercase letters such as X, Y , and Z . The uniform distribution over $\{0, 1\}^t$ is denoted by U_t , and the binomial distribution on n trials with success probability p is denoted by $\text{Bin}(n, p)$. The binary entropy function is denoted by h and all logarithms $\log$ are taken with respect to the base 2.

B. Almost k -wise independent random variables

We start by defining almost k -wise independence and present a related result that are needed for our subsequent derivations.

Definition 8 (ϵ -almost k -wise independent random variable): A random variable $X \in \{0, 1\}^m$ is said to be ϵ -almost k -wise independent if for all sets of k distinct indices $i_1, i_2, \dots, i_k \in \{1, \dots, m\}$ we have

$$|\Pr[X_{i_1} = x_1, \dots, X_{i_k} = x_k] - 2^{-k}| \leq \epsilon$$

for all $(x_1, \dots, x_k) \in \{0, 1\}^k$.

The following result gives an efficient construction of an ϵ -almost k -wise independent random variable which can be generated from few uniformly random bits.

Lemma 9 ([45]): For every m, k , and ϵ there exists a function $g : \{0, 1\}^t \rightarrow \{0, 1\}^m$ with $t = O\left(\log\left(\frac{k \log m}{\epsilon}\right)\right)$ such that $g(U_t)$ is an ϵ -almost k -wise independent random variable over $\{0, 1\}^m$, where U_t denotes the uniform distribution over $\{0, 1\}^t$.

Moreover, $g(z)$ can be computed in time $\text{poly}(m, 1/\epsilon)$ given z . In particular, if $\epsilon = 1/\text{poly}(m)$, then g can be computed in time $\text{poly}(m)$.

C. Nearly-optimal systematic codes for edit errors

We require systematic codes that are robust against edit errors (deletions and insertions). Nearly-optimal systematic codes for adversarial edit errors have been recently constructed using optimal protocols for deterministic document exchange [46], [47]. The following result is relevant to our analysis.

Lemma 10 ([46], [47]): For every m and $t < m$ there exists an efficiently encodable and decodable systematic code $\mathcal{C}_{\text{edit}} \subseteq \{0, 1\}^{m+r}$ with encoder $\text{Enc}_{\text{edit}} : \{0, 1\}^m \rightarrow \{0, 1\}^{m+r}$ and redundancy $r = O(t \log^2 \frac{m}{t} + t)$ that can efficiently correct up to t edit errors. In particular, if $t = \Theta(m)$ then the redundancy is $r = O(m)$.

D. Trace reconstruction

Next, we discuss several results pertaining to the worst-case and average-case trace reconstruction problem that will be useful for our constructions.

Redundancy	Number of traces	Observations
$\frac{n}{\log n}$	$\exp(\log^{2/3} n)$	Any constant deletion probability, balanced GC-content
$\frac{n}{\log \log n}$	$\exp((\log \log n)^{2/3})$	Any constant deletion probability, balanced GC-content
$\frac{n}{\log n}$	$\text{poly}(\log n)$	Small enough constant deletion probability

TABLE I

SUMMARY OF THE PROPERTIES OF OUR MAIN CODE CONSTRUCTIONS. WE HAVE OMITTED HIDDEN CONSTANTS FROM THE PARAMETERS FOR SIMPLICITY.

1) *Worst-case trace reconstruction*: For worst-case reconstruction, the state-of-the-art result used in Section III is summarized below.

Lemma 11 ([4], [5]): For every n and constant deletion probability d there exists an algorithm that reconstructs an arbitrary string $x \in \{0,1\}^n$ with probability at least $1 - \exp(-2n)$ from $\exp(O(n^{1/3}))$ traces in time $\exp(O(n^{1/3}))$.

2) *Trace reconstruction of subsequence-unique strings*: One of the key tools for our constructions in Section IV is a modified version of the efficient trace reconstruction algorithm [2] for what we refer to as *subsequence-unique* strings. This algorithm may also be used for average-case trace reconstruction. We start by defining subsequence-unique strings.

Definition 12 (w -subsequence-unique string): A string $x \in \{0,1\}^n$ is said to be w -subsequence-unique if for every $1 \leq a, b \leq n$ such that $a + w, b + 1.1w \leq n + 1$, and either $a < b$ or $b + 1.1w < a + w$, we have that the substring $x[a, a + w]$ is not a subsequence of $x[b, b + 1.1w]$.

We remark that the constant 1.1 in Definition 12 is arbitrary, and can be replaced by another constant close to 1. Since we make no effort to optimize constants in our results, we use the same constant as in [2]. Intuitively, a string x is w -subsequence-unique if no length w substring of x appears as a *subsequence* of another slightly longer substring of x (except for the trivial case where the longer substring contains the shorter one). Note that these strings have been defined under the name “substring-unique” in [2]. We chose to change this name to avoid confusion with a different definition under the same name described in [22]. The following result about subsequence-unique strings was established in [2].

Lemma 13 ([2, Theorem 2.5]): For $w = 100 \log n$ and a small enough constant deletion probability d , there exists an algorithm that reconstructs every w -subsequence-unique string $x \in \{0,1\}^n$ with probability $1 - 1/\text{poly}(n)$ from $\text{poly}(n)$ traces in time $\text{poly}(n)$.

Since a uniformly random string is w -subsequence-unique with high probability, Lemma 13 applies to average-case trace reconstruction. As we make explicit use of the algorithm behind Lemma 13, for the sake of clarity we provide next a more in-depth discussion of the method. However, before we proceed to the actual description of the algorithm, we briefly introduce some definitions and basic related results.

Given integers i and j and a deletion probability d , we denote the probability that the i -th bit of a string appears as the j -th bit of its trace by $P(i, j)$. Then, we have

$$P(i, j) = \binom{i-1}{j-1} (1-d)^j d^{i-j}.$$

The following lemma states some useful properties of $P(i, j)$.

Lemma 14 ([2, Lemma 2.2]): If $j \leq (1 - 3d)i$, then $P(i, j) \geq 2 \sum_{i' > i} P(i', j)$. Furthermore, if $(1 - 4d)i < j < (1 - 3d)i$, we have $P(i, j) \geq \exp(-6di)$.

Intuitively, the second part of Lemma 14 suggests that we have a good idea of the position of x_i in the trace if i is small. The following result makes use of this observation. It states that we can recover the first $O(\log n)$ bits of an arbitrary string with $\text{poly}(n)$ traces, which is required to bootstrap the trace reconstruction algorithm from [2].

Lemma 15 ([2, Theorem 2.1]): Fix a string $x \in \{0,1\}^n$, and suppose that we know $x_1, \dots, x_{h-1}$. Then, there is an algorithm that recovers x_h from $\exp(O(hd \log(1/d)))$ traces of x with probability $1 - o(1)$, provided that $d < 1/3$.

In the second part of the algorithm, we must look for matchings of certain strings within the traces. To this end, we introduce the following definition.

Definition 16 (Matching): Fix a string $x \in \{0,1\}^n$, and let T denote a trace of x . Then, we say that there is a *matching* of $x[a, b]$ in T if there exists some u such that $T[u - (b - a), u] = x[a, b]$.

Matchings of w -subsequence-unique strings have useful properties, as formalized in the following lemma.

Lemma 17 ([2, Lemma 2.4]): If x is w -subsequence-unique and there is a matching of $x[a, a + w]$ in T , say at $T[u - w, u]$, then the probability that T_{u-1} does not come from $x[a + w, a + 1.1w]$ is at most $nd^{0.001w}$.

We are now in a position to describe the algorithm introduced in [2]. We begin by setting $w = 100 \log n$, $v = w/d$, and $j = (v - 0.1w)(1 - 3d)$. Then, to recover a w -subsequence-unique string x we proceed with two steps: First, we use the algorithm from Lemma 15 to recover the first v bits of x with $\text{poly}(n)$ traces. Now, suppose we have recovered $x_1, \dots, x_{i-1}$ for $i - 1 \geq v$. Our next goal is to recover x_i with $\text{poly}(n)$ traces. Note that if i is relatively large, we cannot use the algorithm from Lemma 15 to recover x_i as it would require more than $\text{poly}(n)$ traces. To achieve our goal, we instead focus on finding matchings of the substring $x[i - v - w, i - v]$ within the trace. Let T denote a trace of x and suppose there is a matching of $x[i - v - w, i - v]$ in T at positions $T[u - w, u]$. Then, we set $V = T[u, \cdot]$, i.e., we let V be the suffix of the trace following the matching. The key property is that $\Pr[V_j = 1]$ satisfies a threshold property depending on the value of x_i . More precisely, there exist two positive values $B_1 > B_0$ sufficiently far apart such that $\Pr[V_j = 1] \leq B_0$ if $x_i = 0$ and $\Pr[V_j = 1] \geq B_1$ if $x_i = 1$. Moreover, all terms in these inequalities can be estimated with a small error from $\text{poly}(n)$ traces of x . As a result, we can reliably estimate x_i by checking whether $\Pr[V_j = 1] \leq B_0$ or $\Pr[V_j = 1] \geq B_1$.

We prove next the threshold property for $\Pr[V_j = 1]$. Let R

denote the position in x of the bit appearing in position $u - 1$ in the trace T of the matching for $x[i - v - w, i - v]$. In other words, R denotes the position in x of the last bit appearing in the matching in T . We may write

$$\begin{aligned} \Pr[V_j = 1] &= \sum_{r=1}^n \Pr[R = r] \Pr[V_j = 1 | R = r] \\ &= \epsilon_i(x) + \sum_{r=i-v}^{i-v+0.1w} \Pr[R = r] \Pr[V_j = 1 | R = r] \\ &= \epsilon_i(x) + \sum_{r=i-v}^{i-v+0.1w} \Pr[R = r] \sum_{\ell=r+1}^n P(\ell - r, j) x_\ell \\ &= \epsilon_i(x) + \sum_{r=i-v}^{i-v+0.1w} \Pr[R = r] \sum_{\ell=r+1}^{i-1} P(\ell - r, j) x_\ell \\ &\quad + \sum_{r=i-v}^{i-v+0.1w} \Pr[R = r] \cdot P(i - r, j) x_i \\ &\quad + \sum_{r=i-v}^{i-v+0.1w} \Pr[R = r] \cdot \sum_{\ell=i+1}^n P(\ell - r, j) x_\ell, \end{aligned}$$

where the second equality follows from Lemma 17 with $0 \leq \epsilon_i(x) \leq nd^{-0.001w}$. Using the first part of Lemma 14, we conclude that $\sum_{\ell=i+1}^n P(\ell - r, j) \leq \frac{1}{2}P(i - r, j)$. As a result, when $x_i = 0$ we have

$$\begin{aligned} \Pr[V_j = 1] &\leq \epsilon_i(x) + \sum_{r=i-v}^{i-v+0.1w} \Pr[R = r] \sum_{\ell=r+1}^{i-1} P(\ell - r, j) x_\ell \\ &\quad + \frac{1}{2} \sum_{r=i-v}^{i-v+0.1w} \Pr[R = r] P(i - r, j), \end{aligned} \quad (2)$$

and when $x_i = 1$ we have

$$\begin{aligned} \Pr[V_j = 1] &\geq \epsilon_i(x) + \sum_{r=i-v}^{i-v+0.1w} \Pr[R = r] \sum_{\ell=r+1}^{i-1} P(\ell - r, j) x_\ell \\ &\quad + \sum_{r=i-v}^{i-v+0.1w} \Pr[R = r] P(i - r, j). \end{aligned} \quad (3)$$

By the second part of Lemma 14, since $i - r \leq v$ and $v = w/d$, we have $P(i - r, j) \geq 2^{-9w}$. Combining this with Lemma 17 for d small enough means that the gap between the right hand side of (2) and (3) is at least $2^{-(9w+1)}$. To finalize the argument, we note that (i) we can efficiently approximate $\Pr[V_j = 1]$ to within an error of, say, 2^{-100w} with high probability from $\text{poly}(n)$ traces of x , and (ii) we can efficiently approximate $\Pr[R = r | R < i]$ to within the same error given that we know $x_1, \dots, x_{i-1}$, provided d is small enough. Since $\Pr[R < i] \geq 1 - nd^{-0.001w}$ by Lemma 17, we can further efficiently approximate $\Pr[R = r]$ to within an error of, say, 2^{-50w} with high probability. From these observations it follows that we can estimate x_i correctly with high probability from $\text{poly}(n)$ traces, where the degree of the polynomial is independent of i , as desired.

III. MARKER-BASED CONSTRUCTIONS

We start with simple constructions of high-rate codes that can be efficiently reconstructed from a few traces. The idea

behind the approach is the following: Each codeword contains markers, consisting of sufficiently long runs of 0's and 1's. Between two consecutive markers, we add a short block containing a codeword from an inner code satisfying a mild constraint.

Intuitively, the runs in the markers will still be long in the trace, and so we hope to be able to correctly identify the positions of all markers in a trace with high probability. After this is done, we can effectively split the trace into many shorter, independent sub-traces corresponding to a block (and possibly some bits from the two markers delimiting it). Then, we can apply worst-case trace reconstruction algorithms to the sub-traces. The savings in the number of traces required for reconstruction stem from the fact that sub-traces are short, and that each trace can be utilized simultaneously (and independently) by all blocks. This idea for reconstruction almost works as is, except that the process of identifying the markers in a trace may be affected by long runs of 0's originating from a block between two markers. However, this can be easily solved by requiring that all runs of 0's in each block are short enough. Many codes, including codes with low redundancy, satisfy the desired property, and hence make for good candidates for the inner code.

We describe and analyze a code based on the idea discussed above in Section III-A. Then, we consider a follow-up construction in Section III-B which requires fewer traces, at the expense of a decrease in the rate. At a high-level, this second code is obtained by introducing two levels of markers and adding some simple error-correction redundancy to the message prior to other encodings. Finally, in Section III-C we extend these ideas to the $\{A, C, G, T\}$ alphabet in order to obtain high-rate codes with desirable properties for use in DNA-based storage. Namely, these codes have balanced GC-content and can be reconstructed from few traces. Such codes are designed by exploiting the fact that the marker-based constructions can be instantiated with a large range of inner codes, and we can make the inner code satisfy stronger constraints than before.

A. A simple construction

Here we provide a precise description of the encoder Enc for our code $\mathcal{C}$ and prove Theorem 2. For simplicity, we consider $d = 1/2$ throughout. This choice of d is arbitrary, and the construction and analysis can be easily generalized to all constant $d \in (0, 1)$.

Let $\ell = 50 \log n$, and define two strings $M_0 = 0^\ell$ and $M_1 = 1^\ell$. Then, a marker M is a string of length 2ℓ of the form $M = M_0 \| M_1 = 0^\ell \| 1^\ell$. We also require an efficiently encodable and decodable inner code $\mathcal{C}' \subseteq \{0, 1\}^{m+r}$ with encoder $\text{Enc}' : \{0, 1\}^m \rightarrow \{0, 1\}^{m+r}$, where $m = \log^2 n$ and r is the redundancy, satisfying the following property.

Property 18: For all $c \in \mathcal{C}'$ and substrings s of c with $|s| = \sqrt{m}$, it holds that $w(s) \geq |s|/3$.

In other words, every codeword of $\mathcal{C}'$ has many 1's in all short enough substrings. Such efficient codes exist with redundancy $r = O(\log m) = O(\log \log n)$, which is enough for our needs. We provide a simple construction in Section III-A1.

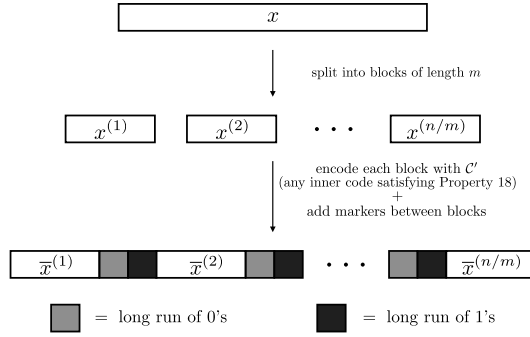


Fig. 1. The general encoder for $\mathcal{C}$. By considering different instantiations of the inner code C' , this construction is used to prove Theorem 2 (see Section III-A), Theorem 4 (with slightly different markers due to the use of a 4-ary alphabet, see Section III-C), and Theorem 7 (see Section IV-B).

Suppose we wish to encode an n -bit message $x \in \{0, 1\}^n$. The encoder Enc on input x proceeds through the following steps:

- 1) Split x into n/m blocks, each of length m
- 2) Encode each block $x^{(i)}$ under the inner code C' to obtain $\bar{x}^{(i)} = \text{Enc}'(x^{(i)}) \in \{0, 1\}^{m+r}$;
- 3) Set the encoding of x , denoted by $\text{Enc}(x)$, to be

$$\text{Enc}(x) = 1^\ell \|\bar{x}^{(1)}\| M \|\bar{x}^{(2)}\| M \dots \|M\| \bar{x}^{(n/m)} \|0^\ell.$$

We remark that the first run 1^ℓ and the last run 0^ℓ are superfluous, and are added only to make the analysis simpler. Computing $\text{Enc}(x)$ from x and decoding x from $\text{Enc}(x)$ can both be done efficiently if the inner code C' is efficiently encodable and decodable. Figure 1 illustrates the encoding procedure for $\mathcal{C}$ with a general inner code C' satisfying Property 18 detailed above.

We now compute the redundancy of $\mathcal{C}$. It is straightforward to see that

$$|\text{Enc}(x)| \leq \frac{n}{m}(|M| + |\bar{x}^{(1)}|) = n + O\left(\frac{n}{\log n}\right) + \frac{nr}{\log^2 n}. \quad (4)$$

As mentioned before, we have $r = O(\log m) = O(\log \log n)$. Therefore, $\mathcal{C}$ can be made to have redundancy $O\left(\frac{n}{\log n}\right)$.

In the remainder of this section, we prove Theorem 2 using $\mathcal{C}$ via a sequence of lemmas. For convenience, we restate the theorem below.

Theorem 19 (Theorem 2, restated): There is an efficient algorithm that recovers every $c \in \mathcal{C}$ from $\exp(O(\log^{2/3} n))$ traces in time $\text{poly}(n)$ with probability $1 - 1/\text{poly}(n)$.

The reconstruction procedure works as follows: First, we show that the markers M still contain long enough runs after they are sent through the deletion channel. Then, we show that no long runs of 0's originate from the sub-traces associated with each block. This implies that we can correctly identify the position of the “01” string of each marker in the trace. As a result, we can split the trace into smaller “sub-traces”, each one associated to a different block $\bar{x}^{(i)}$. Then, we apply a reconstruction algorithm to the set of sub-traces associated to each block in order to reconstruct the blocks, and thus

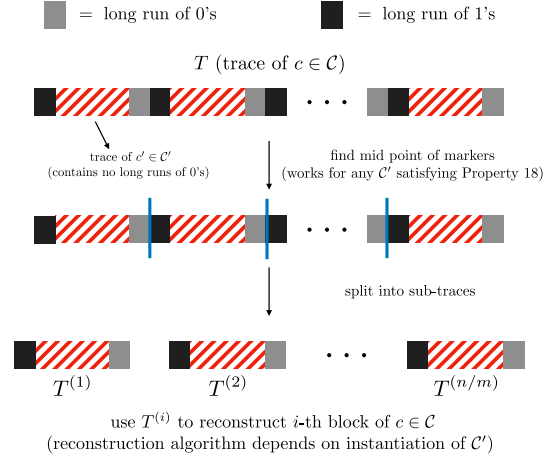


Fig. 2. The reconstruction procedure for $\mathcal{C}$. This high-level approach is used across our work in combination with trace reconstruction algorithms of different levels of sophistication.

the whole codeword. This general reconstruction procedure is illustrated in Figure 2.

To obtain Theorem 2, we show that we can apply the worst-case trace reconstruction algorithm from Lemma 11 to recover each block with high probability and with the desired number of traces. We shall see in Section IV that a careful instantiation of the inner code C' allows us to use more sophisticated trace reconstruction algorithms.

We start by proving that the markers M still contain long runs after they are sent through the deletion channel.

Lemma 20: Let $0^{L_0} 1^{L_1}$ be the output of the deletion channel on input M . Then,

$$\Pr[L_0 > 10 \log n, L_1 > 0] \geq 1 - n^{-3}.$$

Proof: The result follows by a standard application of the Chernoff bound. More precisely, we have $\mathbb{E}[L_0] = 25 \log n$ since $\ell = 50 \log n$ and $d = 1/2$, and hence

$$\begin{aligned} \Pr[L_0 \leq 10 \log n] &= \Pr[L_0 \leq \mathbb{E}[L_0] - 15 \log n] \\ &\leq \exp\left(-\frac{15^2 \log^2 n}{2\mathbb{E}[L_0]}\right) \\ &\leq n^{-4}. \end{aligned}$$

To conclude the proof, we note that $\Pr[L_1 = 0] = 2^{-\ell} = n^{-50}$. A union bound over the two probabilities yields the desired lower bound. ■

We now show that no long runs of 0's originate from the sub-traces associated with each block.

Lemma 21: Let $c \in \mathcal{C}'$. Then, a trace of c does not contain a run of 0's of length at least $10 \log n$ with probability at least $1 - n^{-4}$.

Proof: Since $c \in \mathcal{C}'$, a run of 0's of length at least $10 \log n$ in the trace of c requires that there is a subsequence of c with length at least $10 \times \frac{\log n}{2} - 1 = 5 \log n - 1$ comprised only of *deleted* 1's without any undeleted 1 in between them. The probability that a fixed subsequence of 1's with length $5 \log n - 1$ is completely deleted is at most $n^{-4.5}$. Since there are at most $O(\log^2 n)$ such sequences in c , by the union bound it follows that the desired probability is at most n^{-4} . ■

The next lemma follows immediately by combining Lemmas 20 and 21 with the union bound over the $n/\log^2 n$ blocks.

Lemma 22: Consider the following event E : We correctly identify the separation between the traces of 0^ℓ and 1^ℓ from every marker in the trace of $\text{Enc}(x)$ by looking for all 1's that appear immediately after a run of at least $10 \log n$ 0's.

Then, E happens with probability at least $1 - n^{-2}$ over the randomness of the trace.

We are now ready to prove Theorem 2. Let E denote the event described in Lemma 22. Then, Lemma 22 implies that, conditioned on E , we can split a trace T of $\text{Enc}(x)$ into $n/\log^2 n$ sub-traces $T^{(1)}, \dots, T^{(n/\log^2 n)}$ satisfying the following:

- The sub-traces $T^{(i)}$ are independent for $i = 1, 2, \dots, n/\log^2 n$;
- Each sub-trace $T^{(i)}$ is distributed like a trace of $1^\ell \|\bar{x}^{(i)}\| 0^\ell$ conditioned on the high probability event E .

In fact, each sub-trace $T^{(i)}$ can be identified by looking for the $(i-1)$ -th and i -th runs of 0 of length at least $10 \log n$ in the trace T , and picking every bit in T immediately after the $(i-1)$ -th run up to and including the i -th run.

Observe that $1^\ell \|\bar{x}^{(i)}\| 0^\ell$ has length $O(\log^2 n)$. Suppose that we have $t = \exp(O(\log n)^{2/3})$ independent traces $T_1, \dots, T_t$ of $\text{Enc}(x)$. Let E_{all} denote the event that E holds for all T_i simultaneously. Combining Lemma 22 with the union bound yields

$$\Pr[E_{\text{all}}] \geq 1 - t/n^2 > 1 - 1/n. \quad (5)$$

Fix some trace reconstruction algorithm $\mathcal{A}$, and let $E_{\text{indFail}}^{(i)}$ denote the event that $\mathcal{A}$ fails to recover a fixed string $y^{(i)} = 1^\ell \|\bar{x}^{(i)}\| 0^\ell$ from t independent traces of $y^{(i)}$. Assuming that E_{all} holds, the sub-traces $T_1^{(i)}, \dots, T_t^{(i)}$ (note that $T_j^{(i)}$ denotes the i -th sub-trace of the j -th trace) are distributed as t independent traces of $y^{(i)}$, each also satisfying the conditions that the first run 1^ℓ is not completely deleted, the last run 0^ℓ has length at least $10 \log n$ in the trace, and there is no run of 0's of length at least $10 \log n$ in the trace of $\bar{x}^{(i)}$. We denote the event that these conditions hold for all of the t independent traces of $y^{(i)}$ by $E_{\text{split}}^{(i)}$. Finally, we let E_{fail} denote the event that we fail to recover $\text{Enc}(x)$ from the t i.i.d. traces $T_1, \dots, T_t$. Then, we have

$$\begin{aligned} \Pr[E_{\text{fail}}] &\leq \Pr[E_{\text{fail}}, E_{\text{all}}] + \Pr[\neg E_{\text{all}}] \\ &= \Pr[(\exists i : E_{\text{indFail}}^{(i)}), (\forall i : E_{\text{split}}^{(i)})] + \Pr[\neg E_{\text{all}}] \\ &\leq \Pr[\exists i : E_{\text{indFail}}^{(i)}] + 1/n \end{aligned} \quad (6)$$

$$\leq \sum_{i=1}^{n/\log^2 n} \Pr[E_{\text{indFail}}^{(i)}] + 1/n. \quad (7)$$

The first equality follows from the discussion in the previous paragraph, the second inequality follows from (5), and the third inequality follows by the union bound. Instantiating $\mathcal{A}$ with the worst-case trace reconstruction algorithm from Lemma 11, we conclude from (7) that

$$\Pr[E_{\text{fail}}] \leq n \cdot \exp(-2 \log^2 n) + 1/n < 2/n.$$

As a result, we can successfully recover x from $\exp(O(\log n)^{2/3})$ traces of $\text{Enc}(x)$ with probability at least

$1 - 2/n$. To conclude the proof, we note that we can repeat the process $O(\log n)$ times and take the majority vote to boost the success probability to $1 - 1/p(n)$ for any fixed polynomial p of our choice. The total number of traces required is still $\exp(O(\log^{2/3} n))$. Since recovering each $\bar{x}^{(i)}$ from the associated traces takes time $\exp(O(\log^{2/3} n))$ and the inner code $\mathcal{C}'$ has an efficient decoder, the whole procedure is efficient.

Remark 23: By modifying the inner block length m from $\log^2 n$ to $\log^{1+\gamma} n$ for an arbitrary constant $\gamma \in (0, 1)$, the reasoning above can be adapted to yield efficiently encodable codes with redundancy $O(n/\log^\gamma n)$ which are efficiently reconstructible from $\exp(O(\log^{\frac{1+\gamma}{3}} n))$ traces.

1) Instantiating the inner code: What remains to be done is to instantiate the inner code $\mathcal{C}'$ with the appropriate parameters and properties. To this end, we present a simple construction of an efficiently encodable and decodable inner code $\mathcal{C}'$ with encoder $\text{Enc}' : \{0, 1\}^m \rightarrow \{0, 1\}^{m+r}$ and redundancy $r = O(\log m)$. We can then obtain the desired code by setting $m = \log^2 n$. The starting point is the following result.

Lemma 24: Let $g : \{0, 1\}^t \rightarrow \{0, 1\}^m$ be the function whose existence is guaranteed by Lemma 9 with $k = 3w$ and $\epsilon = 2^{-10w}$ for $w = 100 \log m$ (hence $t = O(\log m)$). Fix some $x \in \{0, 1\}^m$ and consider the random variable $Y = x + g(U_t)$. Then, with probability at least $1 - 2/m$, we have that Y satisfies the following property:

Property 25: One has $w(Y[a, a+w]) \geq 0.4w$ simultaneously for all $1 \leq a \leq m - w + 1$.

Proof: Fix some a . Then, we have

$$\begin{aligned} \Pr[w(Y[a, a+w]) < 0.4w] &= \sum_{y: w(y) < 0.4w} \Pr[Y[a, a+w] = y] \\ &\leq \sum_{y: w(y) < 0.4w} (2^{-w} + 2^{3w} \epsilon) \\ &\leq 2^{wh(0.4)} \cdot 2^{-w+1} \\ &\leq \frac{2}{m^2}. \end{aligned}$$

The first inequality follows because Y is ϵ -almost k -wise independent, and the second inequality follows from a standard bound on the volume of the Hamming ball and the fact that $2^{3w} \epsilon < 2^{-w}$. Since there are at most m choices for a , by the union bound we conclude that Y fails to satisfy the desired property with probability at most $m \cdot 2/m^2 = 2/m$, as desired. ■

Given $x \in \{0, 1\}^m$, we evaluate $\text{Enc}'(x)$ as follows: We iterate over all $z \in \{0, 1\}^t$ until we find a z such that $y = x + g(z)$ satisfies $w(s[a, a+w]) \geq 0.4w$. Such a string z is known to exist by Lemma 24 and can be found in time $\text{poly}(m)$ since $t = O(\log m)$. Then, we set $\text{Enc}'(x) = z \| [x + g(z)]$.

Observe that the redundancy of $\mathcal{C}'$ is exactly $|z| = t = O(\log m)$, and that we have encoders and decoders for $\mathcal{C}'$ running in time $\text{poly}(m)$ since $t = O(\log m)$. To see that $\mathcal{C}'$ satisfies the property required in this section, fix some substring s of $\text{Enc}'(x)$ such that $|s| = \sqrt{m}$. Then, $w(s) \geq 0.4w \cdot |s|/w - t \geq 0.39|s|$ provided that m is large enough.

Finally, we remark that the code used in this marker-based construction is just an example of a viable inner code $\mathcal{C}'$. Any *structured family* of codes satisfying Property 18 may be used instead, and one may envision adding more constraints to $\mathcal{C}'$, depending on the application constraints at hand. We exploit this fact in Sections III-B and III-C. For example, in Section III-C we will require that $\mathcal{C}'$ is a code over $\{A, C, G, T\}$ satisfying an analogue of Property 18 while also having balanced GC -content.

B. Adding a second level of markers

In our next construction, we exploit the fact that the marker-based construction from Section III-A can be instantiated with a large range of inner codes to prove Theorem 3. To do so, we show that we can iterate the marker-based construction so that we can split a trace into even smaller sub-traces with high probability. This leads to a code requiring fewer traces, but with a penalty in the redundancy. We restate Theorem 3 for convenience, where we now use n_0 to denote the message length so as to avoid confusion.

Theorem 26 (Theorem 3, restated): There exists an efficiently encodable code $\mathcal{C}_0 \subseteq \{0, 1\}^{n_0+r_0}$ with encoder $\text{Enc}_0 : \{0, 1\}^{n_0} \rightarrow \{0, 1\}^{n_0+r_0}$ and redundancy $r_0 = O(n_0/\log \log n_0)$ that can be efficiently reconstructed from $\exp(O(\log \log n_0)^{2/3})$ traces with probability at least $1 - 2/n_0$.

As before, for simplicity we set $d = 1/2$ throughout the section. We will use the same construction blueprint as in Section III-A, except for the following differences:

- We assume the n -bit message x belongs to a binary outer code $\mathcal{C}_{\text{out}} \subseteq \{0, 1\}^n$ with encoder $\text{Enc}_{\text{out}} : \{0, 1\}^{n_0} \rightarrow \{0, 1\}^n$ and relative (Hamming) distance³ $30/\log^2 n_0$. In particular, we have $x = \text{Enc}_0(x_0)$ for some $x_0 \in \{0, 1\}^{n_0}$. Such efficiently encodable and decodable codes are known to exist with redundancy $n - n_0 = O\left(n_0 \frac{\log \log n_0}{\log n_0}\right)$ (see Appendix A for a proof). The reasons for using this encoding will be made clear later;
- We consider a different instantiation of the inner code $\mathcal{C}'$ than the one in Section III-A1, which we describe below.

If $\mathcal{C}$ denotes the code obtained via the reasoning of Section III-A (with the different instantiation of the inner code $\mathcal{C}'$ that we describe below) and Enc corresponds to its encoder, then the encoder $\text{Enc}_0 : \{0, 1\}^{n_0} \rightarrow \{0, 1\}^{n_0+r_0}$ for our final code $\mathcal{C}_0$ is obtained by composing the encoders of $\mathcal{C}_{\text{out}}$ and $\mathcal{C}$, i.e.,

$$\text{Enc}_0 = \text{Enc} \circ \text{Enc}_{\text{out}}.$$

We illustrate the encoder detailed above in Figure 3.

a) *Instantiating the inner code $\mathcal{C}'$:* We proceed to describe the encoder Enc' for (our instantiation of) the inner code $\mathcal{C}'$ of $\mathcal{C}$. Essentially, our chosen instantiation of $\mathcal{C}'$ uses the encoder detailed in Section III-A and illustrated in Figure 1 along with the code constructed in Section III-A1, both with different blocklengths.

³The relative Hamming distance of a code is defined as its minimum Hamming distance normalized by its block length.

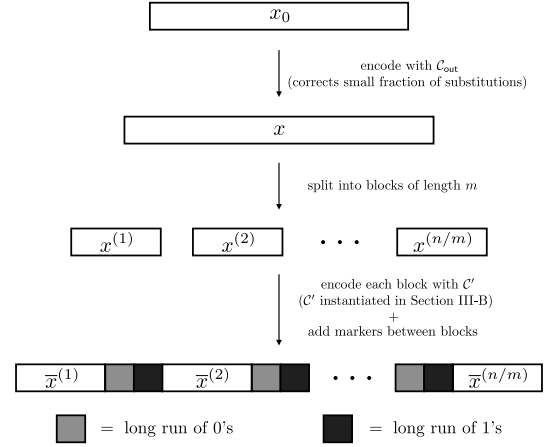


Fig. 3. The encoder for $\mathcal{C}_0$, which is used to prove Theorem 3 (see Section III-B). The only differences from the encoder described in Section III-A consist in the addition of an outer error-correcting step and the fact that we use a specific instantiation of the inner code $\mathcal{C}'$.

Given $y \in \{0, 1\}^m$, where $m = \log^2 n$, we split y into $m/\log^2 m$ blocks of length $\log^2 m$,

$$y = y^{(1)} \| y^{(2)} \| \dots \| y^{(m/\log^2 m)}.$$

Then, we take $\mathcal{C}'' \subseteq \{0, 1\}^{m'+r'}$ with encoder $\text{Enc}'' : \{0, 1\}^{m'} \rightarrow \{0, 1\}^{m'+r'}$ as the efficiently encodable and decodable code constructed in Section III-A1 with message length $m' = \log^2 m$ and redundancy $r' = O(\log m') = O(\log \log m)$. For each i , we define $\bar{y}^{(i)} = \text{Enc}''(y^{(i)})$. Moreover, we let $\ell' = 50 \log m$, and define the marker $M' = 0^{\ell'} \| 1^{\ell'}$. Then, we define $\text{Enc}'(y)$ as

$$\text{Enc}'(y) = M' \| \bar{y}^{(1)} \| M' \| \bar{y}^{(2)} \| M' \| \dots \| M' \| \bar{y}^{(m/\log^2 m)} \| M'.$$

Observe that we can efficiently decode y from $\text{Enc}'(y)$ provided that $\mathcal{C}''$ is efficiently decodable.

We first compute the redundancy of the inner code $\mathcal{C}'$ and the resulting code $\mathcal{C}$ obtained as in Section III-A. We have

$$\begin{aligned} |\text{Enc}'(y)| &= m + \frac{m}{\log^2 m} \cdot (|M'| + O(\log \log m)) \\ &= m + O\left(\frac{m}{\log m}\right). \end{aligned}$$

Thus, $\mathcal{C}'$ has redundancy $r = O(m/\log m)$. Plugging r into (4) and recalling that $m = \log^2 n$, we conclude that $\mathcal{C}$ has redundancy

$$O\left(\frac{n}{\log n}\right) + O\left(\frac{n \log^2 n}{\log^2 n \cdot \log \log n}\right) = O\left(\frac{n}{\log \log n}\right).$$

As a result, since $n = n_0 + O\left(n_0 \frac{\log \log n_0}{\log n_0}\right)$, the code $\mathcal{C}_0$ has redundancy $r_0 = O(n_0/\log \log n_0)$, as desired.

We now show that $\mathcal{C}'$ satisfies Property 18. First, we observe that $\mathcal{C}''$ satisfies Property 25 with m' in place of m . Then, since each M' has weight $0.5|M'|$, we conclude that every substring s of $\text{Enc}'(y)$ such that $|s| \geq \sqrt{m}$ satisfies

$$w(s) \geq 0.4w \cdot |s|/w - \ell' \geq 0.39|s|,$$

provided m is large enough, since $\ell' = O(\log m)$. As a result, Lemma 22 holds for this choice of inner code, and we can

hence focus solely on the trace reconstruction problem for strings of the form

$$1^\ell \|\text{Enc}'(y)\|0^\ell \\ = 1^\ell \|M'\| \bar{y}^{(1)} \|M'\| \dots \|M'\| \bar{y}^{(m/\log^2 m)} \|M'\| 0^\ell, \quad (8)$$

where $\ell = O(\log n) = O(\sqrt{m})$, and provided the number of traces used is significantly smaller than n .

b) The trace reconstruction algorithm.: Since the instantiation of $\mathcal{C}'$ above satisfies Property 18, the results from Section III-A (along with the high-level reconstruction procedure pictured in Figure 2) imply that it now suffices to design a trace reconstruction algorithm for strings of the form $1^\ell \|c'\|0^\ell$, where $c' \in \mathcal{C}'$. With this in mind, we give a reconstruction algorithm for strings of the form (8) that requires $\exp(O(\log^{2/3} m)) = \exp(O(\log \log n_0)^{2/3})$ traces and time, and succeeds with probability at least $1 - 1/\text{poly}(m) = 1 - 1/\text{poly}(\log n_0)$. We have the following two lemmas whose proofs are analogous to those of Lemmas 20 and 21 and hence omitted.

Lemma 27: Let $0^{L_0} 1^{L_1}$ be the output of the deletion channel on input M' . Then,

$$\Pr[L_0 > 10 \log m, L_1 > 0] \geq 1 - m^{-3}.$$

Lemma 28: Let $c \in \mathcal{C}''$. Then, a trace of c does not contain a run of 0's of length at least $10 \log m$ with probability at least $1 - m^{-3}$.

Combining Lemmas 27 and 28 with the union bound leads to the following analogue of Lemma 22.

Lemma 29: Consider the following event E' : We correctly identify the separation between the traces of 0^ℓ and 1^ℓ from every marker in the trace of $\text{Enc}'(x)$ by looking for all 1's that appear immediately after a run of at least $10 \log m$ 0's.

Then, E' happens with probability at least $1 - m^{-2}$ over the randomness of the trace.

As in Section III-A, Lemma 29 implies that, conditioned on E' happening for a trace T of $1^\ell \|\text{Enc}'(y)\|0^\ell$, we can split T into independent sub-traces $T^{(i)}$ each distributed like a trace of $1^\ell \|\text{Enc}'(y^{(i)})\|0^\ell$ conditioned on the high probability event E' .

Let $\mathcal{A}$ denote the worst-case trace reconstruction algorithm from Lemma 11 for strings of length $O(m') = O(\log^2 m)$, with failure probability at most $\exp(-\Omega(\log^2 m))$. A reasoning similar to that preceding (7) with Lemma 29 in place of Lemma 22, and the code $\mathcal{C}'$ designed in this section in place of $\mathcal{C}$ shows that, using algorithm $\mathcal{A}$, we fail to recover $\text{Enc}'(y)$ from $\exp(O(\log^{2/3} m))$ i.i.d. traces of $1^\ell \|\text{Enc}'(y)\|0^\ell$ with probability at most

$$m \cdot \exp(-\Omega(\log^2 m)) + 1/m < 2/m. \quad (9)$$

Let $\mathcal{A}'$ denote the algorithm that recovers $\text{Enc}'(y)$ from $\exp(O(\log^{2/3} m))$ i.i.d. traces of $1^\ell \|\text{Enc}'(y)\|0^\ell$ with failure probability at most $2/m$ as described above. We hope to instantiate (7) directly with $\mathcal{A}'$ to obtain the desired upper bound on the reconstruction failure probability for $\mathcal{C}$. However, this approach does not produce a satisfactory result as the failure probability of $\mathcal{A}'$ is $2/m = 1/\text{poly}(\log n)$, which is too large to be used in the union bound.

Recall from Section III-A (see also Figure 1) that, given $x \in \{0, 1\}^n$, the codeword $\text{Enc}(x)$ of $\mathcal{C}$ is obtained by splitting x into $n/\log^2 n$ blocks $x^{(i)}$ and encoding each block with the encoder Enc' associated with $\mathcal{C}'$. From the discussion in the previous paragraph, a fraction of blocks $x^{(i)}$ will be reconstructed with errors. Below we argue that this fraction is of size at most $10/\log^2 n_0$ with probability at least $1 - 2/n_0$. The reasoning is similar in spirit to that used to derive (6), and it suffices to complete the proof of Theorem 3. In fact, suppose we recovered $\tilde{x}$, which is a guess of x with at most a $(10/\log^2 n_0)$ -fraction of incorrect blocks. In particular, the relative Hamming distance between $\tilde{x}$ and x is at most $10/\log^2 n_0$. Since the relative distance of $\mathcal{C}_{\text{out}}$ is at least $30/\log^2 n_0$ and we assumed that $x \in \mathcal{C}_{\text{out}}$, it follows that $\text{Dec}_{\text{out}}(\tilde{x}) = \text{Dec}_{\text{out}}(x) = x_0$. Therefore, we conclude that we can recover the underlying message x_0 with probability at least $1 - 2/n_0$ from $\exp(O(\log^{2/3} m)) = \exp(O(\log \log n_0)^{2/3})$ i.i.d. traces of $\text{Enc}_0(x)$. This proves Theorem 3.

As the last step, we show that the fraction of bad blocks is small enough with high probability. Suppose that we have access to $t = \exp(O(\log^{2/3} m))$ i.i.d. traces $T_1, \dots, T_t$ of $\text{Enc}(x)$, where Enc is the encoder associated with $\mathcal{C}$. Let E denote the event from Lemma 22, and let E_{all} denote the event that E holds for all T_i simultaneously. As before, assuming that E_{all} holds, the sub-traces $(T_1^{(i)}, \dots, T_t^{(i)})_{1 \leq i \leq n/\log^2 n}$ are independent between all i , and each tuple $T_1^{(i)}, \dots, T_t^{(i)}$ is distributed as t independent traces of $1^\ell \|\text{Enc}'(x^{(i)})\|0^\ell$, each $T_j^{(i)}$ also satisfying the conditions that the first run 1^ℓ is not completely deleted, the last run 0^ℓ has length at least $10 \log n$ in the trace, and no run of 0's has length at least $10 \log n$ in the trace of $\text{Enc}'(x^{(i)})$. Denote the event that both these conditions hold for t independent traces of $1^\ell \|\text{Enc}'(x^{(i)})\|0^\ell$ by $E_{\text{split}}^{(i)}$. Invoking the trace reconstruction algorithm $\mathcal{A}'$ defined above, let $I_{\text{indFail}}^{(i)}$ denote the indicator random variable of the event that $\mathcal{A}'$ fails to recover $1^\ell \|\text{Enc}'(x^{(i)})\|0^\ell$ from t independent traces of $1^\ell \|\text{Enc}'(x^{(i)})\|0^\ell$. Taking into account the previous discussion, we let E_{fail} denote the probability that more than a $(10/\log^2 n_0)$ -fraction of blocks $x^{(i)}$ is recovered with errors. Then, we have

$$\begin{aligned} \Pr[E_{\text{fail}}] &\leq \Pr[E_{\text{fail}}, E_{\text{all}}] + \Pr[\neg E_{\text{all}}] \\ &= \Pr \left[\sum_{i=1}^{n/\log^2 n} I_{\text{indFail}}^{(i)} > \frac{n}{\log^2 n} \cdot \frac{10}{\log^2 n_0}, \forall i : E_{\text{split}}^{(i)} \right] \\ &\quad + \Pr[\neg E_{\text{all}}] \\ &\leq \Pr \left[\sum_{i=1}^{n/\log^2 n} I_{\text{indFail}}^{(i)} > \frac{n}{\log^2 n} \cdot \frac{10}{\log^2 n_0} \right] + 1/n_0. \end{aligned} \quad (10)$$

The first equality follows from the discussion in the previous paragraph, and the second inequality follows from Lemma 22 and the fact that $n > n_0$. Recalling (9), which asserts that the failure probability for $\mathcal{A}'$ is at most $2/m$, shows that

$$\Pr[I_{\text{indFail}}^{(i)}] \leq 2/m = 2/\log^2 n < 2/\log^2 n_0$$

holds for every i . Since the $I_{\text{indFail}}^{(i)}$ are independent for all i , a standard application of the Chernoff bound yields the following lemma.

Lemma 30: We have

$$\Pr \left[\sum_{i=1}^{n/\log^2 n} I_{\text{indFail}}^{(i)} > \frac{n}{\log^2 n} \cdot \frac{10}{\log^2 n_0} \right] \leq n_0^{-10}.$$

We remark that the Chernoff bound yields a stronger upper bound than the one featured in Lemma 30. However, for simplicity we use a weaker upper bound that still suffices for our needs. Combining (10) with Lemma 30 allows us to conclude that $\Pr[E_{\text{fail}}] < 2/n_0$, as desired.

C. A code for DNA-based data storage decodable from a few traces

We describe next how to adapt the ideas from Sections III-A and III-B and combine them with techniques from [48] in order to construct codes over the alphabet $\{A, C, G, T\}$ that have balanced GC -content and provably require few traces for reconstruction. As already pointed out, strings with balanced GC -content are significantly easier to synthesize than their non-balanced counterparts. Therefore, constructions accommodating this constraint are well-suited for use in DNA-based data storage.

The constructions follow those outlined in Sections III-A and III-B (see also Figures 1 and 3). The only modifications are the choice of markers and the definition of the inner code. We focus on discussing these changes and their properties within the setting of Section III-A. The full argument and the extension for the two-level marker-based construction of Section III-B follow in a straightforward manner.

We first describe the modified markers. The marker M used throughout the section is of the form $M = (AC)^\ell \|(TG)^\ell$, where $\ell = 25 \log n$ and n is the message length. Observe that this marker has the same length as the original marker in Section III-A. Moreover, M has balanced GC -content.

In order to proceed as in Section III-A we need to design an efficiently encodable and decodable inner code $\mathcal{C}' \subseteq \{A, C, T, G\}^{m'}$ with balanced GC -content which satisfies a property analogous to Property 18.

Suppose that $\mathcal{C}'$ has encoder $\text{Enc}' : \{0, 1\}^m \rightarrow \{A, C, T, G\}^{m'}$ and that $m' = m/2 + r$, where $m = \log^2 n$ as in Section III-A and r denotes the redundancy to be determined. Given the composition of M , the property we wish $\mathcal{C}'$ to satisfy is the following:

Property 31: For all $c \in \mathcal{C}'$ and substrings s of c with $|s| = \sqrt{m}$, it holds that at least $|s|/3$ symbols of s are T or G .

Similarly to Lemma 21, it can be shown that if $\mathcal{C}'$ satisfies Property 31, then with high probability a trace of $c \in \mathcal{C}'$ will not contain long runs consisting only of symbols A and C . As a result, with high probability we can easily split a trace into multiple sub-traces associated with different blocks as in Section III-A (see the high-level reconstruction procedure in Figure 2). This is accomplished by looking for all long substrings of the trace consisting only of A 's and C 's in the trace. The reason is that, with high probability, each such

substring consists of the trace of an $(AC)^\ell$ substring from a marker M possibly with some extra symbols prepended. In that case we can correctly identify the separation between the traces of $(AC)^\ell$ and $(TG)^\ell$ in all markers by looking for the first T or G after every sufficiently long substring of A 's and C 's.

We proceed to describe the encoder Enc' of the inner code $\mathcal{C}'$ that has redundancy $r = O(\log m)$. We combine a technique from [48] with the code from Section III-A1. As an additional ingredient in the construction, we require an efficiently encodable and decodable binary balanced code $\mathcal{C}_1$ with encoder $\text{Enc}_1 : \{0, 1\}^{m/2} \rightarrow \{0, 1\}^{m/2+r_1}$. Nearly-optimal constructions of such codes are known, and they have redundancy $r_1 = O(\log m)$ [49], [50]. Let $\mathcal{C}_2 \subseteq \{0, 1\}^{m/2+r_2}$ denote the code from Section III-A1 with encoder $\text{Enc}_1 : \{0, 1\}^{m/2} \rightarrow \{0, 1\}^{m/2+r_2}$ and redundancy $r_2 = O(\log m)$. By padding one of $\mathcal{C}_1$ or $\mathcal{C}_2$ appropriately, we may assume that $r_1 = r_2 = r$, i.e., that both codes have the same block length. Similarly to [48], we define the bijection $\Psi : \{0, 1\}^n \times \{0, 1\}^n \rightarrow \{A, C, G, T\}^n$ as

$$\Psi(a, b)_i = \begin{cases} A, & \text{if } (a_i, b_i) = (0, 0), \\ T, & \text{if } (a_i, b_i) = (0, 1), \\ C, & \text{if } (a_i, b_i) = (1, 0), \\ G, & \text{if } (a_i, b_i) = (1, 1). \end{cases}$$

The code $\mathcal{C}'$ is defined via an encoding $\text{Enc}' : \{0, 1\}^m \rightarrow \{A, C, G, T\}^{m/2+r}$ of the form

$$\text{Enc}'(x) = \Psi(\text{Enc}_1(x^{(1)}), \text{Enc}_2(x^{(2)})),$$

where $x = x^{(1)} \| x^{(2)} \in \{0, 1\}^{m/2} \times \{0, 1\}^{m/2}$. It is clear that decoding x from $\text{Enc}'(x)$ can be performed efficiently. We hence have the following lemma.

Lemma 32: The inner code $\mathcal{C}'$ has balanced GC -content and satisfies Property 31.

Proof: Suppose that $c = \Psi(c_1, c_2)$, where $c_1 \in \mathcal{C}_1$ and $c_2 \in \mathcal{C}_2$. To see that c has balanced GC -content, note that the number of C 's and G 's in c equals the weight of c_1 . We have $w(c_1) = |c_1|/2$ since $\mathcal{C}_1$ is a balanced code, and hence c has balanced GC -content. To verify that $\mathcal{C}$ satisfies Property 31, note that the number of T 's and G 's within a substring $c[i, j]$ equals $w(c_2[i, j])$. Since $\mathcal{C}_2$ satisfies Property 18, the proof follows. ■

Given Lemma 32, we can now proceed along the steps described in Section III-A by splitting a trace of $\mathcal{C}$ into many short sub-traces associated with different blocks, and then applying a worst-case trace reconstruction algorithm on each block. We remark that although the algorithm from Lemma 11 works for worst-case trace reconstruction over binary strings, it can be easily adapted for quaternary strings. In fact, if t traces suffice for a worst-case trace reconstruction algorithm to reconstruct a string in $\{0, 1\}^n$ with high probability, then a simple modification of this procedure recovers any quaternary string in $\{A, C, G, T\}^n$ with $2t$ traces. This is achieved by mapping the symbols in the first t traces over $\{A, C, G, T\}$ to traces over $\{0, 1\}$ according to, say, $A \mapsto 0, C \mapsto 0, G \mapsto 1, T \mapsto 1$, and the symbols in the last t traces according to $A \mapsto 0, C \mapsto 1, G \mapsto 0, T \mapsto 1$. We can now run the binary

worst-case algorithm on both sets of t traces, and recover the original string over $\{A, C, G, T\}$ from the two outputs.

Taking into account the previous discussion, applying the reasoning from Section III-A to the marker M and inner code C' defined in this section leads to Theorem 4, which we restate for completeness.

Theorem 33 (Theorem 4, restated): For every deletion probability $d < 1$, there exists an efficiently encodable code $\mathcal{C} \subseteq \{A, C, G, T\}^{n+r}$ with redundancy $r = O(n/\log n)$ and balanced GC -content that can be efficiently reconstructed from $\exp(O(\log n)^{2/3})$ traces.

Following the reasoning from Section III-B with the modified markers and C'' instantiated with the inner code C' we designed in this section proves Theorem 5, which we also restate for completeness.

Theorem 34 (Theorem 5, restated): For every constant deletion probability $d < 1$, there exists an efficiently encodable code $\mathcal{C} \subseteq \{A, C, G, T\}^{n+r}$ with redundancy $r = O(n/\log \log n)$ and balanced GC -content that can be efficiently reconstructed from $\exp(O(\log \log n)^{2/3})$ traces.

Two comments are in place regarding the choice of markers. First, the marker sequence $M = (AC)^\ell \|(TG)^\ell$ may lead to hairpin formations when single stranded DNA is used. Hairpins are double-stranded folds, but may be easily controlled through addition of urea or through temperature management. Second, repeats such as marker repeats are undesirable as they may lead to issues during DNA synthesis. To mitigate this issue, one can alternate marker sequences. For example, two valid marker options are $(AC)^\ell \|(TG)^\ell$ and $(AG)^\ell \|(TC)^\ell$, and any other marker where the sets of symbols used in each side are disjoint and C and G do not appear in the same side is appropriate for use in the construction.

Note that alternating markers in turn requires alternating the inner codes used between markers. This can be accommodated in a straightforward manner. Suppose that the block $x^{(i)}$ precedes an $(AC)^\ell \|(TG)^\ell$ marker. Then, we encode $x^{(i)}$ as usual with Enc' as defined in this section. However, if $x^{(i)}$ precedes an $(AG)^\ell \|(TC)^\ell$ marker, then we encode $x^{(i)}$ by first computing $\text{Enc}'(x^{(i)})$, and then swapping all G 's and C 's in the encoding. Observe that in both cases the encoding has balanced GC -content. Moreover, since C' satisfies Property 31, with high probability the trace of each block's encoding will not have long substrings containing only A 's and C 's (resp. A 's and G 's) before an $(AC)^\ell \|(TG)^\ell$ marker (resp. $(AG)^\ell \|(TC)^\ell$ marker). As before, this means that, with high probability, we can correctly split the full trace into the relevant sub-traces by alternately looking for long substrings composed of A 's and C 's only, and of A 's and A 's and G 's only. In fact, the end of such long substrings corresponds to the beginning of the traces of the $(TG)^\ell$ and $(TC)^\ell$ substrings of the marker, respectively.

IV. REDUCING THE NUMBER OF TRACES FOR SMALL CONSTANT DELETION PROBABILITY

In Section III-A, we gave a construction of marker-based codes that requires a few traces for reconstruction. A simple property of the inner code ensured that we can correctly

identify all markers with high probability, effectively dividing the global trace into many independent, shorter traces (see Figures 1 and 2). After this, we applied the state-of-the-art worst-case trace reconstruction algorithm from Lemma 11 on each short trace in order to obtain the desired codes.

It seems plausible, however, that one could design the inner code more carefully so that many fewer traces are needed to recover the short codewords contained between the markers. This is the main problem we address in what follows. We design a code that, when used as the inner code in the construction from Section III-A, leads to an almost exponential reduction of the number of traces required for reconstruction with only a slight decrease in the code rate, provided that the deletion probability is a sufficiently small constant. The trace reconstruction algorithm we use is a variation of the algorithm for average-case trace reconstruction described in [2, Section 2.3].

Our starting point is a low redundancy code with the property that it can be reconstructed from $\text{poly}(n)$ traces. We discuss this construction in Section IV-A. Then, in Section IV-B we show how to adapt this code so that it can be successfully used as an inner code in the marker-based construction introduced in Section III.

A. Low redundancy codes reconstructible from polynomially many traces

In what follows, we prove Theorem 6. We restate the result for convenience.

Theorem 35 (Theorem 6, restated): For small enough deletion probability d , there exists an efficiently encodable code $\mathcal{C} \subseteq \{0, 1\}^{n+r}$ with encoder $\text{Enc} : \{0, 1\}^n \rightarrow \{0, 1\}^{n+r}$ and redundancy $r = O(\log n)$ that can be efficiently reconstructed from $\text{poly}(n)$ traces with probability at least $1 - \exp(-n)$. The code we construct to prove Theorem 6 will be the starting point for the proof of Theorem 7 in Section IV-B. Roughly speaking, our code encodes n -bit messages into codewords that are *almost* w -subsequence-unique for $w = O(\log n)$, in the sense that all but the first $O(\log n)$ bits of the codeword comprise a w -subsequence-unique string. This is possible because an ϵ -almost k -wise independent random variable over $\{0, 1\}^n$ with the appropriate parameters is w -subsequence-unique with high probability. We make this statement rigorous in the following lemma. We note that the technique in the lemma below has already been used in [47] to obtain strings satisfying related properties, such as substring-uniqueness, with high probability.

Lemma 36: Let $g : \{0, 1\}^t \rightarrow \{0, 1\}^m$ be the function guaranteed by Lemma 9 with $k = 3w$ and $\epsilon = 2^{-10w}$ for $w = 100 \log m$ (hence $t = O(\log m)$). Fix some $x \in \{0, 1\}^m$ and define the random variable $Y = x + g(U_t)$. Then, with probability at least $1 - 1/\text{poly}(m)$ it holds that Y is w -subsequence-unique.

Proof: First, note that Y is ϵ -almost k -wise independent. This proof follows along the same lines as the proof that a random string is w -subsequence-unique with high probability found in [2, Lemma 2.6] with a few simple modifications.

Without loss of generality, fix a and b such that $a < b$, and fix distinct indices $i_1, \dots, i_w \in [b, \dots, b + 1.1w)$. For

convenience, let $\mathcal{S} = \{i_1, \dots, i_w\}$, $\mathcal{S}' = [b, b + 1.1w) \setminus \mathcal{S}$, and $u = \min(a + w, b)$. Then,

$$\begin{aligned} & \Pr[Y_{\mathcal{S}} = Y[a, a + w)] \\ &= \sum_{y, y'} \Pr[Y_{\mathcal{S}} = Y[a, a + w], Y[a, u] = y, Y_{\mathcal{S}'} = y']. \end{aligned} \quad (11)$$

We now show that $Y[a, u]$ and $Y_{\mathcal{S}'}$ completely determine $Y_{\mathcal{S}}$ under the constraint $Y_{\mathcal{S}} = Y[a, a + w]$. This can be seen by induction. First, we must have $Y_{i_1} = Y_a$, and Y_a is determined by $Y[a, u]$ since $a < u$. Now, suppose that $Y_{i_1}, \dots, Y_{i_j}$ are determined by $Y[a, u]$ and $Y_{\mathcal{S}'}$. It must be the case that $Y_{i_{j+1}} = Y_{a+j}$. If $a + j < u$ or $a + j \in \mathcal{S}'$, then $Y_{i_{j+1}}$ is determined by $Y[a, u]$ or $Y_{\mathcal{S}'}$, respectively. On the other hand, if $a + j \geq u$ and $a + j \notin \mathcal{S}'$, then $Y_{a+j} = Y_{i_d}$ for some $d < j + 1$. By the induction hypothesis, Y_{i_d} is determined by $Y[a, u]$ and $Y_{\mathcal{S}'}$, and hence $Y_{i_{j+1}}$ is, too.

As a result, we conclude that there exists a string $\bar{y} = (\bar{y}_1, \dots, \bar{y}_w)$ completely determined by y and y' such that

$$\begin{aligned} & \Pr[Y_{\mathcal{S}} = Y[a, a + w], Y[a, u] = y, Y_{\mathcal{S}'} = y'] \\ &= \Pr[Y_{\mathcal{S}} = \bar{y}, Y[a, u] = y, Y_{\mathcal{S}'} = y']. \end{aligned} \quad (12)$$

Since Y is ϵ -almost $3w$ -wise independent and fewer than $3w$ coordinates are fixed, we have

$$\Pr[Y_{\mathcal{S}} = \bar{y}, Y[a, u] = y, Y_{\mathcal{S}'} = y'] \leq 2^{-1.1w - (u-a)} + 2^{3w\epsilon} \quad (13)$$

for all y and y' . Combining (11), (12), and (13), we conclude that

$$\begin{aligned} \Pr[Y_{\mathcal{S}} = Y[a, a + w]] &\leq 2^{u-a} \cdot 2^{0.1w} (2^{-1.1w - (u-a)} + 2^{3w\epsilon}) \\ &\leq 2^{-w} + 2^{4.1w\epsilon} \\ &\leq 2^{-w+1}, \end{aligned}$$

since $u-a \leq w$ and $\epsilon = 2^{-10w}$. Since there are $\binom{1.1w}{w}$ choices for $\mathcal{S}$ for each pair (a, b) and fewer than m^2 possible pairs (a, b) , the probability that Y is not w -subsequence-unique is at most

$$\begin{aligned} m^2 \binom{1.1w}{w} 2^{-w+1} &= n^2 \binom{1.1w}{w} 2^{-w+1} \\ &\leq m^2 (11e)^{0.1w} 2^{-w+1} \\ &\leq 2m^2 (1.415)^{-w} \\ &\leq m^{-45}, \end{aligned}$$

as desired. $\blacksquare$

Lemma 36 naturally leads to a simple, efficient candidate construction of the encoder Enc : Given $x \in \{0, 1\}^n$, we first iterate over all $z \in \{0, 1\}^t$ until we find z such that $x + g(z)$ is w -subsequence-unique. Most strings z satisfy this, according to Lemma 36. Moreover, since $t = O(\log n)$, we can iterate over all such z in time $\text{poly}(n)$, and verify whether $x + g(z)$ is w -subsequence-unique for each z in $\text{poly}(n)$ time. The latter can be accomplished simply by exhaustive search over all valid pairs of substrings $[a, a + w)$ and $[b, b + 1.1w)$, since $w = O(\log n)$. To recover x from $x + g(z)$ we need to provide z to the receiver. Therefore, the encoder Enc for $\mathcal{C}$ maps a message $x \in \{0, 1\}^n$ to the codeword

$$\text{Enc}(x) = z \parallel [x + g(z)] \in \{0, 1\}^{n+t}, \quad (14)$$

where z is the first string (in lexicographic order) such that $x + g(z)$ is w -subsequence-unique. Observe that the redundancy of $\mathcal{C}$ is exactly $t = O(\log n)$.

1) *The trace reconstruction algorithm*: We describe next an efficient trace reconstruction algorithm for $\mathcal{C}$ that works whenever the deletion probability is a small enough constant, thus proving Theorem 6. This algorithm works very similarly to the one introduced in [2] and described in Section II-D2. As before, we shall set $w = 100 \log n$, $v = w/d = O(\log n)$ and $j = (v - 0.1w)(1 - 3d) = O(\log n)$. Given a codeword $c = \text{Enc}(x) = z \parallel [x + g(z)]$, we proceed as follows: First, we apply the algorithm from Lemma 15 to recover z and the first $2v + w = O(\log n)$ bits of $y = x + g(z)$ with $\text{poly}(n)$ traces (repeating the process $O(n)$ times if necessary) and success probability $1 - \exp(-\Omega(n))$.

Now, suppose that we know $y_1, \dots, y_{i-1}$ for $i-1 \geq 2v + w$. We show how to find y_i with probability $1 - \exp(-\Omega(n))$ from $\text{poly}(n)$ traces, which concludes the proof of Theorem 6.

Let T denote a trace of c . As in Section II-D2, we will look for a matching of $y[i - v - w, i - v)$ within T . However, we shall discard matchings that occur too early in T . More precisely, suppose that $y[i - v - w, i - v)$ is matched with $T[u - w, u)$. We call such a matching *good* if $u - w > |z|$. If T does not contain a good matching of $y[i - v - w, i - v)$, we discard it. Otherwise, if the first good matching occurs at $T[u - w, u)$, we let $V = T[u, \cdot]$ and discard the remaining bits of T . Our observations so far are summarized in the following lemmas.

Lemma 37: For d small enough, the probability that a good matching occurs in T is at least $2^{-(w+1)}$.

Proof: First, observe that the probability that no bit in $y[i - v - w, i - v)$ is deleted is exactly $(1 - d)^w \geq 2^{-w}$. Given this, suppose that $y[i - v - w, i - v)$ shows up in positions $T[P - w, P)$, for some random variable $P \in [n]$. Then, the probability that the given matching is good equals $\Pr[P > |z| + w]$, and $|z| + w \leq Cw$ for a fixed constant $C > 0$, since $|z| = O(\log n)$. Note that we may assume $i - v - w \geq v = w/d$ since we have already learned the first $2v + w$ bits of y . We may also choose $d < 1/10$ small enough such that $v > Cw$. Then, we have

$$\Pr[P \leq |z| + w] \leq \Pr[\text{Bin}(2Cw, 1 - d) \leq Cw] < 1/2,$$

where the last inequality follows from an application of the Chernoff bound. Concluding, the trace T contains a good matching with probability at least $1/2 \cdot 2^{-w} = 2^{-(w+1)}$. $\blacksquare$

Lemma 38: The probability that the last bit of a good matching in T does not come from $y[i - v, i - v + 0.1w)$ is at most $nd^{-w/100} \leq 2^{-100w}$ if d is small enough.

Proof: The probability that the event in question happens is at most the probability that more than $0.1w$ bits are deleted from some substring $y[b, b + 1.1w)$. To see this, first note that the bits in a good matching must come from y . If at most $0.1w$ bits are deleted from every substring $y[b, b + 1.1w)$, then the w bits of the good matching in T for $y[i - v - w, i - v)$ must be a subsequence of $y[b, b + 1.1w)$ for some b , which means $y[i - v - w, i - v)$ appears as a subsequence of $y[b, b + 1.1w)$. Since y is w -subsequence-unique, for this to happen we must have $b \leq i - v - w$ and $b + 1.1w \geq i - v$. Now suppose

that the last bit of the good matching in T does not come from $y[i - v, i - v + 0.1w]$. Then, it must be the case that $y[i - v - w, i - v]$ is a subsequence of $y[b, i - v - 1]$. Since $i - v - 1 < i - v$, this violates the w -subsequence-uniqueness property of y .

For a fixed b , the probability that more than $0.1w$ bits are deleted from $y[b, b + 1.1w]$ is at most $d^{-w/100}$ for d small enough. The result then follows by a union bound, since there are fewer than n choices for b . ■

Let E_{good} denote the event that a good matching occurs in T . From Lemma 37 and the fact that we can efficiently check whether E_{good} occurred for T , it follows that we can efficiently estimate

$$\Pr[V_j = 1|E_{\text{good}}]$$

to within an error of, say, 2^{-100w} from $\text{poly}(n)$ traces, with probability at least $1 - \exp(-\Omega(n))$. Then, we proceed similarly to Section II-D2. Let R be the random variable denoting the coordinate in y of the last bit appearing in the good matching within T . We may then write

$$\begin{aligned} \Pr[V_j = 1|E_{\text{good}}] &= \sum_{r=1}^n \Pr[R = r|E_{\text{good}}] \Pr[V_j = 1|R = r, E_{\text{good}}] \\ &= \epsilon_i(c) + \sum_{r=i-v}^{i-v+0.1w} \Pr[R = r|E_{\text{good}}] \Pr[V_j = 1|R = r] \end{aligned}$$

for $0 \leq \epsilon_i(c) \leq 2^{-100w}$, by Lemma 38. The second equality follows because, once $R = r$ is fixed, V does not depend on whether E_{good} occurs or not, but only depends on the traces of z and $y[1, r]$. Therefore, as in (2) and (3) we have

$$\begin{aligned} y_i = 0 &\implies \Pr[V_j = 1|E_{\text{good}}] \leq \\ &\epsilon_i(c) + \sum_{r=i-v}^{i-v+0.1w} \Pr[R = r|E_{\text{good}}] \sum_{\ell=r+1}^{i-1} P(\ell - r, j) s_\ell \\ &\quad + \frac{1}{2} \sum_{r=i-v}^{i-v+0.1w} \Pr[R = r|E_{\text{good}}] P(i - r, j) \quad (15) \end{aligned}$$

and

$$\begin{aligned} y_i = 1 &\implies \Pr[V_j = 1|E_{\text{good}}] \geq \\ &\epsilon_i(c) + \sum_{r=i-v}^{i-v+0.1w} \Pr[R = r|E_{\text{good}}] \sum_{\ell=r+1}^{i-1} P(\ell - r, j) s_\ell \\ &\quad + \sum_{r=i-v}^{i-v+0.1w} \Pr[R = r|E_{\text{good}}] P(i - r, j). \quad (16) \end{aligned}$$

Similarly to what was done in Section II-D2, since $i - r \leq v$ and $v = w/d$, the second part of Lemma 14 implies that $P(i - r, j) \geq 2^{-9w}$. Combining this result with Lemma 38 shows that the gap between the right hand sides of (15) and (16) is at least $2^{-(9w+1)}$. Each term $\Pr[R = r|E_{\text{good}}]$ can be approximated to within an error of 2^{-90w} with probability at least $1 - \exp(-\Omega(n))$ in time $\text{poly}(n)$. This is accomplished by first using z and the values $y_1, \dots, y_{i-1}$ that we have already recovered to estimate $\Pr[R = r|E_{\text{good}}, R < i]$ to

within a small enough error and with high probability. Then, the fact that $\Pr[R < i|E_{\text{good}}] \geq 1 - 2^{-100w}$ and Lemma 38 imply that

$$|\Pr[R = r|E_{\text{good}}] - \Pr[R = r|E_{\text{good}}, R < i]| \leq 2 \cdot 2^{-100w},$$

which in turn implies a good enough approximation for $\Pr[R = r|E_{\text{good}}]$.

Since we know $y_1, \dots, y_{i-1}$, the discussion above suggests that we can approximate the right hand side of (15) and (16) to within an error of

$$2^{-100w} + n^2 \cdot 2^{-90w} \leq 2^{-80w}$$

with high probability. As already mentioned, we can estimate $\Pr[V_j = 1|E_{\text{good}}]$ to within error 2^{-100w} from $\text{poly}(n)$ traces in time $\text{poly}(n)$ with probability at least $1 - \exp(-\Omega(n))$. Consequently, with probability $1 - \exp(-\Omega(n))$ we can recover y_i correctly from $\text{poly}(n)$ traces, where the degree of this polynomial is independent of i . The success probability can be made at least $1 - \exp(-Cn)$ for any fixed constant C of our choice by repeating the process $O(n)$ times and taking the majority vote. Overall, we must recover fewer than n positions of y , and each position requires $\text{poly}(n)$ traces, where the degree of this polynomial is independent of the position of the bit. As a result, the total number of traces required is $\text{poly}(n)$ and the overall success probability is $1 - 1/\text{poly}(n)$. This proves Theorem 6.

B. Using the code within a marker-based construction

Next, we combine the constructions from Sections III-A and IV-A with some additional modifications in order to prove Theorem 7, which we restate here.

Theorem 39 (Theorem 7, restated): For small enough deletion probability, there exists an efficiently encodable code $\mathcal{C}$ with encoder $\text{Enc} : \{0, 1\}^n \rightarrow \{0, 1\}^{n+r}$ and redundancy $r = O\left(\frac{n}{\log n}\right)$ that can be efficiently reconstructed from $\text{poly}(\log n)$ traces with probability $1 - 1/\text{poly}(n)$.

The basic idea is that we would like to use the code designed in Section IV-A as the instantiation of the inner code $\mathcal{C}'$ in the construction of $\mathcal{C}$ in Section III-A (see Figure 1). Then, we could combine the high-level trace reconstruction procedure illustrated in Figure 2 with the trace reconstruction algorithm from Section IV-A1 on each sub-trace and mitigate the use of worst-case trace reconstruction algorithms. This idea does not work as is, but some modifications to the code from Section IV-A will allow the construction to go through.

The first issue we have to address is for the inner code $\mathcal{C}'$ to satisfy Property 18. If this property holds, then the reasoning of Section III-A (see also Figure 2) implies that we can focus on the trace reconstruction problem for strings of the form

$$1^\ell \|c\| 0^\ell, \quad (17)$$

where $c \in \mathcal{C}'$ has length $O(\log^2 n)$ and $\ell = O(\log n)$, as long as we use a sub-polynomial number of traces in n . From here onwards we focus solely on this setting. If we were to directly apply the trace reconstruction algorithm from Section IV-A1, we would run into a problem. For the aforementioned algorithm to work, we need to bootstrap it by

recovering the first few bits of c using the procedure described in Lemma 15. However, in this case c only appears after a run of length $\ell = O(\log n)$. Even though we know the previous bits, we still require $\text{poly}(n)$ traces to recover the first bit of c in this way, which is not acceptable as we want to use $\text{poly}(\log n)$ traces. Consequently, we need an alternative bootstrapping method. Another issue we need to resolve is that the reconstruction algorithm from Section IV-A1 assumed that all but the first few bits of c lead to a subsequence-unique string. However, this is not the case here, as we must deal with a string of the form $c\|0^\ell$.

Before we proceed to describe a modified version of our code from Section IV-A that avoids the issues raised above, we first prove the following lemma.

Lemma 40: Let $g : \{0,1\}^t \rightarrow \{0,1\}^m$ be the function guaranteed by Lemma 9 with $k = 3w$ and $\epsilon = 2^{-10w}$ for $w = 100 \log m$ (hence $t = O(\log m)$). For arbitrary ℓ and $x \in \{0,1\}^m$, define the random variable $Y = [x + g(U_t)]\|0^\ell$. Then, with probability at least $1 - 1/\text{poly}(m)$ we have that Y satisfies the following property.

Property 41: For any a and b such that $a + w \leq \min(m + 1, b)$, we have that $Y[a, a + w]$ is not a subsequence of $Y[b, b + 1.1w]$.

Proof: Fix a pair (a, b) satisfying $a + w \leq \min(m + 1, b)$ and let $\mathcal{S} \subseteq [b, b + 1.1w]$ be a set of size w . Let $u = \min(m + 1, b + 1.1w)$. Then, we have

$$\Pr[Y[a, a + w] = Y_{\mathcal{S}}] = \sum_y \Pr \left[\begin{matrix} Y[a, a + w] = Y_{\mathcal{S}}, Y[b, u] = y, \\ Y[m + 1, b + 1.1w] = (0, \dots, 0) \end{matrix} \right].$$

Observing that $Y_{\mathcal{S}}$ is completely determined by $Y[b, u]$ and $Y[m + 1, b + 1.1w]$ and that $Y[m + 1, b + 1.1w]$ is fixed, we have

$$\Pr[Y[a, a + w] = Y_{\mathcal{S}}] = \sum_y \Pr \left[\begin{matrix} Y[a, a + w] = y', \\ Y[b, u] = y \end{matrix} \right],$$

for some y' determined by y . Since $x + g(U_t)$ is ϵ -almost $3w$ -wise independent and fewer than $3w$ coordinates are fixed, we have

$$\Pr[Y[a, a + w] = y', Y[b, u] = y] \leq 2^{-w-(u-b)} + 2^{3w}\epsilon.$$

Therefore, it follows that

$$\begin{aligned} \Pr[Y[a, a + w] = Y_{\mathcal{S}}] &\leq 2^{u-b}(2^{-w-(u-b)} + 2^{3w}\epsilon) \\ &\leq 2^{-w} + 2^{4.1w}\epsilon \\ &\leq 2^{-w+1}. \end{aligned}$$

Since there are fewer than m^3 choices for pairs (a, b) and $\binom{1.1w}{w}$ choices for $\mathcal{S}$, from the union bound we conclude similarly to what we did in the proof of Lemma 36 that the probability that the desired event does not occur is at most $m^3 \binom{1.1w}{w} 2^{-w+1} \leq m^{-45}$. ■

Intuitively, Lemma 40 guarantees that $x + g(U_t)$ satisfies a stronger form of subsequence-uniqueness with high probability. In fact, not only is $x + g(U_t)$ w -subsequence-unique with high probability based on Lemma 36, but also it is impossible to find a substring of $x + g(U_t)$ that is a subsequence of $[x + g(U_t)]\|0^\ell$ elsewhere.

We are now ready to describe our modified inner code $\mathcal{C}'$ with encoder $\text{Enc}' : \{0,1\}^m \rightarrow \{0,1\}^{m+r'}$. On an input message $x \in \{0,1\}^m$, Enc' operates as follows:

- 1) Set $x' = 0^{\ell'}\|x$ for $\ell' = 10\ell = O(\sqrt{m})$. Let $m' = |x'|$ and set $w = 100 \log m'$;
- 2) Iterate over all $z \in \{0,1\}^t$ for $t = O(\log m') = O(\log m)$ until a z such that $x' + g(z)$ is w -subsequence-unique and simultaneously satisfies Properties 25 and 41 is found. Such a string z is guaranteed to exist because all such properties hold for $x' + g(U_t)$ with probability $1 - o(1)$ (see Lemmas 24, 36, and 40). Moreover, whether $x' + g(z)$ satisfies all three properties can be checked in time $\text{poly}(m)$;
- 3) Obtain z' from z by setting $z' = \text{Enc}_{\text{edit}}(0\|z)$, where Enc_{edit} is the encoder of the systematic code $\mathcal{C}_{\text{edit}}$ from Lemma 10 robust against $|z|/2$ edit errors and with redundancy $O(|z|) = O(\log m)$. Here, d is assumed to be a small enough constant so that $5d|z'| < |z|/2$, i.e., $\mathcal{C}_{\text{edit}}$ can correct a $5d$ -fraction of edit errors in z' . This is possible because $|z'| = O(|z|)$;
- 4) Define $\text{Enc}'(x) = z'\|x' + g(z) = z'\|y'$.

For a given message $x \in \{0,1\}^m$, we can compute $\text{Enc}'(x)$ in time $\text{poly}(m)$. Furthermore, recalling that $m = \log^2 n$ in the construction of Section III-A, the redundancy of $\mathcal{C}'$ is

$$r' = |z'| + \ell' = O(\log m + \sqrt{m}) = O(\sqrt{m}) = O(\log n).$$

If we use $\mathcal{C}'$ as the inner code in the construction of $\mathcal{C}$ from Section III-A, then according to (4) we obtain an overall redundancy $r = O\left(\frac{n}{\log n}\right)$ for $\mathcal{C}$, as desired. It is also easy to see that $\mathcal{C}'$ satisfies Property 18. By the choice of z , we have $w(y'[a, a + w]) \geq 0.4w$ for every a and $w = 100 \log m'$. Therefore, for any substring s such that $|s| = \sqrt{m}$ we have

$$w(s) \geq 0.4|s| - |z'| \geq 0.39|s|$$

provided that m is large enough, since $|z'| = O(\log m)$. As a result, the reasoning used in Section III-A applies to this choice of $\mathcal{C}'$. To prove Theorem 7, it remains to give a trace reconstruction algorithm to recover strings of the form $1^\ell\|\text{Enc}'(x)\|0^\ell$ from $\text{poly}(m) = \text{poly}(\log n)$ traces with probability, say, $1 - n^{-10}$.

To address the problem, suppose we already have such an algorithm, and call it $\mathcal{A}$. Recall (7) and the definition of the event $E_{\text{indFail}}^{(i)}$ from Section III-A. Instantiating $E_{\text{indFail}}^{(i)}$ with algorithm $\mathcal{A}$ leads to the bound $\Pr[E_{\text{indFail}}^{(i)}] \leq n^{-10}$, for all i . Combining this observation with (7) allows us to conclude that the probability that we successfully recover $c \in \mathcal{C}$ from $\text{poly}(\log n)$ i.i.d. traces of c is at least $1 - 2/n$. Similarly to Section III-A, we can boost the success probability to $1 - 1/p(n)$ for any fixed polynomial of our choice by repeating the process $O(\log n)$ times and by taking a majority vote.

1) *The trace reconstruction algorithm:* Next, we analyze an algorithm for recovering strings of the form $1^\ell\|\text{Enc}'(x)\|0^\ell$ from $\text{poly}(m) = \text{poly}(\log n)$ traces with probability $1 - 1/\text{poly}(n)$. As discussed before, we proceed by adapting the algorithm from Section IV-A1, which in turn is a modified version of the algorithm from [2] described in Section II-D2.

The main difference between the current and the two previously discussed settings is that the original bootstrapping technique cannot be applied, as $\text{Enc}'(x)$ is enclosed by two long runs. We start by showing that the structure of Enc' allows for a simple alternative bootstrapping method.

Recall that $c = \text{Enc}'(x) = z' \| y'$, where $y' = x' + g(z)$ and the first $O(\sqrt{m})$ bits of x' are zero. Therefore, if we can recover z from a few traces of $1^\ell \| c \| 0^\ell$, then we can recover the first $O(\sqrt{m})$ bits of y' , which suffices for bootstrapping, by simply computing $g(z)$. The following lemma states that we can recover z with high probability from $O(\log n)$ traces.

Lemma 42: There is an algorithm that recovers z from $O(\log n)$ traces of $1^\ell \| c \| 0^\ell$ with probability at least $1 - n^{-10}$.

Proof: We begin by recalling that $z' = \text{Enc}_{\text{edit}}(0 \| z)$, and that $\mathcal{C}_{\text{edit}}$ is systematic. This means $z'_1 = 0$, and so with probability $1 - d$ the first 0 appearing in the trace will correspond to z'_1 .

Given a trace T of $1^\ell \| c \| 0^\ell$, we proceed as follows: Let u denote the position of the first 0 in T . Then, we take $\tilde{z} = T[u, u + (1 - d)|z'|]$, feed $\tilde{z}$ into Dec_{edit} , and let the corresponding output be our guess for z . The probability that this procedure fails to yield z is at most the probability that z'_1 was deleted, plus the probability that $\tilde{z}$ is too far away in edit distance from z' given that z'_1 was not deleted. We proceed to bound both probabilities. First, the probability that z'_1 is deleted is exactly d . Second, we assume z'_1 is not deleted and let L denote the length of the trace of $z'[2, \cdot]$ within T . We have $\mathbb{E}[L] = (1 - d)(|z'| - 1)$. Therefore, Chernoff's bound gives

$$\Pr[L \geq (1 - 3d)(|z'| - 1)] \leq \exp\left(-\frac{2d^2}{1 - d}(|z'| - 1)\right).$$

Since d is a constant and $|z'| = \Theta(\log m)$, we conclude that for m large enough we have

$$\Pr[L - (1 - d)(|z'| - 1) \geq 2d(|z'| - 1)] < 1/5.$$

As a result, with probability at least $4/5$ we have that $\tilde{z}$ is within edit distance $5d|z'| < |z|/2$ from z' . If this distance condition holds, then $\text{Dec}_{\text{edit}}(\tilde{z}) = z$.

In summary, the procedure fails to return z with probability at most $d + 1/5 < 1/4$ if d is small enough. Repeating this procedure $O(\log n)$ times and taking the majority vote ensures via the Chernoff bound that we can recover z from $O(\log n)$ traces with success probability at least $1 - 1/p(n)$, for p any choice of a fixed polynomial. ■

Once z has been recovered, the bits of $1^\ell \| c \| 0^\ell = 1^\ell \| z' \| y' \| 0^\ell$ are known up to and including the first $\ell' = O(\sqrt{m})$ bits of y' . Our last task is to recover the remaining bits of y' , and given that we have sufficiently many initial bits from y' we may to this end use the ideas from Section IV-A1. The differences with respect to Section IV-A1 are the following:

- Instead of y , we use $y'' = y' \| 0^\ell$;
- We are only interested in recovering y''_i for $\ell' < i \leq |y'|$, since we already know all other bits of y'' ;
- We change the threshold used to declare that a matching is good: In this case, if T is a trace of $1^\ell \| c \| 0^\ell$ and $y''[i - v - w, i - v]$ is matched with $T[u - w, u]$, then the matching is said to be *good* if $u - w > \ell + |z'|$. This change

ensures that the bits in a good matching always come from $y'' = y' \| 0^\ell$.

Two key lemmas now follow from the previous discussion. Their statements and proofs are similar to the ones of Lemmas 37 and 38 from Section IV-A1, respectively, and we hence only discuss relevant differences. Henceforth, we use T to denote a trace of $1^\ell \| c \| 0^\ell$.

Lemma 43: The probability that a good matching occurs in T is at least $2^{-(w+1)}$.

Lemma 44: For $\ell' < i \leq |y'|$, the probability that the last bit of a good matching in T does not come from $y''[i - v, i - v + 0.1w]$ is at most $nd^{-w/100} \leq 2^{-100w}$ if d is small enough.

Proof: Similarly to the proof of Lemma 38, the probability of the event in the statement of the lemma is upper bounded by the probability that more than $0.1w$ bits are deleted from some substring $y''[b, b + 1.1w]$. We explain next why this is true. First, note that the bits in a good matching must come from y'' . Suppose that at most $0.1w$ bits are deleted from every substring $y''[b, b + 1.1w]$. Then, $y''[i - v - w, i - v]$ must be a subsequence of $y''[b, b + 1.1w]$ for some $1 \leq b \leq |y''| - 1.1w$. We distinguish two cases:

- $b + 1.1w > |y'|$:

Recalling that $v = w/d$, we have $i - v \leq |y'| - w/d \leq |y'| - 1.1w \leq \min(|y'| + 1, b)$, and so Property 41 holds for $y''[i - v - w, i - v]$. Therefore, $y''[i - v - w, i - v]$ cannot be a subsequence of $y''[b, b + 1.1w]$ for any b such that $b + 1.1w > |y'| + 1$. Consequently, we only need to consider values of b such that $b + 1.1w \leq |y'|$;

- $b + 1.1w \leq |y'|$:

Since y' is w -subsequence-unique, we must have $b \leq i - v - w$ and $b + 1.1w \geq i - v$. This implies the desired result as in the proof of Lemma 38;

The remainder of the proof follows along the lines of the proof of Lemma 38. ■

Lemmas 43 and 44 imply that we can recover y''_i with probability $1 - 1/\text{poly}(n)$ via the same reasoning of Section IV-A1 with the small differences described above. The number of traces required to recover y''_i is polynomial in the length of $1^\ell \| c \| 0^\ell$, which equals

$$2\ell + |z'| + |y'| = O(\sqrt{m} + \log m + m) = O(m).$$

Since $m = \log^2 n$, it follows that we can recover y''_i with probability $1 - 1/\text{poly}(n)$ from $\text{poly}(\log n)$ traces. In particular, the success probability can be assumed to be at least $1 - 1/p(n)$ for a fixed polynomial of our choice since we can repeat the process $O(\log n)$ times and take the majority vote while still requiring $\text{poly}(\log n)$ traces. Since Lemma 42 asserts that $O(\log n)$ traces suffice to recover z with high probability, and we need to recover $m = \log^2 n$ bits of y'' , we overall require $\text{poly}(\log n)$ traces to recover $1^\ell \| c \| 0^\ell$ with probability $1 - 1/\text{poly}(n)$. This concludes the proof of Theorem 7.

V. OPEN PROBLEMS

The newly introduced topic of coded trace reconstruction can be extended in many different directions, in which the following problems are of interest:

- Coded trace reconstruction based on a more general set of edit errors. Our constructions only work against i.i.d. deletions. However, errors imposed by the nanopore sequencing process in DNA-based storage also include insertions and substitutions, and the errors may be symbol and context-dependent;
- Narrowing the gap between upper and lower bounds on the number of traces used in reconstruction. The current gap between lower and upper bounds for trace reconstruction is almost exponential in the string length, and there is no study of lower bounds on coded trace reconstruction as a function of the redundancy of the code. We suspect that our current construction of efficient codes are suboptimal;
- The main contributors to the redundancy in all but one of our constructions are the markers placed between blocks. It is of interest to find other constructions that either do not use markers, or that use shorter markers that can still be integrated into the trace reconstruction process;
- The parameters of our codes have hidden constants which depend on some results we invoke (namely, Lemmas 9 and 10), and we have made no effort to optimize constants in this work. It would be interesting to study coded trace reconstruction in the non-asymptotic regime.
- As discussed in Section I, given the portable DNA-based storage system presented in [14], it would be interesting to design codes which simultaneously have good clustering and trace reconstruction properties.
- As discussed in Remark 1, it would be interesting to further investigate the tradeoff between code rate, the number of traces required for coded trace reconstruction, and the reconstruction error probability allowed.

APPENDIX A

TRADEOFF BETWEEN REDUNDANCY AND RELATIVE HAMMING DISTANCE FOR EFFICIENTLY ENCODABLE/DECODABLE BINARY CODES

We establish next the existence of efficiently encodable and decodable codes $\mathcal{C}_{\text{out}} \subseteq \{0,1\}^n$ with encoder $\text{Enc}_{\text{out}} : \{0,1\}^{n_0} \rightarrow \{0,1\}^n$, relative Hamming distance at least $30/\log^2 n_0$, and redundancy $n - n_0 = O\left(n_0 \frac{\log \log n_0}{\log n_0}\right)$. We make use of the Zyablov bound [51, Section 10.2], which states that an efficiently encodable and decodable binary code with rate R and relative Hamming distance δ exists for R and δ satisfying

$$R \geq \max_{0 < r < 1 - h(\delta + \epsilon)} r \left(1 - \frac{\delta}{h^{-1}(1-r) - \epsilon} \right),$$

where $\epsilon > 0$ is arbitrary. This bound is achieved by concatenating a Reed-Solomon outer code with an inner linear code lying on the Gilbert-Varshamov bound. The relevant generator matrices can be constructed in time polynomial in the block length of the concatenated code, and one can efficiently correct substitution errors up to half of its designed distance via generalized minimum distance decoding [51, Section 11.3].

We set $\delta = 30/\log^2 n_0$ and $\epsilon = 1/\log^2 n_0$. Then,

$$h(\delta + \epsilon) = h(31/\log^2 n_0) \leq 62 \cdot \frac{\log \log n_0}{\log^2 n_0}$$

for n_0 large enough. The inequality follows from the fact that $h(p) \leq -2p \log p$ for p small enough. We set $r = 1 - \frac{\log \log n_0}{\log n_0}$ and observe that $r < 1 - h(\delta + \epsilon)$ for n_0 large enough. Moreover, we have

$$h^{-1}(1-r) = h^{-1}\left(\frac{\log \log n_0}{\log n_0}\right) \geq \frac{1}{2 \log n_0}.$$

As a result,

$$\begin{aligned} 1 - \frac{\delta}{h^{-1}(1-r) - \epsilon} &\geq 1 - \frac{\delta}{1/2 \log n_0 - 1/\log^2 n_0} \\ &= 1 - \frac{10/\log^2 n_0}{1/2 \log n_0 - 1/\log^2 n_0} \\ &\geq 1 - \frac{40}{\log n_0}. \end{aligned} \quad (18)$$

Combining (18) with the previously described choice for r , it follows that

$$R \geq r \left(1 - \frac{40}{\log n_0} \right) > 1 - \frac{2 \log \log n_0}{\log n_0},$$

and hence the corresponding redundancy is $O\left(n_0 \frac{\log \log n_0}{\log n_0}\right)$, as desired.

REFERENCES

- [1] T. Batu, S. Kannan, S. Khanna, and A. McGregor, "Reconstructing strings from random traces," in *Proc. 15th Annu. ACM-SIAM Symp. Discrete Algorithms (SODA)*, 2004, pp. 910–918.
- [2] T. Holenstein, M. Mitzenmacher, R. Panigrahy, and U. Wieder, "Trace reconstruction with constant deletion probability and related results," in *Proc. 19th Annu. ACM-SIAM Symp. Discrete Algorithms (SODA)*, 2008, pp. 389–398.
- [3] A. McGregor, E. Price, and S. Vorotnikova, "Trace reconstruction revisited," in *Proc. Eur. Symp. Algorithms (ESA)*, 2014, pp. 689–700.
- [4] A. De, R. O'Donnell, and R. A. Servedio, "Optimal mean-based algorithms for trace reconstruction," in *Proc. 49th Annu. ACM SIGACT Symp. Theory Comput. (STOC)*, 2017, pp. 1047–1056.
- [5] F. Nazarov and Y. Peres, "Trace reconstruction with $\exp(O(n^{1/3}))$ samples," in *Proc. 49th Annu. ACM SIGACT Symp. Theory Comput. (STOC)*, 2017, pp. 1042–1046.
- [6] L. Hartung, N. Holden, and Y. Peres, "Trace reconstruction with varying deletion probabilities," in *Proc. 15th Workshop Analytic Algorithmics Combinatorics (ANALCO)*, 2018, pp. 54–61.
- [7] Y. Peres and A. Zhai, "Average-case reconstruction for the deletion channel: Subpolynomially many traces suffice," in *Proc. IEEE 58th Annu. Symp. Found. Comput. Sci. (FOCS)*, Oct 2017, pp. 228–239.
- [8] N. Holden, R. Pemantle, and Y. Peres, "Subpolynomial trace reconstruction for random strings and arbitrary deletion probability," in *Proc. 31st Conf. Learning Theory (COLT)*, 2018, pp. 1799–1840.
- [9] V. I. Levenshtein, "Efficient reconstruction of sequences," *IEEE Trans. Inf. Theory*, vol. 47, no. 1, pp. 2–22, Jan 2001.
- [10] F. Sala, R. Gabrys, C. Schoeny, and L. Dolecek, "Exact reconstruction from insertions in synchronization codes," *IEEE Trans. Inf. Theory*, vol. 63, no. 4, pp. 2428–2445, April 2017.
- [11] R. Gabrys and E. Yaakobi, "Sequence reconstruction over the deletion channel," *IEEE Trans. Inf. Theory*, vol. 64, no. 4, pp. 2924–2931, April 2018.
- [12] M. Horovitz and E. Yaakobi, "Reconstruction of sequences over non-identical channels," *IEEE Trans. Inf. Theory*, vol. 65, no. 2, pp. 1267–1286, Feb 2019.
- [13] S. H. T. Yazdi, R. Gabrys, and O. Milenkovic, "Portable and error-free DNA-based data storage," *Sci. reports*, vol. 7, no. 1, p. 5011, 2017.
- [14] L. Organick, S. D. Ang, Y.-J. Chen, R. Lopez, S. Yekhanin, K. Makarychev, M. Z. Racz, G. Kamath, P. Gopalan, B. Nguyen *et al.*, "Random access in large-scale DNA data storage," *Nature biotechnology*, vol. 36, no. 3, p. 242, 2018.
- [15] C. Rashtchian, K. Makarychev, M. Racz, S. Ang, D. Jevdjic, S. Yekhanin, L. Ceze, and K. Strauss, "Clustering billions of reads for DNA data storage," in *Proc. Conf. Neural Inf. Process. Syst. (NIPS)*, 2017, pp. 3360–3371.

- [16] T. Shinkar, E. Yaakobi, A. Lenz, and A. Wachter-Zeh, "Clustering-correcting codes," in *Proc. IEEE Int. Symp. Inf. Theory (ISIT)*, July 2019, pp. 81–85.
- [17] Z. Chase, "New lower bounds for trace reconstruction," *arXiv e-prints*, p. arXiv:1905.03031, May 2019.
- [18] N. Holden and R. Lyons, "Lower bounds for trace reconstruction," *arXiv e-prints*, p. arXiv:1808.02336, Aug 2018, to appear in *Ann. Appl. Prob.*
- [19] A. Magner, J. Duda, W. Szpankowski, and A. Grama, "Fundamental bounds for sequence reconstruction from nanopore sequencers," *IEEE Trans. Mol. Biol. Multi-Scale Commun.*, vol. 2, no. 1, pp. 92–106, June 2016.
- [20] H. M. Kiah, G. J. Puleo, and O. Milenkovic, "Codes for DNA sequence profiles," *IEEE Trans. Inf. Theory*, vol. 62, no. 6, pp. 3125–3146, 2016.
- [21] R. Gabrys and O. Milenkovic, "The hybrid k -deck problem: Reconstructing sequences from short and long traces," in *Proc. IEEE Int. Symp. Inf. Theory (ISIT)*, June 2017, pp. 1306–1310.
- [22] —, "Unique reconstruction of coded sequences from multiset substring spectra," in *Proc. IEEE Int. Symp. Inf. Theory (ISIT)*, June 2018, pp. 2540–2544.
- [23] B. Haeupler and M. Mitzenmacher, "Repeated deletion channels," in *Proc. IEEE Inf. Theory Workshop (ITW)*, Nov 2014, pp. 152–156.
- [24] S. R. Srinivasavaradhan, M. Du, S. Diggavi, and C. Fragouli, "On maximum likelihood reconstruction over multiple deletion channels," in *Proc. IEEE Int. Symp. Inf. Theory (ISIT)*, June 2018, pp. 436–440.
- [25] A. Lenz, P. H. Siegel, A. Wachter-Zeh, and E. Yaakobi, "Coding over sets for DNA storage," in *Proc. IEEE Int. Symp. Inf. Theory (ISIT)*, June 2018, pp. 2411–2415.
- [26] W. Song, K. Cai, and K. A. Schouhamer Immink, "Sequence-subset distance and coding for error control for DNA-based data storage," in *Proc. IEEE Int. Symp. Inf. Theory (ISIT)*, July 2019, pp. 86–90.
- [27] J. Sima, N. Raviv, and J. Bruck, "On coding over sliced information," in *Proc. IEEE Int. Symp. Inf. Theory (ISIT)*, July 2019, pp. 767–771.
- [28] A. Lenz, P. H. Siegel, A. Wachter-Zeh, and E. Yaakobi, "Anchor-based correction of substitutions in indexed sets," in *Proc. IEEE Int. Symp. Inf. Theory (ISIT)*, July 2019, pp. 757–761.
- [29] M. Kovačević and V. Y. F. Tan, "Codes in the space of multisets—Coding for permutation channels with impairments," *IEEE Trans. Inf. Theory*, vol. 64, no. 7, pp. 5156–5169, July 2018.
- [30] R. Heckel, I. Shomorony, K. Ramchandran, and D. N. C. Tse, "Fundamental limits of DNA storage systems," in *Proc. IEEE Int. Symp. Inf. Theory (ISIT)*, June 2017, pp. 3130–3134.
- [31] I. Shomorony and R. Heckel, "Capacity results for the noisy shuffling channel," in *Proc. IEEE Int. Symp. Inf. Theory (ISIT)*, July 2019, pp. 762–766.
- [32] N. Goldman, P. Bertone, S. Chen, C. Dessimoz, E. M. LeProust, B. Sipos, and E. Birney, "Towards practical, high-capacity, low-maintenance information storage in synthesized DNA," *Nature*, vol. 494, no. 7435, p. 77, 2013.
- [33] R. N. Grass, R. Heckel, M. Puddu, D. Paunescu, and W. J. Stark, "Robust chemical preservation of digital information on DNA in silica with error-correcting codes," *Angew. Chem. Int. Ed.*, vol. 54, no. 8, pp. 2552–2555.
- [34] S. H. T. Yazdi, Y. Yuan, J. Ma, H. Zhao, and O. Milenkovic, "A rewritable, random-access DNA-based storage system," *Sci. reports*, vol. 5, p. 14138, 2015.
- [35] S. M. H. T. Yazdi, H. M. Kiah, E. Garcia-Ruiz, J. Ma, H. Zhao, and O. Milenkovic, "DNA-based storage: Trends and methods," *IEEE Trans. Mol. Biol. Multi-Scale Commun.*, vol. 1, no. 3, pp. 230–248, Sep. 2015.
- [36] O. Milenkovic, R. Gabrys, H. M. Kiah, and S. H. T. Yazdi, "Exabytes in a test tube," *IEEE Spectrum*, vol. 55, no. 5, pp. 40–45, 2018.
- [37] S. R. Srinivasavaradhan, M. Du, S. Diggavi, and C. Fragouli, "Symbol-wise MAP for multiple deletion channels," in *Proc. IEEE Int. Symp. Inf. Theory (ISIT)*, July 2019, pp. 181–185.
- [38] S. Davies, M. Z. Racz, and C. Rashtchian, "Reconstructing trees from traces," in *Proc. 32nd Conf. Learning Theory (COLT)*, 2019, pp. 961–978.
- [39] F. Ban, X. Chen, A. Freilich, R. A. Servedio, and S. Sinha, "Beyond trace reconstruction: Population recovery from the deletion channel," in *Proc. IEEE 60th Annu. Symp. Found. Comput. Sci. (FOCS)*, Nov 2019.
- [40] F. Ban, X. Chen, R. A. Servedio, and S. Sinha, "Efficient average-case population recovery in the presence of insertions and deletions," in *Proc. APPROX/RANDOM*, 2019, pp. 44:1–44:18.
- [41] A. Krishnamurthy, A. Mazumdar, A. McGregor, and S. Pal, "Trace reconstruction: Generalized and parameterized," in *Proc. 27th Annu. Eur. Symp. Algorithms (ESA)*, 2019, pp. 68:1–68:25.
- [42] M. A. Sini and E. Yaakobi, "Reconstruction of sequences in DNA storage," in *Proc. IEEE Int. Symp. Inf. Theory (ISIT)*, July 2019, pp. 290–294.
- [43] M. Abroshan, R. Venkataramanan, L. Dolecek, and A. Guillén i Fàbregas, "Coding for deletion channels with multiple traces," in *Proc. IEEE Int. Symp. Inf. Theory (ISIT)*, July 2019, pp. 1372–1376.
- [44] J. Brakensiek, R. Li, and B. Spang, "Coded trace reconstruction in a constant number of traces," *arXiv e-prints*, p. arXiv:1908.03996, Aug 2019.
- [45] N. Alon, O. Goldreich, J. Hastad, and R. Peralta, "Simple constructions of almost k -wise independent random variables," *Random Structures & Algorithms*, vol. 3, no. 3, pp. 289–304, 1992.
- [46] B. Haeupler, "An optimal document exchange protocol," in *Proc. IEEE 60th Annu. Symp. Found. Comput. Sci. (FOCS)*, Nov 2019.
- [47] K. Cheng, Z. Jin, X. Li, and K. Wu, "Deterministic document exchange protocols, and almost optimal binary codes for edit errors," in *Proc. IEEE 59th Annu. Symp. Found. Comput. Sci. (FOCS)*, Oct 2018, pp. 200–211.
- [48] S. M. H. Tabatabaei Yazdi, H. M. Kiah, R. Gabrys, and O. Milenkovic, "Mutually uncorrelated primers for DNA-based data storage," *IEEE Trans. Inf. Theory*, vol. 64, no. 9, pp. 6283–6296, Sep. 2018.
- [49] D. Knuth, "Efficient balanced codes," *IEEE Trans. Inf. Theory*, vol. 32, no. 1, pp. 51–53, January 1986.
- [50] K. A. Schouhamer Immink and J. H. Weber, "Very efficient balanced codes," *IEEE J. Sel. Areas Commun.*, vol. 28, no. 2, pp. 188–192, February 2010.
- [51] V. Guruswami, A. Rudra, and M. Sudan, *Essential Coding Theory*, 2018, draft available at <https://cse.buffalo.edu/faculty/atril/courses/coding-theory/book>.

Mahdi Cheraghchi (S'05–M'10–SM'16) is an Assistant Professor of Computer Science and Engineering at the University of Michigan–Ann Arbor. He has been on the faculty of Imperial College London from 2015 to 2019, where he maintains an honorary appointment. After completing his Ph.D. degree, he was affiliated as a post-doctoral researcher with the University of Texas at Austin (2010–11), Carnegie Mellon University (2011–13), MIT (2013–14), and the University of California, Berkeley (2015). He obtained his M.Sc. and Ph.D. degrees in computer science from Ecole Polytechnique Fédérale de Lausanne (EPFL), in 2005 and 2010, respectively, and the B.Sc. degree in computer engineering from Sharif University of Technology in 2004.

Dr. Cheraghchi is broadly interested in theoretical computer science, and his research so far has mainly focused on the interconnections between information and coding theory and theoretical computer science, sparse recovery and high-dimensional geometry, information-theoretic privacy and security, and approximation algorithms.

Ryan Gabrys received the B.S. degree in mathematics and computer science from the University of Illinois at Urbana-Champaign in 2005 and the Ph.D. degree in electrical engineering from the University of California at Los Angeles, in 2014. He is currently a Scientist with the Naval Information Warfare Center. His research interests broadly lie in the areas of theoretical computer science and electrical engineering, including coding theory, combinatorics, and communication theory.

Olgica Milenkovic is a professor of Electrical and Computer Engineering at the University of Illinois, Urbana-Champaign (UIUC), and Research Professor at the Coordinated Science Laboratory. She obtained her Masters Degree in Mathematics in 2001 and PhD in Electrical Engineering in 2002, both from the University of Michigan, Ann Arbor. Her scholarly contributions have been recognized by multiple awards, including the NSF Faculty Early Career Development (CAREER) Award, the DARPA Young Faculty Award, the Deans Excellence in Research Award, and several best paper awards. In 2013, she was elected a UIUC Center for Advanced Study Associate and Willett Scholar while in 2015 she was elected Distinguished Lecturer of the Information Theory Society. In 2018 she became an IEEE Fellow. She has served as Associate Editor of the IEEE Transactions of Communications, the IEEE Transactions on Signal Processing, the IEEE Transactions on Information Theory and the IEEE Transactions on Molecular, Biological and Multi-Scale Communications. In 2009, she was the Guest Editor in Chief of a special issue of the IEEE Transactions on Information Theory on Molecular Biology and Neuroscience.

João Ribeiro (S'18) received the B.Sc. degree in applied mathematics and computation from Instituto Superior Técnico, Lisbon, Portugal, in 2015 and the M.Sc. degree in computer science from the Swiss Federal Institute of Technology (ETH), Zurich, Switzerland, in 2017. Since October 2017 he has been a Ph.D. student at the Department of Computing, Imperial College London, UK.

Mr. Ribeiro is interested in information theory and its connections to coding theory, cryptography, and theoretical computer science.