

Compression of Samplable Sources*

Luca Trevisan[†]
 Computer Science Division
 U.C. Berkeley
 615 Soda Hall
 Berkeley, CA 94720
luca@cs.berkeley.edu

Salil Vadhan[‡]
 Division of Engineering & Applied Sciences
 Harvard University
 33 Oxford Street
 Cambridge, MA 02138
salil@eecs.harvard.edu

David Zuckerman[§]
 Department of Computer Science
 University of Texas
 1 University Station C0500
 Austin, TX 78712
diz@cs.utexas.edu

December 27, 2004

Abstract

We study the compression of polynomially samplable sources. In particular, we give efficient prefix-free compression and decompression algorithms for three classes of such sources (whose support is a subset of $\{0, 1\}^n$).

1. We show how to compress sources X samplable by logspace machines to expected length $H(X) + O(1)$.

Our next results concern flat sources whose support is in $\mathbf{P}$.

2. If $H(X) \leq k = n - O(\log n)$, we show how to compress to length $k + \delta \cdot (n - k)$ for any constant $\delta > 0$; in quasi-polynomial time we show how to compress to length $k + \text{polylog}(n - k)$ even if $k = n - \text{polylog}(n)$.
3. If the support of X is the witness set for a self-reducible $\mathbf{NP}$ relation, then we show how to compress to expected length $H(X) + 5$.

Keywords: expander graphs, arithmetic coding, randomized logspace, pseudorandom generators, approximate counting

*An extended abstract of this paper has appeared in CCC '04 [37].

[†]Supported by NSF Grant CCR-9984783, US-Israel BSF grant 2002246, a Sloan Research Fellowship and an Okawa Foundation Grant.

[‡]Supported by NSF grant CCR-0133096, US-Israel BSF grant 2002246, ONR grant N00014-04-1-0478, and a Sloan Research Fellowship. Work done in part while a fellow at the Radcliffe Institute for Advanced Study at Harvard University.

[§]Supported in part by NSF Grants CCR-9912428 and CCR-0310960, a David and Lucile Packard Fellowship for Science and Engineering, a Radcliffe Institute Fellowship, and a Guggenheim Fellowship.

1 Introduction

Data compression has been studied extensively in the information theory literature (see e.g. [8] for an introduction). In this literature, the goal is to compress a random variable X , which is called a random source. Non-explicitly, the entropy $H(X)$ is both an upper and lower bound on the expected size of the compression (to within an additive $\log n$ term). For explicit (i.e. polynomial-time) compression and decompression algorithms, this bound cannot be achieved for general sources. Thus, existing efficient data-compression algorithms have been shown to approach optimal compression for sources X satisfying various stochastic “niceness” conditions, such as being stationary and ergodic, or Markovian.

In this paper, we focus on the feasibility of data compression for sources satisfying *computational* niceness conditions, most notably efficient samplability. Goldberg and Sipser [10] were the first to study compression of sources satisfying computational, rather than stochastic, conditions. Actually, they did not explicitly discuss random sources, but focused on compressing *languages* in $\mathbf{P}$, and thus implicitly considering sources uniformly distributed on all n -bit strings in such a language.

Samplable Sources with Membership Algorithms

We extend and generalize their study. We focus on sources which are polynomially *samplable*, i.e. can be generated by a probabilistic polynomial-time algorithm. Samplability captures a very general class of sources, and it is arguably a reasonable model for probability distributions generated by various natural and man-made processes. When a distribution is not samplable, the problem of generating the distribution is computationally intractable, and this seems unlikely for “natural” sources.

The languages corresponding to the supports of samplable sources need not be in $\mathbf{P}$. Indeed, while Goldberg & Sipser show that every sparse language in $\mathbf{P}$ can be compressed at least slightly, this is unlikely to be true for all polynomially samplable distributions. In particular, as first observed by Levin, pseudorandom distributions are incompressible and, if pseudorandom generators exist, then there are polynomially samplable pseudorandom distributions. (See Section 3 for more details.)

Therefore, while seeking classes of samplable distributions that can be optimally compressed, we need to impose computational constraints that rule out the possibility of sampling pseudorandom distributions. We do this by considering sources for which membership in the support can be tested in polynomial time. We first study logspace sources, which have this property implicitly, while later we study flat sources with explicit membership algorithms.

Logspace Samplers

We first study sources that can be sampled by a sampling algorithm that uses *logarithmic space*. (As is usual when studying randomized logspace, we only allow the algorithm one-way access to the random tape.) Such sources generalize Markovian sources (which can be thought of as being sampled by a constant-space sampling algorithm). On the other hand, it is known that no such source can be pseudorandom [21].

We show the existence of a universal compression algorithm for such sources that compresses optimally, up to an additive constant, in polynomial time. The compression algorithm is universal

in the sense that it optimally compresses a source X without being given a sampler for X , and just knowing the existence of a sampling algorithm and an upper bound to the space used by the sampler.

If the sampler is known, we use *arithmetic encoding*, a well known optimal compression method that can be used on any source for which it is possible to compute the *cumulative probability* distribution of the source. Our result is then obtained by giving an algorithm for computing cumulative probabilities for sources sampled by logarithmic space algorithms. We also prove a general result showing that if optimal compression is possible for a class of samplable distributions *given the sampler* then, with only a constant additive loss, optimal compression is also possible without being given the sampler, i.e. it is possible to do *universal compression* for this class. Applying this to the result above, we obtain a universal compression algorithm for sources samplable in space $c \log n$ for any constant c .

Flat Sources with Membership Algorithms

We next consider more general samplable sources for which membership in the support can be tested in polynomial time. Without further restrictions, a membership algorithm may not be useful; for example, the support of the source could be $\{0, 1\}^n$ but some strings occur with tiny probability. We therefore require that the source be flat, i.e., uniform on its support. Observe that a membership algorithm rules out the possibility that such a distribution is pseudorandom. Indeed, the membership algorithm gives a way to distinguish the source from any other source of higher entropy.

The case of flat distributions with membership algorithms was studied by Goldberg and Sipser [10] who showed that every such source X on $\{0, 1\}^n$ could be compressed to $k + 3 \log n$ bits provided that the entropy of X is smaller than $k = n - O(\log n)$. We show how to improve the compression length to $k + \delta \cdot (n - k)$ for an arbitrarily small constant $\delta > 0$, which saves more than a constant factor in overhead if $k = n - o(\log n)$. While Goldberg and Sipser use arithmetic encoding, we use a completely different method relying on recent constructions of expander graphs with expansion close to the degree [6]. In addition, our compression algorithm is deterministic, whereas the Goldberg–Sipser algorithm is probabilistic.

In our last main result, we show that if the support of the samplable distribution forms the witness set for a *self-reducible NP* relation, then we can compress almost optimally. As a consequence, we obtain polynomial-time compression algorithms for a wide variety of combinatorial structures for which sampling algorithms are known, e.g. the set of perfect matchings in a bipartite graph [16]. Our compression algorithm computes an “approximate” arithmetic coding, using ideas underlying the proof that sampling implies approximate counting for self-reducible relations [18]. In fact, we show that, for self-reducible relations, near-optimal compression is *equivalent* to almost-uniform sampling (which in turn is known to be equivalent to approximate counting [18, 15]).

Perspective and Open Problems

There are a number of examples where the imposition of complexity-theoretic constraints on traditionally information-theoretic problems has been very fruitful. For example, modern cryptography developed and flourished out of the realization that Shannon’s classic impossibility results [34] could be bypassed via the reasonable assumption that the adversary is computationally bounded [9]. Our restriction to *samplable* sources in particular was motivated by our work in [38], where we consider the somewhat related problem of (deterministic) random extraction, in which one is given a source

of a certain entropy and wants to devise an algorithm that given a sample from the source outputs an almost uniform distribution. This deterministic randomness extraction problem was known to be impossible for general sources [32, 7], and it was known to be possible for very structured sources like Markovian sources (just like the data compression problem). In [38], it is shown that, under certain complexity assumptions, randomness extraction is possible for samplable sources. Another, earlier, work showing the promise of restricting to samplable sources is that of Lipton [23], who showed that if the distribution of errors in a channel is samplable, then it is possible to transmit information reliably even above the capacity bound. As noted above, for data compression, the class of samplable sources is still too general, and thus we have tried to impose sensible additional restrictions that are still computational in nature, yet allow for interesting positive results. However, we have by no means exhausted the possibilities, and there may be other computational constraints that are even more relevant for data compression.

Another motivation for this line of work comes from the general project of understanding information-theoretic aspects of samplable sources. The theory of pseudorandom generators is naturally one major piece of this study. But samplable sources and their information-theoretic properties have also come up in unexpected places, such as in the complete problems for statistical zero knowledge [31, 11]. Understanding the compressibility of samplable sources can contribute to this general study, as it provides another measure of the (computational) randomness in a source. Indeed Yao [41] proposed such a compressibility measure of randomness, and this is one of the several measures of computational randomness recently studied by Barak, Shaltiel, and Wigderson [4]. In the same spirit, a few years ago, Impagliazzo [13] posed an intriguing question about the relationship between the compressibility and another standard measure of computational randomness, pseudentropy. A source has *pseudentropy* at least k if it is computationally indistinguishable from some distribution having entropy at least k . A source of pseudentropy k cannot be compressed to $k - \omega(\log n)$ by an efficient algorithm, and the question is whether the converse is true for samplable distributions. That is, does low pseudentropy imply compressibility for samplable sources? This intriguing question is still an open problem. However, Wee [40] has exhibited an oracle relative to which the answer is no. Specifically, under this oracle there are samplable distributions over $\{0, 1\}^n$ of very low pseudentropy that cannot be compressed to less than $n - O(\log n)$ bits. It would be very interesting to obtain a similar result without oracles, but rather under complexity-theoretic assumptions.

Finally, the notion of compression we study is in some sense the common generalization of two other problems widely studied in the computational complexity literature — specifically “randomness condensers” (or hashing) and “resource-bounded Kolmogorov complexity”. Loosely speaking, in the study of condensers one is interested in efficient compression algorithms, with no bounds on the complexity of decompressing, while in resource-bounded Kolmogorov complexity one is interested in efficient decompression algorithms, with no bounds on the complexity of compressing.

A *lossless condenser* for a source (see e.g. [30, 36]) is a randomized procedure that, with high probability, is injective (or approximately injective) when applied to samples from the source. The output of the condenser is efficiently computable and can be seen as a compression of the sample; however, no efficient decompressing algorithm is required to exist. Condensers have been studied for their applications to randomness *extractors* [27], and no assumption is typically made on the source that they’re applied to, other than the source having bounded “min-entropy”.

Resource-bounded Kolmogorov complexity (cf. [22]) focuses on the following question: for a fixed universal Turing machine U , given a time bound t and a string x , what is the shortest encoding

y of x such that $U(y)$ will output x within t time steps? Thus, here one studies efficient decompression without the requirement that the compressed representation be computable by an efficient algorithm. For example, while the output of a pseudorandom generator is an incompressible source according to our definition, each of the possible outputs of the generator has low resource-bounded Kolmogorov complexity (because the corresponding seed s is an efficiently decompressible representation of the output $G(s)$). The study of *language compression* (see e.g. [5] for recent results and references to earlier work) focuses on the worst-case compressibility (in the above sense) for sources that are flat over an efficiently decidable support (i.e. sources with membership oracles, just as we study).

2 Preliminaries

2.1 Basic definitions

A *source* X is a probability distribution on strings of some length. We write $x \stackrel{R}{\leftarrow} X$ to indicate that x is chosen randomly according to X . We think of X as being a member of a family of distributions (i.e., a probability *ensemble*), in order for asymptotic notions to make sense. The ensemble will usually be of the form $(X_n)_{n \in \mathbb{Z}^+}$, in which case X_n will be distributed on $\{0, 1\}^n$.¹ Sometimes we will consider ensembles $(X_x)_{x \in L}$ indexed by strings in some language $L \subseteq \{0, 1\}^+$, in which case X_x will be distributed over $\{0, 1\}^{p(|x|)}$ for some polynomial p . Here $\Sigma^+ = \Sigma\Sigma^*$ is the set of strings over alphabet Σ , excluding the empty string.

We denote $X(a) = \Pr[X = a]$. The *support* of X is $\text{Sup}(X) = \{x | X(x) > 0\}$. A *flat source* is uniform on its support. U_n is the uniform distribution on $\{0, 1\}^n$.

Definition 2.1. The *entropy* of a distribution X is $H(X) = \mathbb{E}_{x \stackrel{R}{\leftarrow} X} \left[\log \left(\frac{1}{X(x)} \right) \right]$.

Here, and throughout the paper, all logs are to base 2.

2.2 Basics of compression

Definition 2.2. For functions $\text{Enc} : \Sigma^+ \rightarrow \Sigma^+$ and $\text{Dec} : \Sigma^+ \rightarrow \Sigma^+$, we say (Enc, Dec) compresses source X to length m if

1. For all $x \in \text{Sup}(X)$, $\text{Dec}(\text{Enc}(x)) = x$, and
2. $\mathbb{E}[|\text{Enc}(X)|] \leq m$.

We say that the encoding is *prefix-free* if for all $x \neq y$ in $\text{Sup}(X)$, $\text{Enc}(x)$ is not a prefix of $\text{Enc}(y)$.

All of our codes will be prefix-free. It is well known that a prefix-free encoding is “uniquely decodable”; that is, commas are not needed to send multiple samples of X .

Definition 2.3. We say source X is *compressible to length m* if there exists functions Enc and Dec such that (Enc, Dec) compresses X to length m .

¹Note that this differs from the notation used in classical information theory, where one writes X_i for an individual symbol of an infinite stochastic process $X_1, X_2, \dots$ and is concerned with compressing a prefix $(X_1, X_2, \dots, X_n)$ of this process.

The following simple lemma allows us to assume throughout that all encodings are of length at most $n + 1$.

Lemma 2.1. *If a source X_n is compressible to length m , then it is compressible to length $m + 1$ where for all $x \in \text{sup}(X_n)$, $|\text{Enc}(x)| \leq n + 1$.*

Proof. Let (Enc, Dec) compress X_n to length m . Define $\text{Enc}'(x)$ to be $0\text{Enc}(x)$ (0 concatenated with $\text{Enc}(x)$) if $|\text{Enc}(x)| < n$, and $1x$ otherwise. Then $|\text{Enc}'(x)| \leq n + 1$, $E[|\text{Enc}(X_n)|] \leq n + 1$, and there is an efficient inverse Dec' . $\square$

It is well known that a source X is compressible to length $H(X) + 1$ by a prefix-free encoding (see e.g. [8]). If the encoding is required to be uniquely decodable, then X is not compressible to length less than $H(X)$. Although the codes we construct are uniquely decodable, Definition 2.2 above is less restrictive (often called “nonsingular” compression) and allows some random variables X to be compressed to length less than $H(X)$. The biggest gap is obtained by the distribution X_n which chooses i uniformly from 0 to $n-1$ and y uniformly from $\{0, 1\}^{n-i-1}$ and outputs $0^i 1y$. The compressed string is y , which has expected length $H(X_n) - \log n$. We assume the following is known but we do not know a reference.

Lemma 2.2. *A source X_n is not compressible to length less than $H(X_n) - \lceil \log(n + 1) \rceil$*

Proof. Convert any encoding Enc to a prefix-free encoding Enc' as follows. If $|\text{Enc}(x)| \leq n$, then we define $\text{Enc}'(x) = \ell(x)\text{Enc}(x)$, where $\ell(x)$ is the number $|\text{Enc}(x)|$ written in binary, padded to length $\lceil \log(n + 1) \rceil$. If $|\text{Enc}(x)| > n$, then we define $\text{Enc}'(x) = \ell(x)x$, where $\ell(x)$ is the number $n + 1$ written in binary. Then Enc' is prefix-free, and hence uniquely decodable. The new compression length is $E[|\text{Enc}'(X_n)|] \leq E[|\text{Enc}(X_n)| + \lceil \log(n + 1) \rceil]$. By the lower bound for uniquely decodable codes, we have $E[|\text{Enc}'(X_n)|] \geq H(X_n)$. Thus, $E[|\text{Enc}(X_n)|] \geq H(X_n) - \lceil \log(n + 1) \rceil$, as desired. $\square$

We remark that, by a more careful reduction (specifically, encoding $\ell(x)$ in a prefix-free way without padding to length $\lceil \log(n + 1) \rceil$), the loss of $\log(n + 1)$ can be replaced with a term depending logarithmically on only $H(X_n)$ (rather than n).

We mention that for *flat* sources, $H(X) - O(1)$ is again a lower bound.

Lemma 2.3. *A flat source X_n is not compressible to length less than $H(X_n) - 3$.*

Proof. It suffices to show that if X_n is uniform on a set of size 2^k for some integer k , then it is not compressible to length less than $k - 2$. The optimal compression has support uniform on all strings of length less than k , plus some given string of length k . Ignoring the string of length k , this compresses X_n to length greater than

$$\frac{1}{2}(k - 1) + \frac{1}{4}(k - 2) + \frac{1}{8}(k - 3) + \dots \geq k - 2,$$

as needed. $\square$

Since tight non-explicit bounds are known, the interesting issue is *efficient* compressibility, e.g. when Enc and Dec are computed by polynomial-time algorithms. Indeed, much of the field of Data Compression is centered around understanding when this is possible. In order for efficient compressibility to make sense, we must specify how the source is presented. Ideally, the compression

algorithm is only given a random sample from the source, and does not have any global information about the source other than the fact that it comes from some class of sources:

Definition 2.4 (universal compression). Let $\mathcal{C}$ be a class of sources (i.e. class of probability ensembles X_n), and let $m = m(h, n)$ be a function. We say that (Enc, Dec) is a universal compression algorithm for $\mathcal{C}$ with compression length m if for every source $(X_n)_{n \in \mathbb{Z}^+}$ in $\mathcal{C}$, there is a constant c such that (Enc, Dec) compresses X_n to length $m(H(X_n), n) + c$.

For example, the classic Lempel–Ziv method is a universal compression algorithm with compression length $H(X) + o(n)$ for the class of stationary ergodic processes [42]. (That is, the LZ method is guaranteed to effectively compress X_n if there is a stationary ergodic process $Y_1, Y_2, \dots$ such that $X_n = (Y_1, \dots, Y_n)$.)

Since universal compression is only known for a fairly restricted class of sources (and for those, only approaches the optimal compression length asymptotically), it is also of interest to study the case when the compression algorithm may depend on the entire source (rather than a single sample). That is, the compression algorithm is given a description d_n of the source X_n , and we require that $\text{Dec}(\text{Enc}(x, d_n), d_n) = x$ for all $x \in \text{Sup}(X_n)$, and $\mathbb{E}[|\text{Enc}(X_n, d_n)|] \leq m$. In other words, $\text{Enc}'(\cdot) = \text{Enc}(\cdot, d_n)$ and $\text{Dec}'(\cdot) = \text{Dec}(\cdot, d_n)$ should form a compression algorithm for X_n in the sense of Definition 2.2.

When the source is described *explicitly* (e.g. by the list of probability masses assigned to each string x), then standard methods, such as Huffman coding (cf. [8]) compress to length $H(X) + 1$. But here the input size and the running time of the algorithm are both roughly 2^n . Thus, it is more interesting to consider the case when the source is described in some compact, *implicit* form. Then the question is which implicit representations allow for efficient compression.

One general technique for obtaining efficient compression is *arithmetic coding*, which is feasible if computing the cumulative distribution function is feasible.

Lemma 2.4 (arithmetic coding, cf. [8]). Let X be a source on Σ^n and $\prec$ a total order on $\text{Sup}(X)$. Let $F : \Sigma^n \rightarrow [0, 1]$ be the following modification of the cumulative distribution function of X : $F(x) = \sum_{a \prec x} X(a) + X(x)/2$. Define $\text{Enc}(x)$ to be the first $\lceil \log(1/X(x)) \rceil + 1$ bits of $F(x)$. Then Enc is one-to-one and monotone, and $(\text{Enc}, \text{Enc}^{-1})$ compresses X to length $H(X) + 2$. The encoding is prefix-free.

For example, if X is a Markovian source (i.e. the sequence of symbols of X form a Markov chain run for n steps), then it is known that the cumulative distribution function (with respect to the standard lexicographic order) can be computed in polynomial time, and hence so can the arithmetic coding. (See [8].) Note that since Enc is monotone, if Enc can be computed efficiently, then Enc^{-1} can also be computed efficiently by binary search. Several of our positive results will make use of arithmetic coding and variants.

Another useful fact is that it suffices to obtain a decoder which decodes correctly with high probability. Specifically, we say that (Enc, Dec) compresses a source X with *decoding error* ϵ if $\Pr[\text{Dec}(\text{Enc}(X)) \neq X] \leq \epsilon$. The following lemma shows that we can eliminate small decoding error at a small price in compression length and efficiency.

Lemma 2.5. Suppose X_n is a source on $\{0, 1\}^n$ that is compressible to length m with decoding error ϵ , by algorithms (Enc, Dec) computable in time T . Suppose further that for all $x \in \text{Sup}(X_n)$, $|\text{Enc}(x)| \geq m_0$. Then X_n is compressible to length $m + \epsilon(n - m_0) + 1$ (with zero decoding error),

by algorithms $(\text{Enc}', \text{Dec}')$ computable in time $O(T)$. If Enc gives a prefix-free encoding, then so does Enc' .

For example, if X is close (in variation distance) to a source which is highly compressible, then X itself is highly compressible.

Proof of Lemma 2.5. We construct Enc' and Dec' such that for all $x \in \text{Sup}(X_n)$, $\text{Dec}'(\text{Enc}'(x)) = x$. On input x , Enc' first checks if $\text{Dec}(\text{Enc}(x)) = x$. If so, Enc' outputs $0\text{Enc}(x)$ (0 concatenated with $\text{Enc}(x)$). If not, Enc' outputs $1x$. It is easy to see that Enc' and the natural Dec' are as required. $\square$

2.3 Randomized compression

We will also consider compression algorithms that are randomized. Here we consider two variants, one where Enc and Dec have *independent randomness* and one where they have *shared randomness*. In both cases, we measure the compression length as $\mathbb{E}[|\text{Enc}(X, R)|]$, where the expectation is taken over X and the coin tosses R of Enc . They differ in the definition of decoding error. For independent randomness, the decoding error refers to a bound on $\Pr[\text{Dec}(\text{Enc}(X, R_1), R_2) \neq X]$, whereas with shared randomness it refers to a bound on $\Pr[\text{Dec}(\text{Enc}(X, R), R) \neq X]$. Unless otherwise specified, we require that the decoding error is 0 (even with zero error, randomization could be useful in achieving a small compression length with polynomial-time algorithms). Note that zero-error randomized compression algorithms (with either shared or independent randomness) are subject to the same lower bounds on compression length as in Lemmas 2.2 and 2.3. The reason is that for each fixing of the coin tosses, the lower bounds for deterministic compression algorithms apply.

We also note that in the case of shared randomness, decoding error can be eliminated in a manner analogous to Lemma 2.5 (with the same price on compression length and efficiency):

Lemma 2.6. *Suppose X_n is a source on $\{0, 1\}^n$ that is compressible to length m with decoding error ϵ by algorithms (Enc, Dec) with shared randomness, computable in time T . Suppose further that for all $x \in \text{Sup}(X_n)$, $|\text{Enc}(x)| \geq m_0$. Then X_n is compressible to length $m + \epsilon \cdot (n - m_0) + 1$ with shared randomness and zero decoding error, by algorithms $(\text{Enc}', \text{Dec}')$ computable in time $O(T)$. If Enc gives a prefix-free encoding, then so does Enc' .*

For independent randomness, it is not clear how to eliminate decoding error (while maintaining the independence of the randomness used by Enc and Dec). However, it can be made exponentially small:

Lemma 2.7. *Suppose X_n is a source on $\{0, 1\}^n$ that is compressible to length m with decoding error ϵ by algorithms (Enc, Dec) with shared independent randomness, computable in time T . Suppose further that for all $x \in \text{Sup}(X_n)$, $|\text{Enc}(x)| \geq m_0$. Then X_n is compressible to length $m + 3\epsilon \cdot (n - m_0) + 2$ with independent randomness and decoding error 2^{-n} , by algorithms $(\text{Enc}', \text{Dec}')$ computable in time $O(n \cdot T)$. If Enc gives a prefix-free encoding, then so does Enc' .*

Proof of Lemma 2.7. We construct Enc' as follows. On input x , $\text{Enc}'(x)$ computes $y \leftarrow \text{Enc}(x)$. It then runs $O(n)$ independent executions of $\text{Dec}(y)$. If at least a .6 fraction of these executions output x , then it outputs $0y$. Otherwise, it outputs $1x$.

Before describing Dec' , we analyze the compression length of Enc' . Assume without loss of generality that $\epsilon \leq 1/3$. Then by Markov's inequality, with probability at least $1 - 3\epsilon$ over $x \stackrel{R}{\leftarrow} X_n$ and the coin tosses r_1 of Enc , we have $\Pr_{R_2}[\text{Dec}(\text{Enc}(x, r_1), R_2) \neq x] \leq 1/3$. For each such x and r_1 , Enc' will output $1x$ with probability at most 2^{-n} (by a Chernoff bound). Thus, the probability that Enc' outputs $1x$ rather than $0\text{Enc}(x, r_1)$ is at most $3\epsilon + 2^{-n}$. This implies that the average compression length increases by at most $(3\epsilon + 2^{-n}) \cdot (n - m_0) + 1 \leq 3\epsilon \cdot (n - m_0) + 2$.

Now we describe Dec' . On an input of the form $1x$, Dec' outputs x . On an input of the form $0y$, Dec' runs $O(n)$ independent executions of $\text{Dec}(y)$ and outputs the value that appears most often (breaking ties arbitrarily).

Note that decoding errors can only occur when $\text{Enc}'(x)$ outputs a compressed string of the form $0y$, for $y = \text{Enc}(x, r_1)$. For any x, r_1 , we consider two cases. If $\Pr[\text{Dec}(y) = x] \geq .55$, then by a Chernoff bound, Dec' will decode correctly with probability at least $1 - 2^{-n}$. If $\Pr[\text{Dec}(y) = x] \leq .55$, then by a Chernoff bound, Enc' will output $1x$ with probability at least $1 - 2^{-n}$. Thus the decoding error is at most 2^{-n} . $\square$

Finally, we observe that randomized compression algorithms can be converted into deterministic ones at a small cost, under plausible complexity assumptions.

Lemma 2.8. *Suppose there is a function in $\mathbf{E} = \mathbf{DTIME}(2^{O(n)})$ of circuit complexity $2^{\Omega(n)}$. Then for every polynomial-time compression algorithm (Enc, Dec) with shared randomness there exists a deterministic polynomial-time compression algorithm $(\text{Enc}', \text{Dec}')$ such that for every source X_n , if (Enc, Dec) compresses X to length $m = m(H(X_n), n)$, then $(\text{Enc}', \text{Dec}')$ compresses X_n to length $m + O(\log n)$. If Enc gives a prefix-free encoding, then so does Enc' .*

Proof. Let $t(n)$ be a bound on the running time of (Enc, Dec) on inputs of length n . Under the hypothesis, there is a pseudorandom generator $G : \{0, 1\}^{\ell(n)} \rightarrow \{0, 1\}^{t(n)}$ with $\ell(n) = O(\log n)$ such that no circuit of size $t(n)$ can distinguish the output of G from uniform with advantage greater than $\epsilon = 1/t(n)$ [26, 14]. We define $\text{Enc}'(x)$ to be the shortest string in the set $\{s \circ \text{Enc}(x, G(s)) : s \in \{0, 1\}^{\ell(n)}\}$, where $\circ$ denotes concatenation. Now set $\text{Dec}'(s \circ y) = \text{Dec}(y, G(s))$. By inspection, $\text{Dec}'(\text{Enc}'(x)) = x$ for all x .

For the compression length, the pseudorandom property of G implies that for every string $x \in \{0, 1\}^n$,

$$\begin{aligned} \mathbb{E}_S[|\text{Enc}(x, G(S))|] &\leq \mathbb{E}_R[|\text{Enc}(x, R)|] + t(n) \cdot \epsilon \\ &= \mathbb{E}_R[|\text{Enc}(x, R)|] + 1. \end{aligned}$$

Thus,

$$\begin{aligned} \mathbb{E}[|\text{Enc}'(X_n)|] &= \mathbb{E}_{X_n}[\min_s |s \circ \text{Enc}(X_n, G(s))|] \\ &= \mathbb{E}_{X_n}[\min_s |\text{Enc}(X_n, G(s))|] + O(\log n) \\ &\leq \mathbb{E}_{X_n}[\mathbb{E}_S[|\text{Enc}(X_n, G(S))|]] + O(\log n) \\ &\leq \mathbb{E}_{X_n}[\mathbb{E}_R[|\text{Enc}(X_n, R)|] + 1] + O(\log n) \\ &\leq m(H(X_n), n) + O(\log n) \end{aligned}$$

$\square$

3 Samplable Sources

Classical results, such as those mentioned in the previous section, show that data compression is feasible for various classes of sources defined by statistical or information-theoretic constraints (e.g., stationary ergodic sources or Markovian sources). We propose to investigate classes of sources defined by *computational constraints*, specifically samplability:

Definition 3.1. A source X_n is samplable if there is an efficient probabilistic algorithm S such that $S(1^n)$ is distributed according to X_n for every $n \in \mathbb{N}$. “Efficient” can be taken to mean a polynomial-time algorithm, a logarithmic space algorithm, a uniform or nonuniform algorithm, or any other complexity constraint, and will be specified in context.

For sources $(X_x)_{x \in L}$ indexed by strings, we instead require that $S(x)$ is distributed according to X_x for every $x \in L$.

It is natural to consider samplable sources, since any source X which is polynomially compressible to length $H(X)$, and moreover for all $x \in \text{Sup}(X)$, $|\text{Enc}(x)| = H(X)$, is polynomially samplable. This is because $X = \text{Dec}(U_{H(X)})$. Goldberg and Sipser [10] also studied compression of computationally constrained sources, but they focused on the complexity of deciding membership in the support of the source (for flat sources).

We recall that pseudorandom generators yield samplable sources that are incompressible. (In [10], this observation is attributed to L. Levin.)

Proposition 3.1 (Levin). If one-way functions exist, then there exist polynomial-time samplable sources X_n of entropy at most n^ϵ that cannot be compressed to length $n - 3$ by any probabilistic polynomial-time algorithms (Enc, Dec) .

Proof. (sketch) If one-way functions exist, then there exists a pseudorandom generator $G : \{0, 1\}^{n^\epsilon} \rightarrow \{0, 1\}^n$ [12]. Let $X_n = G(U_{n^\epsilon})$. From the pseudorandom property of G , it follows that

$$\Pr[\text{Dec}(\text{Enc}(U_n)) = U_n] \geq \Pr[\text{Dec}(\text{Enc}(X_n)) = X_n] - \text{neg}(n) = 1 - \text{neg}(n),$$

where neg denotes a negligible function. The pseudorandom property of G also implies that

$$\mathbb{E}[|\text{Enc}(U_n)|] \leq \mathbb{E}[|\text{Enc}(X_n)|] + \text{neg}(n) < n - 2.5.$$

That is, (Enc, Dec) compress a $1 - \text{neg}(n)$ fraction of $\{0, 1\}^n$ to average length $n - 2.5$. This is impossible by a counting argument. $\square$

Conversely, it is known that if one-way functions don’t exist, then every polynomial-time samplable flat source X_n can be compressed to length $H(X_n) + O(\log n)$ (for infinitely many n) [40].

Thus, we do not expect to efficiently compress samplable sources in full generality. Instead, we aim to identify natural subclasses of samplable sources for which compression is feasible. We will focus on the case when the compression algorithms are given the sampling algorithm. This is a natural implicit description of the source (like those discussed in Section 2.2). Moreover, efficient compression in this case implies universal compression (for uniform algorithms):

Lemma 3.1. Let $\mathcal{S} \subseteq \Sigma^*$ be a class of sampling algorithms (encoded as strings) and $\mathcal{C}$ be the corresponding class of sources. Suppose that there exist algorithms (Enc, Dec) such that for every

$S \in \mathcal{S}$, $(\text{Enc}(\cdot, S), \text{Dec}(\cdot, S))$ compresses $X_n = S(1^n)$ to length $m = m(H(X_n), n)$ in time $\text{poly}(n) \cdot f(|S|)$ for some function f . Then there exists a polynomial-time universal compression algorithm $(\text{Enc}', \text{Dec}')$ for $\mathcal{C}$ that compresses to length $m + O(1)$. If each encoding $\text{Enc}(\cdot, S)$ is prefix-free, then so is the encoding Enc' .

Proof. Let $\circ$ denote concatenation, and let $\Sigma^* = \{S_1, S_2, S_3, \dots\}$ be an enumeration of all strings in lexicographic order. Let $p(n) \cdot f(|S|)$ be the running time of (Enc, Dec) .

$\text{Enc}'(x)$, on input $x \in \{0, 1\}^n$:

1. For each $i = 1, \dots, n$
 - (a) Run $\text{Enc}(x, S_i)$ for $p(n) \cdot n$ steps, and if it halts, let y_i be the output.
 - (b) Run $\text{Dec}(y_i, S_i)$ for $p(n) \cdot n$ steps. If it outputs x , set $z_i = 0^i 1 \circ y_i$.
 - (c) If either Enc or Dec failed to halt within $p(n) \cdot n$ steps, set $z_i = 1 \circ x$.
2. Output the shortest string among $z_1, z_2, \dots, z_n$.

$\text{Dec}'(i, z)$: If $i = 0$, output z . Otherwise output $\text{Dec}(z, S_i)$.

By inspection, the above algorithms run in polynomial time. For the compression length, suppose X_n is sampled by algorithm $S_k \in \mathcal{S}$. For all $n \geq \max\{k, f(|S_k|)\}$, $\text{Enc}(x, S_k)$ and $\text{Dec}(y_k, S_k)$ will halt within $p(n) \cdot f(n) \leq p(n) \cdot f(|S_k|)$ steps and thus z_k will equal $(k, \text{Enc}(x, S_k))$. Thus, the compression length will be at most

$$\begin{aligned} \mathbb{E}[|\text{Enc}'(X_n)|] &\leq \mathbb{E}[|(k, \text{Enc}(X, S_k))|] \\ &\leq m(H(X_n), n) + O(1), \end{aligned}$$

since k is a constant. For $n \leq \max\{k, f(|S_k|)\}$, the compression length is bounded by a constant. $\square$

Before moving on to our positive results, we observe that good compression of a samplable source implies that the source's entropy can be approximated.

Proposition 3.2. *If a polynomial-time samplable source X_x distributed over $\{0, 1\}^n$ (for $n = n(x)$) is compressible to length $m = m(x)$ by a probabilistic polynomial-time encoding algorithm Enc (even sharing randomness with the decoding algorithm Dec), then there is a probabilistic polynomial-time algorithm A such that $\Pr[A(x) \in [H(X_x) - \log n - 1, m]] \geq 2/3$.*

Proof. By Lemma 2.2 and hypothesis, we have

$$H(X_x) - \log n - 1/2 \leq \mathbb{E}[|\text{Enc}(X_x, R)|] \leq m.$$

A simply estimates the average compression length $\mathbb{E}[|\text{Enc}(X_x, R)|]$ by taking polynomially many independent samples $x_1, \dots, x_k \stackrel{R}{\leftarrow} X_x$ and sequences of coin tosses $r_1, \dots, r_k$ for Enc , and computing the average of the $|\text{Enc}(x_i, r_i)|$'s. Taking $k = O(n^2)$, we obtain an approximation of $\mathbb{E}[|\text{Enc}(X_x, R)|]$ to within $\pm 1/2$ with high probability. $\square$

In particular, one way to show that a family of sources X_x does not have good polynomial-time compression algorithms is to show that it is intractable to approximate the entropy of X_x . For example, it is known that the problem of approximating the entropy of a general polynomial-time samplable source is complete for **SZK**, the class of problems possessing statistical zero knowledge proofs [11]. (More precisely, the problem is the following: given a boolean circuit $C : \{0, 1\}^m \rightarrow \{0, 1\}^n$, approximate the entropy of distribution $X_C = C(U_m)$.) Using this, we obtain the following.

Proposition 3.3. *If $\mathbf{SZK} \neq \mathbf{BPP}$, then there is a family of samplable sources $\{X_x\}_{x \in L}$ that cannot be compressed to length $H(X_x) + n^{1-\alpha}$ by a probabilistic polynomial-time encoding algorithm Enc (even sharing randomness with the decoding algorithm Dec), for any constant $\alpha > 0$.*

This result is incomparable to Proposition 3.1. Proposition 3.3 only requires that the encoding algorithm be efficient, but Proposition 3.1 rules out compression even to length $n - 1$ and uses a qualitatively weaker assumption (Ostrovsky [28] has shown that $\mathbf{SZK} \neq \mathbf{BPP}$ implies the existence of a weak form of one-way functions).

We note that Proposition 3.3 relies on the fact that the compression algorithm is not given a bound on the entropy of X_x . In fact, the literature on *lossless condensers* [30, 36] gives efficient, randomized encoding algorithms that near-optimally compress flat sources (with a small decoding error) when given a bound on the entropy. Condensers do not, however, provide an efficient decoding algorithm. (Indeed, if the decoding algorithm were efficient, then one could eliminate the need to know a bound on the entropy by trying $k = 1, \dots, n$, using the one that gives the smallest encoding length and decodes correctly, and including the value of k used in the compressed string.)

Note that for flat sources, an additive approximation to $H(X_x)$ is equivalent to a multiplicative approximation to $|\text{Sup}(X_x)|$, which is an approximate counting problem in the usual sense. In Section 7, we will exploit this relationship between compression and approximate counting in the opposite direction, using techniques from approximate counting algorithms to develop compression algorithms for a certain class of sources.

4 Sources with Logspace Samplers

In this section we consider sources sampled by logarithmic space randomized algorithms. As usual in the theory of randomized space-bounded algorithms, we consider a model where the space-bounded machine has *one-way* access to a tape containing random bits.

It is known that no pseudorandom generator can be implemented as a log-space machine with one-way access to the seed [21]. (This follows from the fact that deciding if a given string is a possible output of the generator is a problem in non-deterministic log-space, and so it is solvable in polynomial time.)

In the rest of this section we show that optimal compression is possible for sources sampled by one-way log-space algorithms. This complements the result of Goldberg and Sipser [10], who showed optimal compression for flat sources whose support is *decidable* by one-way log-space machines. Moreover, logspace samplers generalize the Markov chain model used often in compression work [42]. This is because a Markov chain with S states can be converted to a machine using space $\log S$. (S is usually viewed as a constant so uniformity issues do not arise.)

Definition 4.1 (Space-bounded Samplable Sources). *We say that a source X_n is samplable in space $s(n)$ if there is a probabilistic Turing machine M such that:*

- $M(1^n)$ has the same distribution as X_n ,
- For every content of the random tape, the computation $M(1^n)$ uses space at most $s(n)$,
- M has one-way access to its random tape,
- M has write-only access to its output.

We say that M is a space- $s(n)$ sampler.

Notice that the bound on the space implies that M runs in time $n2^{O(s(n))}$ and uses at most as many random bits.

The main lemma of this section says that the cumulative probability distributions of logspace-samplable sources can be computed in polynomial time. (A potentially larger class of sources can be handled using the techniques of [2].)

Lemma 4.1. *There is an algorithm A that on input a space- $s(n)$ sampler M and string $x \in \{0, 1\}^n$ runs in time $\text{poly}(n, 2^{s(n)})$ and returns the cumulative probability $\Pr[M(1^n) \preceq x]$, where $\preceq$ denotes lexicographic ordering.*

Proof. Given M , we define a new probabilistic space-bounded machine M' that uses space $O(s(n))$ and with the property that, for every $x \in \{0, 1\}^n$,

$$\Pr[M'(1^n, x) \text{ accepts}] = \Pr[M(1^n) \preceq x]$$

Given $(1^n, x)$, M' simulates $M(1^n)$, and it accepts if and only if the simulated computation outputs a string a such that $a \preceq x$. Since M' does not have enough space to store a , we need to be careful about the way the simulation is performed. Note that if $a \preceq x$ and a and x have the same length, then either $a = x$ or, for some i , a is a string of the form $(x_1, \dots, x_{i-1}, 0, a_{i+1}, \dots, a_n)$, where $x_i = 1$. That is, a starts with a (possibly empty) prefix of x , then it has a zero in a position in which x has a one, and then it continues arbitrarily.

At the beginning of the simulation, the head of M' on the input tape is on the first bit of x . Every time the simulated computation of $M(1^n)$ writes on the output tape, M' compares the bit that $M(1^n)$ is going to write with the current bit of x that it sees on the output tape. If the bits are the same, then M' continues the simulation and moves the input-tape head on to the next symbol of x . If $M(1^n)$ is about to write a one, and the corresponding bit of x is zero, then the simulation halts and M' rejects. If $M(1^n)$ is about to write a zero, and the corresponding bit of x is one, then M' accepts. Also, if the simulation of $M(1^n)$ is completed with the input-tape head moving all the way until the end of x , then also M' accepts. It should be clear that the contents of the random tape for which $M'(1^n, x)$ accepts are precisely those for which $M(1^n)$ outputs a string $\preceq x$.

After constructing M' , it then remains to compute $\Pr[M'(1^n, x) \text{ accepts}]$, which is a standard problem. We enumerate all $S = n \cdot 2^{O(s)}$ possible states of $M'(1^n, x)$, and construct an $S \times S$ matrix P such that $P_{i,j}$ is the probability that $M'(1^n, x)$ goes from state i to state j in one step. We let e be the S -dimensional vector such that $e_i = 1$ if i is the start state of the machine, and $e_i = 0$ otherwise, and we compute the vector eP^S . Then, if A is the set of accepting states of the machine, then $\sum_{a \in A} (eP^S)[a]$ gives the probability that the machine accepts. $\square$

Theorem 1 (Compressing log-space Sources). *Let X_n be a source over $\{0, 1\}^n$ samplable in space $O(\log n)$. Then there are polynomial time algorithms (Enc, Dec) that compress X_n to length $H(X_n) + 2$. The encoding is prefix-free.*

Proof. Combine Lemma 2.4 with Lemma 4.1 □

Corollary 4.1 (Universal Compression of log-space Sources). *For every bound $s(n) = O(\log n)$ there are polynomial-time algorithms (Enc, Dec) such that for every source X_n over $\{0, 1\}^n$ samplable in space $s(n)$, and for every sufficiently large n , (Enc, Dec) compress X_n to length $H(X_n) + O(1)$. The encoding is prefix-free.*

Proof. Combine Theorem 1 with Lemma 3.1. □

5 Sources with Membership Algorithms

In the rest of this paper, we consider an alternative approach to bypassing the impossibility of compressing pseudorandom sources. Here we allow the sampler to be an arbitrary probabilistic polynomial-time algorithm, but explicitly impose the constraint that the source is not pseudorandom.

Definition 5.1. *Let X_n be a flat source. We say that X_n is a source with membership algorithm if there is a polynomial-time algorithm D such that $D(z) = 1 \Leftrightarrow z \in \text{Sup}(X_{|z|})$. For a source X_x indexed by a string x , we require instead that there is a polynomial-time algorithm D such that $D(x, z) = 1 \Leftrightarrow z \in \text{Sup}(X_x)$.*

Note that a source with a membership algorithm cannot be pseudorandom; indeed the algorithm D distinguishes it from all sources of higher entropy.

Are all samplable sources with membership algorithms efficiently compressible? Goldberg and Sipser [10] showed that any source with membership algorithm can be compressed to length $n - \Theta(\log n)$ (provided $H(X_n) < n - (3 + \delta) \log n$). But can they be compressed to length roughly $H(X_n)$? (Think of, say, $H(X_n) = n/2$.) This is an intriguing open question, which we first heard from Impagliazzo [13]. Goldberg and Sipser [10] and Wee [40] provide oracles relative to which the $n - \Theta(\log n)$ bound cannot be improved, and relative to which deterministic compression is impossible.² We know of no other evidence regarding this question without oracles.

In the next two sections, we present two positive results about sources with membership algorithms. In the first, we show how to compress better than Goldberg–Sipser while using *deterministic* compression and decompression algorithms. In particular, if X_n is a source with membership algorithm and $H(X_n) \leq k = n - O(\log n)$, then Goldberg & Sipser showed how to compress X_n to length $k + 3 \log n$ with high probability. We show how to compress to length $k + \delta \cdot (n - k)$ in polynomial-time, for any $\delta > 0$. Thus, for $k = n - o(\log n)$ this is more than a constant factor savings in overhead. In deterministic quasi-polynomial time, we only require $k \geq n - (\log n)^{O(1)}$, and we show how to compress to length $k + \text{polylog}(n - k) \leq k + \text{polylog } \log n$.

Our technique is completely different than that of [10]. Instead of arithmetic coding, we use the recent explicit construction of constant-degree “lossless” expanders [6], together with an idea from distributed algorithms for routing in expander-based networks [3].

In the second result, we show how to compress to length $H(X) + O(1)$ for a large class of sources with membership algorithms, namely those whose supports are self-reducible in the sense of [33].

²We note that Goldberg and Sipser measure compression by the worst-case length (except for a finite number of exceptions, which makes no difference in the Turing machine model), whereas our definitions involve the average-case length, as in [40].

6 Compressing High Entropy Sources

We prove the following theorem.

Theorem 2. *Let X_n be a flat source with membership algorithm and $H(X_n) \leq k$. Then X_n is compressible:*

1. *to length $k + \text{polylog}(n - k)$ in time $n^{O((n-k)/\log(n-k))}$, and*
2. *to length $k + \delta \cdot (n - k) + O(1)$ in polynomial time, for any constant $\delta > 0$, if $k \geq n - O(\log n)$.*

The encodings are prefix-free.

In particular, we get better compression as k approaches n . Yet even for $k = n - \text{polylog}(n)$, we achieve compression to $k + \text{polylog } n$ in quasi-polynomial time.

The starting idea of the proof is that we wish to condense the input distribution, without many collisions of points in the support. Lossless condensers, first defined and constructed in [30, 36], do exactly this. We prove that a good condensing function can be used to compress, and then use the expanders constructed by Capalbo et al. [6] as condensing functions. A simple application of this approach would only compress the source to $k + O(\log n)$ bits. Using an idea from [3], we improve the bound to $k + \delta \cdot (n - k)$.

We begin with the following lemma, which shows how a good condensing function can be used to compress.

Lemma 6.1. *Suppose X_n is a flat source with membership algorithm and $S = \text{Sup}(X)$. Fix a function $f : \{0, 1\}^n \times \{0, 1\}^d \rightarrow \{0, 1\}^m$. Call $z \in \{0, 1\}^m$ S -unique if there is exactly one element $(x, r) \in S \times \{0, 1\}^r$ such that $f(x, r) = z$. Suppose that $\Pr_{x \in X, r \in U_d}[f(x, r) \text{ is } S\text{-unique}] \geq 1 - \epsilon$. Then X_n is compressible to length $m + \epsilon n + 1$ in time $(T_f + T_{f^{-1}}) \cdot \text{poly}(n)$. Here $T_{f^{-1}}$ denotes the time to compute the set $f^{-1}(y)$ on input y . The encoding is prefix-free.*

Proof. Let $\text{Enc}(x)$ be 0 concatenated with the lexicographically first y of the form $f(x, r)$ which is S -unique, or $1x$ if there is no such y . Let $\text{Dec}(1x) = x$ and $\text{Dec}(0y)$ be the $x \in S$ such that $f(x, r) = y$. Then Enc and Dec satisfy the conclusions of the lemma. $\square$

The function f is essentially a disperser. A disperser is a type of expanding graph where the expansion is required only for sets of a particular size. We will need the expansion close to the degree. For our improvements we will need f to represent a true expander, as defined as follows.

Definition 6.1. *A bipartite graph $G = (V, W, E)$ is a (K, A) -expander if for all $T \subseteq V$, $|T| \leq K$, $|\Gamma(T)| \geq A \cdot |T|$.*

The following lemma is self-evident.

Lemma 6.2. *Let $G = (\{0, 1\}^n, \{0, 1\}^m, E)$ be a $(|S|, (1 - \epsilon/2)D_L)$ -expander with left degree D_L and right degree D_R . Assume the edges are labeled from $\{0, 1\}^d$ with unique labels out of a given node in $\{0, 1\}^n$. Define $f(x, r)$ to be the neighbor of x labeled by r . Then f satisfies the conditions of Lemma 6.1.*

We take G to be the expander explicitly constructed by Capalbo et al. [6]:

Theorem 3. [6]³ Let $N = 2^n \geq K = 2^k$, There are explicitly constructible $(K, (1 - \epsilon/2)D_L)$ regular expanders $G = (\{0, 1\}^n, \{0, 1\}^m, E)$ with left degree $D_L = 2^d$ (d to be specified below), and $M = 2^m = O(KD_L/\epsilon)$ such that the set of neighbors of a vertex in $\{0, 1\}^n$ is computable in time $\text{poly}(n, D_L)$ and the neighbors of a vertex in $\{0, 1\}^m$ are computable in time $\text{poly}(n, D_L, N/M)$ with either of the following values of d :

1. $d = \text{poly}(\log(n - k), \log(1/\epsilon))$, or
2. $d = \delta \cdot (n - k) + O(\log(1/\epsilon))$, where δ is an arbitrarily small constant.

Setting $\epsilon = 1/n$, the expanders in Part 2 yield compression length $m = d + k + \log(1/\epsilon) + O(1) = k + \alpha(n - k) + O(\log n) = k + O(\log n)$ for $k = n - O(\log n)$. This compression differs from optimal by an additive $O(\log n)$ term, like in [10]. This was because we had to set $\epsilon = 1/n$. We now describe a method where we can use a larger ϵ .

First note that Hall's theorem implies that in a $(K, 1)$ -expander, there is a matching which matches all vertices in any K -subset of $\{0, 1\}^n$. Such a matching may be used as the compression function. We show that better expansion allows us to construct the matching efficiently.

The idea is as follows. For $S \subseteq V$, let $\text{Uniq}(S) = \{w \in W : |\Gamma(w) \cap S| = 1\}$ be the set of unique neighbors of S . Let $S_0 = S = \text{Sup}(X_n)$. If there is an r such that $y = f(x, r)$ is S -unique, then we encode x by the lexicographically first such y . This gives a maximal matching from S_0 to $\text{Uniq}(S_0)$. If there is no such r , let S_1 denote all the unmatched elements in S , i.e., $S_1 = S_0 \setminus \Gamma(\text{Uniq}(S_0))$. We now encode x by the lexicographically first $y = f(x, r)$ which is S_1 -unique, if it exists. This gives a maximal matching from S_1 to $\text{Uniq}(S_1)$.

In general, in the i th stage, if we haven't encoded x already, we encode x by the lexicographically first $y = f(x, r)$ which is S_i -unique, if it exists. This gives a maximal matching from S_i to $\text{Uniq}(S_i)$. We then set $S_{i+1} = S_i \setminus \Gamma(\text{Uniq}(S_i))$. We do this for at most $t = \lceil (\log n) / \log(1/\epsilon) \rceil$ stages.

The number of stages is chosen so that all but $1/n$ fraction of S have been matched. This is because in each stage, a $1 - \epsilon$ fraction of nodes gets matched, so all but $\epsilon^t \leq 1/n$ nodes remain unmatched. It suffices to compress all but $1/n$ fraction because of Lemma 2.5.

Note that the number of calls to the membership algorithm to check if a node is in S_i is at most $(D_L \cdot D_R)^i$. Here D_R denotes the right degree, $D_R = D_L N/M$.

Using the expanders from Theorem 3 gives compression length

$$m = k + d + \log(1/\epsilon) + O(1)$$

The number of membership calls is

$$\begin{aligned} (D_L \cdot D_R)^t &= (ND_L^2/M)^t \\ &= \text{poly}(1/\epsilon^t, (N/K)^t, D_L^t). \end{aligned}$$

We also need to multiply by the time to compute neighborhoods, $\text{poly}(n, D_L, N/M)$, which gives a total running time of

$$\text{poly}(n, 1/\epsilon^t, (N/K)^t, D_L^t).$$

³Part 2 is not stated in [6], but can be obtained by using the extractor of [43] to construct the “small conductors” used in the zig-zag product there. In [6], the computation time of neighborhoods of right-vertices is also not stated, but it can be deduced from the computation time of neighborhoods of left-vertices and the high-level structure of the construction.

To optimize compression length, we use the expanders from Part 1 and take $\epsilon = 1/(n - k)$. This yields compression length $k + \text{polylog}(n - k)$ and running time

$$\text{poly}(n, 1/(n - k)^t, 2^{(n-k) \cdot t}, 2^{\text{polylog}(n-k) \cdot t}) = n^{O((n-k)/\log(n-k))}.$$

To optimize the running time for $k \geq n - O(\log n)$, we use the expanders from Part 2 and set $\epsilon = 2^{-\delta \cdot (n-k)}$ for an arbitrarily small constant $\delta > 0$. This gives running time $\text{poly}(n)$ and compression length

$$m = k + O(\delta \cdot (n - k)) + O(1) = k + \delta' \cdot (n - k) + O(1).$$

6.1 Generalizing to Non-Flat Sources

We now extend the above techniques to non-flat distributions. For a non-flat distribution X_n , a natural generalization of a membership algorithm is a *probability mass algorithm*: a polynomial-time algorithm D such that for every $z \in \{0, 1\}^n$, $D(z) = \Pr[X_n = z]$. For distributions with a probability mass algorithm, we can extend the above techniques to obtain bounds on compression length in terms of a “truncated” variant of entropy.

Theorem 4. *Let X_n be a source with probability mass algorithm, and let c be any constant. For $x \in \{0, 1\}^n$, define $h(x) = \log(1/\Pr[X_n = x])$,*

$$h'(x) = \begin{cases} n - c \log n & \text{if } h(x) < n - c \log n, \\ h(x) & \text{if } h(x) \in [n - c \log n, n], \\ n & \text{if } h(x) > n, \end{cases}$$

and

$$H'(X_n) = \mathbb{E}_{X_n}[h'(X_n)].$$

Then X_n is compressible:

1. to length $H'(X_n) + \text{polylog}(n - H'(X_n))$ in time $n^{o(\log n)}$, and
2. to length $H'(X_n) + \delta \cdot (n - H'(X_n)) + O(1)$ in polynomial time, for any constant $\delta > 0$.

The encodings are prefix-free.

Proof. First, we observe that the proof of Theorem 2 immediately works for sources that are “nearly flat,” in the sense that every two elements of $\text{Sup}(X_n)$ have probability mass within a factor of, say, 2 of each other. This is the case because if the algorithm correctly compresses all but a $1/n$ fraction of elements of $\text{Sup}(X_n)$, then it will also correctly compress all but a $2/n$ fraction of the distribution X_n . Removing errors via Lemma 2.5 will then increase the compression length by at most 2 bits.

To handle general sources X_n , we bucket the elements of $\text{Sup}(X_n)$ into sets S_i , consisting of elements of probability mass in the interval $[2^{-i}, 2^{-(i+1)})$. Note that given $x \in \text{Sup}(X_n)$, we can determine its bucket $i(x)$ using the probability mass algorithm for X_n .

To compress elements x of S_i (where $i = i(x)$), we use the compression algorithm from the proof of Theorem 2, replacing the parameter k with $k_i = \max\{i + 1, n - (c + 1) \log n\}$. We denote such an encoding of a string $x \in S_i$ by $\text{Enc}_i(x)$. The final encoding $\text{Enc}(x)$ of a string x is as

follows: If $i(x) \leq n$ then we set $\text{Enc}(x) = 0 \circ (n - i(x)) \circ 1 \circ \text{Enc}_{i(x)}(x)$ where $(n - i(x))$ is written as a string of length $2\lceil \log(n - i(x)) \rceil$ by taking its binary expansion and replacing each 0 with 00 and each 1 with 01. If $i(x) \geq n$, then we set $\text{Enc}(x) = 1 \circ x$.

Since the running time of the compression algorithms in Theorem 2 are decreasing in k , the running time of the new compression algorithm can be bounded by substituting $k = n - (c+1) \log n$ into the same expressions, yielding running time $n^{O((c+1) \log n / \log((c+1) \log n))} = n^{o(\log n)}$ in the first case and polynomial time in the second. In the first case, the compression length can be bounded as follows:

$$\begin{aligned}
& \mathbb{E}_{X_n}[|\text{Enc}(X_n)|] \\
& \leq \mathbb{E}_{X_n}[k_{i(X_n)} + \text{polylog}(n - k_{i(X_n)}) + \log(n - i(X_n)) + O(1)] \\
& \leq \mathbb{E}_{X_n}[\max\{i(X_n), n - c \log n\} + \text{polylog}(n - \max\{i(X_n), n - c \log n\}) + O(1)] \\
& \leq \mathbb{E}_{X_n}[\max\{i(X_n), n - c \log n\}] + \text{polylog}(\mathbb{E}_{X_n}[n - \max\{i(X_n), n - c \log n\}]) + O(1) \\
& \leq H'(X_n) + \text{polylog}(n - H'(X_n)) + O(1).
\end{aligned}$$

The second inequality above is obtained by separately considering the cases that $i(X_n) + 1 \leq n - (c+1) \log n$ (in which case the $\log(n - i(X_n))$ term is absorbed into the first term) and $i(X_n) + 1 > n - (c+1) \log n$ (in which case the $\log(n - i(X_n))$ term is absorbed into the second). The third inequality is obtained by applying Jensen's inequality to the function $f(x) = \log^m x$, which is concave when m is constant and x is sufficiently large.

The compression length is bounded similarly in the second case (without needing Jensen's Inequality, since the function $f(x) = \delta \cdot x$ is linear). $\square$

We now deduce two corollaries of the above, giving compression bounds in terms of the actual entropy of the source.

Corollary 6.1. *Let X_n be a source with probability mass algorithm having min-entropy at least $n - c \log n$ for some constant c . Then for any constant $\delta > 0$, X_n is polynomial-time compressible to length $H(X_n) + \delta \cdot (n - H(X_n)) + O(1)$, via a prefix-free encoding.*

Proof. If X_n has min-entropy at least $c - \log n$, then $H'(X_n) \leq H(X_n)$. $\square$

Corollary 6.2. *Let X_n be a source with probability mass algorithm. Then for any constant $c > 0$, X_n is polynomial-time compressible to length $n - (1 - H(X_n)/n) \cdot c \log n + O(1)$.*

Proof. First, we may assume that $H(X_n) \leq n - 2c \log n$; otherwise the previous corollary applies. Call a string x *light* if $h(x) > n - c \log n$. Let L be the set of light strings. By Markov's inequality, $\Pr[X_n \in L] \leq H(X_n)/(n - c \log n)$. So $\Pr[X_n \notin L] \geq 1 - H(X_n)/(n - c \log n) \geq (1 - H(X_n)/n)/2$ for sufficiently large n . Thus,

$$\begin{aligned}
H'(X_n) & \leq \Pr[X_n \in L] \cdot n + \Pr[X_n \notin L] \cdot (n - c \log n) \\
& = n - \Pr[X_n \notin L] \cdot c \log n \\
& \leq n - (1 - H(X_n)/n) \cdot (c \log n)/2
\end{aligned}$$

Setting $\delta = 1/2$ in Theorem 4, we can compress X_n to length

$$n/2 + H'(X_n)/2 + O(1) \leq n - (1 - H(X_n)/n) \cdot (c \log n)/4 + O(1).$$

Increasing c by a factor of 4 yields the desired bound. $\square$

Notice that the results in this section do not require that the source X_n is samplable, but only that X_n has a membership algorithm (in the case of flat sources) or a probability-mass algorithm (in the case of general sources). For flat sources of entropy at least $n - O(\log n)$, a membership algorithm implies samplability: one can sample by randomly picking elements of $\{0, 1\}^n$ and testing if they are in the support of X_n . But our results also apply to sources of entropy smaller than $n - O(\log n)$ (though they will only achieve compression length $n - O(\log n)$). This leads to the question of whether better compression can be achieved based on just the membership algorithm condition. Below we give evidence that membership algorithm condition alone is unlikely to imply near-optimal compression, even for sources of entropy zero.

Proposition 6.1. *Suppose that every family of flat sources $(X_x)_{x \in L}$ of zero entropy with a membership algorithm can be compressed to length $m = m(n)$ by a polynomial-time compression algorithm (Enc, Dec) with shared randomness. Then SAT is in $\mathbf{RTIME}(\text{poly}(n) \cdot 2^{m(n)})$.*

Proof. We show that the hypothesis implies a randomized algorithm for finding satisfying assignments to formulas with a *unique* satisfying assignment, and then apply the Valiant–Vazirani [39] reduction from SAT to UNIQUE-SAT. For a boolean formula φ , consider the source X_φ that is uniform on the satisfying assignments of φ , and let L be the set of formulas φ with exactly one satisfying assignment. Then $(X_\varphi)_{\varphi \in L}$ has a membership algorithm because checking whether an assignment satisfies a formula is easy. Now, if (Enc, Dec) compress X_φ for $\varphi \in L$ to expected length m , then with probability at least $1/(m+1)$ over the coin tosses R of Enc and Dec , the unique satisfying assignment of φ gets compressed to length at most m . In such a case, the satisfying assignment can be found by enumerating all $O(2^m)$ strings z of length at most m and computing $\text{Dec}(z, R)$. Repeating for $O(m)$ independent choices of R amplifies the probability of finding an assignment to $1/2$.

Valiant and Vazirani [39] give a randomized polynomial-time reduction mapping any formula ψ to a formula φ on the same number n of variables such that if ψ is satisfiable, then with constant probability φ has exactly one satisfying assignment, and if ψ is unsatisfiable, then with probability 1 φ is unsatisfiable. Composing this reduction with the above algorithm for finding unique satisfying assignments yields the claimed algorithm for SAT. $\square$

Thus, if SAT requires time $2^{\Omega(n)}$, the above gives a family of zero-entropy sources that cannot be compressed to length $o(n)$. The argument can be modified to give an incompressible family of sources indexed only by input length, under an assumption about “unique nondeterministic exponential time”.

Proposition 6.2. *Suppose that every family of flat sources $(X_n)_{n \in \mathbb{N}}$ of entropy at most 1 with a membership algorithm can be compressed to length $m = m(n)$ by a polynomial-time compression algorithm (Enc, Dec) with shared randomness. Then $\mathbf{UTIME}(2^n) \subseteq \mathbf{RTIME}(2^{O(n)} \cdot 2^{2m(2^{n+O(1)})})$.*

Proof. Let M be a nondeterministic Turing machine running in time $T(\ell) = 2^\ell$ on inputs of length ℓ such that M has zero or one accepting computation on each input. Our aim is to show that, under the hypothesis, $L(M)$ can be decided by a randomized algorithm running in time $2^{O(\ell)} \cdot 2^{2m(2^{\ell+O(1)})}$ on inputs of length ℓ .

Let M' be a nondeterministic TM running in time 2^ℓ that has exactly one more accepting computation than M on each input (by adding a trivial accepting computation). We view each possible input of length ℓ to M' as a binary number n in the interval $[2^{\ell+c}, 2^{\ell+c} + 2^\ell]$, where 2^c upper bounds

the branching factor of M' . Since M' has running time 2^ℓ , computations of M' can be described by strings of length n . We define X_n to be the uniform distribution on the accepting computations of M' . (For n not in an interval $[2^{\ell+1}, 2^{\ell+1} + 2^\ell]$, X_n can be taken to be the distribution that always outputs 0^n .) Notice that X_n has entropy zero or 1, depending on whether M' has 1 or 2 accepting computations on input n . Membership in the support of X_n can be decided in time $2^{O(\ell)} = \text{poly}(n)$.

Now we argue that a good compression algorithm can be used to decide $L(M)$, specifically by yielding an efficient algorithm to find all accepting computations of M . If (Enc, Dec) compress X_n to length m , then with probability at least $1/(m+1)$ over the coin tosses R of Enc, all accepting computations of M' are compressed to length at most $2m$. (At worst, one is compressed to length $2m$ and the other to length zero.) Thus, the accepting computations can be found in time $\text{poly}(n) \cdot 2^{2m(n)} = 2^{O(\ell)} \cdot 2^{2m(2^{\ell+c})}$. $\square$

If we impose the additional condition that X_n is samplable, then we do not know of any evidence suggesting the intractability of near-optimal compression other than the oracle result of Wee [40].

7 Self-Reducible Sets

For a source X_x with membership oracle, the relation $R = \{(x, z) : z \in \text{Sup}(X_x)\}$ is decidable in polynomial time. Thus sources with membership oracles correspond to the uniform distribution on **NP** witness sets. Many natural **NP** witness sets have the following property of self-reducibility:

Definition 7.1 ([33]). A polynomially balanced relation $R \subseteq \Sigma^* \times \Sigma^*$ is self-reducible if there exist polynomial-time computable functions $\ell : \Sigma^* \rightarrow \mathbb{N}$, $\sigma : \Sigma^* \rightarrow \mathbb{N}$, and $\rho : \Sigma^* \times \Sigma^* \rightarrow \Sigma^*$ such that for all $x, w = w_1 \cdots w_m \in \Sigma^*$,

1. $(x, w) \in R \Rightarrow |w| = \ell(x)$.
2. For all x , $\sigma(x) \leq \ell(x)$, and $\ell(x) > 0 \Rightarrow \sigma(x) > 0$.
3. $\sigma(x) = O(\log |x|)$,
4. $(x, w_1 w_2 \cdots w_{\ell(x)}) \in R$ if and only if $(\rho(x, w_1 \cdots w_{\sigma(x)}), w_{\sigma(x)+1} \cdots w_{\ell(x)}) \in R$,
5. $|\rho(x, w_1 w_2 \cdots w_{\sigma(x)})| \leq |x|$.
6. If $\ell(x) = 0$, then R can be decided in polynomial time.

As usual, the language associated with R is $L_R = \{x : \exists w(x, w) \in R\}$.

Intuitively, this definition says that the witness set for a given input can be expressed in terms of witness sets for smaller inputs. Specifically, the witnesses for x which begin with initial segment $w_1 \cdots w_{\sigma(x)}$ are in one-to-one correspondence with the witnesses for the instance $\rho(x, w_1 \cdots w_{\sigma(x)})$. Many natural witness relations are self-reducible in this sense, e.g. satisfying assignments of boolean formulae and perfect matchings in bipartite graphs.

Example 7.1 (perfect matchings). Let R be the relation consisting of pairs (G, M) , where G is a bipartite graph and M is a perfect matching in G . This is self-reducible because the perfect matchings in $G = (V, E)$ that contain some edge $e = (i, j) \in E$ are in one-to-one correspondence

with the perfect matchings in $G' = (V \setminus \{i, j\}, E \setminus \{e\})$, and those that do not contain e are in one-to-one correspondence with the perfect matchings in $G'' = (V, E \setminus \{e\})$.

More formally, we represent G by its $n \times n$ adjacency matrix, and if G has m edges, then M is represented by a bit vector $M_1 M_2 \cdots M_m$ where M_i indicates whether or not edge i is included in the perfect matching. Then we set $\ell(G) = m$, $\sigma(G) = 1$ (unless $m = 0$, in which case $\sigma(G) = 0$), and define $\rho(G, 0)$ to be the graph obtained by removing edge 1 from G (but keeping its endpoints as vertices), and $\rho(G, 1)$ to be the graph obtained by removing edge 1 and its endpoints from G .

Jerrum, Valiant, and Vazirani [18] proved that, for self-reducible relations, witnesses can be generated almost uniformly at random if and only if approximate counting of witnesses can be done in probabilistic polynomial time. And, indeed, there are now many approximate counting algorithms known that have been obtained by first constructing almost-uniform samplers (typically via the Markov chain Monte Carlo method; see the surveys [19, 17, 29]).

The main result of this section adds compression of the witness set to the list of tasks equivalent to sampling and counting.

Theorem 5. *Let R be a self-reducible relation, and for every $x \in L_R$, let X_x be the uniform distribution on $W_x = \{w : (x, w) \in R\}$. If the sources $(X_x)_{x \in L_R}$ are samplable, then they can be efficiently compressed to length $H(X_x) + 5$ with shared randomness and zero decoding error. The encodings are prefix-free.*

Proof. We will show how to compute an “approximate arithmetic encoding” for the sources X_x . A similar approach was used by Goldberg and Sipser [10] in their main result, but as mentioned above they were only able to compress to length $n - O(\log n)$. We use the ideas in the reduction from approximate counting to sampling [18] to obtain an almost-optimal compression length.

The first step is to argue that we can efficiently approximate probabilities of witness prefixes. For an input x and a witness prefix $z = z_1 \cdots z_{\sigma(x)}$, let $p(x, z) = \Pr[X_x|_{\sigma(x)} = z]$, where $a|_t$ denotes the first t bits of a .

Claim 7.1. *There is a probabilistic algorithm $A(x, z, \epsilon, \delta; r)$ (where r are the coin tosses) running in time $\text{poly}(|x|, 1/\epsilon, \log(1/\delta))$ such that*

1. *For every x, z, ϵ, δ , $\Pr[|A(x, z, \epsilon, \delta) - p(x, z)| > \epsilon] \leq \delta$, and*
2. *For every x, ϵ, δ, r , $A(x, \cdot, \epsilon, \delta; r)$ is a probability measure on $\Sigma^{\sigma(x)}$. That is, $\sum_{z \in \Sigma^{\sigma(x)}} A(x, z, \epsilon, \delta; r) = 1$ and for every $z \in \Sigma^{\sigma(x)}$, $A(x, z, \epsilon, \delta; r) \in [0, 1]$.*

The algorithm A simply takes $\text{poly}(1/\epsilon, \log(1/\delta))$ samples from X_x and outputs the fraction that begin with prefix z . The claim follows from a Chernoff Bound.

Fix an input length n , and set $\delta = 2^{-3n}$, $\epsilon = 1/n^{2c}$, for a large constant c to be specified later. For x of length at most n , z of length at most $\sigma(x)$, and a sequence r of $(\text{poly}(n))$ coin tosses for A , define $q_r(x, z) = A(x, z, \epsilon, \delta; r)$. Taking a union bound over all x, z , the following holds with probability at least $1 - 2^{-n}$ over r :

$$|q_r(x, z) - p(x, z)| \leq \epsilon \quad \forall |x| \leq n, |z| = \sigma(x). \quad (1)$$

Our compression and decompression algorithms will choose r at random, so we may assume they have an r that satisfies this condition (the exponentially rare r 's which violate this condition will only increase the expected compression length by at most $\text{poly}(n)/2^{-n}$).

Once r is fixed, the q_r 's induce approximating distributions $\hat{X}_{x,r}$ via self-reducibility:

$\hat{X}_{x,r}$: If $\ell(x) = 0$, output the empty string. Otherwise:

1. Select a prefix $z \in \{0, 1\}^{\sigma(x)}$ according to the distribution $q_r(x, \cdot)$.
2. Recursively sample $z' \leftarrow \hat{X}_{\rho(x,z),r}$.
3. Output zz' .

Moreover, we can recursively compute the cumulative distribution function $\hat{F}_{x,r}(w)$ for $\hat{X}_{x,r}$ with respect to the lexicographic order as follows, writing $w = zz'$ with $|z| = \sigma(x)$:

$$\hat{F}_{x,r}(zz') = \left(\sum_{u < z} q_r(x, u) \right) + q_r(x, z) \cdot \hat{F}_{\rho(x,z),r}(z'). \quad (2)$$

Thus we can compute the arithmetic coding $(\widehat{\text{Enc}}_{x,r}, \widehat{\text{Dec}}_{x,r})$ (Lemma 2.4) for $\hat{X}_{x,r}$ in polynomial time. Our compression algorithms (Enc, Dec) for X_x itself are as follows:

Enc(x, w, r): Let $s = \widehat{\text{Enc}}_{x,r}(w)$. If $|s| \leq \ell(x)$, output $0c$. Otherwise output $1x$.

Dec(x, bs, r): If $b = 0$, output $\widehat{\text{Dec}}_{x,r}(c)$. Otherwise output s .

By inspection, $\text{Dec}(x, \text{Enc}(x, w, r), r) = w$ for all w . Thus, we only need to verify the compression length. To do this, we argue about how well $\hat{X}_{x,r}$ approximates X_x .

Claim 7.2. *With probability at least $1 - 1/(n \cdot \ell)$ over $w \leftarrow X_x$ (where $\ell = \ell(x)$), we have $X_x(w) \leq \sqrt{2} \hat{X}_{x,r}(w)$.*

Proof of claim: To prove this claim, we call a prefix $z \in \Sigma^{\sigma(x)}$ *light* if

$$\Pr [X_x|_{\sigma(x)} = z] \leq 1/(n^c \cdot |\Sigma|^{\sigma(x)}).$$

By a union bound over all $z \in \Sigma^{\sigma(x)}$, the probability that $z \leftarrow X_x|_{\sigma(x)}$ is light is at most $1/n^c$. Thus, if we sample from X_x by first sampling a prefix z and then recursively sampling from $X_{\rho(x,z)}$, we encounter a light prefix somewhere along the way with probability at most $\ell \cdot (1/n^c)$, because there are at most ℓ levels of recursion. For a sufficiently large choice of c , this probability is at most $1/(n \cdot \ell)$.

So we only need to argue that if the sampling of w involves no light prefixes, then $X_x(w) \leq \sqrt{2} \hat{X}_{x,r}(w)$. Let z be the first prefix. By Property (1) of the q 's, we have

$$\begin{aligned} q_r(x, z) &\geq p(x, z) - \epsilon \\ &= p(x, z) - \frac{1}{n^{2c}} \\ &\geq p(x, z) \cdot \left(1 - \frac{|\Sigma|^{\sigma(x)}}{n^c} \right) \\ &\geq p(x, z) \cdot \left(1 - \frac{1}{3\ell} \right), \end{aligned}$$

for a sufficiently large choice of the constant c . By the definition of self-reducibility and the fact that X_x is uniform on W_x , we can expand $X_x(w)$ for any x and $w = z_1 \cdots z_t \in W_x$ as follows:

$$\Pr[X_x = z_1 z_2 \cdots z_t] = p(x_0, z_1) \cdot p(x_1, z_2) \cdot p(x_2, z_3) \cdots p(x_{t-1}, z_t), \quad (3)$$

where $x_0 = x$, $|z_i| = \sigma(x_{i-1})$, $x_i = \rho(x_{i-1}, z_i)$, and $\sigma(x_t) = 0$. Similarly, by the recursive definition of $\hat{X}_{x,r}$, we have:

$$\Pr[\hat{X}_{x,r} = z_1 z_2 \cdots z_t] = q_r(x_0, z_1) \cdot q_r(x_1, z_2) \cdot q_r(x_2, z_3) \cdots q_r(x_{t-1}, z_t), \quad (4)$$

Putting all of the above together, we have $\hat{X}_{x,r}(w) \geq (1 - 1/3\ell)^\ell \cdot X_x(w) \geq X_x(w)/\sqrt{2}$, as desired. $\square$

We can now estimate the compression length of X_x under $(\text{Enc}(x, \cdot, r), \text{Dec}(x, \cdot, r))$. Recall that the arithmetic coding $\widehat{\text{Enc}}_{x,r}(w)$ compresses an individual string w to length $\lceil \log(1/\hat{X}_{x,r}(w)) \rceil + 1$. If r and w satisfy the Inequalities (1) and the conclusion of Claim 7.2, then we can bound this length as

$$|\widehat{\text{Enc}}_{x,r}(w)| = \lceil \log(1/\hat{X}_{x,r}(w)) \rceil + 1 \leq \log(1/X_x(w)) + 5/2.$$

The probability that r and w do not satisfy either the Inequalities (1) or the conclusion of Claim 7.2 is at most $2^{-n} + 1/(n \cdot \ell)$. Thus, the average compression length is at most

$$\begin{aligned} \mathbb{E}_{w \leftarrow X_x, r}[\text{Enc}(x, w, r)] &= \mathbb{E}_{w \leftarrow X_x, r}[\max\{|\widehat{\text{Enc}}_{x,r}(w)|, \ell\}] + 1 \\ &\leq \mathbb{E}_{w \leftarrow X_x}[\log(1/X_x(w)) + 5/2] + (1/(n \cdot \ell) + 2^{-n}) \cdot \ell + 1 \\ &= H(X_x) + 4, \end{aligned}$$

for large enough n , as desired. $\square$

The randomization in the compression algorithms above can be eliminated via Lemma 2.8, under a complexity assumption. However, if we do not care for a full derandomization, and only to eliminate the *shared* randomness, we can use a “random perturbation” trick from [10] to do it without a complexity assumption.

Proposition 7.1. *Let R be a self-reducible relation, and for every x , let X_x be the uniform distribution on $\{w : (x, w) \in R\}$. If the sources X_x are samplable, then they can be compressed by probabilistic polynomial-time algorithms to length $H(X_x) + O(\log n)$ with independent randomness and decoding error 2^{-n} . The encodings are prefix-free.*

Proof. The only use of randomness in the above proof is to compute the approximations q_r satisfying Property (1), and this randomness r needs to be shared so that both the encoder and decoder utilize the same approximations. Thus, it suffices to show how they can compute their approximations independently, yet have the approximations be equal with high probability. Roughly speaking, we do this by perturbing the approximations with random noise η and rounding. It turns out that the noise only needs to be specified to $O(\log n)$ bits and thus can be included as part of the compressed string.

We now proceed with the details. The randomness used by Enc and Dec consists of two parts — r , which is not shared, and η which will be shared (by explicit inclusion in the compressed string). To compute an approximation $q_{r,\eta}(x, z)$, we first use the algorithm $A(x, z, \epsilon, \delta; r)$ from Claim 7.2, setting $\delta = 2^{-3n}$ (as before) and $\epsilon = 1/n^{4c}$ (instead of $1/n^{2c}$). Then we take η , which is a random number in $\{0, 1, \dots, n^c - 1\}$, and set $q'_{r,\eta}(x, z)$ to equal $A(x, z, \epsilon, \delta; r) + \eta/n^{4c}$ rounded to the nearest multiple of $1/n^{3c}$. Note that the noise and rounding increase the error (in approximating $p(x, z)$) by at most $2/n^{3c}$. However, $q'_{r,\eta}(x, \cdot)$ no longer defines a probability measure (because the perturbations have all been positive). Thus we observe that we can (deterministically) convert $q'_{r,\eta}(x, \cdot)$ into a probability measure $q_{r,\eta}(x, \cdot)$, while reducing each entry by at most $2/n^{3c}$.

Notice that with probability at least $1 - 2^{-n}$ over r and η , the $q_{r,\eta}$'s satisfy the following analogue of Property (1):

$$|q_{r,\eta}(x, z) - p(x, z)| \leq \epsilon + O\left(\frac{1}{n^{3c}}\right) < \frac{1}{n^{2c}} \quad \forall |x| \leq n, |z| = \sigma(x).$$

Thus, if the encoding algorithm uses the $q_{r,\eta}$'s in place of the q_r 's, the bound on compression length will hold just as before, except that we add $O(\log n)$ bits to specify the noise $\eta \in \{0, \dots, n^c - 1\}$.

So all that remains is to argue that decoding is correct with high probability. For this, we argue that the encoder and decoder compute the same approximations with high probability. Specifically, we argue that for every x and z ,

$$\Pr_{r_1, r_2, \eta} [q_{r_1, \eta}(x, z) = q_{r_2, \eta}(x, z)] \geq 1 - 2/n^c. \quad (5)$$

First, by Claim 7.2, we know that with probability at least $1 - 2 \cdot 2^{-n}$, both $A(x, z, \epsilon, \delta; r_1)$ and $A(x, z, \epsilon, \delta; r_2)$ differ from $p(x, z)$ by at most $\epsilon = 1/n^{4c}$, so they differ from each other by at most $2/n^{4c}$. Thus there are at most two values of $\eta \in \{0, 1, \dots, n^c - 1\}$ such that $A(x, z, \epsilon, \delta; r_1)$ and $A(x, z, \epsilon, \delta; r_2)$ round to different multiples of $1/n^{3c}$.

To complete the proof, we argue that (whp) both Enc and Dec evaluate the $q_{r,\eta}$'s on some $p(n) = \text{poly}(n)$ inputs (x, z) where $p(n)$ is a fixed polynomial independent of the choice of the constant c , and the sequence of inputs is independent of r and η . Thus, by Inequality (5), that probability that the two algorithms “see” any difference in their approximations (and decoding possibly fails) is at most $p(n) \cdot (2/n^c)$. By Lemma 2.7, we can reduce this decoding error to 2^{-n} while increasing the compression length by at most $(2p(n)/n^c) \cdot \ell + 2 < 3$ bits for a sufficiently large constant c . So we proceed to argue that the number and sequence of evaluations of $q_{r,\eta}$ is indeed fixed (independent of c and the randomness). By inspection, we see that the arithmetic coding $\widehat{\text{Enc}}_{x,r,\eta}(w)$ (Lemma 2.4) only requires evaluating the cumulative distribution function $\hat{F}_{x,r,\eta}$ at w and its predecessor. By Equation (2), we see that evaluating $\hat{F}_{x,r,\eta}$ requires only a fixed polynomial number of evaluations of $q_{r,\eta}$ and the evaluation points are independent of r and η . This handles the encoding algorithm Enc. Now recall that the decoding algorithm decodes $\widehat{\text{Enc}}_{x,r,\eta}(w)$ by using $\hat{F}_{x,r,\eta}$ to do binary search for the sample w . By inspection, if the decoding algorithm were given the *same* function $q_{r,\eta}$ as the encoding algorithm, then the evaluations made in the binary search for w would be independent of r and η (because it would successfully traverse the path down to w). $\square$

The above results actually only require that X_x can be *approximately sampled* in the following sense.

Definition 7.2. A family of sources $(X_x)_{x \in L}$ is approximately samplable if there is a probabilistic algorithm S such that for every $x \in L$ and $\epsilon > 0$, the output $S(x, \epsilon)$ has statistical difference (i.e. variation distance) at most ϵ from X_x , and $S(x, \epsilon)$ runs in time $\text{poly}(|x|, 1/\epsilon)$.

Proposition 7.2. Let R be a self-reducible relation, and for every $x \in L_R$, let X_x be the uniform distribution on $W_x = \{w : (x, w) \in R\}$. If the sources $(X_x)_{x \in L_R}$ are approximately samplable, then they can be efficiently compressed to length $H(X_x) + 6$ with shared randomness and zero decoding error, and to length $H(X_x) + O(\log n)$ with independent randomness and decoding error 2^{-n} . The encodings are prefix-free.

Proof. In the proof of Theorem 5, both the encoding and decoding algorithms use the sampling algorithm for the distributions X_x only as an oracle to obtain samples from the distribution. Since they make only $\text{poly}(n)$ queries to the oracle, if we replace the oracle with a distribution at statistical difference ϵ , the statistical difference of the outcome (i.e. the compressed string, and an indicator for whether or not decoding is successful) will be at most $\epsilon \cdot \text{poly}(n)$. Choosing ϵ to be a sufficiently small polynomial, we can make this statistical difference smaller than $1/(2\ell')$, where ℓ' is the maximum encoding length. This implies that the average encoding length changes by at most $(1/(2\ell')) \cdot \ell' = 1/2$ and the probability of unsuccessful decoding is at most $1/(2\ell')$. Applying Lemma 2.5 completes the proof. $\square$

Thus, we obtain compression algorithms for the wide variety of self-reducible structures for which almost-uniform samplers are known. For example:

Corollary 7.1. The following families of sources X_x can be efficiently compressed to length $H(X_x) + 6$ with shared randomness and zero decoding error, and to length $H(X_x) + O(\log n)$ with independent randomness and decoding error 2^{-n} :

1. $X_G =$ the uniform distribution on all perfect matchings in bipartite graph G [16].
2. $X_G =$ the uniform distribution on all matchings in graph G [15].
3. $X_G =$ the uniform distribution on all independent sets in graph G of degree at most 4 [24].
4. $X_{(a_1, \dots, a_n, b)} =$ the uniform distribution on all “knapsack solutions”, i.e. subsets $S \subseteq [n]$ such that $\sum_{i \in S} a_i \leq b$, where $a_1, \dots, a_n, b$ are positive real numbers [25].
5. $X_\phi =$ the uniform distribution on satisfying assignments of DNF formula ϕ [20, 18].

The citations refer to the papers establishing the approximate samplability of the given distributions. Actually, for DNF formula, the ideas underlying the approximate counting algorithm of [20] directly yields a simple compression algorithm: given a satisfying assignment $w \in \{0, 1\}^t$ of a DNF formula $\phi = C_1 \vee \dots \vee C_m$ with minimum clause length k , we define $\text{Enc}_\phi(w)$ to be $(i, \alpha) \in [m] \times \{0, 1\}^{t-k}$, where C_i is the first clause satisfied by w and α is the restriction of w to the variables outside C_i . It is easy to check that this encoding is efficiently decodable, and compresses to length $\lceil \log m \rceil + t - k \leq \lceil \log m \rceil + H(X_\phi)$. Compressing to length $H(X_\phi) + O(1)$, however, seems less immediate.

The ability to compactly store combinatorial substructures of a graph (as in the above corollary) could be useful, for example, in storing substructures of huge graphs such as the World Wide Web;

indeed, there have been recent efforts at compressing Web graphs [1]. There are many other examples of self-reducible relations to which our technique can be applied; see the surveys [19, 17, 29] and the references therein.

In addition, we can show that compression and almost-uniform sampling are *equivalent*.

Theorem 6. *Let R be a self-reducible relation, and for every x , let X_x be the uniform distribution on $W_x = \{w : (x, w) \in R\}$. Then the following conditions are equivalent:*

1. X_x can be approximately sampled in polynomial time.
2. X_x can be compressed to length $H(X_x) + O(1)$ by probabilistic polynomial-time compression algorithms with shared randomness and zero decoding error.
3. X_x can be compressed to length $H(X_x) + O(\log n)$ by probabilistic polynomial-time compression algorithms with independent randomness and decoding error 2^{-n} .
4. X_x can be compressed to length $H(X_x) + O(\log n)$ by probabilistic polynomial-time compression algorithms with shared randomness and decoding error $1/n$.

Proof. By Proposition 7.2, sampling (Item 1) implies compression in the sense of Items 2 and 3. Each of these latter two items imply Item 4, so we need only argue that Item 4 implies Item 1. So suppose (Enc, Dec) compresses X_x to length $m \leq H(X_x) + c \log n$ with shared randomness. We may assume there is zero decoding error, by Lemma 2.6. By the results of Sinclair and Jerrum [35] (building on [18]), approximate sampling follows if we can approximate $|W_x|$ to within a $\text{poly}(n)$ accuracy factor in polynomial time. This would be easy if we could estimate the average compressed length m ; unfortunately, random sampling from X_x is unavailable to us.

Instead, we use random sampling from the compressed space and decompressing. In particular, we will use random sampling to estimate

$$p_\ell = \mathbb{E}_r \left[\Pr_{y \leftarrow U_{\leq \ell}} [\text{Dec}(x, y, r) \in W_x \& \text{Enc}(x, \text{Dec}(x, y, r), r) = y] \right],$$

where $U_{\leq \ell}$ denotes the uniform distribution on $\{0, 1\}^{\leq \ell}$, the set of strings of length $\leq \ell$. By sampling, with high probability we can find an integer $\hat{m}$ such that $p_{\hat{m}} \geq 1/(8 \cdot n^c \cdot (m+1))$ and $p_i < 1/(4 \cdot n^c \cdot (m+1))$ for all $i > \hat{m}$. (Note that we need only estimate p_i for i up to, say, n times the running time of Enc , because beyond that, p_i is exponentially small.)

We claim that $2^{\hat{m}}$ approximates $|W_x|$ to within a polynomial factor. For one direction, note that when we restrict to y 's satisfying the condition $\text{Enc}(x, \text{Dec}(x, y, r), r) = y$, the mapping $y \mapsto \text{Dec}(x, y, r)$ is injective. Thus,

$$|W_x| \geq p_{\hat{m}} \cdot |\{0, 1\}^{\leq \hat{m}}| \geq \frac{2^{\hat{m}}}{8 \cdot n^c \cdot (m+1)}.$$

For the other direction, note that Markov's inequality implies that

$$\Pr_{w \leftarrow X_x, r} [|\text{Enc}(x, w, r)| \leq m+1] \geq 1 - \frac{m}{m+1} = \frac{1}{m+1}.$$

Therefore, the expected number of encodings of W_x with length at most $m + 1$ is at least

$$\frac{|W_x|}{m+1} \geq \frac{2^m}{n^c} \cdot \frac{1}{m+1} > \frac{|\{0,1\}^{\leq m+1}|}{4 \cdot n^c \cdot (m+1)},$$

Hence $p_{m+1} \geq 1/(4 \cdot n^c \cdot (m+1))$ and thus with high probability, $\hat{m} \geq m+1 \geq H(X_x) - O(\log n)$ (by Lemma 2.2) and thus $2^{\hat{m}} \geq |W_x|/\text{poly}(n)$. $\square$

A final extension we mention is that our results also apply to some non-uniform distributions on the witness set $\{w : (x, w) \in R\}$. Specifically, it applies to sources X_x that are compatible with the self-reduction in the following sense.

Definition 7.3. Let R be a self-reducible NP relation, with corresponding functions $\ell : \Sigma^* \rightarrow \mathbb{N}$, $\sigma : \Sigma^* \rightarrow \mathbb{N}$, and $\rho : \Sigma^* \times \Sigma^* \rightarrow \Sigma^*$ as in Definition 7.1. We say that the sources $(X_x)_{x \in L_R}$ are compatible with R (and ℓ, ρ, σ) if

1. The support of X_x is a subset of $\Sigma^{\ell(x)}$.
2. When $\ell(x) > 0$ (equivalently, $\sigma(x) > 0$), then for every $z \in \Sigma^{\sigma(x)}$ such that X_x has nonzero probability of having prefix z , the distribution of X_x conditioned on having prefix z is precisely $z \circ X_{\rho(x,z)}$.

The above conditions imply that for every $x \in L_R$, the support of X_x is a subset of $W_x = \{w : (x, w) \in R\}$. It can be verified that setting X_x equal to the uniform distribution on W_x is compatible with R . An example of a non-uniform family of sources compatible with a self-reducible relation is the following generalization of Example 7.1:

Example 7.2 (weighted perfect matchings). Let R be the relation consisting of pairs $((G, w), M)$, where G is a bipartite graph with positive real weights $w(e)$ on each edge and M is a perfect matching in G . G and w are encoded by the $n \times n$ weighted adjacency matrix whose (i, j) 'th entry is $w(i, j)$ if (i, j) is an edge, and 0 otherwise. This relation is self-reducible for the same reason as Example 7.1. We define the distribution $X_{G,w}$ to be the one where a perfect matching M is sampled with probability proportional to its weight $w(M) = \prod_{e \in M} w(e)$. (Note that the total weight $\sum_M w(M)$ equals the permanent of the weighted adjacency matrix.) It can be verified that these distributions are compatible with the self-reducibility of the relation R (e.g. when we remove an edge e and its endpoints, every perfect matching M in G that contains e becomes a perfect matching in $G \setminus \{e\}$ with weight $w(M \setminus \{e\}) = w(M)/w(e)$.)

We can also compress such distributions:

Theorem 7. Let R be a self-reducible relation, and let $(X_x)_{x \in L_R}$ be a family of sources compatible with R . If the sources $(X_x)_{x \in L_R}$ are approximately samplable, then they can be efficiently compressed to length $H(X_x) + 6$ with shared randomness and zero decoding error and to length $H(X_x) + O(\log n)$ with independent randomness and decoding error 2^{-n} . The encodings are prefix-free.

Proof. The proof is identical to that of Theorem 5 and Proposition 7.2. The only use of the fact that X_x equals the uniform distribution on W_x is in the proof of Claim 7.2, specifically to establish Equation (3). By inspection, this equation holds for any family of sources compatible with R . $\square$

Many of the known approximate sampling algorithms for self-reducible relations generalize to natural non-uniform distributions that are compatible with the relation. Often, these distributions have interpretations in statistical physics (namely being the “Gibbs distribution” of some physical system). Some examples follow.

Corollary 7.2. *The following families of sources X_x can be efficiently compressed to length $H(X_x) + 6$ with shared randomness and zero decoding error, and to length $H(X_x) + O(\log n)$ with independent randomness and decoding error 2^{-n} :*

1. $X_{G,w}$ = perfect matchings in weighted bipartite graph (G, w) , as in Example 7.2 (a.k.a. the Gibbs distribution on a dimer system) [16].
2. $X_{G,w}$ = matchings on a weighted graph (G, w) , where the weights are presented in unary (a.k.a. the Gibbs distribution on monomer-dimer systems) [15].
3. $X_{G,\lambda}$ = the weighted distribution on independent sets in graph G , where independent set I has weight $\lambda^{|I|}$, and the maximum degree of G is at most $2/\lambda + 2$ (a.k.a. the Gibbs distribution for the hard-core gas model) [24].

Monte Carlo experiments in statistical physics estimate the expectation of various quantities in a physical system (such as the “mean energy”) by randomly sampling configurations of the system (e.g. according to the Gibbs distribution). Compression algorithms such as in Corollary 7.2 could be possible to compactly store the configurations used in such experiments (e.g. for archival purposes, or to reuse the samples later).

Acknowledgements

We thank Boaz Barak, Nenad Dedić, Ned Dimitrov, and Troy Lee for helpful comments and discussions.

References

- [1] M. Adler and M. Mitzenmacher. Toward compressing web graphs. In *Proceedings of the 2001 Data Compression Conference*, 2001.
- [2] Eric Allender, Danilo Bruschi, and Giovanni Pighizzini. The complexity of computing maximal word functions. *Computational Complexity*, 3(4):368–391, 1993.
- [3] S. Arora, F. T. Leighton, and B. M. Maggs. On-line algorithms for path selection in a non-blocking network. *SIAM Journal on Computing*, 25(3):600–625, 1996.
- [4] B. Barak, R. Shaltiel, and A. Wigderson. Computational analogues of entropy. In *11th International Conference on Random Structures and Algorithms*, 2003.
- [5] Harry Buhrman, Troy Lee, and Dieter van Melkebeek. Language compression and pseudo-random generators. In *Proceedings of the 19th Annual IEEE Conference on Computational Complexity*, pages 15–28, 2004.

- [6] M. Capalbo, O. Reingold, S. Vadhan, and A. Wigderson. Randomness conductors and constant-degree lossless expanders. In *Proceedings of the 34th Annual ACM Symposium on Theory of Computing*, pages 659–668, 2002.
- [7] B. Chor and O. Goldreich. Unbiased bits from sources of weak randomness and probabilistic communication complexity. *SIAM Journal on Computing*, 17(2):230–261, 1988.
- [8] T. M. Cover and J. A. Thomas. *Elements of Information Theory*. John Wiley & Sons, Inc., 1991.
- [9] Whitfield Diffie and Martin E. Hellman. New directions in cryptography. *IEEE Transactions in Information Theory*, IT-22(6):644–654, 1976.
- [10] A.V. Goldberg and M. Sipser. Compression and ranking. *SIAM Journal on Computing*, 20:524–536, 1991.
- [11] O. Goldreich and S. Vadhan. Comparing entropies in statistical zero-knowledge with applications to the structure of SZK. In *Proc. of Conference on Computational Complexity*, pages 54–73, 1999.
- [12] J. Håstad, R. Impagliazzo, L. Levin, and M. Luby. A pseudorandom generator from any one-way function. *SIAM Journal on Computing*, 28:1364–1396, 1999.
- [13] R. Impagliazzo, October 1999. Remarks in Open Problem session at the DIMACS Workshop on Pseudorandomness and Explicit Combinatorial Constructions.
- [14] R. Impagliazzo and A. Wigderson. $P = BPP$ if E requires exponential circuits: Derandomizing the XOR lemma. In *Proceedings of the 29th Annual ACM Symposium on Theory of Computing*, pages 220–229, 1997.
- [15] M. Jerrum and A. Sinclair. Approximating the permanent. *SIAM Journal on Computing*, 18(6):1149–1178, 1989.
- [16] M. Jerrum, A. Sinclair, and E. Vigoda. A polynomial-time approximation algorithm for the permanent of a matrix with non-negative entries. In *Proceedings of the 33rd Annual ACM Symposium on Theory of Computing*, pages 712–721, 2001.
- [17] Mark Jerrum and Alistair Sinclair. The Markov chain Monte Carlo method: an approach to approximate counting and integration. In D.S. Hochbaum, editor, *Approximation Algorithms for NP-hard Problems*, chapter 12, pages 482–520. PWS Publishing, 1996.
- [18] M.R. Jerrum, L.G. Valiant, and V.V. Vazirani. Random generation of combinatorial structures from a uniform distribution. *Theoretical Computer Science*, 43:169–188, 1986.
- [19] Ravi Kannan. Markov chains and polynomial time algorithms. In *Proceedings of the 35th Annual IEEE Symposium on Foundations of Computer Science*, pages 656–671. IEEE Comput. Soc. Press, Los Alamitos, CA, 1994.
- [20] Richard M. Karp, Michael Luby, and Neal Madras. Monte Carlo approximation algorithms for enumeration problems. *Journal of Algorithms*, 10(3):429–448, 1989.

- [21] M. Kharitonov, A. V. Goldberg, and M. Yung. Lower bounds for pseudorandom number generators. In *Proceedings of the 30th Annual IEEE Symposium on Foundations of Computer Science*, pages 242–247, 1989.
- [22] Ming Li and Paul Vitanyi. *An introduction to Kolmogorov complexity*. Springer, 1997. 2nd ed.
- [23] R.J. Lipton. A new approach to information theory. In *Proc. of 11th Symposium on Theoretical Aspects of Computer Science*, pages 699–708, 1994.
- [24] Michael Luby and Eric Vigoda. Fast convergence of the Glauber dynamics for sampling independent sets. *Random Structures & Algorithms*, 15(3-4):229–241, 1999. Statistical physics methods in discrete probability, combinatorics, and theoretical computer science (Princeton, NJ, 1997).
- [25] Ben Morris and Alistair Sinclair. Random walks on truncated cubes and sampling 0-1 knapsack solutions (preliminary version). In *Proceedings of the 40th Annual IEEE Symposium on Foundations of Computer Science*, pages 230–240. IEEE, 1999.
- [26] N. Nisan and A. Wigderson. Hardness vs. randomness. *Journal of Computer and System Sciences*, 49:149–167, 1994.
- [27] N. Nisan and D. Zuckerman. Randomness is linear in space. *Journal of Computer and System Sciences*, 52(1):43–52, 1996.
- [28] Rafail Ostrovsky. One-way functions, hard on average problems, and statistical zero-knowledge proofs. In *Proceedings of the Sixth Annual Structure in Complexity Theory Conference*, pages 133–138, Chicago, Illinois, 30 June–3 July 1991. IEEE Computer Society Press,.
- [29] Dana Randall. Mixing. In *Proceedings of the 44th Annual IEEE Symposium on Foundations of Computer Science*, pages 4–15, Cambridge, MA, October 2003. IEEE.
- [30] R. Raz and O. Reingold. On recycling the randomness of states in space bounded computation. In *Proceedings of the 31st Annual ACM Symposium on Theory of Computing*, pages 159–168, 1999.
- [31] Amit Sahai and Salil Vadhan. A complete problem for statistical zero knowledge. *Journal of the ACM*, 50(2):196–249, March 2003. Extended abstract in *FOCS '97*.
- [32] M. Santha and U. V. Vazirani. Generating quasi-random sequences from semi-random sources. *Journal of Computer and System Sciences*, 33:75–87, 1986.
- [33] C.P. Schnorr. Optimal algorithms for self-reducible problems. In *Proceedings of the 3rd International Colloquium on Automata, Languages, and Programming*, pages 322–337, 1976.
- [34] Claude E. Shannon. Communication theory of secrecy systems. *Bell System Technical Journal*, 28:656–715, 1949.
- [35] A.J. Sinclair and M.R. Jerrum. Approximate counting, uniform generation and rapidly mixing Markov chains. *Information and Computation*, 82:93–133, 1989.

- [36] A. Ta-Shma, C. Umans, and D. Zuckerman. Loss-less condensers, unbalanced expanders, and extractors. In *Proceedings of the 33rd Annual ACM Symposium on Theory of Computing*, pages 143–152, 2001.
- [37] Luca Trevisan, Salil Vadhan, and David Zuckerman. Compression of samplable sources. In *Proceedings of the 19th Annual IEEE Conference on Computational Complexity*, pages 1–14, Amherst, MA, 21–24 June 2004. IEEE.
- [38] Luca Trevisan and Salil P. Vadhan. Extracting randomness from samplable distributions. In *Proceedings of the 41st Annual IEEE Symposium on Foundations of Computer Science*, pages 32–42, 2000.
- [39] L. G. Valiant and V. V. Vazirani. NP is as easy as detecting unique solutions. *Theoretical Computer Science*, 47(1):85–93, 1986.
- [40] Hoeteck Wee. On pseudoentropy versus compressibility. In *Proceedings of the 19th Annual IEEE Conference on Computational Complexity*, pages 29–41, 2004.
- [41] A. C. Yao. Theory and applications of trapdoor functions. In *Proceedings of the 23rd Annual IEEE Symposium on Foundations of Computer Science*, pages 80–91, 1982.
- [42] J. Ziv and A. Lempel. Compression of individual sequences by variable rate coding. *IEEE Transactions on Information Theory*, 24:530–536, 1978.
- [43] D. Zuckerman. Randomness-optimal oblivious sampling. *Random Structures and Algorithms*, 11:345–367, 1997.