

One-Shot Learning for k -SAT

Andreas Galanis

University of Oxford, UK

Leslie Ann Goldberg

University of Oxford, UK

Xusheng Zhang

University of Oxford, UK

Abstract

Consider a k -SAT formula Φ where every variable appears at most d times, and let σ be a satisfying assignment of Φ sampled proportionally to $e^{\beta m(\sigma)}$ where $m(\sigma)$ is the number of variables set to true and β is a real parameter. Given Φ and σ , can we learn the value of β efficiently?

This problem falls into a recent line of works about single-sample (“one-shot”) learning of Markov random fields. The k -SAT setting we consider here was recently studied by Galanis, Kandiros, and Kalavasis (SODA’24) where they showed that single-sample learning is possible when roughly $d \leq 2^{k/6.45}$ and impossible when $d \geq (k+1)2^{k-1}$. Crucially, for their impossibility results they used the existence of unsatisfiable instances which, aside from the gap in d , left open the question of whether the feasibility threshold for one-shot learning is dictated by the satisfiability threshold of k -SAT formulas of bounded degree.

Our main contribution is to answer this question negatively. We show that one-shot learning for k -SAT is infeasible well below the satisfiability threshold; in fact, we obtain impossibility results for degrees d as low as k^2 when β is sufficiently large, and bootstrap this to small values of β when d scales exponentially with k , via a probabilistic construction. On the positive side, we simplify the analysis of the learning algorithm and obtain significantly stronger bounds on d in terms of β . In particular, for the uniform case $\beta \rightarrow 0$ that has been studied extensively in the sampling literature, our analysis shows that learning is possible under the condition $d \lesssim 2^{k/2}$. This is nearly optimal (up to constant factors) in the sense that it is known that sampling a uniformly-distributed satisfying assignment is NP-hard for $d \gtrsim 2^{k/2}$.

2012 ACM Subject Classification Mathematics of computing $\rightarrow$ Maximum likelihood estimation; Theory of computation $\rightarrow$ Machine learning theory; Theory of computation $\rightarrow$ Problems, reductions and completeness

Keywords and phrases Computational Learning Theory, k -SAT, Maximum likelihood estimation

Digital Object Identifier 10.4230/LIPIcs.ICALP.2025.84

Category Track A: Algorithms, Complexity and Games

Related Version *Full Version:* <https://arxiv.org/abs/2502.07135> [22]

1 Introduction

A key task that arises in statistical inference is to estimate the underlying parameters of a distribution, frequently based on the assumption that one has access to a sufficiently large number of independent and identically distributed (i.i.d.) samples. However, in many settings it is critical to perform the estimation with substantially fewer samples, driven by constraints in data availability, computational cost, or real-time decision-making requirements. In this

For the purpose of Open Access, the authors have applied a CC BY public copyright licence to any Author Accepted Manuscript version arising from this submission. All data is provided in full in the results section of this paper.



© Andreas Galanis, Leslie Ann Goldberg, and Xusheng Zhang;
licensed under Creative Commons License CC-BY 4.0

52nd International Colloquium on Automata, Languages, and Programming (ICALP 2025).

Editors: Keren Censor-Hillel, Fabrizio Grandoni, Joël Ouaknine, and Gabriele Puppis

Article No. 84; pp. 84:1–84:15



Leibniz International Proceedings in Informatics

Schloss Dagstuhl – Leibniz-Zentrum für Informatik, Dagstuhl Publishing, Germany



paper, we consider the extreme setting where only a single sample is available and investigate the feasibility of parameter estimation in this case. We refer to this setting as “one-shot learning”.

Markov random fields (also known as undirected graphical models) are a canonical framework used to model high-dimensional distributions. The seminal work of [11] initiated the study of one-shot learning for the Ising and spin glass models, a significant class of Markov random fields that includes the well-known Sherrington-Kirkpatrick and Hopfield models. This approach was later explored in greater depth for the Ising model by [3] and subsequently extended to tensor or weighted variants of the Ising model in [25, 35, 15]. Beyond the Ising model, [17, 18] examined one-shot learning in more general settings, notably including logistic regression and higher-order spin systems, obtaining various algorithmic results in “soft-constrained” models, i.e., models where the distribution is supported on the entire state space. [4] showed that efficient parameter estimation using one sample is still possible under the presence of hard constraints which prohibit certain states, relaxing the soft-constrained assumption with “permissiveness”; canonical Markov random fields in this class include various combinatorial models such as the hardcore model (weighted independent sets). Notably, in all these cases, one-shot learning is always feasible with mild average-degree assumptions on the underlying graph (assuming of course access to an appropriate sample).

More recently, [23] investigated one-shot learning for hard-constrained models that are not permissive, focusing primarily on k -SAT and proper colourings; in contrast to soft-constrained models, they showed that one-shot learning is not always possible and investigated its feasibility under various conditions. Their results left however one important question open for k -SAT, in terms of identifying the “right” feasibility threshold. In particular, their impossibility results were based on the existence of unsatisfiable instances for k -SAT, suggesting that it might be the satisfiability threshold that is most relevant for one-shot learning. Here we refute this in a strong way. We show infeasibility well below the satisfiability threshold, and obtain positive results that align closely with the conjectured threshold for sampling satisfying assignments.

1.1 Definitions and Main Results

In the k -SAT model, we consider the state space $\Omega_n := \{\text{TRUE}, \text{FALSE}\}^n$, where each element is an assignment to n Boolean variables. The support of the Markov random field is then restricted to the set of assignments that satisfy a given k -CNF formula. More precisely, we define $\Phi_{n,k,d}$ as the set of CNF formulas with n variables such that each clause has exactly k distinct variables and each variable appears in at most d clauses. For an assignment $\sigma \in \Omega_n$ and a formula $\Psi \in \Phi_{n,k,d}$, we denote by $\sigma \models \Psi$ the event that σ satisfies Ψ and we denote by $m(\sigma)$ the number of variables that are assigned to TRUE in σ . (See Section 1.4 for further details.)

We study the weighted k -SAT model parametrized by β . For a fixed formula $\Psi \in \Phi_{n,k,d}$, the probability for each assignment $\sigma \in \Omega_n$ is given by

$$\Pr_{\Psi,\beta}[\sigma] = \frac{e^{\beta m(\sigma)} \mathbb{1}[\sigma \models \Psi]}{\sum_{\sigma \in \Omega_n} e^{\beta m(\sigma)} \mathbb{1}[\sigma \models \Psi]}, \quad (1)$$

Let $\Omega_n(\Psi) := \{\sigma \in \Omega_n : \sigma \models \Psi\}$ be the support of $\Pr_{\Psi,\beta}$. When $\beta = 0$, this distribution reduces to the uniform distribution over all satisfying assignments $\Omega_n(\Psi)$. For general $\beta \neq 0$, it biases the distribution toward assignments with more TRUE if $\beta > 0$ and biases toward those with more FALSE if $\beta < 0$.

We consider the following one-shot learning task for β . The learner knows parameters d, k and a fixed formula $\Psi \in \Phi_{n,k,d}$. Additionally, the learner has access to a single sample $\sigma \in \Omega_n(\Psi)$ drawn from distribution $\Pr_{\Psi,\beta}[\cdot]$. The learner also knows that β lies within a specified range $|\beta| \leq B$, but it does not know the exact value of β . The goal is to estimate β using these inputs.

To quantify the accuracy of our estimate, we say that $\hat{\beta}$ is an ϵ -estimate if $|\beta - \hat{\beta}| \leq \epsilon$. Typically we want ϵ to decrease as n increases so that $\epsilon \rightarrow 0$ when $n \rightarrow \infty$. In this case we call $\hat{\beta}$ a consistent estimator. On the other hand, if there exists a constant $\epsilon_0 > 0$ such that $\limsup_n |\hat{\beta} - \beta| \geq \epsilon_0$, then $\hat{\beta}$ is not a consistent estimator and we say β is not identifiable by $\hat{\beta}$. Finally, if β is not identifiable by any $\hat{\beta}$, we say it is impossible to estimate β .

Our main algorithmic result is a linear-time one-shot learning algorithm for β in the weighted k -SAT model.

► **Theorem 1.** *Let $B > 0$ be a real number. Let $d, k \geq 3$ be integers such that*

$$d \leq \frac{1}{e^3 \sqrt{k}} \cdot (1 + e^{-B})^{\frac{k}{2}}. \quad (2)$$

There is an estimation algorithm which, for any β^ with $|\beta^*| \leq B$, given any input $\Phi \in \Phi_{n,k,d}$ and a sample from $\sigma \sim \Pr_{\Phi,\beta^*}$, outputs in $O(n + \log(nB))$ time an $O(n^{-1/2})$ -estimate $\hat{\beta}(\sigma)$ such that*

$$\Pr_{\Phi,\beta^*} \left[|\hat{\beta}(\sigma) - \beta^*| = O(n^{-1/2}) \right] = 1 - e^{-\Omega(n)}.$$

Our results improve upon the conditions in [23], which ensure a consistent estimate under the requirement when $d \lesssim (1 + e^{-B})^{k/6.45}$. Based on the corresponding threshold for approximate sampling, the conjectured “true” threshold for d is of the order $(1 + e^{-B})^{\frac{k}{2}}$. Consequently, our improved condition in (2) is only off by a polynomial factor in k relative to this conjectured threshold.

For comparison with the approximate sampling threshold – commonly stated for the uniform k -SAT distribution – we specialize to $B \rightarrow 0$. In that limit, our algorithmic result for single-sample learning holds roughly when $d \lesssim 2^{k/2}$. The best currently known result for efficient sampling, due to [38], holds under the condition $d \lesssim 2^{k/4.82}$, see also the series of works [33, 19, 30, 29]. It is conjectured that the sharp condition for efficient sampling is $d \lesssim 2^{k/2}$, supported by matching hardness results for monotone formulas. It is known in particular that for $d \gtrsim 2^{k/2}$, no efficient sampling algorithm exists (unless $\text{NP} = \text{RP}$).

To complement our algorithmic result, we also present impossibility results, suggesting that conditions like (2) are nearly sharp.

► **Theorem 2.** *Let β^* be a real number such that $|\beta^*| > 1$. Let $k \geq 4$ be an even integer, and let n be a multiple of $k/2$ that is large enough. If*

$$d \geq k^3 \left(1 + \frac{e}{e^{|\beta^*|} - e} \right)^{\frac{k}{2}}, \quad (3)$$

then there exists a formula $\Phi \in \Phi_{n,k,d}$ such that it is impossible to estimate β^ from a sample $\sigma \sim \Pr_{\Phi,\beta^*}$ with high probability.*

For the parameter d around the satisfiability threshold, specifically at the uniquely satisfiable threshold $u(k)$ (see, e.g., [32] on the connection between these two thresholds), if

$$d \geq u(k) = \Theta\left(\frac{2^k}{k}\right), \quad (4)$$

there exists a formula $\Phi \in \Phi_{n,k,d}$ such that it is impossible to estimate β^* from any number of samples $\sigma \sim \Pr_{\Phi,\beta^*}$ because $\Omega_n(\Phi)$ is a deterministic set consisting of a single satisfying assignment that does not depend on β^* . [23] explicitly constructs such a formula Φ , though it requires an additional $O(k^2)$ factor relative to (4), representing the previous best known condition for the impossibility of estimation. Condition (3) in Theorem 2 not only relaxes (4) when $|\beta^*|$ grows large, but it also features the correct $k/2$ exponent, matching that in both (2) and the conjectured threshold. Indeed, when $B \approx |\beta^*| \rightarrow \infty$, conditions (2) and (3) both take the form

$$(1 + O(e^{-|\beta^*|}))^{k/2} \cdot k^{O(1)} \quad (5)$$

These findings partially indicate that, at least for the k -SAT model, the sampling threshold is more relevant to one-shot learning than the satisfiability threshold.

In addition, we find that if we allow β^* to be proportional to k , then learning becomes impossible for a significantly larger range of d . Specifically, unlike condition (3), which requires d to be exponential in k , here we only need d to be quadratic in k , leading to a much sparser formula when k is large.

► **Theorem 3.** *Let $k \geq 4$ be an even integer. For all $\beta^* \in \mathbb{R}$ such that $|\beta^*| \geq k \ln 2$ the following holds. Let n be a multiple of $k/2$ that is large enough. If $d \geq k^2/2$, then there exists a formula $\Phi \in \Phi_{n,k,d}$ such that it is impossible to estimate β^* from a sample $\sigma \sim \Pr_{\Phi,\beta^*}$ with high probability.*

► **Remark 4.** In the regimes where Theorem 2 or Theorem 3 apply, the corresponding formula Φ ensures that there is a single assignment that is the output with all but exponentially-small probability, regardless of the value of β^* . Hence, the proof of this theorem guarantees that it is impossible to learn from exponentially many independent samples. Moreover, for any pair of (β_1, β_2) such that $|\beta_1| > |\beta_2| \geq |\beta^*|$ and $\beta_1\beta_2 \geq 0$, no hypothesis testing for $H_0 : \beta = \beta_1$ versus $H_1 : \beta = \beta_2$ can be done to distinguish $\Pr_{\Phi,\beta_1}$ from $\Pr_{\Phi,\beta_2}$, that is, there exists no sequence of consistent test functions $\phi_n : \Omega_n \rightarrow \{0, 1\}$ such that

$$\lim_{n \rightarrow \infty} \mathbb{E}_{\sigma \sim \Pr_{\Phi,\beta_1}} \phi_n(\sigma) = 0 \text{ and } \lim_{n \rightarrow \infty} \mathbb{E}_{\sigma \sim \Pr_{\Phi,\beta_2}} \phi_n(\sigma) = 1.$$

1.2 Proof overview

We prove Theorem 1 by using the maximum pseudo-likelihood estimator. Establishing the consistency of this estimator – as stated in Theorem 5 – requires demonstrating that the log-pseudo-likelihood function is (strongly) concave; see (11) for the precise formulation. In the k -SAT setting, showing such concavity amounts to showing that with high probability over samples drawn from $\Pr_{\Phi,\beta^*}[\cdot]$, each sample contains a linear number of “flippable variables”. We prove this property in Lemma 6 by applying the Lovász Local Lemma (LLL), which enables us to compare $\Pr_{\Phi,\beta^*}[\cdot]$ to a suitable product distribution – under which the number of flippable variables is guaranteed to be linear. Notably, we apply the LLL directly to a non-local set of variables, in contrast to previous analyses that confined the application of the LLL to local neighborhoods – a restriction that typically imposes stronger constraints on the parameter regime. By circumventing these stronger constraints, our approach achieves its guarantee under the nearly optimal LLL condition.

The main technical novelty of this paper is our negative results. To explain these, we begin by outlining the proof of Theorem 3. For simplicity, we focus on the case where $\beta^* \geq 0$. On a high level, we construct a gadget formula Ψ_0 for which the all-true assignment

σ^+ carries almost all of the probability mass under $\Pr_{\Psi_0, \beta^*}[\cdot]$ provided that $\beta^* \geq k \ln 2$. Consequently, a sample $\sigma \sim \Pr_{\Psi_0, \beta^*}[\cdot]$ drawn from this distribution is nearly deterministic, offering virtually no information about β^* and thus rendering learning impossible. The key property of this gadget is established in Lemma 8. Specifically, the lemma ensures that σ^+ satisfies Ψ_0 and that any other assignment σ with fewer than $2n/k$ variables set to FALSE is not a satisfying assignment of Ψ_0 . We achieve this by incorporating a cyclic structure over the variables that enforces global correlation among the FALSE values in the assignments. In particular, there are no flippable variables in σ^+ .

Towards the proof of Theorem 2, we first leverage the gadget Ψ_0 to show the existence of a stronger gadget Ψ_2 , parametrized by $b > 1$ in Lemma 13, which guarantees that the all-true assignment σ^+ satisfies Ψ_2 and any other assignment with fewer than n/b variables set to FALSE fails to satisfy Ψ_2 . Then we choose b appropriately in terms of β^* to make sure that σ^+ carries nearly all of the probability mass, using some more technical estimates for the corresponding partition function. To build Ψ_2 , we take a finite number of replicas of Ψ_0 on randomly permuted sets of variables. The existence of the desired formula is established using the probabilistic method; specifically, we demonstrate an upper bound on the expectation of $m(\sigma)$ for any satisfying assignment $\sigma \neq \sigma^+$, over the choice of the permutations.

1.3 Related work

Parameter Estimation in Markov Random Fields. A large body of work has focused on parameter estimation under the one-shot learning paradigm (see, e.g., [11, 3, 17, 18, 25, 15, 35]), particularly for Ising-like models in statistical physics and for dependent regression models in statistics. In this work, we follow a similar approach by establishing the consistency of the maximum pseudo-likelihood estimator. Earlier studies (e.g., [26, 14, 13, 24]) have also explored parameter estimation in Markov random fields using the maximum likelihood estimator.

Before our work, the papers [4, 23] were the first to study one-shot learning in hard-constrained models. In particular, the hardcore model analysed in [4] can be viewed as a weighted monotone 2-SAT model, and one natural extension of the hardcore model to k -uniform hypergraphs corresponds to the class of weighted monotone k -SAT models – a special case of the weighted k -SAT models that we consider. Because a typical assignment in these monotone formulas possesses $\Omega(n)$ flippable variables, the pseudo-likelihood estimator remains consistent across all parameter regimes, and no phase transition is expected. The weighted k -SAT problem was analysed in [23], where the authors derived both a consistency condition and an impossibility condition, though a substantial gap remained between them. By tightening the bounds on both ends, our work considerably narrows this gap, nearly closing it entirely.

Related Works in Structural Learning/Testing. An alternative direction in learning Markov Random Fields involves estimating the interaction matrix between variables – a question originally posed by [12]. For the Ising model, this problem has been extensively studied (see, e.g., [6, 37, 9] and the references therein), and subsequent work has extended the results to higher-order models [31, 28, 20]. Recent work [39, 15, 21] has also considered the joint learning of both structure and parameters. Moreover, [36] establishes the information-theoretic limits on what any learner can achieve, and similar analyses have been conducted for hardcore models [8, 7]. While some approaches in this line of work require multiple independent samples, as noted in [15], it is also possible to reduce learning with $O(1)$ samples to a class of special cases within one-shot learning. Related problems in one-shot testing for Markov random fields have also been studied in [10, 16, 34, 2, 5].

1.4 k -SAT Notation

We use standard notation for the k -SAT problem. A formula $\Phi = (V, C)$ denotes a *CNF* (*Conjunctive Normal Form*) formula with variables $V = \{x_1, \dots, x_n\}$ and clauses C . We use $\sigma(x_i)$ and σ_i to denote the truth value of the variable x_i under an assignment $\sigma: V \rightarrow \{\text{TRUE}, \text{FALSE}\}$. For any clause $c \in C$, $\text{var}(c)$ denotes the set of variables appearing in c (negated or not). The *degree* of a variable x_j in Φ is the number of clauses in which x_j or $\neg x_j$ appears, namely $|\{c \in C : x_j \in \text{var}(c)\}|$. The degree of Φ is the maximum, over $x_j \in V$, of the degree of x_j in Φ . As noted in the introduction, $m(\sigma)$ denotes the number of variables that are assigned to TRUE by an assignment σ . Namely, $m(\sigma) := |\{i \in [n] : \sigma_i = \text{TRUE}\}|$.

2 Maximum pseudo-likelihood estimator: the proof of Theorem 1

Section 2.1 introduces the fundamentals of the maximum pseudo-likelihood estimator and analyses its running time for solving the weighted k -SAT problem. In Sections 2.2 and 2.3, we establish the estimator's consistency.

2.1 Overview of maximum (pseudo)-likelihood estimation

For a k -SAT formula Ψ and a satisfying assignment σ , we will use $f(\beta; \sigma)$ to denote the quantity $\Pr_{\Psi, \beta}(\sigma)$ from (1), i.e., $f(\beta; \sigma) = e^{\beta m(\sigma)} / Z(\beta; \Psi)$, where $Z(\beta; \Psi)$ is the normalising constant of the distribution (the partition function).

A standard approach to parameter estimation is to find $\hat{\beta}_{\text{MLE}}(\sigma) := \arg \max_{\beta} f(\beta; \sigma)$, which is commonly referred to as the maximum likelihood estimate (MLE). However, two main obstacles arise when applying MLE directly to the weighted k -SAT problem. First, (approximately) computing $Z(\beta; \Psi)$ is generally intractable because it is an NP-hard computation. Second, even if an approximation algorithm exists for computing $\hat{\beta}_{\text{MLE}}(\sigma)$, there is no guarantee of its consistency, i.e., there is no guarantee that with high probability it is close to β^* . Hence, we take a computationally more tractable variant of MLE from [1] which is called the maximum pseudo-likelihood estimation. Let $f_i(\beta; \sigma)$ be the conditional probability of σ_i , in a distribution with parameter β , conditioned on the value of $\sigma_{-i} := (\sigma_j)_{j \neq i}$. The maximum pseudo-likelihood estimate (MPLE) is defined as

$$\hat{\beta}_{\text{MPLE}}(\sigma) := \arg \max_{\beta} \prod_{i \in V} f_i(\beta; \sigma) = \arg \max_{\beta} \sum_{i \in V} \ln f_i(\beta; \sigma).$$

Here, the objective function $F(\beta; \sigma) := \sum_{i \in V} \ln f_i(\beta; \sigma)$ is the so-called *log-pseudo-likelihood function*. For the weighted k -SAT problem, it is not hard to compute $f_i(\beta; \sigma)$ as

$$f_i(\beta; \sigma) = \frac{e^{\beta \sigma_i}}{e^{\beta} \mathbb{1}[\sigma_{-i} \wedge (\sigma_i \leftarrow \text{TRUE})] + \mathbb{1}[\sigma_{-i} \wedge (\sigma_i \leftarrow \text{FALSE})]},$$

where $\mathbb{1}[\omega]$ is shorthand for $\mathbb{1}[\omega \models \Psi]$. So we can write the log-pseudo-likelihood function for k -SAT as

$$\begin{aligned} F(\beta; \sigma) &= \sum_{i \in V} \ln \left(\frac{e^{\beta \sigma_i}}{e^{\beta} \mathbb{1}[\sigma_{-i} \wedge (\sigma_i \leftarrow \text{TRUE})] + \mathbb{1}[\sigma_{-i} \wedge (\sigma_i \leftarrow \text{FALSE})]} \right), \\ &= \beta m(\sigma) - \sum_{i \in V} \ln (e^{\beta} \mathbb{1}[\sigma_{-i} \wedge (\sigma_i \leftarrow \text{TRUE})] + \mathbb{1}[\sigma_{-i} \wedge (\sigma_i \leftarrow \text{FALSE})]). \end{aligned}$$

For a fixed σ , $F(\cdot; \sigma) : \mathbb{R} \rightarrow \mathbb{R}$ is a function of β . By taking derivative with respect to β , we obtain

$$\frac{\partial F(\beta; \sigma)}{\partial \beta} = m(\sigma) - \sum_{i \in V} \frac{e^\beta \mathbb{1}[\sigma_{-i} \wedge (\sigma_i \leftarrow \text{TRUE})]}{e^\beta \mathbb{1}[\sigma_{-i} \wedge (\sigma_i \leftarrow \text{TRUE})] + \mathbb{1}[\sigma_{-i} \wedge (\sigma_i \leftarrow \text{FALSE})]}. \quad (6)$$

Clearly, $\frac{\partial F(\beta; \sigma)}{\partial \beta}$ is a decreasing function of β , which implies that $F(\beta; \sigma)$ has a unique global maximum that is achieved when $\frac{\partial F(\beta; \sigma)}{\partial \beta} = 0$. Therefore, $\hat{\beta}_{\text{MPLE}}(\sigma)$ can be uniquely defined to be the maximum of $F(\beta; \sigma)$.

Moreover, provided $|\hat{\beta}_{\text{MPLE}}(\sigma)| \leq 2B$, an ϵ -close estimate of $\hat{\beta}_{\text{MPLE}}(\sigma)$ can be computed, using $O(\ln(B/\epsilon))$ steps of binary search for the solution to $\frac{\partial F(\beta; \sigma)}{\partial \beta} = 0$. At each step of the binary search, we evaluate $\frac{\partial F(\beta; \sigma)}{\partial \beta}$ and adjust the binary search interval based on its sign. A naive evaluation of $\frac{\partial F(\beta; \sigma)}{\partial \beta}$ as in (6) would require $\Theta(n)$ operations per step. We can reduce this by exploiting the fact that the summand

$$S_i(\beta) := \frac{e^\beta \mathbb{1}[\sigma_{-i} \wedge (\sigma_i \leftarrow \text{TRUE})]}{e^\beta \mathbb{1}[\sigma_{-i} \wedge (\sigma_i \leftarrow \text{TRUE})] + \mathbb{1}[\sigma_{-i} \wedge (\sigma_i \leftarrow \text{FALSE})]}$$

can take only three values $\{0, 1, e^\beta/(1+e^\beta)\}$. Hence, by grouping the summands according to their values, we obtain

$$m(\sigma) - \sum_{i \in V} S_i(\beta) = m(\sigma) - |\{i \in V : S_i(\beta) = 1\}| - \frac{e^\beta}{1+e^\beta} \cdot \left| \left\{ i \in V : S_i(\beta) = \frac{e^\beta}{1+e^\beta} \right\} \right|. \quad (7)$$

Crucially, the sets $\{i \in V : S_i(\beta) = 1\}$ and $\{i \in V : S_i(\beta) = \frac{e^\beta}{1+e^\beta}\}$ do not depend on β , and thus can be computed only once before the binary search. After this preprocessing, each evaluation of $\frac{\partial F(\beta; \sigma)}{\partial \beta}$ can be done in $O(1)$ time using the decomposition in (7). Overall, to achieve an $O(n^{-1/2})$ -close estimate, the total running time is $O(n + \log(nB))$.

2.2 Consistency of MPLE

In Section 2.1 we gave an algorithm with running time $O(n + \log(nB))$, which takes as input a formula Ψ and an assignment $\sigma \in \Omega_n(\Psi)$ and outputs an estimate $\hat{\beta}(\sigma)$ which satisfies

$$|\hat{\beta}(\sigma) - \hat{\beta}_{\text{MPLE}}(\sigma)| = O(n^{-1/2}),$$

provided $|\hat{\beta}_{\text{MPLE}}(\sigma)| \leq 2B$. Theorem 1 follows immediately from the Theorem 5, which demonstrates $O(\sqrt{n})$ -consistency.

► **Theorem 5.** *Let β^* be a real number, and let $d, k \geq 3$ be integers such that*

$$d \leq \frac{1}{e^3 \sqrt{k}} \cdot (1 + e^{-|\beta^*|})^{\frac{k}{2}}. \quad (8)$$

For any integer n and $\Phi \in \Phi_{n,k,d}$,

$$\Pr_{\Phi, \beta^*} \left[\left| \hat{\beta}_{\text{MPLE}}(\sigma) - \beta^* \right| = O(n^{-1/2}) \right] = 1 - e^{-\Omega(n)}. \quad (9)$$

As in the standard literature [11], the consistency result can be proved using bounds on the derivatives of the log-pseudo-likelihood function f . In particular, following the analysis in [23], (9) is a consequence of these two bounds:

1. A uniform linear upper bound of the second moment of the first derivative of fF : for all $\beta \in \mathbb{R}$,

$$\mathbb{E}_{\Phi, \beta} \left[\left(\frac{\partial F(\beta; \sigma)}{\partial \beta} \right)^2 \right] \leq kdn \quad (10)$$

2. With high probability over $\sigma \sim \Pr_{\Phi, \beta^*}$, there is a uniform lower bound of the second derivative of F :

$$\Pr_{\Phi, \beta^*} \left[\inf_{\beta \in \mathbb{R}} \frac{\partial^2 F(\beta; \sigma)}{\partial \beta^2} = \Omega(n) \right] = 1 - e^{-\Omega(n)}. \quad (11)$$

Moreover, Lemma 7 in [23] proves the bound (10) for all $\beta \in \mathbb{R}$ and all $d, k \geq 3$, and they give a combinatorial expression (equations (9) and (10) in [23]) for $\frac{\partial^2 F(\beta; \sigma)}{\partial \beta^2}$ that will be the useful for proving (11). To state this expression, we introduce the notion of flippable variables. We say a variable v_i is *flippable* in σ if the assignment, obtained by flipping the value of variable v_i while keeping the values of other variables in σ , is still a satisfying assignment, that is, $(\sigma_{-i} \wedge (\neg \sigma_i)) \models \Psi$. We use $e_{v_i}(\sigma)$ to denote the indicator of the event that variable v_i is flippable in a satisfying assignment σ . It was shown that

$$\frac{\partial^2 F(\beta; \sigma)}{\partial \beta^2} = \frac{e^\beta}{(e^\beta + 1)^2} \sum_{v \in V} e_v(\sigma). \quad (12)$$

Hence, proving (11) reduces to establishing a linear lower bound on the number of flippable variables. The main ingredient in the proof of our positive result is Lemma 6, which provides such a lower bound under the condition (8).

► **Lemma 6.** *Let β^* be a real number, and let $d, k \geq 3$ be integers such that*

$$d \leq \frac{1}{e^3 \sqrt{k}} \cdot (1 + e^{-|\beta^*|})^{\frac{k}{2}}.$$

Then for a fixed $\Phi = (V, C) \in \Phi_{n, k, d}$,

$$\Pr_{\Phi, \beta^*} \left[\sum_{v \in V} e_v(\sigma) = \Omega(n) \right] = 1 - e^{-\Omega(n)}. \quad (13)$$

Using Lemma 6 and the identity (12), we derive (11) under the condition (8), thereby completing the proof of Theorem 5. The proof of Lemma 6 will be presented in the next section.

2.3 Proof of Lemma 6: applying the Lovasz's local lemma in a batch

The following Lovasz's local lemma (LLL) from [27] will be useful for our proof of Lemma 6.

► **Lemma 7** (Lemma 26, [27]). *Let $\Phi = (V, C)$ be a CNF formula. Let μ be the product distribution on $\{\text{TRUE}, \text{FALSE}\}^{|V|}$, where each variable is set to TRUE independently with probability $e^\beta / (1 + e^\beta)$. For each $c \in C$, let A_c be the event that clause c is unsatisfied. If there exists a sequence $\{x_c\}_{c \in C}$ such that for each $c \in C$,*

$$\Pr_\mu[A_c] \leq x_c \cdot \prod_{j \in \Gamma(c)} (1 - x_j), \quad (14)$$

where $\Gamma(c) \subseteq C$ is the set of clauses that contain a variable in $\text{var}(c)$, then Φ has a satisfying assignment. Moreover, the distributions $\Pr_{\Phi, \beta}$ and $\Pr_{\mu}$ can be related as follows: for any event E that can be completely determined by the assignment of a set S of variables,

$$\Pr_{\Phi, \beta}[E] \leq \Pr_{\mu}[E] \cdot \prod_{j \in \Gamma(E)} \frac{1}{1 - x_j}, \quad (15)$$

where $\Gamma(E)$ denotes the set of all clauses that contain a variable in $\text{var}(S)$.

In many previous works utilising the LLL for sampling purposes, Lemma 7, or a variant of it, is typically applied to local events, including, for instance, the event that a specific variable is flippable or, more generally, to events happening in the neighbourhood of a vertex. This is present in the approach of [23] which, in the end, imposed more strict conditions on d (relative to k) since it requires “marking” variables appropriately (see [33]). Here, we prove Lemma 6 by applying Lemma 7 directly to a *batch* R of random variables scattered around the graph in one go, removing the need for marking, and relaxing significantly the conditions on d . This simple idea enables our stronger result. In particular, we set $x_c = (d^2k + 1)^{-1}$ and let E be the event $\{\sum_{i=1}^R e_{v_i}(\sigma) < \frac{R}{3}\}$, where $R = \Omega(n)$ and $v_1, \dots, v_R$ are well-separated variables in the hypergraph corresponding to Φ . We refer the reader to the full paper for a detailed proof of this lemma.

3 Impossibility of learning: proofs of Theorems 2 and 3

For the construction of the impossibility instances, we begin with a “gadget” that will serve as a building block.

► **Lemma 8.** *Let $k \geq 4$ be an even integer and let $n \geq k$ be a multiple of $k/2$. If $d \geq k^2/2$, then there is a formula $\Psi_0 \in \Phi_{n, k, d}$ such that if an assignment σ of Ψ_0 is satisfying then either*

1. σ has n TRUEs, or
2. σ has at least $2n/k$ FALSEs.

Similarly (by symmetry), there is a formula $\Psi_1 \in \Phi_{n, k, d}$ such that, for any satisfying assignment σ of Ψ_1 , either σ has n FALSEs, or σ has at least $2n/k$ TRUEs.

The instance Ψ_0 and Ψ_1 will lead to a proof of Theorem 3. For clarity, we denote the assignment to n variables with n TRUEs as σ^+ .

Proof of Theorem 3. Assume $\beta^* \geq k \ln 2$ without loss of generality. Let Ψ_0 be the formula given by Lemma 8 (when $\beta^* \leq -k \ln 2$ we use Ψ_1 instead). Suppose σ is drawn from $\Pr_{\Psi_0, \beta^*}$, and we directly estimate the probability of $\{\sigma = \sigma^+\}$: since σ has at most 2^n possibilities and on event $\{\sigma \neq \sigma^+\}$ the number of TRUEs is at most $n - \frac{2n}{k}$, we have

$$\Pr_{\Psi_0, \beta^*}[\sigma = \sigma^+] \geq \frac{e^{\beta^* n}}{e^{\beta^* n} + 2^n \cdot e^{\beta^* \cdot (n - 2n/k)}} \geq \frac{2^{kn}}{2^{kn} + 2^{n+k \cdot (n - 2n/k)}} = \frac{1}{1 + 2^{-n}}.$$

As the samples from $\Pr_{\Psi_0, \beta^*}$ are insensitive to β^* with high probability, learning β^* from σ is impossible. ◀

We now present a detailed description of the formula that defines Ψ_0 in Lemma 8. Let $k \geq 3$ be an even integer and let $n \geq k$ be a multiple of $k/2$. Let $N = \{0, \dots, n-1\}$. The variables of Ψ_0 are $\{x_i \mid i \in N\}$.

For the construction, it will be helpful to group variables in batches of $k/2$ variables in a cyclic manner. Specifically, for $i \in \mathbb{N}$, consider the i -th batch of indices $\Xi_i = \{i + j \pmod{n} \mid j \in \{0, \dots, k/2 - 1\}\}$ and let $C_i = \{x_\ell \mid \ell \in \Xi_i\}$ be the corresponding variables in the i -th batch. We now introduce two types of length- $k/2$ clauses $W_{i,\ell}$ and Π_i that will be used to form the final length- k clauses. Specifically, for each ℓ in the batch Ξ_i , $W_{i,\ell}$ is the length- $(k/2)$ clause with variable set C_i in which x_ℓ appears positively and all other variables are negated. Let Π_i be the length- $(k/2)$ clause with variable set C_i in which all variables are negated. Finally,

$$\Psi_0 := \bigwedge_{i \in N, \ell \in \Xi_i} (W_{i,\ell} \vee \Pi_{i+k/2}).$$

(so Ψ_0 is the formula with variable set $\{x_i \mid i \in N\}$ and clause set $\{W_{i,\ell} \vee \Pi_{i+k/2} \mid i \in N, \ell \in \Xi_i\}$). The formula Ψ_1 is obtained from Ψ_0 by negating all of the literals.

Note that $\Psi_0, \Psi_1 \in \Phi_{n,d,k}$ for every integer $d \geq k^2/2$, since for each $j \in N$, the literals x_j and $\neg x_j$ occur (together) $k^2/2$ times. To see this, note that for each $j \in N$, x_j or $\neg x_j$ comes up in $W_{i,\ell} \vee \Pi_{i+k/2}$ for all $i \in \{j - k - 1 \pmod{n}, \dots, j\}$ and all $\ell \in \Xi_i$ so this is k different i 's and $k/2$ different ℓ 's. In the proof of Lemma 9 we will demonstrate that Ψ_0 (and analogously, Ψ_1) satisfies the requirements of Lemma 8.

► **Lemma 9.** *Let $k \geq 4$ be an even integer and let $n \geq k$ be a multiple of $k/2$. If $\sigma \neq \sigma^+$ satisfies Ψ_0 then, for all $l \in \{0, 1, \dots, \frac{2n}{k} - 1\}$, σ assigns at least one variable in $C_{lk/2} \cup C_{(l+1)k/2}$ to FALSE and σ has at least $\frac{2n}{k}$ FALSEs in total.*

Proof of Lemma 9. We will use the following claim as a key step of the proof.

▷ **Claim 10.** If $\sigma \neq \sigma^+$ satisfies Ψ_0 , and σ assigns one variable x_a in $C_{\ell k/2}$ to FALSE, then either

1. σ assigns a variable x_b in $C_{(\ell+1)k/2}$ to FALSE, or
2. All variables in $C_{(\ell+1)k/2}$ are assigned to TRUE in σ , and there exist variables $x_c \in C_{(\ell+2)k/2}$ and $x_d \in C_{\ell k/2} \setminus \{x_a\}$ that are assigned to FALSE in σ .

Proof of the Claim. Suppose we are not in the first case, so we will show that σ assigns FALSE to $x_c \in C_{(\ell+2)k/2}$ and $x_d \in C_{\ell k/2} \setminus \{x_a\}$. Since σ satisfies Ψ_0 , it satisfies at least one of $W_{\ell k/2, a}$ and $\Pi_{(\ell+1)k/2}$. Since variables in $C_{(\ell+1)k/2}$ are all assigned to TRUE by the assumption that we are not in the first case, σ does not satisfy $\Pi_{(\ell+1)k/2}$. If σ satisfies $W_{\ell k/2, a}$ then it assigns a variable x_d to FALSE where $d \in \Xi_{\ell k/2} \setminus \{a\}$. Also, the set $\{j \in \Xi_{\ell k/2} : \sigma(x_j) = \text{FALSE}\}$ is not empty. Let $j = \max\{j \in \Xi_{\ell k/2} : \sigma(x_j) = \text{FALSE}\}$. Note that σ does not satisfy $W_{j,j}$, so for σ to satisfy $W_{j,j} \vee \Pi_{j+(k/2)}$, it must satisfy $\Pi_{j+(k/2)}$, which means σ assigns a variable $x_c \in C_{j+(k/2)}$ to FALSE. Since $x_c \notin C_{(\ell+1)k/2}$ and $C_{j+(k/2)} \subseteq C_{(\ell+1)k/2} \cup C_{(\ell+2)k/2}$, we establish that $x_c \in C_{(\ell+2)k/2}$. This concludes the proof of the claim. ◁

We now show how to use the claim to prove the lemma. Fix an assignment $\sigma \neq \sigma^+$ that satisfies Ψ_0 . Fix an index $j(0)$ so that $x_{j(0)}$ is assigned FALSE by σ . By symmetry of Ψ_0 , we could assume $j(0) \in \Xi_0$. Let $\ell(0) = 0$ and $G(0) = \{x_{j(0)}\}$. Consider three sequences $\{j(t)\}_{t \geq 0}$, $\{\ell(t)\}_{t \geq 0}$ and $\{G(t)\}_{t \geq 0}$ defined recursively as follows. For every positive integer t , applying the claim to $a = j(t) \in \Xi_{\ell(t)k/2}$, in the first case we let $\ell(t+1) = \ell(t) + 1$, $j(t+1) = b \in \Xi_{(\ell(t)+1)k/2}$ and $G(t+1) = G(t) \cup \{x_b\}$; in the latter case we let $\ell(t+1) = \ell(t) + 2$, $j(t+1) = c \in \Xi_{(\ell(t)+2)k/2}$ and $G(t+1) = G(t) \cup \{x_c, x_d\}$. By induction on t (with base case $t = 0$) we conclude that (i) $j(t) \in \Xi_{\ell(t)k/2}$ is assigned to FALSE, (ii) σ assigns all variables in $G(t)$ to FALSE, and (iii) $\ell(t) + 2 \geq \ell(t+1) \geq \ell(t) + 1$ for $t < T$, where

$T := \min\{t > 0 : \ell(t) = 2n/k - 1 \text{ or } 2n/k\}$ is a stopping time of $\{j(t), \ell(t), G(t)\}_{t \geq 0}$. By construction, for all $l \in \{0, 1, \dots, \frac{2n}{k} - 1\}$, $G(T) \cap (C_{lk/2} \cup C_{(l+1)k/2})$ is not empty. Hence we have proved the first part of the lemma.

Next we will show that $|G(T)| \geq \frac{2n}{k}$. For this, observe that for $t < T$, $|G(t+1)| = |G(t)| + \ell(t+1) - \ell(t)$. Since $|G(0)| = 1$, $\ell(0) = 0$ and $\ell(T) \geq \frac{2n}{k} - 1$, it holds that $|G(T)| \geq \frac{2n}{k}$. $\blacktriangleleft$

Proof of Lemma 8. In the case of Ψ_0 , first note that for all $i \in N$, σ^+ satisfies all instances of $W_{i,\ell} \vee \Pi_{i+k/2}$, so σ^+ satisfies Ψ_0 . Also, if $\sigma \neq \sigma^+$, then the lemma follows immediately from Lemma 9. The case of Ψ_1 holds analogously. $\blacktriangleleft$

The names of the indices of the variables of Ψ_0 are not very important, and when we generalise the construction in Lemma 13 it will be useful to consider an arbitrary permutation of them. Here is the notation that we will use. Let π be any permutation of $N = \{0, \dots, n-1\}$. We use the notation $\pi(i)$ to denote the element in N that i is mapped to by π . We will construct a formula Ψ_0^π . Taking id to be the identity permutation on N , the formula Ψ_0 that was already defined is Ψ_0^{id} .

For $i \in \mathbb{N}$, let $\Xi_i^\pi = \{\pi(i+j \pmod n) \mid j \in \{0, \dots, k/2-1\}\}$ and let $C_i^\pi = \{x_\ell \mid \ell \in \Xi_i^\pi\}$. For $\ell \in \Xi_i^\pi$, let $W_{i,\ell}^\pi$ be the length- $(k/2)$ clause with variable set C_i^π in which x_ℓ appears positively and all other variables are negated. Let Π_i^π be the length- $(k/2)$ clause with variable set C_i^π in which all variables are negated. Then $\Psi_0^\pi := \bigwedge_{i \in N, \ell \in \Xi_i^\pi} (W_{i,\ell}^\pi \vee \Pi_{i+k/2}^\pi)$. The proof that $\Psi_0^\pi \in \Phi_{n,d,k}$ for any $d \geq k^2/2$ is exactly the same as the case $\pi = \text{id}$. The formula Ψ_1^π is obtained from Ψ_0^π by negating all of the literals.

The following corollary follows immediately from the proof of Lemma 9 (by renaming the indices using π) and the fact that $C_0, C_{k/2}, \dots, C_{2n/k-1}$ are disjoint sets.

► **Corollary 11.** *Let $k \geq 4$ be an even integer and let $n \geq k$ be a multiple of $k/2$. Let π be a permutation of N . If $\sigma \neq \sigma^+$ satisfies Ψ_0^π then there exists a subset $\mathcal{M}^\pi(\sigma) \subseteq \{0, 1, \dots, \frac{2n}{k}-1\}$ of size at least $\frac{n}{k}$ such that for all $\ell \in \mathcal{M}^\pi(\sigma)$, σ assigns at least one variable in $C_{\ell k/2}^\pi$ to FALSE.*

► **Remark 12.** While Corollary 11 does not guarantee the uniqueness of $\mathcal{M}^\pi(\sigma)$, in what follows we define $\mathcal{M}^\pi(\sigma)$ to be the lexicographically smallest set among all the smallest sets satisfying the corollary. Since there are finitely many such sets and they can be lexicographically ordered, $\mathcal{M}^\pi(\sigma)$ is a unique and well-defined set for given π and σ .

Lemma 8 provides formula Ψ_0 with a GAP property applying to the number of FALSEs in its satisfying assignments. In the next lemma, we use Ψ_0 to build a larger formula Ψ_2 that amplifies the GAP property to make the gap arbitrarily large.

► **Lemma 13.** *Let $k \geq 4$ be an even integer, let $b > 1$ be a real number, and let d be an integer satisfying*

$$d \geq k^3 \left(1 - \frac{1}{b}\right)^{-k/2}.$$

Let $n \geq k$ be a sufficiently large multiple of $k/2$. Then there is a formula $\Psi_2 \in \Phi_{n,k,d}$ such that if an assignment σ satisfies Ψ_2 then either

1. σ has n TRUEs, or
2. σ has at least n/b FALSEs.

Similarly (by symmetry), there is a formula $\Psi_3 \in \Phi_{n,k,d}$ such that, for any satisfying assignment σ of Ψ_3 , either σ has n FALSEs, or σ has at least n/b TRUEs.

Specifically, the set of clauses in Ψ_2 is of the form

$$\bigcup_{\pi \in \Gamma} \{W_{i,\ell}^\pi \vee \Pi_{i+k/2}^\pi \mid i \in N, \ell \in \Xi_i^\pi\},$$

where Γ is a set of approximately $\frac{2k}{b} \left(1 - \frac{1}{b}\right)^{-k/2}$ independently-chosen permutations of N . We use the probabilistic method to show that, with positive probability over the choice of Γ , every satisfying assignment $\sigma \neq \sigma^+$ of Ψ_2 has at least n/b variables assigned to FALSE. To facilitate this counting argument, we apply Corollary 11. For a complete proof of this lemma, we direct the reader to the full paper.

We present the following more general result, from which Theorem 2 follows as a corollary. To prove Theorem 2, we apply Lemma 13 to show that the distribution $\Pr_{\Psi_2, \beta^*}$ is sharply concentrated at a single point, σ^+ . This parallels the approach used in Theorem 3, but requires more delicate estimates. The proof of this theorem is also provided in the full version of the paper.

► **Theorem 14.** *For all $\beta \neq 0$, let $\alpha = \alpha(|\beta|)$ be any real in $(0, 1)$ satisfying*

$$|\beta| + \frac{\alpha}{1-\alpha} \ln \alpha + \ln(1-\alpha) > 0. \quad (16)$$

Let $k \geq 4$ be an even integer and let $\beta^ \in \mathbb{R}$ such that $|\beta^*| \geq |\beta|$. Let n be a multiple of $k/2$ that is large enough. If*

$$d \geq k^3 \alpha^{-\frac{k}{2}}, \quad (17)$$

then there exists a formula $\Phi \in \Phi_{n,k,d}$ such that it is impossible to estimate β^ from a sample $\sigma \sim \Pr_{\Phi, \beta^*}$ with high probability.*

Finally, we prove Theorem 2 as a special case of Theorem 14.

Proof. We give an explicit $\alpha : (1, \infty) \rightarrow (0, 1)$ such that $\alpha(|\beta^*|)$ satisfies (16) for any $|\beta^*| > 1$:

$$\alpha(\beta) = 1 - \frac{1}{e^{\beta-1}} = \frac{e^\beta - e}{e^\beta}. \quad (18)$$

Thus, we obtain the explicit lower bound (3) of d by plugging (18) to (17).

Next we verify that the choice of α in (18) satisfies (16). As in the proof of Theorem 14, we define $f(\alpha) := -\frac{\alpha}{1-\alpha} \ln \alpha - \ln(1-\alpha)$. To verify (16), it suffices to show that $f(\alpha(\beta)) < \beta$ for all $\beta > 1$. After some algebraic simplifications, we arrive at

$$\begin{aligned} f(\alpha(\beta)) &= - \left[\frac{e^\beta - e}{e^\beta} \cdot \frac{1}{e^{1-\beta}} \cdot \ln \left(\frac{e^\beta - e}{e^\beta} \right) + \ln \frac{1}{e^{\beta-1}} \right] \\ &= - [(e^{\beta-1} - 1) \cdot (\ln(e^\beta - e) - \beta) + 1 - \beta] \\ &= \ln(e^\beta - e) - 1 + \beta e^{\beta-1} - e^{\beta-1} \ln(e^\beta - e) \\ &= \beta + \ln(1 - e^{1-\beta}) - 1 + \beta e^{\beta-1} - e^{\beta-1} [\beta + \ln(1 - e^{1-\beta})] \\ &= \beta + \ln(1 - e^{1-\beta}) - 1 - e^{\beta-1} \ln(1 - e^{1-\beta}) \\ &= \beta - 1 + (1 - e^{\beta-1}) \ln(1 - e^{1-\beta}) \end{aligned}$$

Notice for all $0 < x < 1$, by Taylor's expansion,

$$\left(1 - \frac{1}{x}\right) \ln(1-x) = -\left(1 - \frac{1}{x}\right) \cdot \sum_{n=1}^{\infty} \frac{x^n}{n} = 1 + \sum_{n=1}^{\infty} \left(\frac{1}{n+1} - \frac{1}{n}\right) x^n = 1 - \sum_{n=1}^{\infty} \frac{x^n}{n(n+1)} < 1.$$

Hence, by setting $x = e^{1-\beta}$, we have shown $f(\alpha(\beta)) < \beta$. ◀

References

- 1 Julian Besag. Spatial interaction and the statistical analysis of lattice systems. *Journal of the Royal Statistical Society. Series B (Methodological)*, 36(2):192–236, 1974.
- 2 Ivona Bezáková, Antonio Blanca, Zongchen Chen, Daniel Štefankovič, and Eric Vigoda. Lower bounds for testing graphical models: colorings and antiferromagnetic Ising models. *J. Mach. Learn. Res.*, 21(1), 2020. URL: <https://jmlr.org/papers/v21/19-580.html>.
- 3 BHASWAR B. BHATTACHARYA and SUMIT MUKHERJEE. Inference in Ising models. *Bernoulli*, 24(1):493–525, 2018.
- 4 Bhaswar B. Bhattacharya and Kavita Ramanan. Parameter estimation for undirected graphical models with hard constraints. *IEEE Transactions on Information Theory*, 67(10):6790–6809, 2021. doi:10.1109/TIT.2021.3094404.
- 5 Antonio Blanca, Zongchen Chen, Daniel Štefankovič, and Eric Vigoda. Hardness of identity testing for restricted boltzmann machines and potts models. *Journal of Machine Learning Research*, 22(152):1–56, 2021. URL: <https://jmlr.org/papers/v22/20-1162.html>.
- 6 Guy Bresler. Efficiently learning Ising models on arbitrary graphs. In *Proceedings of the Forty-Seventh Annual ACM Symposium on Theory of Computing*, STOC '15, pages 771–782, 2015. doi:10.1145/2746539.2746631.
- 7 Guy Bresler, David Gamarnik, and Devavrat Shah. Hardness of parameter estimation in graphical models. In *Proceedings of the 28th International Conference on Neural Information Processing Systems - Volume 1*, NIPS'14, pages 1062–1070, 2014. URL: <https://proceedings.neurips.cc/paper/2014/hash/26dd0dbc6e3f4c8043749885523d6a25-Abstract.html>.
- 8 Guy Bresler, David Gamarnik, and Devavrat Shah. Structure learning of antiferromagnetic Ising models. In *Advances in Neural Information Processing Systems*, volume 27. Curran Associates, Inc., 2014.
- 9 Guy Bresler, David Gamarnik, and Devavrat Shah. Learning graphical models from the glaufer dynamics. *IEEE Transactions on Information Theory*, 64(6):4072–4080, 2018. doi:10.1109/TIT.2017.2713828.
- 10 Guy Bresler and Dheeraj Nagaraj. Optimal single sample tests for structured versus unstructured network data. In *Proceedings of the 31st Conference On Learning Theory*, volume 75 of *Proceedings of Machine Learning Research*, pages 1657–1690, 2018. URL: <http://proceedings.mlr.press/v75/bresler18a.html>.
- 11 Sourav Chatterjee. Estimation in spin glasses: A first step. *The Annals of Statistics*, 35(5):1931–1946, 2007.
- 12 C. Chow and C. Liu. Approximating discrete probability distributions with dependence trees. *IEEE Transactions on Information Theory*, 14(3):462–467, 1968. doi:10.1109/TIT.1968.1054142.
- 13 Francis Comets. On consistency of a class of estimators for exponential families of markov random fields on the lattice. *The Annals of Statistics*, 20(1):455–468, 1992.
- 14 Francis Comets and Basilis Gidas. Asymptotics of Maximum Likelihood Estimators for the Curie-Weiss Model. *The Annals of Statistics*, 19(2):557–578, 1991. doi:10.1214/aos/1176348111.
- 15 Yuval Dagan, Constantinos Daskalakis, Nishanth Dikkala, and Anthimos Vardis Kandiros. Learning Ising models from one or multiple samples. In *Proceedings of the 53rd Annual ACM SIGACT Symposium on Theory of Computing*, STOC 2021, pages 161–168, 2021. doi:10.1145/3406325.3451074.
- 16 Constantinos Daskalakis, Nishanth Dikkala, and Gautam Kamath. Testing Ising models. In *Proceedings of the Twenty-Ninth Annual ACM-SIAM Symposium on Discrete Algorithms*, SODA '18, pages 1989–2007, 2018. doi:10.1137/1.9781611975031.130.
- 17 Constantinos Daskalakis, Nishanth Dikkala, and Ioannis Panageas. Regression from dependent observations. In *Proceedings of the 51st Annual ACM SIGACT Symposium on Theory of Computing*, STOC 2019, pages 881–889, 2019. doi:10.1145/3313276.3316362.

- 18 Constantinos Daskalakis, Nishanth Dikkala, and Ioannis Panageas. Logistic regression with peer-group effects via inference in higher-order ising models. In *AISTATS*, pages 3653–3663, 2020. URL: <http://proceedings.mlr.press/v108/daskalakis20a.html>.
- 19 Weiming Feng, Heng Guo, Yitong Yin, and Chihao Zhang. Fast sampling and counting k -sat solutions in the local lemma regime. *J. ACM*, 68(6), October 2021. doi:10.1145/3469832.
- 20 Jason Gaitonde, Ankur Moitra, and Elchanan Mossel. Bypassing the noisy parity barrier: Learning higher-order markov random fields from dynamics, 2024. arXiv:2409.05284.
- 21 Jason Gaitonde and Elchanan Mossel. A unified approach to learning Ising models: Beyond independence and bounded width. In *Proceedings of the 56th Annual ACM Symposium on Theory of Computing*, STOC 2024, pages 503–514. Association for Computing Machinery, 2024. doi:10.1145/3618260.3649674.
- 22 Andreas Galanis, Leslie Ann Goldberg, and Xusheng Zhang. One-shot learning for k -sat, 2025. doi:10.48550/arXiv.2502.07135.
- 23 Andreas Galanis, Alkis Kalavasis, and Anthimos Vardis Kandiros. Learning hard-constrained models with one sample. *Proceedings of the 2024 Annual ACM-SIAM Symposium on Discrete Algorithms (SODA)*, pages 3184–3196, 2024.
- 24 Charles J. Geyer and Elizabeth A. Thompson. Constrained Monte Carlo maximum likelihood for dependent data. *Journal of the Royal Statistical Society. Series B (Methodological)*, 54(3):657–699, 1992.
- 25 Promit Ghosal and Sumit Mukherjee. Joint estimation of parameters in Ising model. *The Annals of Statistics*, 48(2):785–810, 2020.
- 26 B. Gidas. Consistency of maximum likelihood and pseudo-likelihood estimators for gibbs distributions. In Wendell Fleming and Pierre-Louis Lions, editors, *Stochastic Differential Systems, Stochastic Control Theory and Applications*, pages 129–145, 1988.
- 27 Heng Guo, Mark Jerrum, and Jingcheng Liu. Uniform sampling through the lovász local lemma. *J. ACM*, 66(3), April 2019. doi:10.1145/3310131.
- 28 Linus Hamilton, Frederic Koehler, and Ankur Moitra. Information theoretic properties of Markov random fields, and their algorithmic applications. In *Advances in Neural Information Processing Systems*, volume 30, 2017.
- 29 Kun He, Chunyang Wang, and Yitong Yin. Deterministic counting lovász local lemma beyond linear programming. In *Proceedings of the 2023 Annual ACM-SIAM Symposium on Discrete Algorithms (SODA)*, pages 3388–3425, 2023. doi:10.1137/1.9781611977554.CH130.
- 30 Vishesh Jain, Huy Tuan Pham, and Thuy Duong Vuong. Towards the sampling lovász local lemma. In *2021 IEEE 62nd Annual Symposium on Foundations of Computer Science (FOCS)*, pages 173–183, 2022. doi:10.1109/FOCS52979.2021.00025.
- 31 Adam Klivans and Raghu Meka. Learning graphical models using multiplicative weights. In *2017 IEEE 58th Annual Symposium on Foundations of Computer Science (FOCS)*, pages 343–354, 2017. doi:10.1109/FOCS.2017.39.
- 32 William Matthews and Ramamohan Paturi. Uniquely satisfiable k -sat instances with almost minimal occurrences of each variable. In Ofer Strichman and Stefan Szeider, editors, *Theory and Applications of Satisfiability Testing – SAT 2010*, pages 369–374, 2010. doi:10.1007/978-3-642-14186-7_34.
- 33 Ankur Moitra. Approximate counting, the lovász local lemma, and inference in graphical models. *J. ACM*, 66(2), April 2019. doi:10.1145/3268930.
- 34 Rajarshi Mukherjee, Sumit Mukherjee, and Ming Yuan. Global testing against sparse alternatives under Ising models. *The Annals of Statistics*, 46(5):2062–2093, 2018.
- 35 Somabha Mukherjee, Jaesung Son, and Bhaswar B Bhattacharya. Estimation in tensor Ising models. *Information and Inference: A Journal of the IMA*, 11(4):1457–1500, 2022.
- 36 Narayana P. Santhanam and Martin J. Wainwright. Information-theoretic limits of selecting binary graphical models in high dimensions. *IEEE Transactions on Information Theory*, 58(7):4117–4134, 2012. doi:10.1109/TIT.2012.2191659.

- 37 Marc Vuffray, Sidhant Misra, Andrey Y. Lokhov, and Michael Chertkov. Interaction screening: efficient and sample-optimal learning of ising models. In *Proceedings of the 30th International Conference on Neural Information Processing Systems*, NIPS'16, pages 2603–2611, 2016.
- 38 Chunyang Wang and Yitong Yin. A sampling Lovász local lemma for large domain sizes. In *2024 IEEE 65th Annual Symposium on Foundations of Computer Science (FOCS)*, pages 129–150, 2024. doi:10.1109/FOCS61266.2024.00019.
- 39 Huanyu Zhang, Gautam Kamath, Janardhan Kulkarni, and Zhiwei Steven Wu. Privately learning Markov random fields. In *Proceedings of the 37th International Conference on Machine Learning*, ICML'20, 2020.