



Fast and Perfect Sampling of Subgraphs and Polymer Systems

ANTONIO BLANCA, Pennsylvania State University, USA

SARAH CANNON, Claremont McKenna College, USA

WILL PERKINS, Georgia Institute of Technology, USA

We give an efficient perfect sampling algorithm for weighted, connected induced subgraphs (or *graphlets*) of rooted, bounded degree graphs. Our algorithm utilizes a vertex-percolation process with a carefully chosen rejection filter and works under a percolation subcriticality condition. We show that this condition is optimal in the sense that the task of (approximately) sampling weighted rooted graphlets becomes impossible in finite expected time for infinite graphs and intractable for finite graphs when the condition does not hold. We apply our sampling algorithm as a subroutine to give near linear-time perfect sampling algorithms for polymer models and weighted non-rooted graphlets in finite graphs, two widely studied yet very different problems. This new perfect sampling algorithm for polymer models gives improved sampling algorithms for spin systems at low temperatures on expander graphs and unbalanced bipartite graphs, among other applications.

CCS Concepts: • **Theory of computation** → **Random walks and Markov chains**; **Generating random combinatorial structures**; • **Mathematics of computing** → **Probabilistic algorithms**; **Stochastic processes**;

Additional Key Words and Phrases: Sampling algorithms, subgraphs, polymer models, spin systems, approximate counting

ACM Reference format:

Antonio Blanca, Sarah Cannon, and Will Perkins. 2024. Fast and Perfect Sampling of Subgraphs and Polymer Systems. *ACM Trans. Algor.* 20, 1, Article 5 (January 2024), 30 pages.

<https://doi.org/10.1145/3632294>

1 INTRODUCTION

Sampling is a fundamental computational task: given a specification of a probability distribution on a (large) set of combinatorial objects, output a random object with the specified distribution or with a distribution close to the specified distribution. This task becomes challenging when the specification of the distribution is much more succinct than the set of objects, and one wants to sample using time and space commensurate with the specification. Fundamental examples include

WP is supported in part by NSF grant DMS-2309958 and CCF-2309708. SC is supported in part by NSF grants DMS-1803325 and CCF-2104795. AB is supported in part by NSF grant CCF-2143762.

Authors' addresses: A. Blanca, Pennsylvania State University, Computer Science and Engineering Department, Westgate Building, University Park, PA 16801 USA; e-mail: ablanca@cse.psu.edu; S. Cannon, Claremont McKenna College, Mathematical Sciences Department, 888 Columbia Avenue, Claremont, CA 91711 USA; e-mail: scannon@cmc.edu; W. Perkins, School of Computer Science, Georgia Institute of Technology, 266 Ferst Dr, Atlanta, GA 30332 USA; e-mail: wperkins3@gatech.edu. Permission to make digital or hard copies of all or part of this work for personal or classroom use is granted without fee provided that copies are not made or distributed for profit or commercial advantage and that copies bear this notice and the full citation on the first page. Copyrights for components of this work owned by others than the author(s) must be honored. Abstracting with credit is permitted. To copy otherwise, or republish, to post on servers or to redistribute to lists, requires prior specific permission and/or a fee. Request permissions from [permissions@acm.org](https://permissions.acm.org).

© 2024 Copyright held by the owner/author(s). Publication rights licensed to ACM.

1549-6325/2024/01-ART5 \$15.00

<https://doi.org/10.1145/3632294>

sampling from Markov random fields and probabilistic graphical models and sampling substructures of graphs. We will address both of these examples here and connect them in a new way.

We consider a natural sampling problem: given a bounded-degree graph G , sample a *graphlet* (a connected, vertex-induced subgraph) of G containing a fixed vertex r with probability proportional to an exponential in the size of the subgraph. That is, sample a graphlet S containing vertex r with probability proportional to $\lambda^{|S|}$, where $\lambda > 0$ is a distribution parameter and $|S|$ denotes the number of vertices in S . In the article, we are concerned with small values of λ , where the expected size of a sampled graphlet is much smaller than the size of the graph.

Sampling graphlets is an important task in data science, network analysis, bioinformatics, and sociology, as it allows us to gain information about massive graphs from small sections of it; see, e.g., References [5, 29, 42, 46]. A number of variants of the problem have consequently been studied, including sampling graphlets of a given size uniformly at random or sampling weighted graphlets of all sizes [1, 6, 10–13, 16, 40, 47, 48, 50, 51]. The variant we consider here, i.e., sample a graphlet S with probability proportional to $\lambda^{|S|}$, arises as a key subroutine in recent sampling algorithms for spin systems (hard-core model, Ising model, Potts model, etc.) in the regime of strong interactions via *polymer models* described below in Section 1.2; see References [7, 14, 15, 18, 26, 35, 38, 39, 45].

One major limitation of previous sampling algorithms for graphlets and polymer models (those in, e.g., [17, 25, 35, 45, 50], among others) is the use of exhaustive enumeration of graphlets of a given size; this requires restrictive parameter regimes or large polynomial running times, with the logarithm of the maximum degree Δ of the graph appearing in the exponent of the polynomial. Here, we design a fast perfect sampling algorithm for weighted graphlets based on a vertex percolation process combined with a rejection filter. This method bypasses the enumeration barrier and allows us to design perfect sampling algorithms for a number of applications, substantially improving upon existing algorithms in three ways: (1) our algorithms have considerably faster running times, with no dependence on Δ in the exponent; (2) our algorithms return perfect, rather than approximate, samples from the desired distributions; and (3) our algorithms are conceptually simple and practical to implement.

Our algorithm proceeds as follows. First, run a vertex percolation process on the graph G beginning at vertex r in a **breadth-first search (BFS)** manner, repeatedly adding each adjacent vertex to the graphlet with a carefully-chosen probability p . Once the percolation process terminates, the graphlet is accepted as the random sample with a certain probability that depends on the graphlet and rejected otherwise; if the graphlet is rejected, the algorithm restarts another percolation process from r . Because of the careful way we choose the percolation and rejection probabilities, we can prove the final accepted sample is drawn exactly from the desired distribution and the expected running time is bounded by a constant that depends only on λ and the maximum degree Δ .

For our applications to polymer models, we use this graphlet sampling algorithm as a subroutine to implement a Markov chain on polymer configurations. We then use this Markov chain to devise a perfect sampling algorithm for polymer models based on the **coupling from the past (CFTP)** method from Reference [49] and the notion of bounding chains from Reference [36].

1.1 Sampling Rooted Graphlets

Our key contribution is a new algorithm for perfectly sampling weighted graphlets containing a given vertex r . We start by fixing the model of computation we work with throughout the article. We assume a model that allows for querying the adjacency list of a given vertex in a bounded degree graph in constant time, including in a rooted infinite graph. We also assume that, in a finite graph, we can query a uniformly random vertex in constant time. This is a standard model used in the study of sublinear algorithms [28]. We also assume access to a stream of perfectly random real numbers in $[0, 1]$. The model of computation is fixed for consistency; in particular, our methods

extend to other models, only requiring to adjust the running time to account for any additional computational overhead.

Let $G = (V, E)$ be a finite or infinite graph of maximum degree Δ . For $r \in V$, let $\mathcal{S}(G, r)$ be the set of all connected, vertex-induced subgraphs of G containing r . (The subgraph induced by $U \subseteq V$ has vertex set U and includes all the edges of G with both endpoints in U .) We call r the root of G and the elements of $\mathcal{S}(G, r)$ graphlets rooted at r . For $\lambda > 0$ define the probability distribution $\nu_{G,r,\lambda}$ on $\mathcal{S}(G, r)$ by

$$\nu_{G,r,\lambda}(S) = \frac{\lambda^{|S|}}{Z_{G,r,\lambda}}, \quad \text{where} \quad Z_{G,r,\lambda} = \sum_{\hat{S} \in \mathcal{S}(G, r)} \lambda^{|\hat{S}|}. \quad (1)$$

The distribution is well defined when the normalizing constant $Z_{G,r,\lambda}$, known as the partition function, is finite. This is the case for every graph of maximum degree Δ and every r when λ is below the critical threshold:

$$\lambda_*(\Delta) = \frac{(\Delta - 2)^{\Delta-2}}{(\Delta - 1)^{\Delta-1}}; \quad (2)$$

see Lemma 1.3 below. This threshold was already considered in Reference [50], who gave an ε -approximate sampling algorithm for $\nu_{G,r,\lambda}$ for the class of maximum-degree Δ graphs when $\lambda < \lambda_*(\Delta)$ with running time $\text{poly}(\varepsilon^{-1})$. We give a perfect sampling algorithm for $\nu_{G,r,\lambda}$ for $\lambda < \lambda_*(\Delta)$ with constant expected running time.

THEOREM 1.1. *Fix $\Delta \geq 3$ and let $\lambda < \lambda_*(\Delta)$. There is a randomized algorithm that for any graph $G = (V, E)$ of maximum degree Δ and any $r \in V$ outputs a graphlet distributed according to $\nu_{G,r,\lambda}$ with expected running time bounded by a constant that depends only on Δ and λ .*

Previous algorithms to generate ε -approximate samples from $\nu_{G,r,\lambda}$ (e.g., those in References [17, 25, 35, 50]) exhaustively enumerate all graphlets of size $\leq k$, for some k that depends on the error parameter ε that describes how accurate the sample must be. This results in algorithms with $(1/\varepsilon)^{O(\log \Delta)}$ running times. Applications such as sampling from polymer models require multiple samples from $\nu_{G,r,\lambda}$ and have small error tolerance per sample; in particular, they require $\varepsilon \ll 1/n$, which results in inefficient algorithms with overall running time $n^{O(\log \Delta)}$. The algorithm in Theorem 1.1, on the other hand, is an exact sampler whose expected running time depends only on Δ and λ and thus provides a significant advantage in applications as we detail below.

We also show that Theorem 1.1 is sharp in two ways. First, we establish that there is no polynomial-time approximate sampling algorithm for $\nu_{G,r,\lambda}$ when $\lambda \in (\lambda_*(\Delta), 1)$ for the class of finite graphs of maximum degree at most Δ unless $\text{RP} = \text{NP}$; see Definition 5.3 for the precise definition of a polynomial-time approximate sampler. (We note that a similar hardness result was proved in Reference [50] for the related problem of sampling “unrooted graphlets”; we provide more details about this in Section 1.4.) Second, in the infinite setting, the normalizing constant $Z_{G,r,\lambda}$ may diverge (and consequently, the distribution $\nu_{G,r,\lambda}$ is not well-defined) when $\lambda > \lambda_*(\Delta)$; conversely, we prove that $Z_{G,r,\lambda}$ is finite on every graph of maximum degree Δ when $\lambda \leq \lambda_*(\Delta)$.

LEMMA 1.2. *Fix $\Delta \geq 3$ and $\lambda \in (\lambda_*(\Delta), 1)$. If there is a polynomial-time approximate sampler for $\nu_{G,r,\lambda}$ for finite graphs $G = (V, E)$ of maximum degree Δ and each $r \in V$, then $\text{RP} = \text{NP}$.*

LEMMA 1.3. *The partition function $Z_{G,r,\lambda}$ is finite for every (possibly infinite) graph $G = (V, E)$ of maximum degree Δ and every $r \in V$ if and only if $\lambda \leq \lambda_*(\Delta)$.*

Finally, we mention that the algorithmic result in Theorem 1.1 cannot be extended even to the case $\lambda = \lambda_*(\Delta)$: for the infinite Δ -regular tree, we can show that the expected size of a graphlet sampled from $\nu_{G,r,\lambda}$ is infinite when $\lambda = \lambda_*(\Delta)$, and so it is impossible to have sampling algorithms

with finite expected running time. In summary, the algorithm in Theorem 1.1 for $\lambda < \lambda_*(\Delta)$, combined with the hardness/impossibility results in Lemmas 1.2 and 1.3 for $\lambda > \lambda_*(\Delta)$, provide a resolution to the computational problem of sampling from $\nu_{G,r,\lambda}$ on graphs of maximum degree at most Δ .

As mentioned, our sampling algorithm is based on exploring the connected component of r in a vertex-percolation process. We carefully choose a specific percolation parameter $p \in (0, 1)$ as a function of λ and Δ (see Lemma 2.2). We then perform BFS from r , labeling each new vertex encountered “active” with probability p and “inactive” with probability $1 - p$ independently over all vertices; we continue the BFS exploring only unexplored neighbors of active vertices. In this way, we uncover the “active” component of r , call it γ . We then accept γ with a given probability depending on $\lambda, \Delta, |\gamma|$ and $|\partial\gamma|$, where $\partial\gamma$ denotes the set of vertices outside of γ that are adjacent to γ . If we reject γ , we begin again with a new percolation process. We note that only when $\lambda < \lambda_*(\Delta)$ there exists a suitable percolation probability p that results in a subcritical percolation process, so that the size of the active component has finite expectation and exponential tails. The weighted model we sample from is particularly well suited to this type of exploration algorithm because of the direct connection to a subcritical percolation process.

Random exploration and rejection sampling have been used previously to sample graphlets and other structures, most notably in the recent work of Bressan [10], who uses a novel bucketing scheme in combination with rejection sampling to perfectly sample uniformly random graphlets of size k from a graph, as well as studying the mixing time of the random walk on the set of all such graphlets. See also Reference [3], in which a random growth process and rejection sampling are used to perfectly sample spin configurations.

We prove a more general version of Theorem 1.1 in Section 2, allowing for vertex-labeled graphlets and modifications of the weights by multiplication by a non-negative function bounded by 1. These generalizations are needed for the application to polymer models in Section 1.2.

1.2 Sampling from Polymer Models

We use our algorithm for sampling weighted rooted graphlets to design fast and perfect samplers for polymer models. Polymer models are systems of interacting geometric objects representing defects from pure ground states (i.e., most likely configurations) in spin systems on graphs in classical statistical physics [23, 30, 43]. These geometric objects are most often represented by vertex-labeled graphlets from a given host graph. Recently, polymer models have found application as an algorithmic tool to sample from spin systems on various classes of graphs in strong interaction regimes; see, e.g., References [7, 14, 15, 17, 18, 21, 25, 26, 34, 35, 38, 39, 45]. In these applications, the problem of sampling weighted vertex-labeled rooted graphlets emerged as a significant computational barrier.

We will work with *subset polymer models* in which all polymers are vertex-labeled graphlets from a host graph $G = (V, E)$. These models were defined in Reference [30] and generalized in Reference [43]. Such a polymer model consists of:

- A set $C = C(G)$ of polymers, each of which is a graphlet in G with the vertices of the graphlet labeled with colors from a set Σ of size q .
- Weights $w_\gamma \geq 0$ for each $\gamma \in C$. We assume without loss of generality that all vertex-labeled graphlets of G , including each individual vertex $v \in V$, are elements of C , by setting $w_\gamma = 0$ when necessary.
- An incompatibility relation $\sim$ defined by connectivity. We say two polymers $\gamma, \gamma' \in C$ are *incompatible* and write $\gamma \sim \gamma'$ if the union of their corresponding vertices induces a connected subgraph in G . Otherwise they are *compatible* and we write $\gamma \not\sim \gamma'$.

Table 1. Comparison of Conditions and Running Times of Known Polymer Sampling Algorithms

Condition	Bound on exponential decay of weights	Running Time	Type of Sampler
Kotecký–Preiss [14, 35, 38]	$w_\gamma \leq (e^2 q \Delta)^{- \gamma }$	$n^{O(\log \Delta)}$	approximate
Polymer Sampling [17, 25]	$w_\gamma \leq (e^5 q^3 \Delta^3)^{- \gamma }$	$O(n \log n)$	approximate
Clique dynamics [24]	$w_\gamma \leq (eq \Delta)^{- \gamma }$	$n^{O(\log \Delta)}$	approximate
This work (Theorem 1.4)	$w_\gamma \leq (eq \Delta)^{- \gamma }$	$O(n \log n)$	perfect

Let $\Omega(C)$ denote the set of all sets of pairwise compatible polymers from C . The polymer model is the Gibbs distribution μ on $\Omega(C)$ defined by

$$\mu(X) = \frac{\prod_{\gamma \in X} w_\gamma}{Z(C)}, \quad \text{where} \quad Z(C) = \sum_{X \in \Omega(C)} \prod_{\gamma \in X} w_\gamma$$

is the polymer model partition function. The size $|\gamma|$ of a polymer γ is the number of vertices in the corresponding graphlet. We let C_v be all polymers containing vertex v .

We will often work with a family of polymer models corresponding to an infinite family of host graphs G . We say the weights of a family of polymer models are *computationally feasible* if w_γ can be computed in time polynomial in $|\gamma|$ uniformly over the polymer models in the family.

Algorithms for sampling polymer models fall into two classes: those based on truncating the cluster expansion of a polymer model to approximate a partition function and using self-reducibility to sample, and those based on Markov chains on the set of collections of compatible polymers. The cluster expansion approach, while giving polynomial-time algorithms, is relatively inefficient in general, with the degree of the polynomial bound on the running time growing with the degree of the underlying graph; e.g., running time $n^{O(\log \Delta)}$ for n -vertex graph of maximum degree Δ . A Markov chain approach based on adding and remove polymers from a polymer configuration in principle can be much faster (near linear time in the size of the graph) but runs into one hitch: a much stricter condition on the parameters of the model is needed to perform one update step of the Markov chain efficiently (the “polymer sampling condition” in References [17, 25]). We solve this problem by adapting our rooted graphlet sampler to sample polymers models, leading to a near linear-time perfect sampling algorithm for polymer models under the least restrictive conditions known (see Table 1).

THEOREM 1.4. *Fix $\Delta \geq 3$, $q \geq 1$, $\theta \in (0, 1)$, and $\lambda < \lambda_*(\Delta, q) := \frac{(\Delta-2)^{\Delta-2}}{q(\Delta-1)^{\Delta-1}}$. There is a perfect sampling algorithm for μ with expected running time $O(|V| \log |V|)$ for any family of subset polymer models on maximum degree Δ graphs $G = (V, E)$ with computationally feasible weights satisfying:*

$$w_\gamma \leq \lambda^{|\gamma|} \text{ for all } \gamma \in C; \text{ and} \quad (3)$$

$$\sum_{\gamma \ni v} |\gamma| w_\gamma \leq \theta \text{ for all } v \in V. \quad (4)$$

The threshold defined in Equation (3) is the generalization of the critical threshold for rooted graphlet sampling to the labeled case (taking $q = 1$ recovers the definition in Equation (2)). Theorem 1.4 improves upon the known results for sampling from polymer models in two ways. For a very general class of polymer models, our algorithm simultaneously provides *perfect sampling* with *near-linear running time* under weak conditions on the polymer weights. We now review previous works to illustrate these improvements; see the accompanying Table 1.

A number of conditions on polymer weights have been used to provide efficient sampling algorithms. The first papers in this direction (including References [14, 35, 38]) used the Kotecký–Preiss condition for convergence of the cluster expansion of the polymer model partition function [43]: $\sum_{Y' \neq Y} w_{Y'} e^{|Y'|} \leq |Y| \forall Y \in C$. This condition is typically verified by ensuring that:

$$\sum_{Y \ni v} w_Y e^{|Y|} \leq 1 \quad \forall v \in V. \quad (5)$$

Since the number of vertex-labeled rooted graphlets of size k in a maximum degree Δ graph grows roughly like $(eq\Delta)^{k-1}$ (see Reference [8]), weights of polymers of size k must decay roughly like $(e^2q\Delta)^{-k}$ for the polymer model to satisfy Equation (5), with the extra factor of e coming from the exponential in the left hand side of the condition Equation (5).

The major downside to the algorithms based on the cluster expansion, i.e., those using Equation (5) or the Kotecký–Preiss condition, is that the running times obtained are of the form $n^{O(\log \Delta)}$. Subsequent works, namely [17, 25], addressed this downside but at the cost of a significantly stricter condition on the polymer weights.

In Reference [17], the authors devised a new Markov chain algorithm for sampling from polymer models. The condition on the polymer weights for rapid mixing of this chain is somewhat less restrictive than the Kotecký–Preiss condition; it is the Polymer Mixing condition:

$$\sum_{Y' \neq Y} |Y'| w_{Y'} \leq \theta |Y| \quad \forall Y \in C \text{ for some } \theta \in (0, 1). \quad (6)$$

This requires weights of polymers of size k to decay like $(eq\Delta)^{-k}$, a savings of factor e in the base of the exponent over (5). However, to implement a single step of this Markov chain in constant expected time, a stronger condition (the Polymer Sampling condition) was required:

$$w_Y \leq (e^5 \Delta^3 q^3)^{-|Y|}. \quad (7)$$

This is a significant loss of a factor $e^3 \Delta^2 q^2$ in the base of the exponent compared with Equation (5), but the resulting sampling algorithm does run in near linear time.

In Reference [24], the authors use a different Markov chain condition, the Clique Dynamics condition, similar to Equation (6), which requires weights of polymers of size k to decay like $(eq\Delta)^{-k}$, saving the same factor e over (5). Their running times, though, are again of the form $n^{O(\log \Delta)}$ since implementing one step of their Markov chain involves enumerating rooted polymers of size $O(\log n)$.

Our results are a “best-of-both-worlds” for polymer sampling: under the conditions Equations (3) and (4) that both require polymer weights to decay like $(eq\Delta)^{-k}$ (this is shown later; see, e.g., the proof of Corollary 1.6), we obtain a near linear time algorithm. Moreover, unlike any of the previous results, our algorithm is a perfect sampler.

To conclude this section, we comment briefly on the algorithm we design to sample from μ . Our starting point is the polymer dynamics Markov chain from Reference [17]. We use it to implement a CFTP algorithm (see Reference [49]). To do so efficiently (in terms of the number of steps of the Markov chain), we design a new “bounding Markov chain” for the polymer dynamics, a method pioneered in References [31, 36], and to implement each step of the Markov chain efficiently, we turn to our sampler for weighted rooted graphlets from Theorem 1.1.

1.3 Applications to Spin Systems

Our new algorithm for sampling subset polymer models can be used as a subroutine in essentially all previous applications of polymer models for spin system sampling at low temperatures, including those in [14, 15, 17, 18, 21, 25, 26, 34, 38, 45]. This results in faster sampling algorithms under

less restrictive conditions on model parameters in all those settings. As examples, we fleshed out here the details in two of these applications; more details are provided in Section 4.

Hard-core model on bipartite graphs. The hard-core model on a graph G is the probability distribution μ_G^{hc} on $\mathcal{I}(G)$, the set of all independent sets of G , with

$$\mu_G^{hc}(I) = \frac{\lambda^{|I|}}{Z_G^{hc}(\lambda)}, \quad \text{where} \quad Z_G^{hc}(\lambda) = \sum_{I \in \mathcal{I}(G)} \lambda^{|I|}. \quad (8)$$

The complexity of approximate counting and sampling from μ_G^{hc} on bounded-degree graphs is well understood: there is a computational threshold at some $\lambda_c(\Delta)$, with efficient algorithms for $\lambda < \lambda_c(\Delta)$ [4, 19, 20, 54] and hardness above the threshold (no polynomial-time algorithms unless $\text{NP} = \text{RP}$) [27, 52, 53]. However, on bipartite graphs, the complexity of these problems is unresolved and is captured by the class #BIS (approximately counting independent sets in bipartite graphs) defined by Dyer, Goldberg, Greenhill, and Jerrum [22].

Theorem 1.4 implies the existence of a fast perfect sampling algorithm for the hard-core model in a certain class of bipartite graphs called *unbalanced bipartite graphs*, considered in References [14, 24].

COROLLARY 1.5. *There is a perfect sampling algorithm for μ_G^{hc} running in expected time $O(n \log n)$ for n -vertex bipartite graphs G with bipartition (L, R) , with maximum degree Δ_L in L , maximum degree Δ_R in R , and minimum degree δ_R in R if*

$$\lambda(1 + (1 + e)(\Delta_L - 1)\Delta_R) < (1 + \lambda)^{\delta_R/\Delta_L}. \quad (9)$$

Approximate sampling algorithms with large polynomial run times were previously given for this problem when $6\lambda\Delta_L\Delta_R < (1 + \lambda)^{\delta_R/\Delta_L}$ in Reference [14] and when $3.3353\lambda\Delta_L\Delta_R < (1 + \lambda)^{\delta_R/\Delta_L}$ in [24]. Our result applies to a comparable parameter range: inequality (9) holds, for instance, when $(1 + e)\lambda\Delta_L\Delta_R < (1 + \lambda)^{\delta_R/\Delta_L}$, or when $3\lambda\Delta_L\Delta_R < (1 + \lambda)^{\delta_R/\Delta_L}$ and $\Delta_L < 6$. More importantly, our algorithm is the first to achieve perfect sampling and near-linear running time.

Potts model on expander graphs. The Q -color ferromagnetic Potts model on a graph $G = (V, E)$ is the probability distribution μ_G^{Potts} on the set of Q -colorings of the vertices of G ; i.e., $\{1, \dots, Q\}^V$. Each Q -coloring σ is assigned probability $\mu_G^{\text{Potts}}(\sigma) \propto e^{\beta m(G, \sigma)}$, where $m(G, \sigma)$ is the number of monochromatic edges of G under the coloring σ and $\beta > 0$ is a model parameter. When the parameter β is large, and G has some structure (e.g., G is an expander graph), typical configurations drawn from μ_G^{Potts} are dominated by one of the Q colors; that is, there is phase coexistence in the model. This enables sampling using subset polymer models.

Recall that an n -vertex graph $G = (V, E)$ is an α -expander if for all subsets $S \subseteq V$ with $|S| \leq n/2$, the number of edges in E with exactly one endpoint in S is at least $\alpha|S|$.

COROLLARY 1.6. *Consider the Q -color ferromagnetic Potts model on an α -expander n -vertex graph of maximum degree Δ . Suppose*

$$\beta \geq \frac{1 + \log\left(\frac{\Delta+1}{e\Delta} + 1\right) + \log((Q-1)\Delta)}{\alpha}. \quad (10)$$

Then, there is a sampling algorithm with expected running time $O(n \log n)$ that outputs a sample σ with distribution $\hat{\mu}$ so that $\|\hat{\mu} - \mu_G^{\text{Potts}}\|_{TV} \leq e^{-\Omega(n)}$.

Previously, Reference [17] provided a ε -approximate sample for μ_G^{Potts} in time $O(n \log(n/\varepsilon) \log(1/\varepsilon))$ whenever $\beta \geq \frac{5+3\log((Q-1)\Delta)}{\alpha}$. Condition (10) holds when $\beta \geq \frac{1.2+\log((Q-1)\Delta)}{\alpha}$,

so our algorithm applies to a larger range of parameters and removes the dependence on ε from the running time. We do not achieve perfect sampling in this application only because the subset polymer models used give approximations of μ_G^{Potts} rather than describing μ_G^{Potts} exactly.

1.4 Sampling Unrooted Graphlets in Finite Graphs

As another application of our algorithm for sampling weighted rooted graphlets, we consider next the problem of sampling weighted *unrooted* graphlets in a finite graph. Given a finite graph G , let $\mathcal{S}(G)$ be the set of all graphlets of G . Define the distribution $\nu_{G,\lambda}$ on $\mathcal{S}(G)$ by

$$\nu_{G,\lambda}(\gamma) = \frac{\lambda^{|\gamma|}}{Z_{G,\lambda}}, \quad \text{where } Z_{G,\lambda} = \sum_{\gamma \in \mathcal{S}(G)} \lambda^{|\gamma|}.$$

Read-McFarland and Štefankovič [50] gave a polynomial-time approximate sampling algorithm for $\nu_{G,\lambda}$ for the class of maximum-degree Δ graphs when $\lambda < \lambda_*(\Delta)$ and prove that there is no such algorithm for $\lambda \in (\lambda_*(\Delta), 1)$ unless $\text{NP} = \text{RP}$.¹ We give a new algorithm for this problem, covering the entire $\lambda < \lambda_*(\Delta)$ regime, and improving on the result of Reference [50] in two ways: (i) our running time is constant in expectation (with no dependence on n), while the running time of the ε -approximate sampler in Reference [50] is $n \cdot (1/\varepsilon)^{O(\log \Delta)}$; and (ii) our algorithm outputs a perfect sample instead of an approximate one (and thus, the running time has no dependence on any approximation parameter).

THEOREM 1.7. *Fix $\Delta \geq 3$ and let $\lambda < \lambda_*(\Delta)$. Then for the class of finite graphs of maximum degree Δ there is a randomized algorithm running in constant expected time that outputs a perfect sample from $\nu_{G,\lambda}$. The expected running time is bounded as a function of Δ and λ .*

The algorithm we use for this theorem is a modification of the one for sampling rooted graphlets. We pick a uniformly random $v \in V$, run the same BFS percolation exploration, and accept the connected component of v with an adjusted probability (to account for the fact that a graphlet can be generated from any of its vertices). The acceptance probability is bounded away from 0 and so the algorithm runs in constant expected time. As mentioned earlier, the ε -approximate sampling algorithm from Reference [50] is based on the exhaustive enumeration of all subgraphs of size $\leq k$, for some k that depends on ε . Our new algorithm entirely bypasses this enumeration barrier.

2 GRAPHLET SAMPLING: ALGORITHMS

In this section, we present our efficient perfect sampling algorithm for weighted, vertex-labeled graphlets containing a fixed vertex r from a maximum degree Δ graph; in particular, in Section 2.1, we prove a generalized version of Theorem 1.1 from the introduction. We also provide in Section 2.2, our algorithm for sampling weighted graphlets (i.e., the unrooted, unlabeled case) and establish Theorem 1.7. Our hardness results, that is, Lemmas 1.2 and 1.3, are proved later in Section 5.

2.1 Sampling Rooted Vertex-Labeled Graphlets

Let $G = (V, E)$ be a (possibly infinite) graph of maximum degree Δ . For $U \subseteq V$, let $G[U]$ denote the corresponding vertex-induced subgraph of G ; specifically, $G[U] = (U, E(U))$, where $E(U) \subseteq E$ is the set of edges of G with both endpoints in U . A vertex-induced subgraph is a *graphlet* if it is connected. For $r \in V$, let $\mathcal{S}(G, r)$ be the set of all graphlets of G that contain vertex r . We call the graphlets in $\mathcal{S}(G, r)$ the graphlets rooted at r .

¹In Reference [50], the threshold is incorrectly stated as $\lambda < \lambda_*(\Delta + 1)$; this is due to a minor error interchanging the infinite Δ -regular tree with the infinite Δ -ary tree; with this small correction their analysis goes through with the bound on λ as stated here.

Let $\Sigma = \{1, \dots, q\}$ be a set of vertex labels or colors, and let $\mathcal{S}(G, r, q) = \bigcup_{(S, E(S)) \in \mathcal{S}(G, r)} \Sigma^S$ be the set of all vertex-labeled graphlets rooted at r . Given a real parameter $\lambda > 0$, we assign to each rooted vertex-labeled graphlet $\gamma \in \mathcal{S}(G, r, q) \cup \{\emptyset\}$ with $|\gamma|$ vertices the weight $w_\gamma = \lambda^{|\gamma|} f(\gamma)$, where $f : \mathcal{S}(G, r, q) \cup \{\emptyset\} \rightarrow [0, 1]$. Note that $0 \leq w_\gamma \leq \lambda^{|\gamma|}$, which will be important for later analysis.

Define the probability distribution $\nu_{G, r, \lambda}$ on $\mathcal{S}(G, r, q) \cup \{\emptyset\}$ by setting

$$\nu_{G, r, \lambda}(\gamma) = \frac{w_\gamma}{Z(G, r, \lambda)}, \quad (11)$$

where $Z(G, r, \lambda) = \sum_{\gamma' \in \mathcal{S}(G, r, q) \cup \{\emptyset\}} w_{\gamma'}$. We assume that G , f , q and λ are such that $Z(G, r, \lambda)$ is finite, so that this distribution is well defined. When $q = 1$ and $f(\gamma) = \mathbb{1}(\gamma \neq \emptyset)$, $\nu_{G, r, \lambda}$ corresponds exactly to the distribution defined in Equation (1) over the unlabeled graphlets of G rooted at r .

We consider the problem of sampling from $\nu_{G, r, \lambda}$; this more general version of the sampling problem is later used as a subroutine for sampling polymer systems in Section 3. Let

$$\lambda_*(\Delta, q) := \frac{(\Delta - 2)^{\Delta-2}}{q(\Delta - 1)^{\Delta-1}};$$

cf., (2). Our main algorithmic result for sampling colored rooted graphlets is the following.

THEOREM 2.1. *Fix $\Delta \geq 3$, $q \geq 1$, and suppose $0 < \lambda < \lambda_*(\Delta, q)$. There is a randomized algorithm to exactly sample from $\nu_{G, r, \lambda}$ for graphs G of maximum degree Δ and functions $f : \mathcal{S}(G, r, q) \cup \{\emptyset\} \rightarrow [0, 1]$ where $f(\gamma)$ is computable in time polynomial in $|\gamma|$; this randomized algorithm has expected running time $O(Z_{G, r, \lambda}^{-1})$.*

Theorem 1.1 from the introduction corresponds to the special case when $q = 1$ and $f(\gamma) = \mathbb{1}(\gamma \neq \emptyset)$ (in this case $Z_{G, r, \lambda} \geq \lambda$). Other mild assumptions on the function f , e.g., $f(\emptyset) = 1$ or $f(r) = 1$, ensure that $Z_{G, r, \lambda}$ is bounded away from 0 and, consequently, that the sampling algorithm in the theorem also has constant expected running time in those cases.

As a warm-up, let us consider first our algorithm for sampling labeled rooted graphlets on a finite graph $G = (V, E)$ with $f = 1$, and purposely omit certain non-essential implementation details for clarity. First, we find $p \in (0, 1)$ such that $\frac{p}{q}(1-p)^{\Delta-2} = \lambda$; this choice of p will be justified in what follows. The algorithm then repeats the following process until a vertex-labeled graphlet is accepted:

- (1) Each vertex of the graph is independently assigned with probability p a uniform random color from $\{1, \dots, q\}$, or it is marked as “not colored” with the probability $1 - p$.
- (2) Let $\tilde{\gamma}$ be the vertex-labeled graphlet from $\mathcal{S}(G, r, q) \cup \{\emptyset\}$ corresponding the colored connected component of r ; i.e., the set of vertices connected to r by at least one path of colored vertices.
- (3) Observe that the probability that $\tilde{\gamma} = \gamma$ is $(p/q)^{|\gamma|} (1-p)^{|\partial\gamma|}$, where $\partial\gamma$ denotes to set of vertices in G that are not in γ but are adjacent to a vertex in γ (with a slight abuse of notation, we let $|\partial\emptyset| = 1$). When $\tilde{\gamma} = \gamma$, our aim is to output γ with probability $\propto \lambda^{|\gamma|}$ which has no dependence on $\partial\gamma$. Therefore, we use a “rejection filter” and only accept γ with probability $(1-p)^{(\Delta-2)|\gamma|+2-|\partial\gamma|}$, so that the probability that γ is the output becomes:

$$\left(\frac{p}{q}\right)^{|\gamma|} (1-p)^{|\partial\gamma|} (1-p)^{(\Delta-2)|\gamma|+2-|\partial\gamma|} = (1-p)^2 \left(\frac{p}{q}(1-p)^{\Delta-2}\right)^{|\gamma|} = (1-p)^2 \lambda^{|\gamma|}. \quad (12)$$

From Equation (12), the choice of p such that $\frac{p}{q}(1-p)^{\Delta-2} = \lambda$ is apparent. We will prove that only when $\lambda < \lambda_*(\Delta, q)$ there exists $p \in (0, 1)$ such that $\frac{p}{q}(1-p)^{\Delta-2} = \lambda$. In the actual implementation of the algorithm, it will in fact suffice to find an approximation for p .

We comment briefly on the intuition for the rejection filter. The acceptance probability must include a factor of $(1-p)^{-|\partial\gamma|}$, so that the final acceptance probability depends on $|\gamma|$ but not on $|\partial\gamma|$. However, $(1-p)^{-|\partial\gamma|} > 1$ is not a valid probability, so we use instead $(1-p)^{(\Delta-2)|\gamma|+2-|\partial\gamma|}$, which is at most 1 since $(\Delta-2)|\gamma|+2 \geq |\partial\gamma|$. This bound on $|\partial\gamma|$ is best possible since it is tight for the Δ -regular tree. We note that using looser bounds for $|\partial\gamma|$ affects the range of the parameter λ for which we can find $p \in (0, 1)$ so that $\frac{p}{q}(1-p)^{\Delta-2} = \lambda$.

Finally, we mention that the algorithm as described requires $\Omega(|V|)$ time per iteration and cannot be extended to infinite graphs. This is easily corrected by assigning colors starting from r and revealing only the colored component of r in a breadth-first fashion. The threshold $\lambda_*(\Delta, q)$ is sharp in the sense that only when $\lambda < \lambda_*(\Delta, q)$ is the value of p such that the revealing process is a sub-critical process that creates a small component with high probability. This ensures the algorithm can be implemented efficiently. In particular, we stress that our algorithm avoids exhaustively enumerating labeled graphlets, as done in previous methods [17].

Before giving the implementation details of our algorithm and proving Theorem 2.1, we consider the problem of finding $p \in (0, 1)$ such that $\frac{p}{q}(1-p)^{\Delta-2} = \lambda$. For $\Delta \geq 3$ and $q \geq 1$, consider the real function $g(x) = \frac{x}{q}(1-x)^{\Delta-2}$. It can be readily checked that the function g is continuous and differentiable in $[0, 1]$, has a unique maximum at $x = \frac{1}{\Delta-1}$ with $g(\frac{1}{\Delta-1}) = \lambda_*(\Delta, q)$, is increasing in $[0, \frac{1}{\Delta-1}]$, and decreasing in $[\frac{1}{\Delta-1}, 1]$. This implies that only when $\lambda < \lambda_*(\Delta, q)$, there exists a value of $p \in [0, \frac{1}{\Delta-1})$ such that $g(p) = \lambda$. In particular, when $\lambda > \lambda_*(\Delta, q)$, there is no value of p for which $g(p) = \lambda$ and when $\lambda = \lambda_*(\Delta, q)$, the only possible value is $p = \frac{1}{\Delta-1}$. The latter case would result in a critical percolation process, corresponding to the fact that the expected size of a graphlet from $\nu_{G,r,\lambda}$ has no uniform upper bound in the class of graphs of maximum degree Δ ; in fact, it is infinite on the Δ -regular tree. We can find a suitable approximation for p when $\lambda < \lambda_*(\Delta, q)$ via a simple (binary search) procedure.

LEMMA 2.2. *For any $\lambda \in [0, \lambda_*(\Delta, q))$ we can find rational numbers $\hat{p} \in [0, \frac{1}{\Delta-1})$ and $\hat{\lambda} \in [\lambda, \lambda_*(\Delta, q)]$ such that $g(\hat{p}) = \hat{\lambda}$ in $O(|\log \frac{1}{\Delta q(\lambda_* - \lambda)}|)$ time.*

The proof of this lemma appears after the proof of Theorem 2.1. We now prove Theorem 2.1, including giving a more detailed version of the algorithm outlined above that includes the previously omitted implementation details and allows for general functions $f : \mathcal{S}(G, r, q) \cup \{\emptyset\} \rightarrow [0, 1]$.

PROOF OF THEOREM 2.1. For ease of notation, let $\lambda_* = \lambda_*(\Delta, q)$. Our algorithm to sample from $\nu_{G,r,\lambda}$ when $\lambda < \lambda_*$ explores from r in a breadth-first manner and stops once it has revealed the colored connected component of r . It proceeds as follows:

- (1) Find $\hat{p} \in [0, \frac{1}{\Delta-1})$ and $\hat{\lambda} \in [\lambda, \lambda_*)$ such that $g(\hat{p}) = \hat{\lambda}$. This can be done in time $O(|\log \frac{1}{\Delta q(\lambda_* - \lambda)}|)$ per Lemma 2.2.
- (2) Let Q be a queue. With probability $1 - \hat{p}$ do not add r to Q ; otherwise, assign r a color uniformly at random from $\{1, \dots, q\}$ and add r to Q . Mark r as explored.
- (3) While $Q \neq \emptyset$, repeat the following:
 - 3.1) Pop a vertex v from Q .
 - 3.2) For each unexplored neighbor w of v , with probability $1 - \hat{p}$ do not add w to Q ; otherwise, assign w a color uniformly at random from $\{1, \dots, q\}$ and add w to Q . Mark w as explored (regardless of whether it was added to Q or not).
- (4) Let γ be the vertex-labeled graphlet from $\mathcal{S}(G, r, q) \cup \{\emptyset\}$ corresponding the colored connected component of r . Accept γ with probability:

$$f(\gamma) \cdot (1 - \hat{p})^{(\Delta-2)|\gamma|+2-|\partial\gamma|} \left(\frac{\lambda}{\hat{\lambda}} \right)^{|\gamma|}.$$

(5) If γ is rejected, go to Step 2 and repeat.

The probability of obtaining $\gamma \in \mathcal{S}(G, r, q) \cup \{\emptyset\}$ in an iteration of the algorithm is

$$\left(\frac{\hat{p}}{q}\right)^{|\gamma|} (1 - \hat{p})^{|\partial\gamma|} \cdot f(\gamma)(1 - \hat{p})^{(\Delta-2)|\gamma|+2-|\partial\gamma|} \left(\frac{\lambda}{\hat{\lambda}}\right)^{|\gamma|} = (1 - \hat{p})^2 f(\gamma) \lambda^{|\gamma|} = (1 - \hat{p})^2 w_\gamma,$$

and thus the overall acceptance probability in an iteration is

$$\rho := (1 - \hat{p})^2 \sum_{\gamma \in \mathcal{S}(G, r, q) \cup \{\emptyset\}} w_\gamma = (1 - \hat{p})^2 Z_{G, r, \lambda}.$$

Then,

$$\Pr[\gamma \in \mathcal{S}(G, r, q) \cup \{\emptyset\} \text{ is the output}] = \sum_{t \geq 1} (1 - \hat{p})^2 w_\gamma (1 - \rho)^{t-1} = \frac{(1 - \hat{p})^2 w_\gamma}{\rho} = v_{G, r, \lambda}(\gamma).$$

We next bound the expected running time of the algorithm. We claim first that expected running per iteration is, at most, a constant that depends only on a , Δ and q . If γ is the configuration generated in an iteration, it is discovered in $O(|\gamma| + |\partial\gamma|) = O(|\gamma|)$ time and, by assumption, $f(\gamma)$ can be computed in at most $O(|\gamma|^a)$ time, for suitable a constant $a > 0$. Let $\hat{\mu}$ the output distribution of Step 3 of the algorithm. Then, there exists a constant $C = C(q, \Delta) > 0$ such that the expected running time of each iteration is at most:

$$C \sum_{\gamma \in \mathcal{S}(G, r, q) \cup \emptyset} |\gamma|^{\max\{a, 1\}} \Pr_{\hat{\mu}}[\gamma] = C \cdot \mathbb{E}_{\hat{\mu}}[|\gamma|^{\max\{1, a\}}]. \quad (13)$$

We show next that $|\gamma|$ (under $\hat{\mu}$) is stochastically dominated by a random variable $W = X + Y$ (i.e., $|\gamma| < W$), where X and Y are i.i.d. random variables corresponding to the cluster size of a homogeneous Galton–Watson process with offspring distribution $\text{Bin}(\Delta - 1, \hat{p})$. We recall that this is the branching process that starting from a single vertex (or individual) N_0 , adds $Z_1 \sim \text{Bin}(\Delta - 1, \hat{p})$ descendants to N_0 . The process is then repeated for each new descendant. The process can either die out or go on forever; the cluster size is the number of descendants of N_0 . When different offspring of N_0 use different distributions to generate its descendants, the process is called heterogeneous (see, e.g., Reference [37] for additional background).

To see that $|\gamma| < W = X + Y$, first note that $|\gamma| < L$, where L is the cluster size of a heterogeneous Galton–Watson process, in which the root vertex has offspring distribution $\text{Bin}(\Delta, \hat{p})$ and every other vertex has offspring distribution $\text{Bin}(\Delta - 1, \hat{p})$. This is because the branching process generating γ includes the root only with probability $\hat{p}$ (the root is always present in the Galton–Watson process), and, in addition, it considers at most Δ (from the root) or $\Delta - 1$ (from any other vertex) potential branches (or descendants). In turn, we can bound the cluster size L by $L < X + Y$, since, in the branching process corresponding to L , we can couple the first $\Delta - 1$ branches of the root N_0 with X (starting at the root) and the remaining branch with Y (starting at the child of the root not coupled with X).

It is well-known that X and Y have finite moments when $(\Delta - 1)\hat{p} < 1$ (see, e.g., Reference [37]). In particular, there exists a constant $A = A(a, \Delta, \hat{p}) > 0$ such that

$$\mathbb{E}_{\hat{\mu}}[|\gamma|^a] \leq \mathbb{E}[L^a] \leq \mathbb{E}[(X + Y)^a] \leq 2^a (\mathbb{E}[X^a] + \mathbb{E}[Y^a]) \leq A. \quad (14)$$

This, together with Equation (13), shows that the expected running time in each iteration of the algorithm is bounded by $C \cdot A$.

Now, let R be the number of times Steps 2–5 are repeated, let T be the overall running time of the algorithm. Then:

$$\mathbb{E}[T] = \sum_{t \geq 1} \mathbb{E}[T \mid R = t] \Pr[R = t] \leq C \cdot A \cdot \sum_{t \geq 1} t(1 - \rho)^{t-1} \rho \leq \frac{CA}{\rho}, \quad (15)$$

and the result follows. $\square$

We conclude this section with the proof of Lemma 2.2.

PROOF OF LEMMA 2.2. It suffices to find $\hat{p} \in [p, \frac{1}{\Delta-1}]$. This can be done via binary search in t steps, provided $t \geq 0$ is such that $\frac{1}{\Delta-1} \cdot \frac{1}{2^t} \leq \frac{1}{\Delta-1} - p$. Since $g' \leq \frac{1}{q}$, it follows from the mean value theorem that $q(\lambda_* - \lambda) \leq \frac{1}{\Delta-1} - p$. Thus, for the binary search to require at most t steps, it is sufficient to pick t so that $\frac{1}{\Delta-1} \cdot \frac{1}{2^t} \leq q(\lambda_* - \lambda)$, and the result follows. $\square$

2.2 Sampling Unrooted Graphlets

We consider next the problem of sampling weighted graphlets from a finite graph $G = (V, E)$ of maximum degree Δ ; specifically, in this variant of the sampling problem, we consider unrooted, unlabeled, weighted graphlets of G . Let $\mathcal{S}(G)$ be the set of all graphlets of G . We define the probability distribution $v_{G,\lambda}$ on $\mathcal{S}(G)$ by setting

$$v_{G,\lambda}(S) = \frac{\lambda^{|S|}}{Z_{G,\lambda}},$$

where $Z_{G,\lambda} = \sum_{S' \in \mathcal{S}(G) \cup \{\emptyset\}} \lambda^{|S'|}$. The problem of (approximately) sampling from $v_{G,\lambda}$ is quite natural. In Reference [50], it was established that this problem is computationally hard when $\lambda > \lambda_*(\Delta) = \frac{(\Delta-2)^{\Delta-2}}{(\Delta-1)^{\Delta-1}}$; an ε -approximate sampling algorithm was also given in Reference [50] for the case when $\lambda < \lambda_*(\Delta)$ with running time $n \cdot (1/\varepsilon)^{O(\log \Delta)}$. We establish the following:

THEOREM 2.3. *Suppose $\Delta \geq 3$ and $\lambda > 0$ are such that $\lambda < \lambda_*(\Delta)$. There is a randomized algorithm to exactly sample from $v_{G,\lambda}$ with $O(1)$ expected running time for finite graphs G of maximum degree Δ .*

PROOF. For ease of notation, we set $\lambda_* = \lambda_*(\Delta)$ throughout this proof. Our algorithm to sample from $v_{G,\lambda}$ is based on the algorithm to sample from $v_{G,r,\lambda}$ (the rooted, vertex-labeled, weighted case). The idea is to pick a root uniformly at random and run the algorithm for the rooted case from this random vertex with the rejection filter adjusted to account for the fact that a graphlet can be generated from any of its vertices. It proceeds as follows:

- (1) Find $\hat{p} \in [0, \frac{1}{\Delta-1})$ and $\hat{\lambda} \in [\lambda, \lambda_*)$ such that $g(\hat{p}) = \hat{\lambda}$ using the method from Lemma 2.2.
- (2) Pick a vertex $r \in V$ uniformly at random.
- (3) Let Q be a queue. With probability $\hat{p}$ add r to Q and mark it as colored. Mark r as explored.
- (4) While $Q \neq \emptyset$, repeat the following:
 - 3.1) Pop a vertex v from Q .
 - 3.2) For each unexplored neighbor w of v , with probability $\hat{p}$ add w to Q and mark w as colored. Mark w as explored.
- (5) Let $S \in \mathcal{S}(G)$ be the graphlet corresponding to the colored connected component of v . Accept S with probability:

$$\frac{1}{|S|} \cdot (1 - \hat{p})^{(\Delta-2)|S|+2-|\partial S|} \left(\frac{\lambda}{\hat{\lambda}} \right)^{|S|}.$$

- (6) If S is rejected, go back to Step 2 and repeat.

The analysis of this algorithm is similar to that in the proof of Theorem 2.1. Let $n = |V|$. The probability that the algorithm outputs S in an iteration is

$$\sum_{v \in S} \frac{1}{n} \cdot \hat{p}^{|S|} (1 - \hat{p})^{|\partial S|} \cdot \frac{1}{|S|} \cdot (1 - \hat{p})^{(\Delta-2)|S|+2-|\partial S|} \left(\frac{\lambda}{\hat{\lambda}} \right)^{|S|} = \frac{(1 - \hat{p})^2 \lambda^{|S|}}{n}. \quad (16)$$

Hence, conditioned on acceptance, the probability of obtaining $S \in \mathcal{S}(G)$ is thus $\nu_{G,\lambda}(S)$, and so the output distribution of the algorithm is $\nu_{G,\lambda}$.

For the running time of the algorithm, we note that Step 4 of the algorithm is analogous to Step 3 of the algorithm in the proof of Theorem 2.1, and so the expected running time of each round is also bounded by a constant $C = C(\Delta, \hat{p}) > 0$. Let T be the overall the running time of the algorithm. From Equation (16), we have that the overall acceptance probability in a round is $\rho = \frac{(1-\hat{p})^2 Z(G,\lambda)}{n}$. Then, as in Equation (15), we deduce that $\mathbb{E}[T] = O(nZ(G,\lambda)^{-1})$. Since $Z(G,\lambda) \geq n\lambda$, we have $\mathbb{E}[T] = O(1)$. $\square$

3 APPLICATIONS TO POLYMER MODELS

In this section, we show how to use our algorithm for sampling rooted vertex-labeled graphlets from Section 2 to sample from subset polymer models and prove Theorem 1.4.

Consider a subset polymer model on an n -vertex graph $G = (V, E)$; see Section 1.2 for the definition. Recall that we use C_v for the set of all polymers containing vertex $v \in V$, and let $\gamma_\emptyset$ denote the empty polymer. Define the distribution ν_v on $C_v \cup \{\gamma_\emptyset\}$ where $\nu_v(\gamma) = \frac{w_\gamma}{40}$ and

$$\nu_v(\gamma_\emptyset) = 1 - \frac{1}{40} \sum_{\hat{\gamma} \in C_v \cup \{\gamma_\emptyset\}} w_{\hat{\gamma}}.$$

We note that ν_v is well-defined when condition (3) holds, since under this condition we have $\sum_{\hat{\gamma} \in C_v \cup \{\gamma_\emptyset\}} w_{\hat{\gamma}} \leq 40$; this is proved later in Lemma 5.1. The following Markov chain on $\Omega(C)$ is similar to the one introduced in Reference [17].

Polymer dynamics. Given a configuration $X_t \in \Omega(C)$, form X_{t+1} as follows:

- (1) Pick $v \in V$ uniformly at random and let $S_v = \{\gamma \in X_t : v \in \gamma\}$ (note that S_v is either empty or contains 1 polymer).
- (2) With probability $1/41$, let $X_{t+1} = X_t \setminus S_v$.
- (3) Otherwise, with the remaining probability $40/41$, sample γ from ν_v . Let $X_{t+1} = X_t \cup \{\gamma\}$ if $X_t \cup \{\gamma\} \in \Omega(C)$ and let $X_{t+1} = X_t$ otherwise.

We first note that this Markov chain is irreducible (every configuration can reach and can be reached from the empty set of polymers with positive probability), and aperiodic (there is positive probability of remaining on the same state). It is also reversible with respect to μ : letting P be the transition matrix for this chain, $\Gamma \in \Omega(C)$, and a polymer $\gamma \notin \Gamma$ be such that $\Gamma \cup \{\gamma\} \in \Omega(C)$, we have $\mu(\Gamma \cup \{\gamma\})/\mu(\Gamma) = w_\gamma$ and for $\gamma \neq \gamma_\emptyset$:

$$\frac{P(\Gamma, \Gamma \cup \{\gamma\})}{P(\Gamma \cup \{\gamma\}, \Gamma)} = \frac{\sum_{v \in \gamma} \frac{40}{41n} \nu_v(\gamma)}{\frac{|\gamma|}{41n}} = w_\gamma.$$

Now, to implement a single update step of the polymer dynamics, one must sample from ν_v in Step 3. We give a fast perfect sampler for ν_v .

THEOREM 3.1. *Fix $\Delta \geq 3, q \geq 1$, and let $\lambda < \lambda_*(\Delta, q)$. Consider a family of subset polymer models defined on graphs of maximum degree Δ with computationally feasible weights that satisfy $w_\gamma \leq \lambda^{|\gamma|}$.*

There is a randomized algorithm to sample perfectly from ν_v for any $v \in V(G)$ with a constant expected running time.

PROOF. For each vertex v , let $\hat{\nu}_v$ be the distribution over $C_v \cup \{\gamma_0\}$ where $\hat{\nu}_v(\gamma) = w_\gamma/Z_v$ with $Z_v = \sum_{\gamma \in C_v \cup \{\gamma_0\}} w_\gamma$ and $w_{\gamma_0} = 1$. By Theorem 2.1, we can perfectly from $\hat{\nu}_v$ in constant expected time, with the constant depending λ, Δ , and q . Our algorithm to sample from ν_v draws a perfect sample from $\hat{\nu}_v$ first and then outputs the sample with probability $Z_v/40$; otherwise, it outputs the empty polymer γ_0 . The output distribution of this algorithm is exactly ν_v , and so all that remains is for us to show how to sample exactly from a Bernoulli distribution with parameter $Z_v/40$, denoted $\text{Ber}(Z_v/40)$.

We can sample exactly from $\text{Ber}(1/Z_v)$ in constant expected time by drawing an exact sample from $\hat{\nu}_v$ (using the algorithm from Theorem 2.1) and outputting 1 if the sample is γ_0 and 0 otherwise. With this, we can then use a “Bernoulli factory” to obtain a perfect sample from $\text{Ber}(Z_v/40)$. Specifically, we know how to simulate a Bernoulli coin with parameter $p = 1/Z_v$ and would like to obtain a Bernoulli coin with parameter $f(p) = 1/(40p) = Z_v/40$. This is possible in constant expected time (independent of p) since $1 + \lambda \leq Z_v \leq 40$, and the function $f(p) = 1/(40p)$ is real analytic in the closed interval $[1/40, 1 + \lambda]$; see Theorem 2 from Reference [55]. $\square$

Note that Bernoulli factories have been used in a similar fashion to design perfect sampling algorithms for CSP solutions and spin models in References [2, 32, 33].

Using this theorem, we give next a perfect sampling algorithm that works whenever a new condition (4) is satisfied (our algorithm also requires the assumptions in Theorem 3.1).

3.1 Perfect Sampling for Polymer Systems: Proof of Theorem 1.4

We propose here an algorithm to output a perfect sample from μ . Our algorithm is based on the polymer dynamics and the CFTP method [49], using the notion of bounding Markov chains [31, 36] to efficiently implement it.

We proceed with the proof of Theorem 1.4. We start with the description of a grand coupling for the polymer dynamics, which is then used to implement a CFTP algorithm. For an n -vertex graph $G = (V, E)$, let $\{X_t^\Gamma\}$ denote an instance of the polymer dynamics started from the polymer configuration $\Gamma \in \Omega(C)$. For all $\Gamma \in \Omega(C)$, the chains $\{X_t^\Gamma\}$ are coupled by choosing the same uniform random vertex $v \in V$, the same polymer γ sampled from ν_v , and the same uniform random number in $[0, 1]$ to decide whether to remove S_v (Step 2) or to add γ (Step 3). A CFTP algorithm will find a time $-T$ such the grand coupling started from all possible states at time $-T$ coalesces to a single state by time 0. This guarantees that the output of the algorithm, that is, the state at time 0, has distribution μ (see Theorem 1 from Reference [49]). Such a T can be found with a binary search procedure. Unfortunately, implementing the CFTP algorithm in this manner for the polymer dynamics Markov chain is infeasible, since it requires simulating an exponential number of copies of the polymer dynamics, one from each $\Gamma \in \Omega(C)$.

To work around this, we consider a bounding Markov chain for the polymer dynamics rather than the polymer dynamics chain itself. Bounding Markov chains were pioneered in References [31, 36] as a method for efficiently implementing CFTP. The bounding chain for the polymer dynamics has state space $\Omega(C) \times 2^C$ and will be denoted by $\{B_t, D_t\}$, where $B_t \in \Omega(C)$ and $D_t \subseteq C$ are sets of polymers. The chain will maintain throughout that all polymers in B_t are compatible and that every polymer in B_t is compatible with every polymer in D_t . The polymers in D_t do *not* need to be compatible with each other. A step of the bounding Markov chain is defined next.

Polymer Dynamics Bounding Chain. Given $\{B_t, D_t\}$, the chain generates $\{B_{t+1}, D_{t+1}\}$ by:

- (1) Uniformly at random, select $v \in V$.

- (2) With probability $1/41$, remove all polymers containing v by setting $B_{t+1} = B_t \setminus C_v$ and $D_{t+1} = D_t \setminus C_v$.
- (3) With the remaining probability $40/41$, draw a sample γ according to ν_v and:
 - (a) If γ is compatible with B_t and γ is compatible with $D_t \setminus \{\gamma\}$, let $B_{t+1} = B_t \cup \{\gamma\}$ and let $D_{t+1} = D_t \setminus \{\gamma\}$.
 - (b) Else if γ is compatible with B_t but γ is not compatible with D_t , let $B_{t+1} = B_t$ and let $D_{t+1} = D_t \cup \{\gamma\}$.
 - (c) If γ is incompatible with B_t , do nothing: $B_{t+1} = B_t$ and $D_{t+1} = D_t$.

Observe that polymers are only added to B_t if they are compatible with all other polymers in B_t ; hence, if B_0 is a valid polymer configuration, so is B_t for all $t \geq 0$.

To implement a step of the polymer dynamics bounding chain it suffices to pick a vertex $v \in V$ uniformly at random, a uniform random number in $[0, 1]$, and a polymer γ from ν_v , just like for the polymer dynamics. Hence, we can couple the evolution of $\{B_t, D_t\}$ with the grand coupling of the polymer dynamics described earlier. If we set $B_0 = \emptyset$ and $D_0 = C$, it can be checked that for all $\Gamma \in \Omega(C)$ and all $t \geq 0$:

$$B_t \subseteq X_t^\Gamma \subseteq B_t \cup D_t.$$

Indeed, this holds initially for $t = 0$, and the grand coupling ensures that whenever a polymer is removed from X_t^Γ it is also removed from B_t , and whenever a polymer is added to X_t^Γ it is also added to B_t or D_t . Consequently, $\{B_t, D_t\}$ is a bounding chain for the polymer dynamics. In particular, the first time $B_t = B_t \cup D_t$, all instances X_t^Γ have necessarily coalesced to the same configuration. This bounding chain allows us to implement the CFTP algorithm efficiently, as follows.

Coupling from the Past. Set $k = 1$.

- (A) For $t = -2^k, -2^k + 1, \dots, -2^{k-1}$ generate $\rho_t = (v_t, \gamma_t, r_t)$ by choosing $v_t \in V$ uniformly at random, $r_t \in [0, 1]$ uniformly at random, and by sampling $\gamma_t \in C_{v_t}$ from ν_{v_t} .
- (B) Set $B_{-2^k} = \emptyset$ and $D_{-2^k} = C$.
- (C) Simulate the polymer dynamics bounding chain from time -2^k to time 0 using $\rho_{-2^k}, \dots, \rho_{-1}$.
- (D) If $B_0 = B_0 \cup D_0$, then output B_0 ; otherwise, set $k \rightarrow 2k$ and repeat the process from Step (A).

This implementation of the CFTP algorithm provides a perfect sample from μ ; see Reference [49]. It remains for us to show that it is efficient. For this, we show first that the expected number of steps of the polymer dynamics bounding chain throughout the execution of the algorithm is $O(n \log n)$. Afterward, we will show how to implement steps so that they can be executed in amortized constant expected time.

LEMMA 3.2. *Suppose a subset polymer model on an n -vertex graph satisfies condition (4). Then, the expected number of steps of the polymer dynamics bounding chain in the coupling from past algorithm is $O(n \log n)$.*

PROOF. Let us first consider an evolution of the bounding chain $\{B_t, D_t\}$ from initial state $\{B_0, D_0\} = \{\emptyset, C\}$. Let T be the first time $B_t = B_t \cup D_t$; that is, the time until $D_t = \emptyset$. We bound first the expected value of T . Let $\phi_t = |(\cup_{\gamma \in D_t} \gamma) \cap V|$, the number of vertices of G that are included in at least one polymer of D_t . Note that $\phi_t \leq n$. We analyze the expected value of $\phi_{t+1} - \phi_t$.

With probability $\phi_t/(41n)$, a vertex in a polymer of D_t is picked and all polymers in D_t containing that vertex are removed, in which case $\phi_{t+1} \leq \phi_t - 1$. (A polymer may also be removed from D_t in Step 3(a), but this case is omitted from our calculations as it never increases ϕ_t .)

The potential ϕ_t may increase in Step 3(b), when a polymer γ may be added to D_t . Such a polymer is only added if it is incompatible with D_t . The probability this occurs is at most:

$$\sum_{v \in D_t} \frac{40}{41n} \sum_{\gamma \neq v} v_v(\gamma)$$

For each such polymer, the expected increase in the size of ϕ_t is at most $|\gamma|$. We can calculate:

$$\mathbb{E}[\phi_{t+1} - \phi_t \mid B_t, D_t] \leq -\frac{\phi_t}{41n} + \sum_{v \in D_t} \frac{40}{41n} \sum_{\gamma \neq v} |\gamma| v_v(\gamma) = -\frac{\phi_t}{41n} + \frac{\phi_t}{41n} \sum_{\gamma \neq v} |\gamma| w_\gamma.$$

By Equation (4), there exists a constant $\theta < 1$ such that $\sum_{\gamma \neq v} |\gamma| w_\gamma \leq \theta$. It follows that

$$\mathbb{E}[\phi_{t+1} - \phi_t \mid B_t, D_t] \leq -\frac{\phi_t}{41n} (1 - \theta).$$

This shows that $\{\phi_t\}$ is a stochastic process with variable multiplicative drift. Since T is also the first time $\phi_t < 1$, using standard hitting time estimates (see, e.g., Theorem 10 in Reference [44]), we get:

$$\mathbb{E}[T] \leq \frac{41n}{1 - \theta} + \frac{41n}{1 - \theta} \int_1^n \frac{1}{z} dz = O(n \log n).$$

Now, let $-M$ be the first time in past such that, after setting $B_{-M} = \emptyset$ and $D_{-M} = C$, we have $B_0 = B_0 \cup D_0$. The CFTP algorithm simulates at most $2M + M + M/2 + \dots + 1 = O(M)$ steps of the bounding chain. The result follows by noting that $\mathbb{E}[M] = \mathbb{E}[T] = O(n \log n)$. $\square$

It remains to consider how to efficiently implement the steps of the polymer dynamics bounding chain. This is subtle because D_t may initially contain an exponentially large number of polymers, and care is thus needed in how D_t is represented and stored. We describe and analyze efficient data structures for B_t and D_t in the following lemma.

LEMMA 3.3. *Suppose a subset polymer model on an n -vertex graph of maximum degree $\Delta \geq 3$ with computationally feasible weights that satisfy $w_\gamma \leq \lambda^{|\gamma|}$ for some $\lambda < \lambda_*(\Delta, q)$. There exists a compact representation of B_t and D_t that uses $O(n + t)$ space in expectation. Using this representation, each iteration of the Polymer Dynamics Bounding Chain can be executed in amortized constant expected time, and the termination condition $B_t = B_t \cup D_t$ can be checked in constant time.*

PROOF. First, by Lemma 3.1, there is a randomized algorithm to sample perfectly from v_v for any $v \in V$ with constant expected running time. Therefore, the expected size of the polymer γ produced by the sampling procedure is constant. This will imply many of our data structure operations, which take (amortized) time $O(|\gamma|)$ for some γ sampled from v_v , take constant expected time.

For simplicity, we assume time begins at $t = 0$ and at some point in the future we wish to check whether $B_t = B_t \cup D_t$. Shifting the indices appropriately, this data structure can be applied in each round of CFTP from each negative starting time. Note because all polymers in B_t are compatible with all polymers in D_t , the condition $B_t = B_t \cup D_t$ is equivalent to the condition $D_t = \emptyset$.

Data Structure for B_t : Note B_t is a collection of compatible polymers. The operations to be performed on B_t include addition of a polymer (Step 3a), deletion of any polymer containing a particular vertex (Step 2), and comparison of a polymer to B_t to determine whether it is compatible with B_t (Step 3). B_t will be stored as a length n array, and we use $\bar{B}_t$ to denote the state of the array at time t . Initially, $\bar{B}_0$ has only 0 entries. Each entry may also have a pointer to a node of a doubly linked list. Each polymer in B_t is stored as a doubly linked list. This requires $O(n)$ space.

When γ is added to B_t , for each $v \in \gamma$ we set $\bar{B}_t(v) = 1$. We also create a doubly linked list $\mathcal{L}^\gamma$, which has a node for each $v \in \gamma$ that notes the color v is assigned in γ and contains a forward

pointer, a backward pointer, and a pointer back to $\bar{B}_t(v)$. Each entry of $\bar{B}_t$ for $v \in \gamma$ also points to the corresponding node of $\mathcal{L}^\gamma$. Doing these updates to the data structure for B_t takes time $O(|\gamma|)$.

The polymer removal step only occurs in Step 2 when any polymers containing v are removed from B_t . There is at most one such polymer γ in B_t containing v . The other vertices in γ can quickly be found using the doubly linked list $\mathcal{L}^\gamma$ and its pointers back to $\bar{B}_t$. We can set $B_t(w) = 0$ for each $w \in \gamma$ and remove $\mathcal{L}^\gamma$. This removal of γ from B_t occurs in time $O(|\gamma|)$.

To compare a polymer γ to B_t to determine whether it is compatible with B_t , one looks up every vertex w that is in or adjacent to γ in $\bar{B}_t$ to see if it is part of a polymer of B_t , that is, if $\bar{B}_t(w) = 1$. If at least one of these vertices w has $\bar{B}_t(w) = 1$, γ is incompatible with B_t ; if all these w have $\bar{B}_t(w) = 0$, then γ is compatible with B_t . This takes time $O(\Delta|\gamma|)$, which because we assume Δ is constant is $O(|\gamma|)$.

Note each of these operations can be performed in time $O(|\gamma|)$, where γ is a polymer that was at some (possibly earlier) time step drawn from v_v for some v . When each polymer γ is sampled in Step 3 in some iteration of the Polymer Dynamics Bounding Chain, all $O(|\gamma|)$ operations that can potentially be performed on it in the future are immediately charged to it. As argued in the first paragraph above, the expected size of a polymer γ drawn in Step 3 is constant, so the amount $O(|\gamma|)$ charged to each sampled polymer γ is constant in expectation. We conclude updating B_t takes expected amortized constant time. This data structure requires $O(n)$ space.

Data Structure for D_t : Much more care is needed with how D_t is stored and accessed. Initially $D_0 = C$, which contains an exponential number of polymers. The operations we need to perform on D_t include deleting all polymers containing a single vertex v (Step 2), checking whether a polymer is compatible with D_t (Step 3b), adding a polymer to D_t (Step 3b), and checking whether a polymer γ is compatible with $D_t \setminus \{\gamma\}$ and if so deleting γ from D_t (Step 3b).

The key observation is that once a particular vertex v of G is selected in Step 1 and the coin flip is such that Step 2 is performed, no polymers containing v remain in C and the size of D_t has been reduced dramatically. Because of this, it makes sense to keep track of which vertices in D_t have been the subject of a deletion in Step 2 at least once. We let D_t^* be (the state at time t) of an array of length n with an entry corresponding to each vertex v of G , where $D_t^*(v) = 0$ if Step 2, deleting C_v from D_t , has been performed for v at least once, and $D_t^*(v) = 1$ if it has not. We also let N_t^* be the number of 1's in D_t^* . Initially, $D_0^*(v) = 1$ for all v and $N_0^* = n$. At any time step, we know any polymer γ where $D_t^*(v) = 1$ for all $v \in \gamma$ is in D_t . Because of this, initially D_0^* completely describes D_0 .

Amidst these deletions, polymers are also added to D_t , and a polymer may contain both vertices where $D_t^*(v) = 0$ and vertices where $D_t^*(v) = 1$. All polymers added to D_t in Step 3b will be stored separately from D_t^* . The data structure used here will have the same idea as that for B_t . There will be an array $\bar{D}_t$ with an entry for each vertex where $\bar{D}_t(v)$ gives the number of polymers containing v that are currently stored in D_t . We also keep track of $\bar{N}_t$, the total number of polymers that are currently in D_t . Adding a polymer γ to D_t involves incrementing N_t , incrementing the corresponding entries in $\bar{D}_t$, and adding a doubly linked list $\mathcal{L}^\gamma$ connecting the vertices of the polymer together. However, because polymers in D_t need not be compatible, there may be more than one polymer in D_t containing a given vertex v . Because of this, instead of each vertex being in at most one doubly linked list $\mathcal{L}^\gamma$, it may be in many such doubly linked lists. To maintain all such pointers for vertex v (two for each $\mathcal{L}^\gamma$ where $v \in \gamma$), these pointers are themselves stored in another doubly linked list $\mathcal{L}^v$.

The necessary operations that must be performed on D_t can be implemented as follows. Note that initially $D_0^*(v) = 1$ for all v , $N_0^* = n$, $\bar{D}_0(v) = 0$ for all v , $\bar{N}_0 = 0$, there are no doubly linked lists $\mathcal{L}^\gamma$, and each doubly linked list $\mathcal{L}^v$ is empty. These steps are presented in a logical

order for explaining the data structures for D_t , rather than in the order in which they occur in the description of the Polymer Dynamics Bounding Chain.

- *Adding polymer γ to D_t (Step 3b):* For each v in γ , increment $\overline{D}_t(v)$. Increment $\overline{N}_t$. Create a doubly linked list $\mathcal{L}^\gamma$, which has a node for each $v \in \gamma$ that notes the color v is assigned in γ and containing a forward pointer and a backward pointer. For each v , this node containing the pointers of $\mathcal{L}^\gamma$ is inserted at the front of the doubly linked list $\mathcal{L}^v$. This takes time $O(|\gamma|)$.
- *Deleting all polymers containing v from D_t (Step 2):* If $D_t^*(v) = 1$, decrement N_t^* and set $D_t^*(v) = 0$. If $\overline{D}_t(v) > 0$, set $\overline{N}_t \leftarrow \overline{N}_t - \overline{D}_t(v)$ and set $\overline{D}_t(v) = 0$. Additionally, for each node in $\mathcal{L}^v$, we explore the corresponding linked list $\mathcal{L}^\gamma$ using the pointers in the node. We delete every other node of $\mathcal{L}^\gamma$, possible in time $O(|\gamma|)$ because for list $\mathcal{L}^w$ also containing a node of $\mathcal{L}^\gamma$, the nodes before and after this node in $\mathcal{L}^w$ can easily be connected to each other because $\mathcal{L}^w$ is doubly linked. For each node in $\mathcal{L}^w$ that is deleted, we also decrement $\overline{D}_t(w)$. This takes time $O(\sum_{\gamma \ni v} |\gamma|)$.
- *Checking if polymer γ is compatible with D_t (Step 3b):* For each w that is in or adjacent to γ , check whether $D_t^*(w) = 0$ and $\overline{D}_t(w) = 0$. If both are 0 for all such w , then γ is compatible with D_t . If at least one is nonzero, then γ is not compatible with D_t . This takes time $O(\Delta|\gamma|)$.
- *Checking whether a polymer γ is compatible with $D_t \setminus \{\gamma\}$ and if so, deleting γ from D_t (Step 3a):* We only need to check this case when γ is compatible with B_t . We first check whether γ is compatible with D_t , as described above. If it is, we add γ to B_t . If γ is compatible with D_t then it cannot be in D_t , so D_t and $D_t \setminus \{\gamma\}$ are the same and no further steps are needed. If γ is incompatible with D_t , the next step is to check if it is also incompatible with $D_t \setminus \{\gamma\}$. We check if $D_t^*(v) = 0$ for all v in or adjacent to γ , $\overline{D}_t(v) = 0$ for all v adjacent to γ , and $\overline{D}_t(v) = 1$ for all $v \in \gamma$; if at least one of these does not hold, γ cannot be compatible with $D_t \setminus \{\gamma\}$ and so we are not in Step 3a and we proceed to Step 3b. If all these hold, we check further to see whether all $\overline{D}_t(v) = 1$ for $v \in \gamma$ because of the presence of the single polymer γ or due to the presence of other polymers, which we can detect by looking at the doubly linked list(s) connecting the vertices in γ , including the colors assigned to each vertex. If there are polymers other than γ here, either multiple smaller polymers or a polymer with the same vertices as γ but differently assigned colors, then γ is incompatible with $D_t \setminus \{\gamma\}$ and we are not in Step 3a. If instead, we find exactly polymer γ , then γ is compatible with $D_t \setminus \{\gamma\}$ and we add γ to B_t and delete it from D_t as previously described. This takes time $O(\Delta|\gamma|)$.

Because any polymer added to D_t was drawn from the distribution v_v for some v , as argued above its expected size is constant. Therefore all of the implementations above, with the possible exception of deleting all polymers containing a single vertex, takes expected constant time. Setting $D_t^*(v) = 0$ takes constant time, but it may take longer to remove any polymers that were added as linked lists $\mathcal{L}^\gamma$. However, we can amortize the cost of deletion if we pay for the cost of deleting a polymer when we add it, as all doubly linked lists $\mathcal{L}^\gamma$ must be added before they are deleted. Doing this amortization makes the cost of adding a polymer to D_t be $O(|\gamma|) + O(|\gamma|) = O(|\gamma|)$, while the amortized cost of deleting all polymers containing a vertex v is now $O(1)$. Thus, all necessary operations for D_t can be performed in amortized expected constant time.

Finally, we note that the termination condition $B_t = B_t \cup D_t$, equivalent to $D_t = \emptyset$, can be checked in constant time by verifying $N_t^* = 0$ and $\overline{N}_t = 0$. The first condition $N_t^* = 0$ verifies that every vertex v has had all polymers containing v deleted at least once, and the second condition $\overline{N}_t = 0$ verifies that there are currently no additional polymers in D_t .

The total space used for these data structures after t steps of the algorithm is in expectation at most $O(n+t)$: $O(n)$ for $\overline{B}_t$, D_t^* and $\overline{D}_t$, and $O(|\gamma|)$ for each polymer γ added to D_t , which is constant

in expectation. There are at most t polymers in D_t , one added each step, so this is $O(t)$ space at most in expectation. $\square$

Together, Theorem 3.1, Lemma 3.2, and Lemma 3.3 imply Theorem 1.4.

4 APPLICATIONS TO LOW-TEMPERATURE SPIN SYSTEMS

We present in this section the details of the applications of spin system sampling given in Section 1.3.

4.1 The Hard-Core Model on Unbalanced Bipartite Graphs

We start by proving Corollary 1.5 for which we use Theorem 1.4.

PROOF OF COROLLARY 1.5. Let G be an n -vertex bipartite graph with partite sets L and R . Suppose that the vertices in L have maximum degree Δ_L , and that the vertices in R have maximum degree Δ_R and minimum degree δ_R .

We can sample from the hardcore model distribution μ_G^{hc} (see Equation (8) for its the definition) by considering an auxiliary distribution $\bar{\mu}$, which is a distribution on subsets $S \subseteq R$. For a subset $S \subseteq R$ with neighbors $N(S)$ in L , this distribution is given by

$$\bar{\mu}(S) = \frac{(1 + \lambda)^{|L \setminus N(S)|} \cdot \lambda^{|S|}}{Z},$$

where $Z = \sum_{S \subseteq R} (1 + \lambda)^{|L \setminus N(S)|} \cdot \lambda^{|S|}$. One can sample from μ_G^{hc} by first sampling $S \subseteq R$ according to $\bar{\mu}$, adding S to the independent set, and then for each vertex $v \in L \setminus N(S)$ include v in the independent set with probability $\lambda/(1 + \lambda)$ independently. This results in exactly the desired distribution μ_G^{hc} .

To sample from $\bar{\mu}$, we use a polymer model, which we define next. Let R^2 be the graph whose vertices are the vertices of R , where two vertices are adjacent if they are at distance two in G . Note the maximum degree in R^2 is $\Delta_R(\Delta_L - 1)$. We define a polymer γ to be a connected subset of R^2 , its neighborhood $N(\gamma)$ to be all vertices in L adjacent to a vertex in γ , and its weight as

$$w_\gamma = \frac{\lambda^{|\gamma|}}{(1 + \lambda)^{|N(\gamma)|}} \leq \left(\frac{\lambda}{(1 + \lambda)^{\delta_R/\Delta_L}} \right)^{|\gamma|}.$$

Polymers will not be labeled; that is, $q = 1$. Compatible polymer configurations in R^2 correspond exactly to subsets $S \subseteq R$, and the weights w_γ mean a compatible polymer configuration corresponding to subset S has probability exactly $\bar{\mu}(S)$.

To sample from this subset polymer model in R^2 , we show that when condition (9) holds, the conditions Equations (3) and (4) of Theorem 1.4 hold, implying the existence of a perfect sampling algorithm with expected running time $O(n \log n)$.

First, when Equation (9) holds, it follows that $e\lambda(\Delta_L - 1)\Delta_R \leq (1 + \lambda)^{\delta_R/\Delta_L}$. Then

$$w_\gamma \leq \left(\frac{\lambda}{(1 + \lambda)^{\delta_R/\Delta_L}} \right)^{|\gamma|} \leq \left(\frac{1}{e(\Delta_L - 1)\Delta_R} \right)^{|\gamma|},$$

and since $\Delta_R(\Delta_L - 1)$ is the maximum degree in the host graph R^2 , condition (3) holds.

From Equation (9), it also follows that

$$\frac{e\lambda\Delta_R(\Delta_L - 1)}{(1 + \lambda)^{\delta_R/\Delta_L}} < \frac{e\Delta_R(\Delta_L - 1)}{1 + (1 + e)\Delta_R(\Delta_L - 1)}.$$

Because all quantities in this equation are constants that do not depend on the size of the graph, we conclude there exists $\theta \in (0, 1)$ such that

$$\frac{e\lambda\Delta_R(\Delta_L - 1)}{(1 + \lambda)^{\delta_R/\Delta_L}} \leq \theta \frac{e\Delta_R(\Delta_L - 1)}{1 + (1 + e)\Delta_R(\Delta_L - 1)}. \quad (17)$$

Using Lemma 5.2 (which precisely counts the number of rooted subtrees of size k of the Δ -regular tree), the number of polymers of size k that are incompatible with v is at most $(\Delta_R(\Delta_L - 1) + 1)(e\Delta_R(\Delta_L - 1))^{k-1}/k$; recall that $\Delta_R(\Delta_L - 1)$ is the maximum degree of R^2 . It follows from this and Equation (17) that

$$\begin{aligned} \sum_{\gamma \neq v} |\gamma| w_\gamma &\leq \sum_{k=1}^{\infty} k \cdot \frac{(\Delta_R(\Delta_L - 1) + 1)(e\Delta_R(\Delta_L - 1))^{k-1}}{k} \cdot \left(\frac{\lambda}{(1 + \lambda)^{\delta_R/\Delta_L}} \right)^k \\ &= \frac{\Delta_R(\Delta_L - 1) + 1}{e\Delta_R(\Delta_L - 1)} \sum_{k=1}^{\infty} \left(\frac{e\lambda\Delta_R(\Delta_L - 1)}{(1 + \lambda)^{\delta_R/\Delta_L}} \right)^k \\ &\leq \frac{\Delta_R(\Delta_L - 1) + 1}{e\Delta_R(\Delta_L - 1)} \sum_{k=1}^{\infty} \theta \left(\frac{e\Delta_R(\Delta_L - 1)}{1 + (1 + e)\Delta_R(\Delta_L - 1)} \right)^k \\ &= \theta \cdot \frac{\Delta_R(\Delta_L - 1) + 1}{e\Delta_R(\Delta_L - 1)} \frac{\frac{e\Delta_R(\Delta_L - 1)}{1 + (1 + e)\Delta_R(\Delta_L - 1)}}{1 - \frac{e\Delta_R(\Delta_L - 1)}{1 + (1 + e)\Delta_R(\Delta_L - 1)}} \\ &= \theta. \end{aligned}$$

Thus, condition Equation (4) also holds, and Theorem 1.4 supplies the perfect sampling algorithm with the desired running time. $\square$

4.2 Potts Model on Expander Graphs

The Q -color Potts model on G at inverse temperature β is a distribution over all (not necessarily proper) Q -colorings of V . Let $\Omega_{G,Q}$ be all colorings $\omega : V \rightarrow [Q]$. For a coloring $\omega \in \Omega_{G,Q}$ with $m(G, \omega)$ monochromatic edges, this distribution has

$$\mu_G^{\text{Potts}}(\omega) = \frac{e^{\beta m(G, \omega)}}{Z},$$

where $Z = \sum_{\omega \in \Omega_{G,Q}} e^{\beta m(G, \omega)}$.

For $j \in [Q]$, let $\Omega_{G,Q,j}$ be all colorings in $\Omega_{G,Q}$ such that strictly more than $n/2$ vertices are assigned color j . Let $\overline{\Omega}_{G,Q} := \bigcup_{j \in [Q]} \Omega_{G,Q,j}$ and consider the distribution $\overline{\mu}^{\text{Potts}}$ given by

$$\overline{\mu}^{\text{Potts}}(\omega) = \frac{e^{\beta m(G, \omega)}}{\hat{Z}} \mathbb{1}(\omega \in \overline{\Omega}_{G,Q}), \text{ where } \hat{Z} = \sum_{\omega \in \overline{\Omega}_{G,Q}} e^{\beta m(G, \omega)}.$$

It follows from Reference [38] that, under suitable conditions, (as we detail next) a perfect sample from $\overline{\mu}^{\text{Potts}}$ is an e^{-n} -approximate sample from μ^{Potts} .

Let $\overline{\mu}_j^{\text{Potts}}$ be the distribution $\overline{\mu}^{\text{Potts}}$ conditioned on being close to the ground state that is entirely color j , that is, conditioned on being in $\Omega_{G,Q,j}$. Because all Q ground states are symmetric, one can sample from $\overline{\mu}^{\text{Potts}}$ by first picking a uniformly random $j \in [Q]$, and then sampling from $\overline{\mu}_j^{\text{Potts}}$. Sampling from $\overline{\mu}_j^{\text{Potts}}$ can be done using the polymer model we define next.

For $\omega \in \Omega_{G,Q,j}$, let $\Gamma(\omega) = \{v \in V : \omega(v) \neq j\}$. Consider a subset polymer model whose host graph is G where a polymer is a graphlet of G whose vertices have colors other than j (there are thus $Q - 1$ colors available to color the vertices of a polymer). For a polymer γ , we let $w_\gamma = \exp(-\beta B(\gamma))$

where $B(\gamma)$ is the number of bichromatic edges of the colored subgraph γ plus the number of boundary edges of γ .

There is a bijection between Potts configurations in $\Omega_{G,Q,j}$ and compatible polymer configurations consisting of at most $n/2$ total vertices, where vertices assigned a color other than j are identified. A sampling algorithm for this polymer model can give a sampling algorithm for $\bar{\mu}_j^{\text{Potts}}$, where if the polymer configuration produced has more than $n/2$ total vertices, that configuration is rejected and resampling occurs. As we will see, the probability of needing to resample is small.

Using this polymer model representation, Reference [38] gives an efficient ε -approximate sampling algorithm with for μ_G^{Potts} , derived from an approximate sampling algorithm for $\bar{\mu}_j^{\text{Potts}}$. This algorithm applies to all α -expanding graphs G with maximum degree Δ whenever $\alpha > 0$, $\Delta \geq 3$, $Q \geq 2$, and $\beta > 4 \log((Q-1)\Delta)/\alpha$. However, it involves a polymer enumeration step, and so its runtime is (omitting the dependence on other parameters) of the form $n^{O(\log \Delta)}$.

In Reference [17], the authors take a Markov chain approach and give an ε -approximate sampling algorithm for μ^{Potts} , again via approximate sampling from $\bar{\mu}_j^{\text{Potts}}$, on α -expanding graphs G with maximum degree Δ whenever $\beta \geq \frac{5+3 \log((Q-1)\Delta)}{\alpha}$. This has running time $O(n \log(n/\varepsilon) \log(1/\varepsilon))$. In both of these prior works, it must hold that $\varepsilon \geq Qe^{-n}$.

Our algorithm gives an even larger range of β in which $O(n \log n)$ sampling is possible, and removes the dependence on ε . It produces an exact sample for $\bar{\mu}^{\text{Potts}}$ rather than an approximate sample, although this sample is still only an e^{-n} -approximate sample for μ_G^{Potts} .

PROOF OF COROLLARY 1.6. From the discussion above, it suffices to generate a perfect sample from $\bar{\mu}_j^{\text{Potts}}$ for any $j \in [Q]$ using the subset polymer model described above. We will show that conditions Equations (3) and (4) of Theorem 1.4 hold, implying the existence of a perfect sampling algorithm for $\bar{\mu}_j^{\text{Potts}}$ with expected running time $O(n \log n)$.

When $\beta \geq \frac{1+\log((Q-1)\Delta)}{\alpha}$, since G is an α expander, $w_\gamma = \exp(-\beta B(\gamma)) \leq (\frac{1}{e\Delta(Q-1)})^{|\gamma|} \leq \lambda^{|\gamma|}$ for a suitable $\lambda < \lambda_*(\Delta, Q-1)$, and thus condition (3) is met. Moreover, when (10) holds, then rearranging terms shows

$$e(Q-1)\Delta e^{-\alpha\beta} < \frac{e\Delta}{e\Delta + \Delta + 1}.$$

Because all quantities in this equation are constants that do not depend on the size of the graph, we conclude there exists $\theta \in (0, 1)$ such that

$$e(Q-1)\Delta e^{-\alpha\beta} \leq \theta \cdot \frac{e\Delta}{e\Delta + \Delta + 1}.$$

Using the bound $(e\Delta)^{k-1}/k$ for the number of graphlets of size k containing a given vertex, we deduce that the number of polymers incompatible with v of size k is at most $(\Delta+1)(e\Delta)^{k-1}(Q-1)^k/k$. (This bound for the number of graphlets appears in Reference [9] but can also be deduced by direct computation from Lemma 5.2 below, which provides a tighter bound.) Therefore, we get

$$\begin{aligned} \sum_{\gamma \not\ni v} |\gamma| w_\gamma &\leq \sum_{k=1}^{\infty} k \cdot \frac{(\Delta+1)(e\Delta)^{k-1}(Q-1)^k}{k} e^{-\beta\alpha k} = \frac{\Delta+1}{e\Delta} \sum_{k=1}^{\infty} (e\Delta(Q-1)e^{-\beta\alpha})^k \\ &\leq \frac{\Delta+1}{e\Delta} \sum_{k=1}^{\infty} \theta \cdot \left(\frac{e\Delta}{e\Delta + \Delta + 1} \right)^k = \theta \cdot \frac{\Delta+1}{e\Delta} \left(\frac{\frac{e\Delta}{e\Delta + \Delta + 1}}{1 - \frac{e\Delta}{e\Delta + \Delta + 1}} \right) = \theta. \end{aligned}$$

Thus, condition Equation (4) holds, and Theorem 1.4 provides a perfect sampling algorithm with the desired running time. $\square$

5 GRAPHLET SAMPLING: HARDNESS

In this section, we establish the sharpness of the threshold $\lambda_*(\Delta, q)$. In particular, we prove Lemmas 1.2 and 1.3 from the introduction. We first prove the following more general variant of Lemma 1.3, which corresponds to the $q = 1$ case.

LEMMA 5.1. *The partition function $Z_{G,r,\lambda}$ is finite for every (possibly infinite) graph $G = (V, E)$ of maximum degree Δ and every $r \in V$ if and only if $\lambda \leq \lambda_*(\Delta, q)$. Moreover, $Z_{G,r,\lambda} \leq 40$ when $\lambda \leq \lambda_*(\Delta, q)$.*

The following combinatorial fact will be used in the proof of Lemma 5.1.

LEMMA 5.2. *Consider the infinite Δ -regular tree rooted at ρ , and let T_k denote the number of subtrees of size k rooted at ρ . Then for $k \geq 1$:*

$$T_k = \frac{\Delta}{(\Delta - 1)k + 1} \binom{(\Delta - 1)k + 1}{k - 1}.$$

PROOF. For positive integers a and b , let

$$A_k(a, b) = \frac{a}{bk + a} \binom{bk + a}{k}.$$

The numbers $A_k(a, b)$ are a generalization of the Catalan numbers and satisfy the recurrence

$$A_k(a + c, b) = \sum_{i=0}^k A_i(a, b) A_{k-i}(c, b); \quad (18)$$

see Reference [41]. It is known (see, e.g., Lemma 1 in Reference [50]) that the number of subtrees of size k containing the root of the infinite $(\Delta - 1)$ -ary tree is

$$\frac{1}{(\Delta - 2)k + 1} \binom{(\Delta - 1)k}{k} = \frac{1}{(\Delta - 1)k + 1} \binom{(\Delta - 1)k + 1}{k},$$

which is also equal to $A_k(1, \Delta - 1)$. Hence, for $k \geq 2$, we have

$$T_k = \sum_{k_1=0}^{k-1} \sum_{k_2=0}^{k-1-k_1} \cdots \sum_{k_{\Delta-1}=0}^{k-1-(k_1+\cdots+k_{\Delta-2})} A_{k-1-(k_1+\cdots+k_{\Delta-1})}(1, \Delta - 1) \cdot \prod_{i=1}^{\Delta-1} A_{k_i}(1, \Delta - 1).$$

Using Equation (18), we see that

$$\sum_{k_{\Delta-1}=0}^{k-1-(k_1+\cdots+k_{\Delta-2})} A_{k-1-(k_1+\cdots+k_{\Delta-1})}(1, \Delta - 1) A_{k_{\Delta-1}}(1, \Delta - 1) = A_{k-1-(k_1+\cdots+k_{\Delta-2})}(2, \Delta - 1),$$

and using Equation (18) repeatedly, we get for $k \geq 2$

$$T_k = A_{k-1}(\Delta, \Delta - 1) = \frac{\Delta}{(\Delta - 1)(k - 1) + \Delta} \binom{(\Delta - 1)(k - 1) + \Delta}{k - 1} = \frac{\Delta}{(\Delta - 1)k + 1} \binom{(\Delta - 1)k + 1}{k - 1}.$$

The claim also holds trivially for $k = 1$ and the result follows. $\square$

We are now ready to provide the proof of Lemma 5.1.

PROOF OF LEMMA 5.1. Let $C_{v,k}$ be the number of graphlets of G with k vertices that contain v . Then, since $f \leq 1$ by assumption, we have

$$Z_{G,r,\lambda} = \sum_{\gamma \in \mathcal{S}(G,r,q) \cup \{\emptyset\}} \lambda^{|\gamma|} f(\gamma) \leq \sum_{k \geq 0} \sum_{\gamma \in \mathcal{S}(G,r,q) \cup \{\emptyset\}: |\gamma|=k} \lambda^k = \sum_{k \geq 0} C_{v,k} \cdot \lambda^k q^k.$$

Consider the infinite Δ -regular tree rooted at ρ , and let T_k be the number of subtrees of size k rooted at ρ so that $C_{v,k} \leq T_k$. (The latter bound captures the key idea in the proof: the infinite Δ -regular tree is the worst case among graphs of maximum degree Δ and the threshold $\lambda_*(\Delta, q)$ arises as the threshold at which there is a change in the number of fixed points of the corresponding tree recurrence.)

Then, it follows from Lemma 5.2 that

$$Z_{G,r,\lambda} \leq 1 + \sum_{k \geq 1} \frac{\Delta}{(\Delta-1)k+1} \binom{(\Delta-1)k+1}{k-1} \lambda^k q^k.$$

We apply the ratio test for the summation on the right hand side and consider the limit:

$$L := \lim_{k \rightarrow \infty} \frac{\frac{\Delta}{(\Delta-1)(k+1)+1} \binom{(\Delta-1)(k+1)+1}{k} q^{k+1} \lambda^{k+1}}{\frac{\Delta}{(\Delta-1)k+1} \binom{(\Delta-1)k+1}{k-1} q^k \lambda^k} = q\lambda \lim_{k \rightarrow \infty} \frac{(\Delta-1)k+1}{(\Delta-1)(k+1)+1} \frac{\binom{(\Delta-1)(k+1)+1}{k}}{\binom{(\Delta-1)k+1}{k-1}}.$$

We have that

$$A_k := \frac{\binom{(\Delta-1)(k+1)+1}{k}}{\binom{(\Delta-1)k+1}{k-1}} = \frac{1}{k} \cdot \frac{[(\Delta-1)(k+1)+1]! [(\Delta-2)k+2]!}{[(\Delta-2)k+\Delta]! [(\Delta-1)k+1]!}.$$

Using the inequality $h(n)e^{1/(12n+1)} \leq n! \leq h(n)e^{1/12n}$ where $h(n) = \sqrt{2\pi n} n^{n+1/2} e^{-n}$, and setting

$$B_k = \frac{h((\Delta-1)(k+1)+1)h((\Delta-2)k+2)}{h((\Delta-2)k+\Delta)h((\Delta-1)k+1)}$$

we get

$$\frac{B_k}{k} \frac{e^{\frac{1}{12[(\Delta-1)(k+1)+1]} + \frac{1}{12[(\Delta-2)k+2]}}}{e^{\frac{1}{12[(\Delta-2)k+\Delta]} + \frac{1}{12[(\Delta-1)k+1]}}} \leq A_k \leq \frac{B_k}{k} \cdot \frac{e^{\frac{1}{12[(\Delta-1)(k+1)+1]} + \frac{1}{12[(\Delta-2)k+2]}}}{e^{\frac{1}{12[(\Delta-2)k+\Delta]} + \frac{1}{12[(\Delta-1)k+1]}}}.$$
 (19)

From a direct calculation, it can be checked that

$$\begin{aligned} B_k &= \frac{1}{e} \left(1 + \frac{\Delta-1}{(\Delta-1)k+1}\right)^{(\Delta-1)k+1} \left(1 - \frac{\Delta-2}{(\Delta-2)k+\Delta}\right)^{(\Delta-2)k+\Delta} \\ &\quad \times \frac{[(\Delta-1)(k+1)+1]^{\Delta-1}}{[(\Delta-2)k+2]^{\Delta-2}} \frac{[(\Delta-2)k+2]^{1/2}}{[(\Delta-1)k+1]^{1/2}} \frac{[(\Delta-1)(k+1)+1]^{1/2}}{[(\Delta-2)k+\Delta]^{1/2}}. \end{aligned}$$

Hence,

$$\lim_{k \rightarrow \infty} \frac{B_k}{k} = \frac{1}{e} e^{\Delta-1} \frac{1}{e^{\Delta-2}} \frac{(\Delta-1)^{\Delta-1}}{(\Delta-2)^{\Delta-2}} = \frac{(\Delta-1)^{\Delta-1}}{(\Delta-2)^{\Delta-2}}.$$

From this and Equation (19) we deduce that $\lim_{k \rightarrow \infty} A_k = \frac{(\Delta-1)^{\Delta-1}}{(\Delta-2)^{\Delta-2}}$, which implies that

$$L = q\lambda \frac{(\Delta-1)^{\Delta-1}}{(\Delta-2)^{\Delta-2}}.$$

From the ratio test, we can then conclude that the series converges when $\lambda < \lambda_*(\Delta, q)$ and diverges when $\lambda > \lambda_*(\Delta, q)$.

It remains for us to consider the $\lambda = \lambda_*(\Delta, q)$ case, where

$$Z_{G,r,\lambda} \leq A + \sum_{k \geq 2} \frac{\Delta}{(\Delta-1)k+1} \binom{(\Delta-1)k+1}{k-1} \frac{(\Delta-2)^{(\Delta-2)k}}{(\Delta-1)^{(\Delta-1)k}}$$

with $A = 1 + \frac{(\Delta-2)^{(\Delta-2)}}{(\Delta-1)^{(\Delta-1)}} \leq 5/4$. We also have for $k \geq 2$ that

$$\binom{(\Delta-1)k+1}{k-1} \leq \frac{h((\Delta-1)k+1)}{h(k-1)h((\Delta-2)k+2)} \times \frac{e^{\frac{1}{12((\Delta-1)k+1)}}}{e^{\frac{1}{12(k-1)+1}} e^{\frac{1}{12((\Delta-2)k+2)+1}}} \leq \frac{h((\Delta-1)k+1)}{h(k-1)h((\Delta-2)k+2)}.$$

Therefore,

$$\begin{aligned} Z_{G,r,\lambda} &\leq \frac{5}{4} + \sum_{k \geq 2} \frac{\Delta}{(\Delta-1)k+1} \frac{h((\Delta-1)k+1)}{h(k-1)h((\Delta-2)k+2)} \frac{(\Delta-2)^{(\Delta-2)k}}{(\Delta-1)^{(\Delta-1)k}} \\ &= \frac{5}{4} + \frac{1}{\sqrt{2\pi}} \sum_{k \geq 2} \frac{\Delta}{(\Delta-1)k+1} \frac{[(\Delta-1)k+1]^{(\Delta-1)k+3/2}}{(k-1)^{k-1/2}[(\Delta-2)k+2]^{(\Delta-2)k+2+1/2}} \frac{(\Delta-2)^{(\Delta-2)k}}{(\Delta-1)^{(\Delta-1)k}} \\ &= \frac{5}{4} + \frac{1}{\sqrt{2\pi}} \sum_{k \geq 2} \frac{\Delta[(\Delta-1)k+1]^{1/2}(k-1)^{1/2}}{[(\Delta-2)k+2]^{5/2}} \left[\frac{((\Delta-1)k+1)^{\Delta-1}(\Delta-2)^{\Delta-2}}{(k-1)((\Delta-2)k+2)^{\Delta-2}(\Delta-1)^{\Delta-1}} \right]^k. \end{aligned}$$

Now, letting $L_1(\Delta, k) = \frac{\Delta[(\Delta-1)k+1]^{1/2}(k-1)^{1/2}}{[(\Delta-2)k+2]^{5/2}}$, $L_2(\Delta, k) = \frac{((\Delta-1)k+1)^{\Delta-1}(\Delta-2)^{\Delta-2}}{(k-1)((\Delta-2)k+2)^{\Delta-2}(\Delta-1)^{\Delta-1}}$, and noting that L_1 is a decreasing function of Δ , we have

$$L_1(\Delta, k) \leq \frac{3(2k+1)^{1/2}(k-1)^{1/2}}{(k+2)^{5/2}} \leq \frac{3\sqrt{2}(k+2)^{1/2}(k+2)^{1/2}}{(k+2)^{5/2}} \leq \frac{3\sqrt{2}}{(k+2)^{3/2}},$$

and,

$$\begin{aligned} L_2(\Delta, k) &= \left(1 - \frac{\Delta}{(\Delta-1)(\Delta-2)k+2(\Delta-1)}\right)^{\Delta-2} \left(1 + \frac{\Delta}{(\Delta-1)(k-1)}\right) \\ &\leq \exp \left[\frac{\Delta}{\Delta-1} \left(-\frac{\Delta-2}{(\Delta-2)k+2} + \frac{1}{k-1} \right) \right] \leq \exp \left[\frac{\Delta}{(\Delta-1)(k-1)} \right] \leq \exp \left[\frac{3}{2(k-1)} \right]. \end{aligned}$$

Hence,

$$Z_{G,r,\lambda} \leq \frac{5}{4} + \frac{3}{\sqrt{\pi}} \sum_{k \geq 2} \frac{1}{(k+2)^{3/2}} \exp \left[\frac{3k}{2(k-1)} \right] \leq \frac{5}{4} + \frac{3e^3}{\sqrt{\pi}} \sum_{k \geq 2} \frac{1}{(k+2)^{3/2}} \leq 40$$

when $\lambda = \lambda_*(\Delta, q)$. □

We proceed next with the proof of Lemma 1.2. Let us first formally define the notion of a polynomial-time approximate sampler.

Definition 5.3. An algorithm $\mathcal{A}$ that takes as input a graph $G = (V, E)$, $\lambda > 0$, a vertex $r \in V$, and an accuracy parameter $\varepsilon \in (0, 1]$ is a polynomial-time approximate sampler for $\nu_{G,r,\lambda}$ if it returns a sample from a distribution $\mu_{\mathcal{A}}$ such that $\|\mu_{\mathcal{A}} - \nu_{G,r,\lambda}\|_{TV} \leq \varepsilon$ and has a running time that is polynomial in $|V|$ and $1/\varepsilon$. A polynomial-time approximate sampler for $\nu_{G,\lambda}$ is defined analogously.

We will require the following result.

LEMMA 5.4. Fix $\lambda < 1$, and let $q = 1$ and $f = 1$. There is an algorithm with the following guarantees. The algorithm takes as input a finite graph $G = (V, E)$, a vertex $v \in V$, and parameters $\varepsilon, \delta \in (0, 1)$. The algorithm has access to an exact sampling algorithm for $\nu_{G,v,\lambda}$. With probability at least $1 - \delta$ the algorithm outputs a value $\tilde{Z}_{G,v,\lambda}$ such that $(1 - \varepsilon)\tilde{Z}_{G,v,\lambda} \leq Z_{G,v,\lambda} \leq (1 + \varepsilon)\tilde{Z}_{G,v,\lambda}$. The algorithm's sample complexity (queries to $\nu_{G,v,\lambda}$) and running time are $\text{poly}(|V|, \varepsilon, \log(1/\delta))$.

PROOF. Let $u_1, u_2, \dots, u_{n-1}$ be the vertices in $V \setminus \{v\}$ in a fixed order we define later. Let G_i be the graph that results from removing vertices $u_1, \dots, u_i$ from G and let $Z_i = Z_{G_i, v, \lambda}$. We let $Z_0 = Z_{G, v, \lambda}$ and note that $Z_{n-1} = 1 + \lambda$. Then:

$$\frac{1 + \lambda}{Z_0} = \prod_{i=0}^{n-2} \frac{Z_{i+1}}{Z_i}.$$

Let $p_i = \frac{Z_{i+1}}{Z_i}$. For each p_i , we will find $\tilde{p}_i$ such that with probability at least $1 - \delta/n$:

$$\left(1 - \frac{\varepsilon}{4n}\right) \tilde{p}_i \leq p_i \leq \left(1 + \frac{\varepsilon}{4n}\right) \tilde{p}_i. \quad (20)$$

This implies that

$$\left(1 - \frac{\varepsilon}{2n}\right) \frac{1}{\tilde{p}_i} \leq \frac{1}{p_i} \leq \left(1 + \frac{\varepsilon}{2n}\right) \frac{1}{\tilde{p}_i},$$

and so letting $\tilde{Z}_{G, v, \lambda} = (1 + \lambda) \prod_{i=0}^{n-2} \frac{1}{\tilde{p}_i}$, we get that

$$(1 - \varepsilon) \tilde{Z}_{G, v, \lambda} \leq \left(1 - \frac{\varepsilon}{2n}\right)^n \tilde{Z}_{G, v, \lambda} \leq Z_0 \leq \left(1 + \frac{\varepsilon}{2n}\right)^n \tilde{Z}_{G, v, \lambda} \leq (1 + \varepsilon) \tilde{Z}_{G, v, \lambda},$$

with probability at least $1 - \delta$ by a union bound.

To obtain Equation (20), we define the sequence of vertices $u_1, u_2, \dots, u_{n-1}$ such that u_i is a leaf in the BFS tree of G_{i-1} rooted at v . This way, we can conveniently claim that $1 \geq p_i \geq \frac{1}{2}$. To see this, note that $S \in \mathcal{S}(G_{i+1})$ if and only if $S \in \mathcal{S}(G_i) : u_{i+1} \notin S$, so

$$Z_{i+1} = \sum_{S \in \mathcal{S}(G_{i+1})} \lambda^{|S|} = \sum_{S \in \mathcal{S}(G_i) : u_{i+1} \notin S} \lambda^{|S|}. \quad (21)$$

Now, if $S \in \mathcal{S}(G_i)$ and S contains u_{i+1} , then, since u_{i+1} is a leaf in the BFS tree rooted at v , $S \setminus \{u_{i+1}\}$ is a connected subgraph that contains v . That is, every $S \in \mathcal{S}(G_i)$ that contains u_{i+1} can be mapped to a unique subgraph in $\mathcal{S}(G_i)$ that does not contain u_{i+1} of size $|S| - 1$. Hence, since $\lambda < 1$,

$$\sum_{S \in \mathcal{S}(G_i) : u_{i+1} \in S} \lambda^{|S|} \leq \sum_{S \in \mathcal{S}(G_i) : u_{i+1} \notin S} \lambda^{|S|},$$

and so

$$Z_i = \sum_{S \in \mathcal{S}(G_i) : u_{i+1} \in S} \lambda^{|S|} + \sum_{S \in \mathcal{S}(G_i) : u_{i+1} \notin S} \lambda^{|S|} \leq 2 \sum_{S \in \mathcal{S}(G_i) : u_{i+1} \notin S} \lambda^{|S|}.$$

Combined with Equation (21), this gives $p_i = \frac{Z_{i+1}}{Z_i} \geq \frac{1}{2}$. Therefore, to deduce Equation (20), it suffices to find $\tilde{p}_i$ such that

$$\tilde{p}_i - \frac{\varepsilon}{16n} \leq p_i \leq \tilde{p}_i + \frac{\varepsilon}{16n}.$$

For this, we draw L samples from $\nu_{G_i, v, \lambda}$ and let X_j be the indicator random variable for the event that the j th sample contains u_{i+1} . Letting $\tilde{p}_i = \frac{1}{L} \sum_{j=1}^L X_j$, we get from a Chernoff bound that

$$\Pr[|\tilde{p}_i - p_i| \geq \rho p_i] \leq 2e^{-\frac{L\rho^2 p_i}{3}}$$

and setting $\rho = \frac{\varepsilon}{16np_i}$ and $L = 384 \frac{n^2}{\varepsilon^2} \log(2n/\delta)$:

$$\Pr\left[|\tilde{p}_i - p_i| \geq \frac{\varepsilon}{16n}\right] \leq 2e^{-\frac{L\rho\varepsilon}{24n}} \leq 2e^{-\frac{L\varepsilon^2}{384n^2}} = \frac{\delta}{n}.$$

In summary, we have provided an algorithm that computes $\tilde{p}_i$ such that Equation (20) holds with probability at least $1 - \delta/n$. The algorithm has sample complexity and running time $\text{poly}(n, 1/\varepsilon, \log(1/\delta))$. $\square$

We can now provide the proof of Lemma 1.2.

PROOF OF LEMMA 1.2. It suffices to prove the result for $f = 1$. Let $G = (V, E)$ be a finite graph of maximum degree Δ , and for ease of notation let $\lambda_* = \lambda_*(\Delta, 1)$. We show that if there is a polynomial-time approximate sampler for $\nu_{G,v,\lambda}$ for each $v \in V$ (see Definition 5.3), then there is also a polynomial-time approximate sampler for $\nu_{G,\lambda}$. In Reference [50], it was shown that there is no polynomial-time approximate sampler for $\nu_{G,\lambda}$ when $\lambda \in (\lambda_*, 1)$ unless $\text{NP} = \text{RP}$, and thus our result follows.²

First, given exact sampling algorithms for $\nu_{G,v,\lambda}$ for each $v \in V$ with $\text{poly}(n)$ running times, and a desired accuracy parameter $\hat{\varepsilon} \in (0, 1)$, we provide a sampling algorithm whose output distribution is within $\hat{\varepsilon}$ total variation distance from $\nu_{G,\lambda}$. The running time of the algorithm will be $\text{poly}(n, 1/\hat{\varepsilon})$. We then modify the algorithm to use instead the polynomial-time approximate samplers for $\nu_{G,v,\lambda}$ (instead of the exact samplers) and show that this has a negligible effect on its running time and output distribution.

The algorithm is the following:

- (1) Let $\varepsilon = \min\{\frac{1}{8}, \frac{\hat{\varepsilon}}{6} - \frac{1}{4n}\}$. For each $v \in V$, using the exact sampler for $\nu_{G,v,\lambda}$ and the algorithm from Lemma 5.4, obtain $\tilde{Z}_{G,v,\lambda}$ such that, with probability at least $1 - \frac{1}{n^2 2^n}$,

$$(1 - \varepsilon)\tilde{Z}_{G,v,\lambda} \leq Z_{G,v,\lambda} \leq (1 + \varepsilon)\tilde{Z}_{G,v,\lambda} \quad (22)$$

- (2) Set $t = 0$ and choose a vertex $v \in V$ with probability $\frac{\tilde{Z}_{G,v,\lambda}}{\sum_w \tilde{Z}_{G,w,\lambda}}$;
- (3) Using the exact sampler for $\nu_{G,v,\lambda}$, draw a sample S from $\nu_{G,v,\lambda}$;
- (4) Accept and output S with probability $1/|S|$;
- (5) If S is rejected, increase t , and if $t \leq 2n^2$ go to Step 2 and repeat; otherwise output $\emptyset$.

First, we note that by Lemma 5.4, Step 1 can be implemented in $\text{poly}(n, 1/\varepsilon)$. Steps 2 to 5 of the algorithm are repeated at most $O(n^2)$ times, and each of those step takes $\text{poly}(n)$ time. Therefore, the overall running time of the algorithm is $\text{poly}(n, 1/\varepsilon) = \text{poly}(n, 1/\hat{\varepsilon})$. Moreover, (22) holds for every $v \in V$ with probability at least $1 - \frac{1}{n^2 2^n}$ by a union bound.

Let μ_{ALG} be the output distribution of the algorithm, and let μ_{ALG}^* be the output distribution when (i) Step 1 succeeds in finding the required approximations and (ii) a subgraph is accepted in Step 4 before $t > 2n^2$. We show next that $\|\mu_{\text{ALG}}^* - \nu_{G,\lambda}\|_{\text{TV}} \leq \hat{\varepsilon}$. Let $\tilde{Z} = \sum_{v \in V} \tilde{Z}_{G,v,\lambda}$ and let $\mathcal{S} = \mathcal{S}(G) \cup \{\emptyset\}$. The probability that the algorithm outputs $S \in \mathcal{S}$ in Step 4 in a give iteration is

$$\sum_{v \in V} \frac{\tilde{Z}_{G,v,\lambda}}{\tilde{Z}} \nu_{G,v,\lambda}(S) \frac{1}{|S|} =: \frac{\phi(S)}{\tilde{Z}}. \quad (23)$$

Then, $\mu_{\text{ALG}}^*(S) = \frac{\phi(S)}{\sum_{S \in \mathcal{S}} \phi(S)}$, and

$$\begin{aligned} \|\mu_{\text{ALG}}^* - \nu_{G,\lambda}\|_{\text{TV}} &= \frac{1}{2} \sum_{S \in \mathcal{S}} \left| \frac{\phi(S)}{\sum_{W \in \mathcal{S}} \phi(W)} - \frac{\lambda^{|S|}}{Z_{G,\lambda}} \right| = \frac{1}{2} \sum_{S \in \mathcal{S}} \left| \frac{\sum_{v \in V} \nu_{G,v,\lambda}(S) \tilde{Z}_{G,v,\lambda}}{|S| \sum_{W \in \mathcal{S}} \phi(W)} - \frac{\lambda^{|S|}}{Z_{G,\lambda}} \right| \\ &\leq \frac{1}{2} \sum_{S \in \mathcal{S}} \sum_{v \in V} \left| \frac{\nu_{G,v,\lambda}(S) \tilde{Z}_{G,v,\lambda}}{|S| \sum_{W \in \mathcal{S}} \phi(W)} - \frac{\lambda^{|S|}}{|S| Z_{G,\lambda}} \right| \end{aligned}$$

²We note that the hardness result in Reference [50] is stated for the harder problem of producing approximate samples with a running time depending polynomially on $\log(1/\varepsilon)$, but the proofs in Reference [50] extend to our setting without modification.

$$\begin{aligned}
&= \frac{1}{2} \sum_{S \in \mathcal{S}} \sum_{v \in S} \frac{Z_{G,v,\lambda}}{|S|Z_{G,\lambda}} \left| \frac{v_{G,v,\lambda}(S) \cdot \tilde{Z}_{G,v,\lambda}}{Z_{G,v,\lambda}} \cdot \frac{Z_{G,\lambda}}{\sum_{W \in \mathcal{S}} \phi(W)} - \frac{\lambda^{|S|}}{Z_{G,v,\lambda}} \right| \\
&= \frac{1}{2} \sum_{S \in \mathcal{S}} \sum_{v \in S} \frac{v_{G,v,\lambda}(S) \cdot Z_{G,v,\lambda}}{|S|Z_{G,\lambda}} \left| \frac{\tilde{Z}_{G,v,\lambda}}{Z_{G,v,\lambda}} \cdot \frac{Z_{G,\lambda}}{\sum_{W \in \mathcal{S}} \phi(W)} - 1 \right|. \tag{24}
\end{aligned}$$

From Equation (22) we get

$$1 - 2\varepsilon \leq \frac{1}{1 + \varepsilon} \leq \frac{\tilde{Z}_{G,v,\lambda}}{Z_{G,v,\lambda}} \leq \frac{1}{1 - \varepsilon} \leq 1 + 2\varepsilon,$$

since $\varepsilon \leq 1/8$. Moreover,

$$\sum_{W \in \mathcal{S}} \phi(W) = \sum_{W \in \mathcal{S}} \sum_{v \in W} \frac{v_{G,v,\lambda}(W) \tilde{Z}_{G,v,\lambda}}{|W|} \leq (1 + 2\varepsilon) \sum_{W \in \mathcal{S}} \sum_{v \in W} \frac{\lambda^{|W|}}{|W|} = (1 + 2\varepsilon) Z_{G,\lambda},$$

and similarly we can obtain that $\sum_{W \in \mathcal{S}} \phi(W) \geq (1 - 2\varepsilon) Z_{G,\lambda}$. Combining these bounds we get:

$$\left| \frac{\tilde{Z}_{G,v,\lambda}}{Z_{G,v,\lambda}} \cdot \frac{Z_{G,\lambda}}{\sum_{W \in \mathcal{S}} \phi(W)} - 1 \right| \leq \frac{4\varepsilon}{1 - 2\varepsilon} \leq 8\varepsilon$$

Plugging this into Equation (24) we deduce that:

$$\|\mu_{\text{ALG}}^* - v_{G,\lambda}\|_{\text{TV}} \leq 4\varepsilon \sum_{S \in \mathcal{S}} \sum_{v \in S} \frac{v_{G,v,\lambda}(S) \cdot Z_{G,v,\lambda}}{|S|Z_{G,\lambda}} = \frac{4\varepsilon}{Z_{G,\lambda}} \sum_{S \in \mathcal{S}} \sum_{v \in S} \frac{\lambda^{|S|}}{|S|} \leq 4\varepsilon. \tag{25}$$

Let $\mathcal{E}_1$ be the event that Step 1 succeeds in finding the approximations, let $\mathcal{E}_2$ the event that the algorithm accepts and outputs a graphlet in Step 4, and let $\mathcal{E}$ be the event that both $\mathcal{E}_1$ and $\mathcal{E}_2$ occur. Then, the triangle inequality and Equation (25) imply:

$$\|\mu_{\text{ALG}} - v_{G,\lambda}\|_{\text{TV}} \leq \|\mu_{\text{ALG}}^* - v_{G,\lambda}\|_{\text{TV}} + \|\mu_{\text{ALG}}^* - \mu_{\text{ALG}}\|_{\text{TV}} \leq 4\varepsilon + \|\mu_{\text{ALG}}^* - \mu_{\text{ALG}}\|_{\text{TV}}. \tag{26}$$

We proceed to bound $\|\mu_{\text{ALG}}^* - \mu_{\text{ALG}}\|_{\text{TV}}$.

$$\begin{aligned}
\|\mu_{\text{ALG}}^* - \mu_{\text{ALG}}\|_{\text{TV}} &= \|\mu_{\text{ALG}} - \mu_{\text{ALG}}(\cdot \mid \mathcal{E})\|_{\text{TV}} \leq \mu_{\text{ALG}}(\neg \mathcal{E}) \\
&\leq \mu_{\text{ALG}}(\neg \mathcal{E}_1) + \mu_{\text{ALG}}(\neg \mathcal{E}_2 \mid \mathcal{E}_1) \leq \frac{1}{2n} + \mu_{\text{ALG}}(\neg \mathcal{E}_2 \mid \mathcal{E}_1). \tag{27}
\end{aligned}$$

So it remains for us to bound $\mu_{\text{ALG}}(\neg \mathcal{E}_2 \mid \mathcal{E}_1)$. For this, note that the probability that a subgraph is accepted in Step 4 in an iteration is

$$\sum_{S \in \mathcal{S}} \sum_{v \in S} \frac{\tilde{Z}_{G,v,\lambda}}{\tilde{Z}} v_{G,v,\lambda}(S) \frac{1}{|S|} \geq \frac{1 - 2\varepsilon}{\tilde{Z}} \sum_{S \in \mathcal{S}} \sum_{v \in S} \frac{\lambda^{|S|}}{|S|} \geq \frac{(1 - 2\varepsilon) Z_{G,\lambda}}{\tilde{Z}}.$$

Since

$$\tilde{Z} = \sum_{v \in V} \tilde{Z}_{G,v,\lambda} \leq (1 + 2\varepsilon) \sum_{v \in V} Z_{G,v,\lambda} \leq (1 + 2\varepsilon) n Z_{G,\lambda},$$

we get the acceptance probability in Step 4 in an iteration is at least

$$\frac{1 - 2\varepsilon}{(1 + 2\varepsilon)n} \geq \frac{1 - 4\varepsilon}{n} \geq \frac{1}{2n},$$

provided $\mathcal{E}_1$ occurs. Therefore, if X is geometric random variable with parameter $1/(2n)$, we have

$$\mu_{\text{ALG}}(\neg \mathcal{E}_2 \mid \mathcal{E}_1) \leq \Pr[X \geq 2n^2] \leq \left(1 - \frac{1}{2n}\right)^{2n^2} \leq \frac{1}{e^n},$$

and plugging this bound into Equations (27) and (26) we obtain:

$$\|\mu_{\text{ALG}} - \nu_{G,\lambda}\|_{\text{TV}} \leq 4\varepsilon + \frac{3}{2n}.$$

We have established so far the hardness of exactly sampling from $\nu_{G,v,\lambda}$ for $\lambda \in (\lambda_*, 1)$. The same reduction (i.e., algorithm) with minor adjustments works for approximate sampling. Suppose that in Step 3 of the algorithm, we instead generate $T = 2n^2$ samples from a distribution μ_v such that $\|\mu_v - \nu_{G,v,\lambda}\|_{\text{TV}} \leq \frac{\varepsilon}{n^2}$. Let $\mu_v^{\otimes T}$ and $\nu_{G,v,\lambda}^{\otimes T}$ be the product distributions corresponding to T independent samples from μ_v and $\nu_{G,v,\lambda}$ respectively. We have

$$\|\mu_v^{\otimes T} - \nu_{G,v,\lambda}^{\otimes T}\|_{\text{TV}} \leq T\|\mu_v - \nu_{G,v,\lambda}\|_{\text{TV}} \leq \frac{T\varepsilon}{n^2} \leq 2\varepsilon. \quad (28)$$

Let $\tilde{\mu}_{\text{ALG}}$ be the output distribution of the algorithm when using samples from μ_v in Step 3. Then:

$$\|\tilde{\mu}_{\text{ALG}} - \nu_{G,\lambda}\|_{\text{TV}} \leq \|\tilde{\mu}_{\text{ALG}} - \mu_{\text{ALG}}\|_{\text{TV}} + \|\mu_{\text{ALG}} - \nu_{G,\lambda}\|_{\text{TV}} \leq \|\tilde{\mu}_{\text{ALG}} - \mu_{\text{ALG}}\|_{\text{TV}} + 4\varepsilon + \frac{3}{2n}.$$

Consider the following coupling between $\tilde{\mu}_{\text{ALG}}$ and μ_{ALG} : use the same randomness for Steps 1, 2 and 5 and the optimal coupling for $\mu_v^{\otimes T}$ and $\nu_{G,v,\lambda}^{\otimes T}$ for Step 3. Then, from Equation (28), we get

$$\|\tilde{\mu}_{\text{ALG}} - \nu_{G,\lambda}\|_{\text{TV}} \leq 6\varepsilon + \frac{3}{2n} \leq \hat{\varepsilon}.$$

That is, we obtain a polynomial-time approximate sampler for $\nu_{G,\lambda}$ which completes the reduction. $\square$

ACKNOWLEDGEMENTS

This work was carried out as part of the AIM SQuARE workshop ‘Connections between computational and physical phase transitions.’ We thank Tyler Helmuth, Alexandre Stauffer, and Izabella Stuhl for many helpful conversations. We also thank the anonymous referees for their many insightful corrections and suggestions.

REFERENCES

- [1] Matteo Agostini, Marco Bressan, and Shahrzad Haddadan. 2019. Mixing time bounds for graphlet random walks. *Inform. Process. Lett.* 152 (2019), 105851.
- [2] Konrad Anand, Andreas Göbel, Marcus Pappik, and Will Perkins. 2023. Perfect sampling for hard spheres from strong spatial mixing. In *Approximation, Randomization, and Combinatorial Optimization. Algorithms and Techniques (APPROX/RANDOM’23)*, Vol. 275. Schloss Dagstuhl – Leibniz-Zentrum für Informatik, Dagstuhl, 38:1–38:18.
- [3] Konrad Anand and Mark Jerrum. 2022. Perfect sampling in infinite spin systems via strong spatial mixing. *SIAM J. Comput.* 51, 4 (2022), 1280–1295.
- [4] Nima Anari, Kuikui Liu, and Shayan Oveis Gharan. 2021. Spectral independence in high-dimensional expanders and applications to the hardcore model. *SIAM J. Comput.* 0 (2021), FOCS20–1.
- [5] Kim Baskerville, Peter Grassberger, and Maya Paczuski. 2007. Graph animals, subgraph sampling, and motif search in large networks. *Physical Review E* 76, 3 (2007), 036107.
- [6] Mansurul A. Bhuiyan, Mahmudur Rahman, Mahmuda Rahman, and Mohammad Al Hasan. 2012. Guise: Uniform sampling of graphlets for large graph analysis. In *Proceedings of the 2012 IEEE 12th International Conference on Data Mining*. IEEE, 91–100.
- [7] Christian Borgs, Jennifer Chayes, Tyler Helmuth, Will Perkins, and Prasad Tetali. 2020. Efficient sampling and counting algorithms for the Potts model on $\mathbb{Z}^d$ at all temperatures. In *Proceedings of the 52nd Annual ACM SIGACT Symposium on Theory of Computing (STOC’20)*. 738–751.
- [8] Christian Borgs, Jennifer Chayes, Jeff Kahn, and László Lovász. 2013. Left and right convergence of graphs with bounded degree. *Random Structures & Algorithms* 42, 1 (2013), 1–28.
- [9] Christian Borgs and John Z. Imbrie. 1989. A unified approach to phase diagrams in field theory and statistical mechanics. *Communications in Mathematical Physics* 123, 2 (1989), 305–328.
- [10] Marco Bressan. 2021. Efficient and near-optimal algorithms for sampling connected subgraphs. In *Proceedings of the 53rd Annual ACM SIGACT Symposium on Theory of Computing (STOC’21)*. 1132–1143.

- [11] Marco Bressan, Flavio Chierichetti, Ravi Kumar, Stefano Leucci, and Alessandro Panconesi. 2017. Counting graphlets: Space vs time. In *Proceedings of the 10th ACM International Conference on Web Search and Data Mining (WSDM'17)*. 557–566.
- [12] Marco Bressan, Flavio Chierichetti, Ravi Kumar, Stefano Leucci, and Alessandro Panconesi. 2018. Motif counting beyond five nodes. *ACM Transactions on Knowledge Discovery from Data (TKDD)* 12, 4 (2018), 1–25.
- [13] Marco Bressan, Stefano Leucci, and Alessandro Panconesi. 2021. Faster motif counting via succinct color coding and adaptive sampling. *ACM Transactions on Knowledge Discovery from Data (TKDD)* 15, 6 (2021), 1–27.
- [14] Sarah Cannon and Will Perkins. 2020. Counting independent sets in unbalanced bipartite graphs. In *Proceedings of the Thirty-First Annual ACM-SIAM Symposium on Discrete Algorithms (SODA'20)*. 1456–1466.
- [15] Charles Carlson, Ewan Davies, and Alexandra Kolla. 2020. Efficient algorithms for the Potts model on small-set expanders. arXiv:2003.01154. Retrieved from <https://arxiv.org/abs/2003.01154>
- [16] Xiaowei Chen, Yongkun Li, Pinghui Wang, and John C. S. Lui. 2016. A general framework for estimating graphlet statistics via random walk. *Proceedings of the VLDB Endowment* 10, 3 (2016), 253–264.
- [17] Zongchen Chen, Andreas Galanis, Leslie A Goldberg, Will Perkins, James Stewart, and Eric Vigoda. 2021. Fast algorithms at low temperatures via Markov chains. *Random Structures & Algorithms* 58, 2 (2021), 294–321.
- [18] Zongchen Chen, Andreas Galanis, Daniel Štefankovič, and Eric Vigoda. 2022. Sampling colorings and independent sets of random regular bipartite graphs in the non-uniqueness region. In *Proceedings of the 2022 Annual ACM-SIAM Symposium on Discrete Algorithms (SODA'22)*. SIAM, 2198–2207.
- [19] Zongchen Chen, Kuikui Liu, and Eric Vigoda. 2020. Rapid mixing of Glauber dynamics up to uniqueness via contraction. In *2020 IEEE 61st Annual Symposium on Foundations of Computer Science (FOCS'20)*. IEEE, 1307–1318.
- [20] Zongchen Chen, Kuikui Liu, and Eric Vigoda. 2021. Optimal mixing of Glauber dynamics: Entropy factorization via high-dimensional expansion. In *Proceedings of the 53rd Annual ACM SIGACT Symposium on Theory of Computing (STOC'21)*. 1537–1550.
- [21] Matthew Coulson, Ewan Davies, Alexandra Kolla, Viresh Patel, and Guus Regts. 2020. Statistical physics approaches to unique games. In *Proceedings of the 35th Computational Complexity Conference (CCC'20)*.
- [22] Martin Dyer, Leslie Ann Goldberg, Catherine Greenhill, and Mark Jerrum. 2004. The relative complexity of approximate counting problems. *Algorithmica* 38, 3 (2004), 471–500.
- [23] Sacha Friedli and Yvan Velenik. 2017. *Statistical Mechanics of Lattice Systems: A Concrete Mathematical Introduction*. Cambridge University Press.
- [24] Tobias Friedrich, Andreas Göbel, Martin S Krejca, and Marcus Pappik. 2020. Polymer dynamics via cliques: New conditions for approximations. arXiv:2007.08293. Retrieved from <https://arxiv.org/abs/2007.08293>
- [25] Andreas Galanis, Leslie Ann Goldberg, and James Stewart. 2021. Fast algorithms for general spin systems on bipartite expanders. *ACM Transactions on Computation Theory (TOCT)* 13, 4 (2021), 1–18.
- [26] Andreas Galanis, Leslie Ann Goldberg, and James Stewart. 2022. Fast mixing via polymers for random graphs with unbounded degree. *Information and Computation* 285, Part B (2022), 104894.
- [27] Andreas Galanis, Daniel Štefankovič, and Eric Vigoda. 2016. Inapproximability of the partition function for the anti-ferromagnetic Ising and hard-core models. *Combinatorics, Probability and Computing* 25, 4 (2016), 500–559.
- [28] Oded Goldreich and Dana Ron. 1997. Property testing in bounded degree graphs. In *Proceedings of the 29th Annual ACM Symposium on Theory of Computing (STOC'97)*. 406–415.
- [29] Joshua A. Grochow and Manolis Kellis. 2007. Network motif discovery using subgraph enumeration and symmetry-breaking. In *Annual International Conference on Research in Computational Molecular Biology*. Springer, 92–106.
- [30] Christian Gruber and Hervé Kunz. 1971. General properties of polymer systems. *Communications in Mathematical Physics* 22, 2 (1971), 133–161.
- [31] Olle Haggstrom and Karin Nelander. 1999. On exact simulation of Markov random fields using coupling from the past. *Scandinavian Journal of Statistics* 26, 3 (1999), 395–411.
- [32] Kun He, Chunyang Wang, and Yitong Yin. 2022. Sampling Lovász local lemma for general constraint satisfaction solutions in near-linear time. In *Proceedings of the 2022 IEEE 63rd Annual Symposium on Foundations of Computer Science (FOCS'22)*. IEEE, 147–158.
- [33] Kun He, Kewen Wu, and Kuan Yang. 2023. Improved bounds for sampling solutions of random CNF formulas. In *Proceedings of the 2023 Annual ACM-SIAM Symposium on Discrete Algorithms (SODA'23)*. SIAM, 3330–3361.
- [34] Tyler Helmuth, Matthew Jenssen, and Will Perkins. 2023. Finite-size scaling, phase coexistence, and algorithms for the random cluster model on random graphs. In *Annales de l'Institut Henri Poincaré (B) Probabilités et statistiques*, Vol. 59. Institut Henri Poincaré, 817–848.
- [35] Tyler Helmuth, Will Perkins, and Guus Regts. 2020. Algorithmic Pirogov–Sinai theory. *Probability Theory and Related Fields* 176, 3 (2020), 851–895.
- [36] Mark Huber. 2004. Perfect sampling using bounding chains. *The Annals of Applied Probability* 14, 2 (2004), 734–753.
- [37] Svante Janson, Andrzej Ruciński, and Tomasz Łuczak. 2011. *Random Graphs*. John Wiley & Sons.

- [38] Matthew Jenssen, Peter Keevash, and Will Perkins. 2020. Algorithms for #BIS-hard problems on expander graphs. *SIAM J. Comput.* 49, 4 (2020), 681–710.
- [39] Matthew Jenssen, Aditya Potukuchi, and Will Perkins. 2022. Approximately counting independent sets in bipartite graphs via graph containers. In *Proceedings of the 2022 Annual ACM-SIAM Symposium on Discrete Algorithms (SODA'22)*. SIAM, 499–516.
- [40] Madhav Jha, C. Seshadhri, and Ali Pinar. 2015. Path sampling: A fast and provable method for estimating 4-vertex subgraph counts. In *Proceedings of the 24th International Conference on World Wide Web (WWW'15)*. 495–505.
- [41] Reza Kahkeshani. 2013. A generalization of the Catalan numbers. *Journal of Integer Sequences* 16, 2 (2013), 3.
- [42] Nadav Kashtan, Shalev Itzkovitz, Ron Milo, and Uri Alon. 2004. Efficient sampling algorithm for estimating subgraph concentrations and detecting network motifs. *Bioinformatics* 20, 11 (2004), 1746–1758.
- [43] Roman Kotecký and David Preiss. 1986. Cluster expansion for abstract polymer models. *Communications in Mathematical Physics* 103, 3 (1986), 491–498.
- [44] Timo Kötzing and Martin S. Krejca. 2019. First-hitting times under drift. *Theoretical Computer Science* 796 (2019), 51–69.
- [45] Chao Liao, Jiabao Lin, Pinyan Lu, and Zhenyu Mao. 2019. Counting independent sets and colorings on random regular bipartite graphs. In *Proceedings of the Approximation, Randomization, and Combinatorial Optimization. Algorithms and Techniques (APPROX/RANDOM'19)*. Schloss Dagstuhl-Leibniz-Zentrum fuer Informatik.
- [46] Xuesong Lu and Stéphane Bressan. 2012. Sampling connected induced subgraphs uniformly at random. In *Proceedings of the International Conference on Scientific and Statistical Database Management*. Springer, 195–212.
- [47] Ryuta Matsuno and Aristides Gionis. 2020. Improved mixing time for k -subgraph sampling. In *Proceedings of the 2020 SIAM International Conference on Data Mining*. SIAM, 568–576.
- [48] Kirill Paramonov, Dmitry Shemetov, and James Sharpnack. 2019. Estimating graphlet statistics via lifting. In *Proceedings of the 25th ACM SIGKDD International Conference on Knowledge Discovery & Data Mining*. 587–595.
- [49] James Gary Propp and David Bruce Wilson. 1996. Exact sampling with coupled Markov chains and applications to statistical mechanics. *Random Structures & Algorithms* 9, 1-2 (1996), 223–252.
- [50] Andrew Read-McFarland and Daniel Štefankovič. 2021. The hardness of sampling connected subgraphs. In *Proceedings of the Latin American Symposium on Theoretical Informatics*. Springer, 464–475.
- [51] Nino Shervashidze, SVN Vishwanathan, Tobias Petri, Kurt Mehlhorn, and Karsten Borgwardt. 2009. Efficient graphlet kernels for large graph comparison. In *Proceedings of the Artificial Intelligence and Statistics*. PMLR, 488–495.
- [52] Allan Sly. 2010. Computational transition at the uniqueness threshold. In *Proceedings of the 2010 IEEE 51st Annual Symposium on Foundations of Computer Science (FOCS'10)*. IEEE, 287–296.
- [53] Allan Sly and Nike Sun. 2012. The computational hardness of counting in two-spin models on d -regular graphs. In *Proceedings of the 2012 IEEE 53rd Annual Symposium on Foundations of Computer Science*. IEEE, 361–369.
- [54] Dror Weitz. 2006. Counting independent sets up to the tree threshold. In *Proceedings of the 38th Annual ACM Symposium on Theory of Computing*. 140–149.
- [55] Nacu Șerban and Peres Yuval. 2005. Fast simulation of new coins from old. *The Annals of Applied Probability* 15, 1A (2005), 93–115.

Received 1 June 2022; revised 15 May 2023; accepted 23 October 2023