

Authors are encouraged to submit new papers to INFORMS journals by means of a style file template, which includes the journal title. However, use of a template does not certify that the paper has been accepted for publication in the named journal. INFORMS journal templates are for the exclusive purpose of submitting to an INFORMS journal and should not be used to distribute the papers in print or online or to submit the papers to another publication.

Hidden Integrality and Semi-random Robustness of SDP Relaxation for Sub-Gaussian Mixture Model

Yingjie Fei

School of Operations Research and Information Engineering, Cornell University, yf275@cornell.edu

Yudong Chen

Department of Computer Sciences, University of Wisconsin-Madison, yudong.chen@wisc.edu

We consider the problem of estimating the discrete clustering structures under the Sub-Gaussian Mixture Model. Our main results establish a *hidden integrality* property of a semidefinite programming (SDP) relaxation for this problem: while the optimal solution to the SDP is not integer-valued in general, its estimation error can be upper bounded by that of an idealized integer program. The error of the integer program, and hence that of the SDP, are further shown to decay *exponentially* in the signal-to-noise ratio. In addition, we show that the SDP relaxation is robust under the semi-random setting in which an adversary can modify the data generated from the mixture model. In particular, we generalize the hidden integrality property to the semi-random model and thereby show that SDP achieves the optimal error bound in this setting. These results together highlight the “global-to-local” mechanism that drives the performance of the SDP relaxation.

To the best of our knowledge, our result is the first exponentially decaying error bound for convex relaxations of mixture models. A corollary of our results shows that in certain regimes the SDP solutions are in fact integral and exact. More generally, our results establish sufficient conditions for the SDP to correctly recover the cluster memberships of $(1 - \delta)$ fraction of the points for any $\delta \in (0, 1)$. As a special case, we show that under the d -dimensional Stochastic Ball Model, SDP achieves non-trivial (sometimes exact) recovery when the center separation is as small as $\sqrt{1/d}$, which improves upon previous exact recovery results that require constant separation.

Key words: Exponential rates, hidden integrality, mixture models, SDP relaxation, semi-random robustness, stochastic ball model

MSC2000 subject classification: 62H30

OR/MS subject classification: Primary: Statistics; secondary: Cluster analysis

History:

1. Introduction We consider the Sub-Gaussian Mixture Model (SGMM), in which one is given n random points drawn from a mixture of k sub-Gaussian distributions with different means/centers. SGMM, particularly its special case the Gaussian Mixture Model (GMM), is widely used in a broad range of applications including speaker identification, background modeling and online recommendation. In these applications, one is typically interested in two types of statistical inference problems under SGMM:

- **Clustering:** (approximately) identify the cluster membership of each point, that is, which of the k mixture components generates a given point;

- **Parameter estimation:** estimate the parameters (e.g., means/centers) of the k components, or the density of the entire mixture.

Standard approaches to these problems, such as k-means clustering, typically lead to integer programming formulations that are non-convex and NP-hard to optimize [5, 30, 43]. Consequently, much work has been devoted to developing computationally tractable algorithms for SGMM; examples include expectation maximization [19], Lloyd’s algorithm [41], spectral methods [58], the method of moments [53], and many more. Among them, the convex relaxation methods, including those based on linear programming (LP) and semidefinite programming (SDP), have emerged as a promising approach for SGMM. This approach has several attractive properties: (a) it is solvable in polynomial time, and does not require a good initial solution to be provided; (b) it has the flexibility to incorporate different quality metrics and additional constraints; (c) it is not restricted to specific forms of SGMM (such as Gaussian distribution), and is robust against model misspecification [55, 54, 52]; (d) it can provide a certificate for optimality [29].

Theoretical performance guarantees for convex relaxation methods have been investigated in a body of both old and recent work. As will be discussed in greater details in the related work section (Section 2), these existing results often come in one of two forms:

1. How well the (rounded) solution of a relaxation optimizes a particular objective function (e.g., the k-means or k-medians objective) compared to the original integer program, as studied in classical work on *approximation factors* [13, 33, 54, 38];
2. When the solution of a relaxation exactly recovers the ground-truth clustering, a phenomenon known as *exact recovery* and studied in a more recent line of work [52, 8, 47, 29, 40].

In many applications, optimizing a particular objective function, and designing approximation algorithms for doing so, are often only a means to an end, whereas the ultimate goal is to solve the two statistical inference problems above, both of which involve learning the true underlying model that generates the observed data. Results on exact recovery are more directly relevant to this goal. However, such results often require very stringent conditions on the separation or signal-to-noise ratio (SNR) of the model. In practice, convex relaxation solutions are rarely exact, even when the data are generated from the assumed model. On the other hand, researchers have observed that the solutions, while not exact or integer-valued, are often a good approximation to the desired solution that represents the true model [47]. Such a phenomenon is not captured by the results on exact recovery.

In this paper, we aim to strengthen our understanding of the convex relaxation approach for clustering SGMM. In particular, we study the regime where solutions of convex relaxations are not exact, and directly characterize the *estimation errors* of the solutions—namely, their distance to desired (integer) solution corresponding to the true underlying model.

1.1. Our Contributions For a class of SDP relaxations for SGMM, our results reveal a perhaps surprising property thereof: while the SDP solutions are not integral in general, their estimation errors can be controlled by that of the solutions of an idealized integer program (IP), in which one tries to estimate cluster memberships when an oracle reveals the *true centers* of the SGMM. We refer to the latter program as the *Oracle Integer Program*. In particular, we show that, in a precise sense to be formalized later, the estimation errors of the SDP and Oracle IP satisfy the relationship (Theorem 1)

$$\text{error}(\text{SDP}) \lesssim \text{error}(\text{IP}) \quad (1)$$

under certain conditions. We refer to this property as *hidden integrality* of the SDP relaxation; its proof in fact involves showing that certain intermediate linear optimization

problems are integral. We then further upper bound the error of the Oracle IP and show that it decays *exponentially* in terms of the SNR (Theorem 2):

$$\text{error(IP)} \lesssim \exp[-\Omega(\text{SNR}^2)], \quad (2)$$

where the SNR is defined as the ratio of the center separation and the standard deviation of the mixture components. Combining these two results immediately leads to explicit bounds on the error of the SDP solutions (Corollary 1).

1.1.1. Robustness under Semi-random Model Our results can be generalized to a so-called *semi-random* version of SGMM. In this setting, an adversary is allowed to modify the data points generated from SGMM in an arbitrary and potentially adversarial way (subject to certain monotonicity constraints). This semi-random setting captures unpredictable deviations from the nominal SGMM—which is common in real data—and it is well recognized to be much more challenging than the original, purely random model. In fact, many existing algorithms provably fail in the semi-random setting [36, 10].

We show that SDP relaxation has an inherent robustness property under the semi-random model. In particular, we establish the following error bounds:

$$\text{error(SDP)} \lesssim \text{error(IP)} + \varepsilon \quad \text{and} \quad \text{error(IP)} \lesssim \exp[-\Omega(\text{SNR}^2)], \quad (3)$$

where ε denotes the additional error induced by the adversary; see Theorems 4 and 5 for the precise statement of this result and the expression for ε . In certain regimes, the error ε is dominated by $\exp[-\Omega(\text{SNR}^2)]$, the error of the Oracle IP, hence the error of the SDP is (order-wise) unaffected under the semi-random model. In other regimes, the additional error ε can be shown to be fundamentally unavoidable. Note that the adversary only affects the error of the SDP relative to the Oracle IP, but not the error of the Oracle IP itself.

1.1.2. Consequences When the SNR is sufficiently large, the above results imply that the SDP solutions are integral and exact up to numerical errors, hence recovering existing results on exact recovery as a special case. If the SNR is low and the SDP solutions are fractional, one may obtain an explicit clustering from the SDP solutions via a simple, optimization-free rounding procedure. We show that the error of this explicit clustering (in terms of the fraction of points misclassified) is also bounded by the error of the Oracle IP and hence also decays exponentially in the SNR (Theorem 3). As a consequence, we obtain sufficient conditions for misclassifying at most δ fraction of the points for any given $\delta \in [0, 1]$.

Significantly, our results often match and sometimes improve upon state-of-the-art performance guarantees in settings for which known results exist, and lead to new guarantees in other less studied settings of SGMM. For instance, a corollary of our results shows that under the Stochastic Ball Model, SDP achieves non-trivial (sometimes exact) cluster recovery even when the center separation Δ is as small as $O(\sqrt{1/d})$, where d is the dimension (Section 4.4). For high dimensional settings, this bound goes beyond existing results that only consider exact recovery and require $\Delta = \Omega(1)$. We discuss the implications of our results in details and compare with existing ones after we state the main theorems.

1.1.3. The “Global-to-Local” Phenomenon Our results above are obtained in two steps: (a) relating the SDP to the Oracle IP, and (b) bounding the Oracle IP errors. This two-step approach allows us to decouple the two types of mechanisms that determine the performance of the SDP relaxation:

- On the one hand, step (a) is done by leveraging the structures of the *entire* dataset with n points. In particular, certain global spectral properties of the data ensure that the error of the SDP is non-trivially bounded (in terms of the Oracle IP error). This step is relatively insensitive to the specific structure of the SGMM.
- On the other hand, as shall become clear in the sequel, the Oracle IP essentially reduces to n independent clustering problems, one for each data point. Knowing the true cluster centers, the Oracle IP is optimal in terms of the clustering errors. No other algorithms (including SDP relaxations) can achieve a strictly better error, due to the inherent randomness of *individual* data points. Step (b) above hence captures the local mechanism that determines fine-grained error rates as a function of the SNR.

Our analysis establishes the hidden integrality property as the bridge between these two mechanisms. Our result hence highlights the power of the SDP approach: it is able to capture the global and local structures of the data simultaneously, without requiring a good initial solution or sophisticated pre-processing/post-processing steps.

In the context of the semi-random model, the error bound (3) shows that the effect of the adversary is restricted to the global regime in step (a) and does not play a role in the local step (b). We emphasize that clustering under semi-random models is a highly challenging problem by its own, and an entire line of work is devoted to this problem [16, 25, 36, 10, 44]. Our two-step, hidden integrality-based approach allows for a streamlined analysis of this setting, as it isolates precisely how the adversary can impact the clustering error of the SDP.

1.2. Paper Organization The remainder of the paper is organized as follows. In Section 2, we discuss related work on SGMM and its special cases. In Section 3, we describe the problem setup for SGMM and the SDP relaxation approach. In Section 4, we present our main theoretical results, discuss their consequences and compare with existing results. We prove our main theorems in Sections 5 and 6. The paper is concluded with a discussion of future directions in Section 7. The proofs of some technical results are deferred to the Appendix.

2. Related Work The study of SGMM has a long history and is still an active area of research. Here we review the most relevant work on algorithms for SGMM with theoretical performance guarantees.

The work [17] is among the first to obtain performance guarantees for GMM. Subsequent work has developed a variety of algorithms including spectral methods, which achieved improved guarantees. These results establish sufficient conditions, in terms of the separation between the cluster centers (or equivalently the SNR), for achieving (near)-exact recovery of the cluster memberships. One of the best results is obtained in [58] and requires $\text{SNR} \gtrsim (k \ln n)^{1/4}$, which is later extended in a long line of work including [4, 37, 9]. We compare these results with ours in Section 4.

Expectation-Maximization (EM) and the closely related Lloyd’s algorithm are two of the most popular methods for GMM. Despite their empirical effectiveness, non-asymptotic statistical guarantees are established only recently; see the work in [11, 34, 61, 42]. All these results assume that one has access to a sufficiently good initial solution, typically obtained by spectral methods. Although some recent work shows that EM converges from random initialization under certain restricted settings of GMM [18, 60], for general settings EM is known to suffer from the existence of bad local minima [31]. Robustness of the Lloyd’s algorithm under a semi-random GMM is studied in the work [10].

Most relevant to us is work on convex relaxation methods for GMM and k-means/median problems. A class of SDP relaxations, often called the *Peng-Wei relaxations*, are developed in the seminal work in [55, 54]. Thanks to convexity, these methods

do not suffer from the issues of bad local minima faced by EM/Lloyd’s, though it is far from trivial how to round their (typically fractional) solutions into a valid clustering solution with provable quality guarantees. In this direction, the work in [8, 29, 40] establish sufficient conditions for LP/SDP relaxations to achieve exact recovery, and the work in [47] proves approximate recovery error bounds for SDP. These results are directly comparable to ours. We discuss them in more details in Section 4 after presenting our main theorems.

In the last decade, the related Stochastic Block Model (SBM) has also witnessed significant progress on understanding convex relaxation methods; see [1, 39] for a comprehensive survey. In particular, much work has been done on exact recovery guarantees for SDP relaxations of SBM [36, 6, 15, 7]. A more recent line of work establishes approximate recovery guarantees of the SDPs [27, 50] in the low-SNR regime. Particularly relevant to us is the work in [24, 22, 23], who also establish exponentially decaying error bounds. Interestingly, these results also highlight a so-called “local-to-global amplification” phenomenon [1, 3, 2], which is related to what we discussed in Section 1.1.3. Despite the apparent similarity, however, our results on SGMM require very different analytical techniques, due to the fundamental difference between the geometric and probabilistic structures of SGMM and SBM. Moreover, our results establish a general hidden integrality property of SDP relaxations, which we believe holds more broadly beyond specific models like SBM and SGMM.

3. Models and Algorithms In this section, we formally set up the clustering problem under SGMM and describe our SDP relaxation approach.

3.1. Notations Vectors and matrices are denoted by bold letters such as $\mathbf{u}$ and $\mathbf{M}$. For a vector $\mathbf{u}$, we denote by u_i its i -th entry. For a matrix $\mathbf{M}$, we use $\text{Tr}(\mathbf{M})$ to denote its trace, M_{ij} its (i, j) -th entry, $\text{diag}(\mathbf{M})$ the vector of its diagonal entries, $\|\mathbf{M}\|_1 := \sum_{i,j} M_{ij}$ its entry-wise ℓ_1 norm, $\|\mathbf{M}\|_{\text{op}}$ its spectral norm (maximum singular value), $\mathbf{M}_{i\bullet}$ its i -th row and $\mathbf{M}_{\bullet j}$ its j -th column. We write $\mathbf{M} \succeq 0$ if $\mathbf{M}$ is symmetric positive semidefinite (psd). The trace inner product between two matrices $\mathbf{M}$ and $\mathbf{Q}$ of the same dimension is denoted by $\langle \mathbf{M}, \mathbf{Q} \rangle := \text{Tr}(\mathbf{M}^\top \mathbf{Q})$. For a number a , $\mathbf{M} \geq a$ means that $M_{ij} \geq a$ for all entries (i, j) . We denote by $\mathbf{1}_m$ the all-one column vector of dimension m . For a positive integer i , let $[i] := \{1, 2, \dots, i\}$. For two non-negative sequences $\{a_i\}$ and $\{b_i\}$, we write $a_i \lesssim b_i$ if there exists a universal constant $C > 0$ such that $a_i \leq Cb_i$ for all i , and write $a_i \asymp b_i$ if $a_i \lesssim b_i$ and $b_i \lesssim a_i$.

We recall that the sub-Gaussian norm of a random variable X is defined as

$$\|X\|_{\psi_2} := \inf \{t > 0 : \mathbb{E} \exp(X^2/t^2) \leq 2\},$$

and X is called sub-Gaussian if $\|X\|_{\psi_2} < \infty$. Note that Gaussian and bounded random variables are sub-Gaussian. Denote by $\mathbb{S}^{m-1}$ the m -dimensional unit ℓ_2 sphere. A random vector $\mathbf{x} \in \mathbb{R}^m$ is sub-Gaussian if for all fixed vectors $\mathbf{u} \in \mathbb{R}^m$, the one dimensional marginal $\langle \mathbf{x}, \mathbf{u} \rangle$ is a sub-Gaussian random variable. The sub-Gaussian norm of $\mathbf{x}$ is defined as $\|\mathbf{x}\|_{\psi_2} := \sup_{\mathbf{u} \in \mathbb{S}^{m-1}} \|\langle \mathbf{x}, \mathbf{u} \rangle\|_{\psi_2}$. With this notation, a random vector $\mathbf{x} \sim \mathcal{N}(\mathbf{0}, \Sigma)$ from the multivariate Gaussian distribution is sub-Gaussian with $\|\mathbf{x}\|_{\psi_2} = \|\Sigma\|_{\text{op}}$.

3.2. Sub-Gaussian Mixture Model We focus on the Sub-Gaussian Mixture Model (SGMM) with balanced clusters.

MODEL 1 (Sub-Gaussian Mixture Model). Let $\boldsymbol{\mu}_1, \dots, \boldsymbol{\mu}_k \in \mathbb{R}^d$ be k unknown cluster centers. We observe n random points in $\mathbb{R}^d$ of the form

$$\mathbf{h}_i := \boldsymbol{\mu}_{\sigma^*(i)} + \mathbf{g}_i, \quad i \in [n]$$

where $\sigma^*(i) \in [k]$ is the unknown cluster label of the i -th point, and $\{\mathbf{g}_i\}$ are zero-mean independent sub-Gaussian random vectors with sub-Gaussian norms $\|\mathbf{g}_i\|_{\psi_2} \leq \tau$.¹ We assume that the ground-truth clusters have equal sizes, that is, $|\{i \in [n] : \sigma^*(i) = a\}| = \frac{n}{k}$ for each $a \in [k]$.

Let $\boldsymbol{\sigma}^* \in [k]^n$ be the vector of the true cluster labels; that is, its i -th coordinate is $\sigma_i^* \equiv \sigma^*(i)$ (we use these two notations interchangeably in the paper.) The task is to estimate the underlying clustering $\boldsymbol{\sigma}^*$ given the observed data $\{\mathbf{h}_i : i \in [n]\}$.

Note that in Model 1 we do not require $\{\mathbf{g}_i\}$ to be isotropic or identically distributed. This model includes several important mixture models as special cases:

- Spherical GMM, where $\{\mathbf{g}_i\}$ are Gaussian with the covariance matrix $\tau^2 \mathbf{I}$.
- More general GMM where the k clusters have non-identical and non-diagonal covariance matrices $\{\boldsymbol{\Sigma}_a\}_{a \in [k]}$.
- The Stochastic Ball Model [52], where $\{\mathbf{g}_i\}$ are bounded and rotationally invariant random variables; we discuss this model in details in Section 4.4

Throughout the paper we assume $n \geq 4$ and $k \geq 2$ to avoid degeneracy. Denote by $\Delta_{ab} := \|\boldsymbol{\mu}_a - \boldsymbol{\mu}_b\|_2$ the separation of the centers of clusters a and b , and $\Delta := \min_{a \neq b \in [k]} \|\boldsymbol{\mu}_a - \boldsymbol{\mu}_b\|_2$ the minimum separation of the centers. Playing a crucial role in our results is the quantity

$$s := \frac{\Delta}{\tau}, \quad (4)$$

which is a measure of the signal-to-noise ratio of the SGMM.

3.3. Semidefinite Programming Relaxation We now describe our SDP relaxation for clustering SGMM. To begin, note that each candidate clustering of n points into k clusters can be represented by an *assignment matrix* $\mathbf{F} \in \{0, 1\}^{n \times k}$ with

$$F_{ia} = \begin{cases} 1 & \text{if point } i \text{ is assigned to cluster } a, \\ 0 & \text{otherwise.} \end{cases}$$

Let $\mathcal{F} := \{\mathbf{F} \in \{0, 1\}^{n \times k} : \mathbf{F}\mathbf{1}_k = \mathbf{1}_n\}$ be the set of all possible assignment matrices. To cluster the points $\{\mathbf{h}_i\}$, a natural approach is to find an assignment $\mathbf{F}$ that minimizes the total within-cluster pairwise distance. Arranging the pairwise squared distance as a matrix $\mathbf{A} \in \mathbb{R}^{n \times n}$ with

$$A_{ij} = \|\mathbf{h}_i - \mathbf{h}_j\|_2^2, \quad \text{for each } (i, j) \in [n] \times [n],$$

we can express the above objective as

$$\sum_{i,j} A_{ij} \mathbb{I}\{\text{points } i \text{ and } j \text{ are assigned to the same cluster}\} = \sum_{i,j} A_{ij} (\mathbf{F}\mathbf{F}^\top)_{ij} = \langle \mathbf{F}\mathbf{F}^\top, \mathbf{A} \rangle.$$

Therefore, the approach described above can be formulated as the integer program (5) below:

¹ More explicitly, the sub-Gaussian assumption is equivalent to $\mathbb{E} \exp(\langle \mathbf{g}_i, \mathbf{v} \rangle) \leq \exp(\tau^2 \|\mathbf{v}\|_2^2 / 2), \forall \mathbf{v} \in \mathbb{R}^d$.

$$\begin{aligned}
 \min_{\mathbf{F}} \langle \mathbf{F}\mathbf{F}^\top, \mathbf{A} \rangle & \quad \min_{\mathbf{Y}} \langle \mathbf{Y}, \mathbf{A} \rangle \\
 \text{s.t. } \mathbf{F} \in \mathcal{F}, & \quad \text{s.t. } \mathbf{Y}\mathbf{1}_n = \frac{n}{k}\mathbf{1}_n, \\
 \mathbf{1}_n^\top \mathbf{F} = \frac{n}{k}\mathbf{1}_k^\top. & \quad \text{diag}(\mathbf{Y}) = \mathbf{1}_n, \\
 & \quad \mathbf{Y} \succeq 0, \\
 & \quad \mathbf{Y} \in \{0, 1\}^{n \times n}, \text{rank}(\mathbf{Y}) = k.
 \end{aligned} \tag{5} \tag{6}$$

In program (5) the additional constraint $\mathbf{1}_n^\top \mathbf{F} = \frac{n}{k}\mathbf{1}_k^\top$ enforces that all k clusters have the same size $\frac{n}{k}$, as we are working with SGMM whose true clusters are balanced. Under this balanced model, it is not hard to see that the program (5) is equivalent to the classical k-means formulation. With a change of variable $\mathbf{Y} = \mathbf{F}\mathbf{F}^\top$, we may lift the program (5) to the space of $n \times n$ matrices and obtain the equivalent formulation (6).

Both programs (5) and (6) involve non-convex combinatorial constraints and are computationally hard to solve. To obtain a tractable formulation, we drop the non-convex rank constraint in (6) and replace the integer constraint with the box constraint $0 \leq \mathbf{Y} \leq 1$ (the constraint $\mathbf{Y} \leq 1$ is in fact redundant). Doing so leads to the following SDP relaxation:

$$\begin{aligned}
 \hat{\mathbf{Y}} \in \arg \min_{\mathbf{Y} \in \mathbb{R}^{n \times n}} \langle \mathbf{Y}, \mathbf{A} \rangle \\
 \text{s.t. } \mathbf{Y}\mathbf{1}_n = \frac{n}{k}\mathbf{1}_n, \\
 \text{diag}(\mathbf{Y}) = \mathbf{1}_n, \\
 \mathbf{Y} \succeq 0, \\
 \mathbf{Y} \leq 1.
 \end{aligned} \tag{7}$$

Let $\mathbf{F}^* \in \mathcal{F}$ be the assignment matrix associated with the true underlying clustering of the SGMM; that is, $F_{ja}^* = \mathbb{I}\{\sigma^*(j) = a\}$ for each $j \in [n]$ and $a \in [k]$. The performance of the SDP is measured against the true *cluster matrix* $\mathbf{Y}^* := \mathbf{F}^*(\mathbf{F}^*)^\top \in \{0, 1\}^{n \times n}$, which has the explicit form

$$Y_{ij}^* = \begin{cases} 1 & \text{if } \sigma^*(i) = \sigma^*(j), \text{ i.e., points } i \text{ and } j \text{ are in the same cluster,} \\ 0 & \text{if } \sigma^*(i) \neq \sigma^*(j), \text{ i.e., points } i \text{ and } j \text{ are in different clusters,} \end{cases}$$

with the convention that $Y_{ii}^* = 1, \forall i \in [n]$. The matrix $\mathbf{Y}^*$ hence encodes the ground-truth clustering labels σ^* . If we order the rows and columns of $\mathbf{Y}^*$ according to σ^* , then $\mathbf{Y}^*$ takes the block-diagonal form

$$\mathbf{Y}^* = \begin{pmatrix} \mathbf{J}_{n/k} & & \\ & \ddots & \\ & & \mathbf{J}_{n/k} \end{pmatrix}, \tag{8}$$

where $\mathbf{J}_{n/k}$ denotes the $\frac{n}{k}$ -by- $\frac{n}{k}$ all-one matrix. From Equation (8) it is clear that $\mathbf{Y}^*$ is feasible to the SDP (7). We view an optimal solution $\hat{\mathbf{Y}}$ to (7) as an estimate of the true clustering $\mathbf{Y}^*$. Our goal is to characterize the estimation error $\|\hat{\mathbf{Y}} - \mathbf{Y}^*\|_1$ in terms of the number of points n , the number of clusters k , the data dimension d and the SNR s defined in (4). Note that here we measure the error of $\hat{\mathbf{Y}}$ in the ℓ_1 metric. As we shall see later, this metric is directly related to the clustering error, i.e., the fraction of misclassified points.

We remark that the SDP (7) is somewhat different from the more classical SDP relaxation of k-means proposed in [54], as the latter involves normalizing the variable by the cluster sizes. The SDP (7) has been studied under SGMM in [40], and is closely related to the SDP formulation considered in [7] for the Stochastic Block Model.

3.4. Explicit Clustering In general the SDP solution $\hat{\mathbf{Y}}$ is not in the block-diagonal form of (8) and hence does not directly correspond to an explicit clustering. Nevertheless, we may easily extract cluster memberships from $\hat{\mathbf{Y}}$ using a simple greedy procedure that runs in time linear in the size of $\hat{\mathbf{Y}}$.

The procedure consists of two steps. In the first step, as given in Algorithm 1, we treat the rows of $\hat{\mathbf{Y}}$ as points in $\mathbb{R}^n$, and consider the ℓ_1 balls centered at each row with a certain radius. The ball that contains the most rows is identified, and the indices of the rows in this ball are output as a set. The process continues iteratively with the remaining rows of $\hat{\mathbf{Y}}$. This procedure outputs a collection of index sets whose sizes are no larger than $\frac{n}{k}$ but may not equal to each other.

Algorithm 1 Initial Grouping

Input: data matrix $\hat{\mathbf{Y}} \in \mathbb{R}^{n \times n}$, number of points n , target number of clusters k .

1. Initialize $B_0 \leftarrow \emptyset$, $t \leftarrow 0$, $V \leftarrow [n]$.
 2. While $V \setminus \bigcup_{i=0}^t B_i \neq \emptyset$:
 - (a) Set $t \leftarrow t + 1$ and $V_t \leftarrow V \setminus \bigcup_{i=0}^{t-1} B_i$.
 - (b) For each $u \in V_t$, set $B(u) \leftarrow \left\{ w \in V_t : \|\hat{\mathbf{Y}}_{u\bullet} - \hat{\mathbf{Y}}_{w\bullet}\|_1 \leq \frac{n}{4k} \right\}$.
 - (c) Set $B_t \leftarrow \arg \max_{B(u): u \in V_t} |B(u)|$.
 - (d) If $|B_t| > \frac{n}{k}$, then remove arbitrary elements in B_t so that $|B_t| = \frac{n}{k}$.
- Output: sets $\{B_t\}_{t \geq 1}$.
-

In the second step, we convert the sets output above into k equal-size clusters. This is done by identifying the k largest sets among them, and distributing points in the remaining sets across the chosen k sets so that each of these sets contains exactly $\frac{n}{k}$ points. Combining these two steps gives our final algorithm, **cluster**, for extracting an explicit clustering from the SDP solution $\hat{\mathbf{Y}}$. This procedure is given as Algorithm 2.

Algorithm 2 cluster

Input: data matrix $\hat{\mathbf{Y}} \in \mathbb{R}^{n \times n}$, number of points n , target number of clusters k .

1. Run Algorithm 1 with $\hat{\mathbf{Y}}$, n and k as input and obtain the sets $\{B_t\}_{t \geq 1}$.
 2. Choose the k largest sets among $\{B_t\}_{t \geq 1}$; call the chosen sets $\{U_t\}_{t \in [k]}$.
 3. Arbitrarily distribute elements of $\{B_t\}_{t \geq 1} \setminus \{U_t\}_{t \in [k]}$ among $\{U_t\}_{t \in [k]}$ so that each U_t has exactly $\frac{n}{k}$ elements.
 4. For each $i \in [n]$: set $\hat{\sigma}_i \leftarrow t$, where t is the unique index in $[k]$ such that $U_t \ni i$.
- Output: clustering assignment $\hat{\sigma} \in [k]^n$.
-

The output of the above procedure

$$\hat{\sigma} := \text{cluster}(\hat{\mathbf{Y}}, n, k)$$

is a vector in $[k]^n$ such that point i is assigned to the $\hat{\sigma}_i$ -th cluster. We are interested in controlling the clustering error of $\hat{\sigma}$ relative to the ground-truth clustering σ^* . Let $\mathcal{S}_k$ denote the symmetric group consisting of all permutations of $[k]$. The clustering error is defined by

$$\text{err}(\hat{\sigma}, \sigma^*) := \min_{\pi \in \mathcal{S}_k} \frac{1}{n} |\{i \in [n] : \hat{\sigma}_i \neq \pi(\sigma_i^*)\}|, \quad (9)$$

which is the proportion of points that are misclassified, modulo permutation of the cluster labels.

Variants of the above `cluster` procedure have been considered before in [45, 47]. Our results in the next section establish that the clustering error $\text{err}(\hat{\boldsymbol{\sigma}}, \boldsymbol{\sigma}^*)$ is always upper bounded by the ℓ_1 error $\|\hat{\mathbf{Y}} - \mathbf{Y}^*\|_1$ of the SDP solution.

4. Main Results In this section, we establish the connection between the estimation error of the SDP relaxation (7) and that of what we call the Oracle Integer Program. Using this connection, we derive explicit bounds on the error of the SDP, and explore the implications of these results under different settings.

4.1. Oracle Integer Program Consider an idealized setting where an oracle reveals the true cluster centers $\{\boldsymbol{\mu}_a\}_{a \in [k]}$. Moreover, we are given the data points $\{\bar{\mathbf{h}}_i\}_{i \in [n]}$, where $\bar{\mathbf{h}}_i := \boldsymbol{\mu}_{\sigma^*(i)} + (2\kappa)^{-1}\mathbf{g}_i$ with $\kappa := 1/8$ and $\{\mathbf{g}_i\}$ are the same realizations of the random variables in the original SGMM. In other words, $\{\bar{\mathbf{h}}_i\}$ are the same as the original data points $\{\mathbf{h}_i\}$, except that the sub-Gaussian norm of the noise $\{\mathbf{g}_i\}$ is scaled by $(2\kappa)^{-1} = 4$.

To cluster the points $\{\bar{\mathbf{h}}_i\}$ in this setting, a natural approach is to simply assign each point to the closest cluster center, so that the total distance of the points to their assigned centers is minimized. Representing each candidate clustering assignment using an assignment matrix $\mathbf{F} \in \mathcal{F}$ as before, we see that the quantity

$$\eta(\mathbf{F}) := \sum_{j \in [n]} \sum_{a \in [k]} \|\bar{\mathbf{h}}_j - \boldsymbol{\mu}_a\|_2^2 F_{ja} \quad (10)$$

is exactly the sum of the distances of each point to its assigned center. The idealized clustering procedure above thus amounts to solving the following *Oracle Integer Program (IP)*:

$$\min_{\mathbf{F}} \eta(\mathbf{F}), \quad \text{s.t. } \mathbf{F} \in \mathcal{F}. \quad (11)$$

Note that this program is separable across the rows of $\mathbf{F}$ and can be reduced to n independent optimization problems, one for each data point $\bar{\mathbf{h}}_i$.

A priori, there is no obvious connection between the estimation error of a solution to the Oracle IP and that of a solution to the SDP. In particular, the SDP is oblivious to the true centers and in general has fractional optimal solutions. Nevertheless, we are able to establish a formal relationship between these two programs, as is done below.

4.2. Errors of SDP Relaxation and Oracle IP We begin by making the following observations. Recall that $\mathbf{F}^* \in \mathcal{F}$ is the true assignment matrix as defined in Section 3.3. For each assignment matrix $\mathbf{F} \in \mathcal{F}$, it is easy to see that the quantity $\frac{1}{2}\|\mathbf{F} - \mathbf{F}^*\|_1$ equals the number of points that are assigned differently in $\mathbf{F}$ and $\mathbf{F}^*$. Therefore, this quantity measures the clustering error of $\mathbf{F}$. On the other hand, for a solution $\mathbf{F} \in \mathcal{F}$ to potentially be an optimal solution of the Oracle IP (11), it must satisfy $\eta(\mathbf{F}) \leq \eta(\mathbf{F}^*)$ since $\mathbf{F}^*$ is feasible to (11). Consequently, the quantity

$$\max \left\{ \frac{1}{2} \|\mathbf{F} - \mathbf{F}^*\|_1 : \mathbf{F} \in \mathcal{F}, \eta(\mathbf{F}) \leq \eta(\mathbf{F}^*) \right\} \quad (12)$$

represents the worst-case error of a potentially optimal solution to the Oracle IP.

The quantity in (12) in fact gives an *upper bound* of the error of the optimal solution $\hat{\mathbf{Y}}$ to the SDP relaxation, as is shown in the theorem below.

THEOREM 1 (IP bounds SDP). *Under Model 1, there exist some universal constants $C_s > 0, C \geq 1$ for which the following holds. If the SNR satisfies*

$$s^2 \geq C_s \left(\sqrt{\frac{kd}{n} \log n} + \frac{kd}{n} + k \right), \quad (13)$$

then we have

$$\frac{\|\hat{\mathbf{Y}} - \mathbf{Y}^*\|_1}{\|\mathbf{Y}^*\|_1} \leq 2 \cdot \max \left\{ \frac{\|\mathbf{F} - \mathbf{F}^*\|_1}{\|\mathbf{F}^*\|_1} : \eta(\mathbf{F}) \leq \eta(\mathbf{F}^*), \mathbf{F} \in \mathcal{F} \right\}$$

with probability at least $1 - n^{-C} - 2e^{-n}$.

We prove this theorem in Section 5. The proof has two main steps: (i) showing that with high probability the SDP error is upper bounded by the objective value of a linear program (LP), and (ii) showing that the LP admits an *integral* optimal solution and relating this solution to the Oracle IP error in (12). We note that the key step (ii), which establishes a hidden integrality property, is completely deterministic. The SNR condition (13) is required only in the probabilistic step (i). As we elaborate in Sections 4.2.1–4.4, our SNR condition holds even in the regime where exact recovery is impossible, and is often milder than existing results on convex relaxations.

To obtain an explicit bound on the SDP error, it suffices to upper bound the error (12) of the Oracle IP. This turns out to be a relatively easy task compared to directly controlling the error of the SDP: the Oracle IP has only finitely many feasible solutions, allowing one to use a union-bound-like argument. Our analysis establishes that the error of the Oracle IP decays *exponentially* in the SNR, as summarized in the theorem below.

THEOREM 2 (Exponential rate of IP). *Under Model 1, there exist universal constants $C_s, C_g, C_e > 0$ for which the following holds. If $s^2 \geq C_s k$, then we have*

$$\max \left\{ \frac{\|\mathbf{F} - \mathbf{F}^*\|_1}{\|\mathbf{F}^*\|_1} : \eta(\mathbf{F}) \leq \eta(\mathbf{F}^*), \mathbf{F} \in \mathcal{F} \right\} \leq C_g \exp \left[-\frac{s^2}{C_e} \right]$$

with probability at least $1 - \frac{3}{2}n^{-1}$.

The proof is given in Section 6. Combining Theorems 1 and 2, we immediately deduce that the SDP (7) also achieves an exponentially decaying error rate.

COROLLARY 1 (Exponential rate of SDP). *Under Model 1 and the SNR condition (13), there exist universal constants $C_m, C_e > 0$ such that*

$$\frac{\|\hat{\mathbf{Y}} - \mathbf{Y}^*\|_1}{\|\mathbf{Y}^*\|_1} \leq C_m \exp \left[-\frac{s^2}{C_e} \right]$$

with probability at least $1 - 2n^{-1}$.

Our last theorem concerns the explicit clustering $\hat{\sigma}$ extracted from $\hat{\mathbf{Y}}$ using the `cluster` procedure described in Section 3.4. In particular, we show that the clustering error of $\hat{\sigma}$ is upper bounded by the ℓ_1 error of $\hat{\mathbf{Y}}$; consequently, $\hat{\sigma}$ inherits the exponential error bound of $\hat{\mathbf{Y}}$.

THEOREM 3 (Clustering error). *The clustering $\hat{\sigma}$ extracted from $\hat{\mathbf{Y}}$ satisfies the deterministic inequality*

$$\text{err}(\hat{\sigma}, \sigma^*) \lesssim \frac{\|\hat{\mathbf{Y}} - \mathbf{Y}^*\|_1}{\|\mathbf{Y}^*\|_1}.$$

Consequently, under Model 1 and the SNR condition (13), there exist universal constants $C_m, C_e > 0$ such that

$$\text{err}(\hat{\sigma}, \sigma^*) \leq C_m \exp \left[-\frac{s^2}{C_e} \right]$$

with probability at least $1 - 2n^{-1}$.

The proof is given in Appendix A. Note that the above clustering error bound is information-theoretically optimal (up to a constant in the exponent) in view of the minimax results in [42].

We would like to mention that the SNR condition (13) for Theorems 1–3 and Corollary 1 can be improved to

$$s^2 \geq C_s \left(\sqrt{\frac{kd}{n} \log n} + k \sqrt{\frac{d}{n}} + k \right),$$

by adopting the proof strategies in the conference version of this paper [21]. We adopt the current proof as it allows us to streamline the analysis for both SGMM (Model 1) and the semi-random SGMM (Model 2 to be introduced below).

4.2.1. Consequences We explore the consequences of our error bounds in Corollary 1 and Theorem 3.

Exact recovery: If the SNR s satisfies the condition (13) and moreover $s^2 \gtrsim \log n$, then Theorem 3 guarantees that $\text{err}(\hat{\sigma}, \sigma^*) < \frac{1}{n}$, which means that $\text{err}(\hat{\sigma}, \sigma^*) = 0$ and hence the true underlying clusters are recovered exactly. In fact, in this case Corollary 1 guarantees the SDP solution satisfies the bound $\|\hat{\mathbf{Y}} - \mathbf{Y}^*\|_1 < \frac{1}{4}$, so simply rounding $\hat{\mathbf{Y}}$ *element-wise* produces the ground-truth cluster matrix $\mathbf{Y}^*$. Note that the SNR conditions above can be simplified to $s^2 \gtrsim k + \log n$ when $n \gtrsim d$.

Recovery with δ -error: Our results are applicable even in the regime with a lower SNR, for which exact recovery of the clusters is impossible due to the overlap between clusters. In this regime, Theorem 3 implies the following *approximate recovery* guarantees: for any number $\delta \in (0, 1)$, if s satisfies the condition (13) and $s^2 \gtrsim \log \frac{1}{\delta}$, then $\text{err}(\hat{\sigma}, \sigma^*) \leq \delta$ and hence SDP correctly recovers the cluster memberships of at least $(1 - \delta)$ fraction of the points.

In Section 4.5 to follow we compare the above results with existing ones. As a passing note, the above results on *clustering error* further imply error bounds on estimating the *cluster centers*. We do not delve into the details and refer the readers to our conference paper [21] for such a result.

4.3. Robustness under Semi-random Model In this section we extend the above results to the so called semi-random setting, in which an omniscient adversary is allowed to modify, in a coordinated way, the data generated from a probabilistic model. In the literature [25, 44, 36, 48, 10], semi-random models have been recognized as a more flexible model for non-ideal, real-world data, and are used as a benchmark for evaluating algorithmic robustness. Algorithms that over-exploit the idealized structures in a purely random model (e.g., independence, homogeneity or Gaussianity), often fail completely in the semi-random version of the model [25, 36, 48].

In the context of SGMM, we consider the following semi-random model proposed in [10].

MODEL 2 (Semi-random Sub-Gaussian Mixture Model). Let $\{\mathbf{h}_i\}$ be points generated from Model 1. An adversary can arbitrarily move each point $\mathbf{h}_i$ towards its cluster center $\boldsymbol{\mu}_{\sigma^*(i)}$ to produce a new point $\tilde{\mathbf{h}}_i$; that is

$$\tilde{\mathbf{h}}_i = \boldsymbol{\mu}_{\sigma^*(i)} + \tilde{\mathbf{g}}_i, \quad \text{where } \tilde{\mathbf{g}}_i = \alpha_i \mathbf{g}_i \text{ for some } \alpha_i \in [0, 1].$$

Here, $\{\alpha_i\}$ are chosen arbitrarily from $[0, 1]$ and may be correlated with $\{\mathbf{g}_i\}$ and with each other.

Our goal is to estimate the true clustering σ^* using the SDP (7) with the modified data $\{\tilde{\mathbf{h}}_i\}$ as the input. Before proceeding, we make two remarks on the above model. (i) Although the adversary shrinks the noise $\{\mathbf{g}_i\}$, it does not necessarily make the clustering problem easier, as the adversary can calibrate its output to create spurious local structures in the data. (ii) The adversary must move the original point $\mathbf{h}_i$ in the direction of $\mathbf{g}_i = \mathbf{h}_i - \boldsymbol{\mu}_{\sigma^*(i)}$. If the adversary were allowed to move in other directions (without increasing the distance $\|\mathbf{h}_i - \boldsymbol{\mu}_{\sigma^*(i)}\|_2$), then the clustering structure may be completely lost, since most points $\tilde{\mathbf{h}}_i$ may become closer to a different cluster center than to its own center $\boldsymbol{\mu}_{\sigma^*(i)}$. See [10, Section 1] for a detailed discussion of these two points.

Under Model 2, we consider the same Oracle IP (11) but with $\{\boldsymbol{\mu}_{\sigma^*(j)} + 8\tilde{\mathbf{g}}_j\}$ (that is, $\{\tilde{\mathbf{h}}_i\}$ with variance augmented by 8) as the input. That is, the objective function η therein is replaced by

$$\tilde{\eta}(\mathbf{F}) := \sum_{j \in [n]} \sum_{a \in [k]} \|(\boldsymbol{\mu}_{\sigma^*(j)} + 8\tilde{\mathbf{g}}_j) - \boldsymbol{\mu}_a\|_2^2 F_{ja}.$$

The following theorem is an analogue of Theorem 1 and bounds the SDP error in terms of the Oracle IP error.

THEOREM 4 (IP bounds SDP, semi-random). Under Model 2, there exist some universal constants $C_s, C > 0$ for which the following holds. If the SNR satisfies

$$s^2 \geq C_s k, \tag{14}$$

and $n \gtrsim k(d + \log n)$, then with probability at least $1 - 4n^{-1} - 2^{-d}$, we have

$$\frac{\|\hat{\mathbf{Y}} - \mathbf{Y}^*\|_1}{\|\mathbf{Y}^*\|_1} \leq 2 \cdot \max \left\{ \frac{\|\mathbf{F} - \mathbf{F}^*\|_1}{\|\mathbf{F}^*\|_1} : \tilde{\eta}(\mathbf{F}) \leq \tilde{\eta}(\mathbf{F}^*), \mathbf{F} \in \mathcal{F} \right\} + \varepsilon(n, k, d, s) \tag{15}$$

where

$$\varepsilon(n, k, d, s) := \tilde{O} \left(\frac{kd}{ns^4} + e^{-s^4/C} \right)$$

and $\tilde{O}(\cdot)$ hides a multiplicative factor of $\log(d + \log n)$.

The proof is given in Appendix B and closely follows that of Theorem 1. We see that the bound is similar to non-semi-random setting except for the additional, polynomial error term $\varepsilon(n, k, d, s)$. This term captures the “global” effect (cf. Section 1.1.3) of the adversary, who can make some points from two different clusters closer to each other and thus increase the clustering error.

The next theorem shows that the Oracle IP obeys the same error bound as in Theorem 2.

THEOREM 5 (Exponential rate of IP, semi-random). *Under Model 2, there exist universal constants $C_s, C_g, C_e > 0$ for which the following holds. If $s^2 \geq C_s k$, then we have*

$$\max \left\{ \frac{\|\mathbf{F} - \mathbf{F}^*\|_1}{\|\mathbf{F}^*\|_1} : \tilde{\eta}(\mathbf{F}) \leq \tilde{\eta}(\mathbf{F}^*), \mathbf{F} \in \mathcal{F} \right\} \leq C_g \exp \left[-\frac{s^2}{C_e} \right]$$

with probability at least $1 - \frac{3}{2}n^{-1}$.

The proof is provided in Section 6. Theorem 5 shows that the adversary has essentially no “local effect” (cf. Section 1.1.3) and does not change the error of the Oracle IP. This is intuitive: since the Oracle IP knows the true cluster centers and the adversary can only move points towards their centers, the error of the Oracle IP can only improve or stay the same.

Combining Theorems 4 and 5 gives the following explicit error bound for the SDP relaxation.

COROLLARY 2 (Exponential rate of SDP, semi-random). *Under Model 2, the SNR condition (14) and the sample complexity condition $n \gtrsim k(d + \log n)$, there exists some universal constant $C_e > 0$ such that*

$$\text{err}(\hat{\boldsymbol{\sigma}}, \boldsymbol{\sigma}^*) \stackrel{(i)}{\lesssim} \frac{\|\hat{\mathbf{Y}} - \mathbf{Y}^*\|_1}{\|\mathbf{Y}^*\|_1} \stackrel{(ii)}{\lesssim} \tilde{O} \left(\frac{kd}{ns^4} + e^{-s^2/C_e} \right) \stackrel{(iii)}{\lesssim} \begin{cases} \tilde{O} \left(\frac{kd}{ns^4} \right), & \text{if } s^2 \gtrsim \log \left(\frac{n}{kd} \right), \\ \tilde{O} \left(e^{-s^2/C_e} \right), & \text{otherwise,} \end{cases}$$

with probability at least $1 - 6n^{-1} - 2^{-d}$.

Proof. The inequality (i) follows from part 1 of Theorem 3. The inequality (ii) holds by combining Theorems 4 and 5 and noting that $e^{-s^4/C} \leq e^{-s^2/C}$ as $s \geq 1$ by assumption. The inequality (iii) holds since $\frac{kd}{ns^4} \geq e^{-s^2/C_e} \iff s^2 \geq c_0 \log \left(\frac{n}{kd} \right)$ for a sufficiently large constant $c_0 > 0$. $\square$

The result in Corollary 2 demonstrates a phase transition phenomenon for the semi-random model. In the low-SNR regime, the error bound decays exponentially in the SNR s and is the same as in the non-semi-random setting. As mentioned, such an error is unavoidable even in standard SGMM [42]. In the high-SNR regime, the effect of the adversary becomes dominant, in which case the error decays at a slower, polynomial rate. Interestingly, a lower bound result is established in [10, Theorem 4.1], which shows that any k -means based algorithm must incur this polynomial error. Therefore, we believe that the bound in Corollary 2 is unimprovable. To the best of our knowledge, this is the first phase transition result for semi-random SGMM for any algorithm.

4.3.1. Comparison The best result for Semi-random SGMM is given in [10]. They analyze the Lloyd’s k -means algorithm and show that its output $\hat{\boldsymbol{\sigma}}_{\text{lloyd}}$ satisfies [10, Theorem 3.1]:

$$\text{err}(\hat{\boldsymbol{\sigma}}_{\text{lloyd}}, \boldsymbol{\sigma}^*) = \tilde{O} \left(\frac{kd}{ns^4} \right), \quad \text{if } s^2 \gtrsim \min\{k, d\} \cdot \log n.$$

Assuming $k \leq d$, one sees that our result in Corollary 2 strictly generalizes theirs under the setting of equal-sized clusters. In particular, in the high-SNR regime with $s^2 \gtrsim \log n$, both results provide the same polynomial error bound $\tilde{O}(kd/(ns^4))$. Our result further applies to the low-SNR regime with $s^2 \lesssim \log n$, while theirs does not. In fact, it is not clear whether the Lloyd’s algorithm can achieve the same robust error bound as the SDP relaxation in this more challenging regime.

When $k \geq d$, our result requires the SNR condition $s^2 \gtrsim k$ whereas their result require $s^2 \gtrsim d \log n$, which are incomparable. We suspect that both conditions can be improved, though we do not have a formal proof.

4.4. Stochastic Ball Model In this section, we illustrate the power of our main theorems by deriving several new results for the Stochastic Ball Model, formally described below. This model was introduced in [52] and has recently attracted attention in the computer science and mathematical programming communities [8, 29, 40].

MODEL 3 (Stochastic Ball Model). *Under Model 1, we assume in addition that each $\mathbf{g}_i$ is sampled from a rotationally invariant distribution supported on the unit ℓ_2 ball in $\mathbb{R}^d$.*

Under the above model, each data point $\mathbf{h}_i = \boldsymbol{\mu}_{\sigma^*(i)} + \mathbf{g}_i$ is sampled from the unit ball around its cluster center $\boldsymbol{\mu}_{\sigma^*(i)}$. It is not hard to see that this model is a special case of SGMM, with its sub-Gaussian norm given below:

FACT 1. *Under Model 3, each $\mathbf{g}_i$ has sub-Gaussian norm $\|\mathbf{g}_i\|_{\psi_2} \leq \tau = C\sqrt{\frac{1}{d}}$ for some universal constant $C > 0$.*

For completeness we prove this standard fact in Appendix C.2. Specializing Corollary 1 and Theorem 3 to the Stochastic Ball Model, we obtain the following sufficient conditions on the minimum center separation Δ for exact and approximate recovery:

$$\Delta^2 \gtrsim \sqrt{\frac{k}{nd} \log n} + \frac{k}{n} + \frac{k}{d} + \begin{cases} \frac{\log n}{d}, & \text{for exact recovery,} \\ \frac{\log \delta^{-1}}{d}, & \text{for recovery of } (1 - \delta) \text{ fraction of the points.} \end{cases}$$

The state-of-the-art results for Stochastic Ball Models are given in [8, 29, 40], which establish that SDP achieves exact recovery when n is sufficiently large and $\Delta^2 \geq 4 + \theta(k, d)$ for some non-negative function $\theta(\cdot)$. Regardless of the values of k and d , these results all require the separation to be at least a constant and thus the balls to be disjoint. In contrast, our results above are applicable to a *small-separation* regime that is not covered by these existing results. In particular, when n is large and $k = O(1)$, our results guarantee that SDP achieves approximate recovery when $\Delta^2 \gtrsim \frac{1}{d}$, which can be arbitrarily smaller than a constant when the dimension d grows. Moreover, the recovery is exact if $\Delta^2 \gtrsim \frac{\log n}{d}$, which can again be arbitrarily small as long as n does not grow exponentially fast (i.e., $n = e^{o(d)}$). Therefore, in the high dimensional setting, our results guarantee strong performance of the SDP *even when the centers are very close and the balls overlap with each other*.

It may appear a bit counter-intuitive that exact/approximate recovery is achievable when the separation is so small. Such a result is a manifestation of the geometry in high dimensions: the relative volume of the intersection of two balls vanishes as the dimension grows. As a passing note, the result in [40, Corollary 4.3] establishes a *necessary* condition $\Delta \geq 1 + \sqrt{1 + \frac{2}{d+2}}$ for exact recovery. Our result above does *not* contradict this condition, as the latter allows n to grow arbitrarily fast, in which case with high probability some points will land in the intersection.

4.5. Comparison with Existing Results In this section we compare our results above with the ones in the literature on clustering SGMM. Our focus is on results that provide explicit clustering error bounds in the regime where exact cluster recovery is impossible.

4.5.1. Clustering Error Bounds The work of [47] considers the Peng-Wei SDP relaxation introduced in [54]. An intermediate result of theirs, after appropriate rescaling, establishes the polynomial error bound $\|\hat{\mathbf{Y}} - \mathbf{Y}^*\|_F^2 \lesssim \frac{n^2}{s^2}$ under the setting of balanced clusters and $s^2 \gtrsim k$. In comparison, our exponential error bound $\|\hat{\mathbf{Y}} - \mathbf{Y}^*\|_F^2 \leq \|\hat{\mathbf{Y}} - \mathbf{Y}^*\|_1 \lesssim$

$\frac{n^2}{k}e^{-s^2}$ is strictly better. The work in [42] proves an exponentially decaying error rate similar to ours, but for a different algorithm (Lloyd’s algorithm). Their results require the SNR condition $s^2 \gg k^2 + k^3 \frac{d}{n}$ and $k^3 \ll \frac{n}{\log n}$ as $n \rightarrow \infty$. Our SNR condition in (13) has milder dependency on k , though dependency on n and d are a bit more subtle. We do note that under their more restricted SNR condition, their results provide tight constants in the exponent of the error rate.

Several papers [26, 14, 51] appeared around or after the conference version [21] of this manuscript was published. All these papers establish an exponential error bound of the form $\mathbf{err}(\hat{\boldsymbol{\sigma}}, \boldsymbol{\sigma}^*) \lesssim \exp[-s^2/C]$ that is similar to ours, though the details differ. Restricting to the spherical Gaussian mixture model with $k=2$ component, the work [51] provides a refinement of the above exponent error bound with a precise coefficient C . In particular, they establish a sharp bound $\mathbf{err}(\hat{\boldsymbol{\sigma}}, \boldsymbol{\sigma}^*) \leq \Phi^c[(1 \pm o(1))r]$, where Φ^c is the complementary cdf of the standard Gaussian distribution and $r := s^2/\sqrt{s^2 + d/n}$ is a form of SNR. This bound is achieved using an iterative algorithm that is different from our SDP approach. The work [14] generalizes the Peng-Wei SDP relaxation to the non-Euclidean setting. When specialized to the (Euclidean) SGMM, their error bound is essentially identical to ours. The work [26] also considers the Peng-Wei SDP relaxation and allows for imbalanced clusters. Their bound involves an alternative definition for the SNR $s_0^2 := \min(s^2, \ell s^4 R^{-1})$, where R represents the effective rank of the mixture model.² Since $s_0 \leq s$, their exponential error bound decays no faster than ours in Corollary 1. On the other hand, their bound holds under the SNR condition $s^2 \gtrsim (1 + \sqrt{R/n})k$, which is weaker than ours when the effective rank is low. In terms of the algorithm, they propose to extract explicit clustering from SDP solutions using an approximate k -medoids algorithm. This algorithm itself involves solving a linear program and is more complicated than our greedy extraction procedure in Algorithm 2.

We note that none of the above work provides results for the semi-random model and the Stochastic Ball Model, unlike what we do in Sections 4.3 and 4.4.

4.5.2. Conditions for Exact Recovery As discussed in Section 4.2.1, a corollary of our results provides sufficient condition, in terms of the SNR s^2 , for exact recovery of the cluster. While exact recovery is not our focus, we nevertheless summarize and compare with several most representative results in Table 1. It can be seen that our result is comparable with, and sometimes better than, the other results in the table. Only the paper [58] proves a strictly better SNR condition, which remains the best to date. Their result is, however, only established for the special case of mixture of spherical Gaussians; in fact, their analysis makes substantial use of the structure of this special case. In comparison, our result is more general and applies to non-spherical and sub-Gaussian mixtures. Their result also requires a higher sample complexity $n = \Omega(k^2 d^3)$ as compared to ours $n = \Omega(kd)$.

To conclude this section, we mention that there is a large body of work on estimating the cluster *centers* of Gaussian mixture model, or estimating the density of the entire mixture distribution. See, for example, the work in [20, 28, 32, 35, 49, 56]. Our work instead focuses on the problem of *clustering* the data points. Results for these two problems are not directly comparable in general. There exist settings in which center estimation can be done while clustering is impossible; for example, when the cluster centers are identical.

² Explicitly, $R := \frac{\max_{a \in [k]} \|\boldsymbol{\Sigma}_a\|_F^2}{\max_{a \in [k]} \|\boldsymbol{\Sigma}_a\|_{\text{op}}^2}$, where $\boldsymbol{\Sigma}_a$ is the covariance matrix of the a -th cluster.

Paper	SNR Condition on s^2	Algorithm
[58]	$\Omega(\sqrt{k \log n})$ for spherical Gaussian mixture	Spectral
[4]	$\Omega(k \log n + k^2)$	Spectral
[37]	$\Omega(k^2 \cdot \text{polylog}(n))$	Spectral
[9]	$\Omega(k \cdot \text{polylog}(n))$	Spectral
[42]	$\Omega(k^2 + \log n)$	Spectral + Lloyd's
[40, 57, 14, 26]	$\Omega(k + \log n)$	SDP
This paper	$\Omega(k + \log n)$	SDP

TABLE 1. Summary of existing results on exact cluster recovery for GMM.

5. Proof of Theorem 1 In this section, we prove Theorem 1, which relates the errors of SDP (7) and Oracle IP formulation (11). Some additional notations are used in the proof and the rest of this paper. We define the shorthand $\gamma := \|\hat{\mathbf{Y}} - \mathbf{Y}^*\|_1$ for the ℓ_1 error of the SDP solution. For a matrix $\mathbf{M}$, we write $\|\mathbf{M}\|_\infty := \max_{i,j} |M_{ij}|$ as its entry-wise ℓ_∞ norm. We let $\mathbf{I}$ and $\mathbf{J}$ be the $n \times n$ identity matrix and all-one matrix, respectively. For a real number x , $\lceil x \rceil$ denotes its ceiling. We denote by $C_a^* := \{i \in [n] : \sigma^*(i) = a\}$ the set of indices of points in cluster a , and we define $\ell := |C_a^*| = \frac{n}{k}$ to be the cluster size. Throughout the proof, we use $i, j \in [n]$ to index the data points, and $a, b \in [k]$ to index the clusters. We sometimes omit the ranges of these variables, i.e., $[n]$ and $[k]$, to avoid cluttered notation.

Our proof consists of three main steps:

- **Step 1:** Use optimality of $\hat{\mathbf{Y}}$ to derive a *basic inequality* satisfied by the error matrix $\hat{\mathbf{Y}} - \mathbf{Y}^*$;
- **Step 2:** Relate the basic inequality to a linear program (LP) parameterized by the SDP error $\gamma := \|\hat{\mathbf{Y}} - \mathbf{Y}^*\|_1$;
- **Step 3:** Show that the LP is integral and thereby bound γ by the error of the Oracle IP.

Without loss of generality, assume that the SDP error $\gamma > 0$. For ease of understanding the proof, it is convenient to think of the n data points as appropriately ordered and hence $\mathbf{Y}^*$ takes the block diagonal form as in (8), though the proof does not actually rely on this ordering.

Step 1: basic inequality To streamline the proof, we first record several basic structural properties of the error matrix $\hat{\mathbf{Y}} - \mathbf{Y}^*$ in the following proposition, whose proof is given in Section 5.1.

FACT 2. Define the shorthand $\mathbf{B} := \hat{\mathbf{Y}} - \mathbf{Y}^*$. We have

- (a) (zero diagonal) $\forall i \in [n] : B_{ii} = 0$;
- (b) (signs of diagonal and off-diagonal blocks) $B_{ij} \in \begin{cases} [-1, 0], & \text{if } \sigma_i^* = \sigma_j^*, \\ [0, 1], & \text{otherwise;} \end{cases}$
- (c) (zero row sum) $\forall i \in [n] : \sum_{j: \sigma^*(j) \neq \sigma^*(i)} B_{ij} = -\sum_{j: \sigma^*(j) = \sigma^*(i)} B_{ij}$;
- (d) (zero block row sum) $\forall a \in [k] : \sum_{i: \sigma^*(i) = a} \sum_{j: \sigma^*(j) \neq a} B_{ij} = -\sum_{i: \sigma^*(i) = a} \sum_{j: \sigma^*(j) = a} B_{ij}$;
- (e) (blockwise decomposition of SDP error) $\gamma = \sum_{i,j: \sigma^*(i) \neq \sigma^*(j)} B_{ij} - \sum_{i,j: \sigma^*(i) = \sigma^*(j)} B_{ij}$;
- (f) (diagonal and off-diagonal blocks equally divide SDP error) $\frac{\gamma}{2} = \sum_{i,j: \sigma^*(i) \neq \sigma^*(j)} B_{ij} - \sum_{i,j: \sigma^*(i) = \sigma^*(j)} B_{ij}$.

We next decompose the input matrix of pairwise squared distances as

$$\mathbf{A} = \mathbf{C} + \mathbf{C}^\top - 2\mathbf{H}\mathbf{H}^\top,$$

where $\mathbf{H} \in \mathbb{R}^{n \times d}$ is the matrix whose i -th row is the data point $\mathbf{h}_i$, and $\mathbf{C} \in \mathbb{R}^{n \times n}$ is the matrix where the entries in the i -th row are identical and equal to $\|\mathbf{h}_i\|_2^2$. By Fact 2(c), we have $\langle \hat{\mathbf{Y}} - \mathbf{Y}^*, \mathbf{C} \rangle = \langle \hat{\mathbf{Y}} - \mathbf{Y}^*, \mathbf{C}^\top \rangle = 0$, which implies $\langle \hat{\mathbf{Y}} - \mathbf{Y}^*, \mathbf{C} + \mathbf{C}^\top \rangle = 0$.

Let $\mathbf{G} := \mathbf{H} - \mathbb{E}\mathbf{H}$ be the centered version of $\mathbf{H}$. We can compute

$$\mathbf{H}\mathbf{H}^\top = (\mathbf{G} + \mathbb{E}\mathbf{H})(\mathbf{G} + \mathbb{E}\mathbf{H})^\top = \mathbf{G}\mathbf{G}^\top + \mathbf{G}(\mathbb{E}\mathbf{H})^\top + (\mathbb{E}\mathbf{H})\mathbf{G}^\top + (\mathbb{E}\mathbf{H})(\mathbb{E}\mathbf{H})^\top$$

and

$$\mathbb{E}\mathbf{H}\mathbf{H}^\top = \mathbb{E}\mathbf{G}\mathbf{G}^\top + (\mathbb{E}\mathbf{H})(\mathbb{E}\mathbf{H})^\top.$$

Therefore, we have the expression

$$\mathbf{H}\mathbf{H}^\top - \mathbb{E}\mathbf{H}\mathbf{H}^\top = (\mathbf{G}\mathbf{G}^\top - \mathbb{E}\mathbf{G}\mathbf{G}^\top) + \mathbf{G}(\mathbb{E}\mathbf{H})^\top + (\mathbb{E}\mathbf{H})\mathbf{G}^\top.$$

Let $\mathbf{U} := \frac{1}{\sqrt{\ell}}\mathbf{F}^*$ so that it takes the form

$$U_{ia} = \frac{1}{\sqrt{\ell}}F_{ia}^* = \frac{1}{\sqrt{\ell}} \cdot \mathbb{I}\{\sigma^*(i) = a\}.$$

Note that the columns of $\mathbf{U}$ are the left singular vectors of $\mathbf{Y}^*$. For each matrix $\mathbf{M} \in \mathbb{R}^{n \times n}$, define the projection $\mathcal{P}_T(\mathbf{M}) := \mathbf{U}\mathbf{U}^\top\mathbf{M} + \mathbf{M}\mathbf{U}\mathbf{U}^\top - \mathbf{U}\mathbf{U}^\top\mathbf{M}\mathbf{U}\mathbf{U}^\top$ and its orthogonal complement $\mathcal{P}_{T^\perp}(\mathbf{M}) := \mathbf{M} - \mathcal{P}_T(\mathbf{M})$.

Now, recall that $\hat{\mathbf{Y}}$ is optimal and $\mathbf{Y}^*$ is feasible to the SDP (7). Combining with the above decomposition of $\mathbf{A}$ and $\mathbf{H}$, we obtain the following basic inequality:

$$\begin{aligned} 0 &\leq -\frac{1}{2} \langle \hat{\mathbf{Y}} - \mathbf{Y}^*, \mathbf{A} \rangle \\ &= \langle \hat{\mathbf{Y}} - \mathbf{Y}^*, \mathbf{H}\mathbf{H}^\top - \mathbb{E}\mathbf{H}\mathbf{H}^\top \rangle + \langle \hat{\mathbf{Y}} - \mathbf{Y}^*, \mathbb{E}\mathbf{H}\mathbf{H}^\top \rangle \\ &= \langle \hat{\mathbf{Y}} - \mathbf{Y}^*, \mathbf{G}\mathbf{G}^\top - \mathbb{E}\mathbf{G}\mathbf{G}^\top + \mathbf{G}(\mathbb{E}\mathbf{H})^\top + (\mathbb{E}\mathbf{H})\mathbf{G}^\top \rangle + \langle \hat{\mathbf{Y}} - \mathbf{Y}^*, \mathbb{E}\mathbf{H}\mathbf{H}^\top \rangle \\ &\stackrel{(i)}{=} \underbrace{\langle \hat{\mathbf{Y}} - \mathbf{Y}^*, \mathbf{G}\mathbf{G}^\top \rangle}_{S_1} + \underbrace{\langle \hat{\mathbf{Y}} - \mathbf{Y}^*, \mathbf{G}(\mathbb{E}\mathbf{H})^\top + (\mathbb{E}\mathbf{H})\mathbf{G}^\top \rangle}_{S_2} + \underbrace{\langle \hat{\mathbf{Y}} - \mathbf{Y}^*, (\mathbb{E}\mathbf{H})(\mathbb{E}\mathbf{H})^\top \rangle}_{S_3} \\ &= \underbrace{\langle \hat{\mathbf{Y}} - \mathbf{Y}^*, \mathcal{P}_T(\mathbf{G}\mathbf{G}^\top) \rangle}_{S_1} + \underbrace{\langle \hat{\mathbf{Y}} - \mathbf{Y}^*, \mathcal{P}_{T^\perp}(\mathbf{G}\mathbf{G}^\top) \rangle}_{S_2} + 2 \cdot \underbrace{\langle \hat{\mathbf{Y}} - \mathbf{Y}^*, \mathbf{G}(\mathbb{E}\mathbf{H})^\top \rangle}_{S_3} \\ &\quad + \underbrace{\langle \hat{\mathbf{Y}} - \mathbf{Y}^*, (\mathbb{E}\mathbf{H})(\mathbb{E}\mathbf{H})^\top \rangle}_{S_4}, \end{aligned} \tag{16}$$

where in step (i) we use the identities $\langle \hat{\mathbf{Y}} - \mathbf{Y}^*, \mathbb{E}\mathbf{G}\mathbf{G}^\top \rangle = 0$ and $\langle \hat{\mathbf{Y}} - \mathbf{Y}^*, \mathbb{E}\mathbf{H}\mathbf{H}^\top \rangle = \langle \hat{\mathbf{Y}} - \mathbf{Y}^*, (\mathbb{E}\mathbf{H})(\mathbb{E}\mathbf{H})^\top \rangle$, which follow from Fact 2(a) and the fact that $\mathbb{E}\mathbf{G}\mathbf{G}^\top$ is a diagonal matrix. Here, $\hat{\mathbf{Y}} - \mathbf{Y}^*$ takes the role of the error matrix, and S_1 and S_2 are the products of $\mathbf{G}\mathbf{G}^\top$ with the error matrix projected to the column space of $\mathbf{Y}^*$ and its orthogonal complement, respectively. The quantities S_3 and S_4 can be equivalently written as $S_3 = \sum_{i,j} (\hat{\mathbf{Y}} - \mathbf{Y}^*)_{ij} \langle \mathbf{g}_i, \boldsymbol{\mu}_{\sigma^*(j)} \rangle$ and $S_4 = \sum_{i,j} (\hat{\mathbf{Y}} - \mathbf{Y}^*)_{ij} \langle \boldsymbol{\mu}_{\sigma^*(i)}, \boldsymbol{\mu}_{\sigma^*(j)} \rangle$.

Step 2: from SDP to LP The following propositions provide high probability bounds on the terms S_1 , S_2 and S_4 that appear on the RHS of the basic inequality (16).

PROPOSITION 1. *If $s^2 \geq C \left(\sqrt{\frac{kd}{n} \log n} + \sqrt{\frac{k}{n} \log n} + \frac{kd}{n} \right)$ for some universal constant $C > 0$, then $S_1 \leq \frac{1}{100} \Delta^2 \gamma$ with probability at least $1 - n^{-8}$.*

PROPOSITION 2. We have $S_2 \leq \frac{\gamma}{\ell} \|\mathbf{G}\|_{\text{op}}^2$. Moreover, if $s^2 \geq Ck \left(\frac{d}{n} + 1\right)$ for some universal constant $C > 0$, then $\frac{\gamma}{\ell} \|\mathbf{G}\|_{\text{op}}^2 \leq \frac{1}{100} \Delta^2 \gamma$ with probability at least $1 - e^{-n/2}$.

PROPOSITION 3. We have $S_4 = -\frac{1}{2} \sum_{a,b \in [k]: a \neq b} \Delta_{ab}^2 [\sum_{i \in C_a^*, j \in C_b^*} (\hat{\mathbf{Y}} - \mathbf{Y}^*)_{ij}]$. Furthermore, $S_4 \leq -\frac{1}{4} \Delta^2 \gamma$.

We prove the above propositions in Sections 5.2, 5.3 and 5.4, respectively, using appropriate concentration inequalities. The above bounds on S_1, S_2 and S_4 together imply that $S_1 + S_2 \leq -\frac{1}{2} S_4$ with probability at least $1 - (2n)^{-C'} - 2e^{-n}$ for some universal constant $C' > 0$. Combining with the basic inequality (16), we obtain that with the same probability, there holds

$$0 \leq S_3 + \frac{1}{4} S_4. \quad (17)$$

The rest of the proof is completely deterministic, in which we exploit the above inequality (17) and the structures of S_3 and S_4 while treating $\{\mathbf{g}_j\}$ therein as fixed quantities. For each $j \in [n]$ and $a \in [k]$, define the variable

$$\beta_{ja} := \langle \boldsymbol{\mu}_a - \boldsymbol{\mu}_{\sigma^*(j)}, \mathbf{g}_j \rangle - \kappa \Delta_{\sigma^*(j),a}^2, \quad (18)$$

where we recall that $\kappa := 1/8$. It is not hard to show that the quantity $S_3 + \frac{1}{4} S_4$ can be expressed a sum of $\{\beta_{ja}\}$ weighted by entries of the error matrix $\hat{\mathbf{Y}} - \mathbf{Y}^*$. This is the content of the next lemma, whose proof is given in Section 5.5.

LEMMA 1. We have that

$$\frac{1}{\ell} \left(S_3 + \frac{1}{4} S_4 \right) = \sum_j \sum_{a \neq \sigma^*(j)} \beta_{ja} \left(\frac{1}{\ell} \sum_{i \in C_a^*} (\hat{\mathbf{Y}} - \mathbf{Y}^*)_{ji} \right). \quad (19)$$

Note that the RHS of Equation (19) is linear in the quantities $\hat{X}_{ja} := \frac{1}{\ell} \sum_{i \in C_a^*} (\hat{\mathbf{Y}} - \mathbf{Y}^*)_{ji}$, $(j, a) \in [n] \times [k]$. With this observation in mind, we proceed to control the RHS of (19) with a linear program (LP). In particular, consider the following LP parameterized by a number $R \in [0, n]$:

$$V(R) := \left\{ \begin{array}{l} \max_{\mathbf{X}} \sum_j \sum_{a \neq \sigma^*(j)} \beta_{ja} X_{ja} \\ \text{s.t. } 0 \leq X_{ja} \leq 1, \quad \forall j \in [n], a \neq \sigma^*(j) \\ \sum_{a \neq \sigma^*(j)} X_{ja} \leq 1, \quad \forall j \in [n] \\ \sum_j \sum_{a \neq \sigma^*(j)} X_{ja} = R, \end{array} \right\}. \quad (20)$$

Note that this LP is always feasible. Now, Facts 2(b) and 2(f) imply that

$$\frac{\gamma}{2} = - \sum_{i,j: \sigma^*(i) = \sigma^*(j)} (\hat{\mathbf{Y}} - \mathbf{Y}^*)_{ij} \in (0, n\ell]$$

and thus

$$\sum_{j \in [n]} \sum_{a \neq \sigma^*(j)} \left(\frac{1}{\ell} \sum_{i \in C_a^*} (\hat{\mathbf{Y}} - \mathbf{Y}^*)_{ji} \right) = \frac{1}{\ell} \sum_{i,j: \sigma^*(i) \neq \sigma^*(j)} (\hat{\mathbf{Y}} - \mathbf{Y}^*)_{ij} = \frac{\gamma}{2\ell} \in (0, n].$$

Together with Facts 2(b) and 2(c), we conclude that the variables $\hat{X}_{ja} := \frac{1}{\ell} \sum_{i \in C_a^*} (\hat{\mathbf{Y}} - \mathbf{Y}^*)_{ji}$, $(j, a) \in [n] \times [k]$ are feasible to the LP (20) with $R = \frac{\gamma}{2\ell}$, hence the RHS of (19) is upper bounded by $V\left(\frac{\gamma}{2\ell}\right)$. Combining with Equation (17), we obtain the inequality

$$0 \leq \frac{1}{\ell} \left(S_3 + \frac{1}{4} S_4 \right) \leq V\left(\frac{\gamma}{2\ell}\right), \quad (21)$$

Step 3: from LP to Oracle IP The inequality (21) immediately implies a bound on the SDP error: $\frac{\gamma}{2\ell} \leq \max \{R \in (0, n] : V(R) \geq 0\}$. It is not hard to see that the last RHS is itself an LP. In fact, up to a factor of 2, we can restrict to integer values for the variable R . Indeed, inspecting the LP (20) we see that it satisfies $V\left(\frac{\gamma}{2\ell}\right) \leq \max \{V\left(\lceil \frac{\gamma}{2\ell} \rceil\right), V\left(\lceil \frac{\gamma}{2\ell} \rceil - 1\right)\}$. Combining with Equation (21), we obtain the bound

$$\begin{aligned} \frac{\gamma}{2\ell} \leq \left\lceil \frac{\gamma}{2\ell} \right\rceil &\leq \max \{R \in \{1, 2, \dots, n\} : V(R) \vee V(R-1) \geq 0\} \\ &\leq 1 + \max \{R \in \{1, 2, \dots, n\} : V(R) \geq 0\} \\ &\leq 2 \max \{R \in \{1, 2, \dots, n\} : V(R) \geq 0\} \\ &= 2 \max \{R \in \{0, 1, \dots, n\} : V(R) \geq 0\}. \end{aligned} \quad (22)$$

As the next and crucial step, we observe that for each integer $R \in \{0, 1, \dots, n\}$, the LP (20) defining $V(R)$ is in fact *integral*, that is, it has an optimal solution $\{X_{ja}\}$ satisfying $X_{ja} \in \{0, 1\}$, $\forall (j, a) \in [n] \times [k]$. Therefore, the value of $V(R)$ is unchanged if we replace the constraint $0 \leq X_{ja} \leq 1$ in the LP (20) with the integer constraint $X_{ja} \in \{0, 1\}$. With this replacement, we can expand the RHS of the bound (22) into a single IP:

$$\begin{aligned} &\max \{R \in \{0, 1, \dots, n\} : V(R) \geq 0\} \\ &= \left\{ \begin{array}{l} \max_{R, \mathbf{X}} R \\ \text{s.t. } R \in \{0, 1, \dots, n\} \\ \sum_j \sum_{a \neq \sigma^*(j)} \beta_{ja} X_{ja} \geq 0 \\ X_{ja} \in \{0, 1\}, \quad \forall j, a \neq \sigma^*(j) \\ \sum_{a \neq \sigma^*(j)} X_{ja} \leq 1, \quad \forall j \\ \sum_j \sum_{a \neq \sigma^*(j)} X_{ja} = R, \end{array} \right\} = \left\{ \begin{array}{l} \max_{\mathbf{X}} \sum_j \sum_{a \neq \sigma^*(j)} X_{ja} \\ \text{s.t. } \sum_j \sum_{a \neq \sigma^*(j)} \beta_{ja} X_{ja} \geq 0 \\ X_{ja} \in \{0, 1\}, \quad \forall j, a \neq \sigma^*(j) \\ \sum_{a \neq \sigma^*(j)} X_{ja} \leq 1, \quad \forall j \end{array} \right\} =: \text{IP}_1. \end{aligned} \quad (23)$$

Note that the above argument is what underlies the hidden integrality property of the SDP (7): even though the *solution* of the SDP is not integral, the *error* of the solution can be bounded by an LP that is integral.

As the last step, we show that the IP_1 in (23) above is in fact equivalent to the Oracle IP error as defined in Equation (12). In particular, recalling that $\eta(\mathbf{F})$ and $\mathcal{F}$ are the objective function and feasible set of the Oracle IP, we have the following identity:

LEMMA 2. *We have that*

$$\text{IP}_1 = \max \left\{ \frac{1}{2} \|\mathbf{F} - \mathbf{F}^*\|_1 : \eta(\mathbf{F}) \leq \eta(\mathbf{F}^*), \mathbf{F} \in \mathcal{F} \right\}.$$

We prove this lemma in Section 5.6 by a careful change-of-variable argument. Combining Equations (22), (23) and Lemma 2, we obtain that

$$\gamma \leq 4\ell \cdot \text{IP}_1 = 2\ell \cdot \max \{ \|\mathbf{F} - \mathbf{F}^*\|_1 : \eta(\mathbf{F}) \leq \eta(\mathbf{F}^*), \mathbf{F} \in \mathcal{F} \}.$$

Dividing both sides by $\|\mathbf{Y}^*\|_1 = n\ell = \|\mathbf{F}^*\|_1 \ell$ proves Theorem 1.

5.1. Proof of Fact 2 Fact 2(a) follows from the fact that $\hat{\mathbf{Y}}$ obeys the diagonal constraint in the SDP (7).

Fact 2(b) follows from the fact that $\hat{\mathbf{Y}}$ is feasible to the SDP (7) and hence satisfies $\hat{Y}_{ij} \in [-1, 1]$.

Fact 2(c) follows from the fact that both $\hat{\mathbf{Y}}$ and $\mathbf{Y}^*$ obey the row-sum constraint of the SDP (7).

Fact 2(d) follows from summing over $\{i : \sigma^*(i) = a\}$ for both sides of the equation in Fact 2(c).

Fact 2(e) follows from Fact 2(b) and the definition of γ .

Fact 2(f) follows from Facts 2(d) and 2(e).

5.2. Proof of Proposition 1 To bound S_1 , we begin with the decomposition

$$\begin{aligned} S_1 &= \left\langle \hat{\mathbf{Y}} - \mathbf{Y}^*, \mathbf{U}\mathbf{U}^\top (\mathbf{G}\mathbf{G}^\top) \right\rangle + \left\langle \hat{\mathbf{Y}} - \mathbf{Y}^*, (\mathbf{G}\mathbf{G}^\top) \mathbf{U}\mathbf{U}^\top \right\rangle - \left\langle \hat{\mathbf{Y}} - \mathbf{Y}^*, \mathbf{U}\mathbf{U}^\top (\mathbf{G}\mathbf{G}^\top) \mathbf{U}\mathbf{U}^\top \right\rangle \\ &\leq 2 \cdot \underbrace{\left| \left\langle \hat{\mathbf{Y}} - \mathbf{Y}^*, \mathbf{U}\mathbf{U}^\top (\mathbf{G}\mathbf{G}^\top) \right\rangle \right|}_{T_1} + \underbrace{\left| \left\langle \hat{\mathbf{Y}} - \mathbf{Y}^*, (\mathbf{G}\mathbf{G}^\top) \mathbf{U}\mathbf{U}^\top \right\rangle \right|}_{T_2}. \end{aligned}$$

By the generalized Holder's inequality, we have the bounds

$$T_1 \leq \gamma \cdot \|\mathbf{U}\mathbf{U}^\top (\mathbf{G}\mathbf{G}^\top)\|_\infty$$

and

$$\begin{aligned} T_2 &= \left| \left\langle \hat{\mathbf{Y}} - \mathbf{Y}^*, \mathbf{U}\mathbf{U}^\top (\mathbf{G}\mathbf{G}^\top) \mathbf{U}\mathbf{U}^\top \right\rangle \right| \\ &= \left| \left\langle (\hat{\mathbf{Y}} - \mathbf{Y}^*) \mathbf{U}\mathbf{U}^\top, \mathbf{U}\mathbf{U}^\top (\mathbf{G}\mathbf{G}^\top) \right\rangle \right| \\ &\leq \gamma \cdot \|\mathbf{U}\mathbf{U}^\top (\mathbf{G}\mathbf{G}^\top)\|_\infty, \end{aligned}$$

where the last inequality holds since the definition $\mathbf{U} = \frac{1}{\sqrt{\ell}} \mathbf{F}^*$ implies that $\mathbf{U}\mathbf{U}^\top = \frac{1}{\ell} \mathbf{Y}^*$, which further leads to $\|(\hat{\mathbf{Y}} - \mathbf{Y}^*) \mathbf{U}\mathbf{U}^\top\|_1 \leq \|\hat{\mathbf{Y}} - \mathbf{Y}^*\|_1 = \gamma$. Combining the above bounds, we obtain that

$$S_1 \leq 3\gamma \cdot \|\mathbf{U}\mathbf{U}^\top (\mathbf{G}\mathbf{G}^\top)\|_\infty.$$

Note that there are $m = nk$ distinct random variables in $\mathbf{U}\mathbf{U}^\top (\mathbf{G}\mathbf{G}^\top)$ and let us call them $X_1, \dots, X_m$. For each $i \in [m]$, one can see that X_i takes the form $\left\langle \mathbf{g}_{u(i)}, \frac{1}{|\mathcal{M}_i|} \sum_{j \in \mathcal{M}_i} \mathbf{g}_j \right\rangle$ for some index $u(i) \in [n]$ and some set $\mathcal{M}_i \subset [n]$ with cardinality $|\mathcal{M}_i| = \ell$. Applying the concentration inequality in Lemma 12 with $\mathcal{M} = \mathcal{M}_i$, we obtain that for a fixed $i \in [m]$, with probability at least $1 - n^{-10}$ there holds

$$X_i \leq C\tau^2 \left(\sqrt{\frac{d \log n}{\ell}} + \frac{\log n}{\sqrt{\ell}} + \frac{d}{\ell} \right)$$

for some universal constant $C > 0$. Applying a union bound gives that

$$S_1 \leq 3C\gamma\tau^2 \left(\sqrt{\frac{kd \log n}{n}} + \sqrt{\frac{k}{n}} \log n + \frac{kd}{n} \right)$$

with probability at least $1 - n^{-8}$. The result follows from the condition of the proposition.

5.3. Proof of Proposition 2 In this section we control the term S_2 . Since the matrix $\mathcal{P}_{T^\perp}(\hat{\mathbf{Y}} - \mathbf{Y}^*) = \mathcal{P}_{T^\perp}(\hat{\mathbf{Y}}) = (\mathbf{I} - \mathbf{U}\mathbf{U}^\top)\hat{\mathbf{Y}}(\mathbf{I} - \mathbf{U}\mathbf{U}^\top)$ is positive semidefinite, we can apply the generalized Holder's inequality to obtain

$$S_2 = \left\langle \mathcal{P}_{T^\perp}(\hat{\mathbf{Y}} - \mathbf{Y}^*), \mathbf{G}\mathbf{G}^\top \right\rangle \leq \text{Tr} \left[\mathcal{P}_{T^\perp}(\hat{\mathbf{Y}} - \mathbf{Y}^*) \right] \cdot \|\mathbf{G}\mathbf{G}^\top\|_{\text{op}}.$$

For the second RHS term above, we have the equality $\|\mathbf{G}\mathbf{G}^\top\|_{\text{op}} = \|\mathbf{G}\|_{\text{op}}^2$, which can be established by taking the Singular Value Decomposition (SVD) of $\mathbf{G}$. For the first RHS term, we have the identity

$$\begin{aligned} \text{Tr} \left[\mathcal{P}_{T^\perp}(\hat{\mathbf{Y}} - \mathbf{Y}^*) \right] &= \text{Tr} \left[(\mathbf{I} - \mathbf{U}\mathbf{U}^\top) (\hat{\mathbf{Y}} - \mathbf{Y}^*) (\mathbf{I} - \mathbf{U}\mathbf{U}^\top) \right] \\ &\stackrel{(i)}{=} \text{Tr} \left[(\mathbf{I} - \mathbf{U}\mathbf{U}^\top) (\hat{\mathbf{Y}} - \mathbf{Y}^*) \right] \\ &\stackrel{(ii)}{=} \text{Tr} \left[-\mathbf{U}\mathbf{U}^\top (\hat{\mathbf{Y}} - \mathbf{Y}^*) \right] \\ &= \left\langle -\frac{1}{\ell} \mathbf{Y}^*, \hat{\mathbf{Y}} - \mathbf{Y}^* \right\rangle \\ &\stackrel{(iii)}{=} \frac{\gamma}{2\ell}, \end{aligned}$$

where step (i) holds since the trace is invariant under cyclic permutation and the matrix $\mathbf{I} - \mathbf{U}\mathbf{U}^\top$ is idempotent, step (ii) holds by Fact 2(a), and step (iii) follows from Fact 2(f). Combining the above results proves that $S_2 \leq \frac{\gamma}{\ell} \|\mathbf{G}\|_{\text{op}}^2$, the first inequality in the proposition.

We further note the identity $\|\mathbf{G}\|_{\text{op}}^2 = \|\mathbf{G}^\top \mathbf{G}\|_{\text{op}}$, which again follows from the SVD of $\mathbf{G}$. The spectral norm of $\mathbf{G}^\top \mathbf{G} = \sum_{i \in [n]} \mathbf{g}_i \mathbf{g}_i^\top$ can be controlled using the following standard result.

LEMMA 3 (Lemma A.2 in [42]). *We have $\|\sum_{i=1}^n \mathbf{g}_i \mathbf{g}_i^\top\|_{\text{op}} \leq 6\tau^2(n+d)$ with probability at least $1 - e^{-n/2}$.*

Applying Lemma 3, we obtain that with probability at least $1 - e^{-n/2}$, there holds the inequality

$$S_2 \leq \frac{\gamma}{\ell} \cdot 6\tau^2(n+d) = 6\gamma\tau^2 k \left(1 + \frac{d}{n}\right).$$

The second part of the proposition then follows from the assumption that $s^2 \geq Ck \left(\frac{d}{n} + 1\right)$.

5.4. Proof of Proposition 3 It can be seen that

$$((\mathbb{E}\mathbf{H})(\mathbb{E}\mathbf{H})^\top)_{ij} = \begin{cases} \|\boldsymbol{\mu}_{\sigma^*(i)}\|_2^2 & \text{if } \sigma^*(i) = \sigma^*(j), \\ \langle \boldsymbol{\mu}_{\sigma^*(i)}, \boldsymbol{\mu}_{\sigma^*(j)} \rangle & \text{otherwise.} \end{cases}$$

In view of Equation (8), we may partition the matrix $\hat{\mathbf{Y}} - \mathbf{Y}^*$ into k^2 blocks of size $\ell \times \ell$. Define $T_{ab} := \sum_{i \in C_a^*, j \in C_b^*} (\hat{\mathbf{Y}} - \mathbf{Y}^*)_{ij}$ to be the sum of entries within the (a, b) -th block. We have

$$\begin{aligned} S_4 &= \sum_{a \in [k]} T_{aa} \|\boldsymbol{\mu}_a\|_2^2 + 2 \sum_{a, b \in [k]: a < b} T_{ab} \langle \boldsymbol{\mu}_a, \boldsymbol{\mu}_b \rangle \\ &\stackrel{(i)}{=} \sum_{a, b \in [k]: a < b} T_{ab} \Delta_{ab}^2 \\ &\stackrel{(ii)}{=} \frac{1}{2} \sum_{a, b \in [k]: a \neq b} T_{ab} \Delta_{ab}^2, \end{aligned}$$

where step (i) follows from Fact 2(d), and step (ii) holds since $T_{ab} = T_{ba}$ as implied by the symmetry of $\hat{\mathbf{Y}} - \mathbf{Y}^*$. This proves the first part of the proposition. The second part of the proposition, $S_4 \leq -\frac{1}{4}\Delta^2\gamma$, follows from combining Fact 2(f), the fact that $\Delta_{ab} \geq \Delta$ for $a \neq b$ by definition, and the fact that $T_{ab} \geq 0$ for $a \neq b$ as implied by Fact 2(b).

5.5. Proof of Lemma 1 Let $\mathbf{B} := \hat{\mathbf{Y}} - \mathbf{Y}^*$. Recall the definitions of S_3 and S_4 from Equation (16). We have

$$\begin{aligned} S_3 &= \sum_j \sum_a \sum_{i \in C_a} B_{ji} \langle \boldsymbol{\mu}_a, \mathbf{g}_j \rangle \\ &= \ell \sum_j \sum_a \langle \boldsymbol{\mu}_a, \mathbf{g}_j \rangle \left(\frac{1}{\ell} \sum_{i \in C_a^*} B_{ji} \right) \\ &= \ell \sum_j \sum_{a \neq \sigma^*(j)} \langle \boldsymbol{\mu}_a - \boldsymbol{\mu}_{\sigma^*(j)}, \mathbf{g}_j \rangle \left(\frac{1}{\ell} \sum_{i \in C_a^*} B_{ji} \right), \end{aligned}$$

where the last step holds by Fact 2(c). On the other hand, by Proposition 3 we have

$$S_4 = -\ell \sum_j \sum_{a \neq \sigma^*(j)} \frac{1}{2} \Delta_{\sigma^*(j),a}^2 \left(\frac{1}{\ell} \sum_{i \in C_a^*} B_{ji} \right).$$

Summing up the last two equations gives

$$\frac{1}{\ell} \left(S_3 + \frac{1}{4} S_4 \right) = \sum_j \sum_{a \neq \sigma^*(j)} \left(\langle \boldsymbol{\mu}_a - \boldsymbol{\mu}_{\sigma^*(j)}, \mathbf{g}_j \rangle - \frac{1}{8} \Delta_{\sigma^*(j),a}^2 \right) \left(\frac{1}{\ell} \sum_{i \in C_a^*} B_{ji} \right).$$

Note that the quantity inside the first parenthesis above is exactly β_{ja} by definition in Equation (18). This completes the proof of the lemma.

5.6. Proof of Lemma 2 Recall that $\eta(\mathbf{F})$ defined in Equation (10) is the objective of the Oracle IP, the set $\mathcal{F} := \{\mathbf{F} \in \{0,1\}^{n \times k} : \mathbf{F}\mathbf{1}_k = \mathbf{1}_n\}$ contains all assignment matrices feasible to the Oracle IP, and $\mathbf{F}^* \in \mathcal{F}$ is the true assignment matrix, that is, $F_{ja}^* = \mathbb{I}\{a = \sigma^*(j)\}$ for all $j \in [n], a \in [k]$.

Let us reparameterize the integer program IP_1 in (23) by a change of variable. Consider any feasible solution $\mathbf{X} \in \{0,1\}^{n \times k}$ of IP_1 . For each $j \in [n]$, we may fix $X_{j,\sigma^*(j)} = -\sum_{a \neq \sigma^*(j)} X_{ja}$. Doing so does not affect the feasibility and objective value of $\mathbf{X}$ with respect to IP_1 . Define the new variable $\mathbf{F} := \mathbf{F}^* + \mathbf{X}$. The objective function and constraints of the original variable $\mathbf{X}$ can be mapped to those of $\mathbf{F}$. In particular, for the objective we have the identity

$$\sum_j \sum_{a \neq \sigma^*(j)} X_{ja} = \sum_j \sum_{a \neq \sigma^*(j)} (F_{ja} - F_{ja}^*) = \frac{1}{2} \|\mathbf{F} - \mathbf{F}^*\|_1.$$

For the constraints we have the equivalency

$$\left. \begin{aligned} X_{ja} &\in \{0,1\}, \quad \forall j \in [n], a \neq \sigma^*(j) \\ \sum_{a \neq \sigma^*(j)} X_{ja} &\leq 1, \quad \forall j \in [n] \\ X_{j,\sigma^*(j)} &= -\sum_{a \neq \sigma^*(j)} X_{ja}, \quad \forall j \in [n] \end{aligned} \right\} \iff \mathbf{F} \in \mathcal{F};$$

and

$$\sum_j \sum_{a \neq \sigma^*(j)} \beta_{ja} X_{ja} \stackrel{(i)}{=} \sum_j \sum_a \beta_{ja} X_{ja} = \sum_j \sum_a \beta_{ja} F_{ja} - \sum_j \sum_a \beta_{ja} F_{ja}^* \stackrel{(ii)}{=} \sum_j \sum_a \beta_{ja} F_{ja},$$

where steps (i) and (ii) both follow from the fact that $\beta_{j,\sigma^*(j)} = 0, \forall j$. It follows that IP_1 has the same optimal value as another integer program with the variable $\mathbf{F}$:

$$\text{IP}_1 = \max \left\{ \frac{1}{2} \|\mathbf{F} - \mathbf{F}^*\|_1 : \sum_j \sum_a \beta_{ja} F_{ja} \geq 0, \mathbf{F} \in \mathcal{F} \right\} := \text{IP}_2. \quad (24)$$

We simplify the first constraint in IP_2 . Recall that $\Delta_{\sigma^*(j),a}^2 = \|\boldsymbol{\mu}_{\sigma^*(j)} - \boldsymbol{\mu}_a\|_2^2$ is the separation between clusters $\sigma^*(j)$ and a , and $\bar{\mathbf{h}}_j := \boldsymbol{\mu}_{\sigma^*(j)} + (2\kappa)^{-1} \mathbf{g}_j, j \in [n]$ are data points generated from SGMM with augmented variance. By definition of β_{ja} , we have

$$\begin{aligned} \beta_{ja} &= \langle \boldsymbol{\mu}_a - \boldsymbol{\mu}_{\sigma^*(j)}, \mathbf{g}_j \rangle - \kappa \Delta_{\sigma^*(j),a}^2 \\ &= \kappa \left(-\|\bar{\mathbf{h}}_j - \boldsymbol{\mu}_a\|_2^2 + \|(2\kappa)^{-1} \mathbf{g}_j\|_2^2 \right), \end{aligned}$$

where the last step can be verified by plugging the expressions for $\Delta_{\sigma^*(j),a}^2$ and $\bar{\mathbf{h}}_j$. It follows that for each $\mathbf{F} \in \mathcal{F}$, we have

$$\begin{aligned} \sum_j \sum_a \beta_{ja} F_{ja} &= \kappa \sum_j \sum_a \left(-\|\bar{\mathbf{h}}_j - \boldsymbol{\mu}_a\|_2^2 + \|(2\kappa)^{-1} \mathbf{g}_j\|_2^2 \right) F_{ja} \\ &\stackrel{(i)}{=} \kappa \left(-\sum_j \sum_a \|\bar{\mathbf{h}}_j - \boldsymbol{\mu}_a\|_2^2 F_{ja} + \sum_j \|(2\kappa)^{-1} \mathbf{g}_j\|_2^2 \sum_a F_{ja}^* \right) \\ &= \kappa \left(-\sum_j \sum_a \|\bar{\mathbf{h}}_j - \boldsymbol{\mu}_a\|_2^2 F_{ja} + \sum_j \sum_a \|\bar{\mathbf{h}}_j - \boldsymbol{\mu}_{\sigma^*(j)}\|_2^2 F_{ja}^* \right) \\ &\stackrel{(ii)}{=} \kappa \left(-\sum_j \sum_a \|\bar{\mathbf{h}}_j - \boldsymbol{\mu}_a\|_2^2 F_{ja} + \sum_j \sum_a \|\bar{\mathbf{h}}_j - \boldsymbol{\mu}_a\|_2^2 F_{ja}^* \right), \end{aligned}$$

where step (i) holds because $\sum_a F_{ja} = 1 = \sum_a F_{ja}^*, \forall j$, and step (ii) holds because $F_{ja}^* \neq 0$ only if $a = \sigma^*(j)$. Recalling the definition $\eta(\mathbf{F}) := \sum_j \sum_a \|\bar{\mathbf{h}}_j - \boldsymbol{\mu}_a\|_2^2 F_{ja}$ given in Section 4.1, we have the compact expression

$$\sum_j \sum_a \beta_{ja} F_{ja} = \kappa (\eta(\mathbf{F}^*) - \eta(\mathbf{F})), \quad \forall \mathbf{F} \in \mathcal{F}. \quad (25)$$

Therefore, the first constraint in IP_2 is satisfied if and only if $\eta(\mathbf{F}) \leq \eta(\mathbf{F}^*)$. Substituting into Equation (24), we complete the proof of the lemma.

6. Proof of Theorems 2 and 5 We first prove Theorem 2. The proof of Theorem 5 follows from a simple extension.

We define the following shorthand for the error of the Oracle IP:

$$\widehat{\delta} := \max \left\{ \frac{1}{2} \|\mathbf{F} - \mathbf{F}^*\|_1 : \eta(\mathbf{F}) \leq \eta(\mathbf{F}^*), \mathbf{F} \in \mathcal{F} \right\}.$$

Note that $\widehat{\delta}$ takes integer values in $[0, n]$. If $\widehat{\delta} = 0$ then the theorem follows trivially. We therefore assume that $\widehat{\delta} \in \{1, 2, \dots, n\}$. Let $\widehat{\mathbf{F}} \in \{0, 1\}^{n \times k}$ be an optimal solution to the

above maximization problem. Define the matrix $\widehat{\mathbf{M}} \in \{0, 1\}^{n \times k}$ via $\widehat{M}_{ja} := \widehat{F}_{ja}(1 - F_{ja}^*)$. We have

$$\begin{aligned} 0 &\leq \kappa \left(\eta(\mathbf{F}^*) - \eta(\widehat{\mathbf{F}}) \right) \\ &\stackrel{(i)}{=} \sum_{j \in [n]} \sum_{a \in [k]} \beta_{ja} \widehat{F}_{ja} \stackrel{(ii)}{=} \sum_{j \in [n]} \sum_{a \in [k]} \beta_{ja} \widehat{M}_{ja}, \end{aligned} \quad (26)$$

where step (i) holds due to the expression (25) of $\eta(\cdot)$, and step (ii) holds since $\widehat{F}_{ja} = \widehat{M}_{ja}$ if $a \neq \sigma^*(j)$ and $\beta_{ja} = 0$ if $a = \sigma^*(j)$. Note that the random variable β_{ja} defined in (18) has mean $-\kappa \Delta_{\sigma^*(j),a}^2$ and sub-Gaussian norm $\|\beta_{ja}\|_{\psi_2} \leq \tau \Delta_{\sigma^*(j),a}$, and that $\{\beta_{ja}\}$ are independent across j . To control the RHS of (26), we make use of the following uniform concentration result, which is proved in Section 6.1.

LEMMA 4. *Let $\mathbf{Z} \in \mathbb{R}^{n \times k}$ be a matrix with independent rows, such that for each $(j, a) \in [n] \times [k]$, Z_{ja} is a zero-mean sub-Gaussian random variable with sub-Gaussian norm no larger than ρ_{ja} . Then for some universal constant $C > 0$, we have with probability at least $1 - \frac{1.5}{n}$,*

$$\begin{aligned} \sum_j \sum_a |Z_{ja}| M_{ja} &\leq C \sqrt{t \left(\sum_j \sum_a \rho_{ja}^2 M_{ja} \right) \log(3nk/t)}, \\ &\quad \forall t \in \{1, 2, \dots, n\}; \forall \mathbf{M} \in \{0, 1\}^{n \times k} : \mathbf{M} \mathbf{1}_k \leq \mathbf{1}_n, \|\mathbf{M}\|_1 = t. \end{aligned} \quad (27)$$

It is easy to verify that $\widehat{\mathbf{M}} \mathbf{1}_k \leq \mathbf{1}_n$ and $\|\widehat{\mathbf{M}}\|_1 = \frac{1}{2} \|\widehat{\mathbf{F}} - \mathbf{F}^*\|_1 = \widehat{\delta}$. Therefore, we can apply Lemma 4 with $Z_{ja} = \beta_{ja} + \kappa \Delta_{\sigma^*(j),a}^2$ and $\rho_{ja} = \tau \Delta_{\sigma^*(j),a}$ to bound the RHS of (26). Doing so gives that with probability at least $1 - \frac{1.5}{n}$, we have

$$0 \leq C \sqrt{\widehat{\delta} \tau^2 \left(\sum_j \sum_a \Delta_{\sigma^*(j),a}^2 \widehat{M}_{ja} \right) \log(3nk/\widehat{\delta})} - \kappa \sum_j \sum_a \Delta_{\sigma^*(j),a}^2 \widehat{M}_{ja}.$$

Now, for the sake of deriving a contradiction, assume that $\widehat{\delta} > 3nke^{-s^2/C_0^2}$ for a fixed constant $C_0 > C/\kappa$. Continuing from the RHS of the last display equation, we obtain that

$$\begin{aligned} 0 &\leq C \sqrt{\widehat{\delta} \tau^2 \left(\sum_j \sum_a \Delta_{\sigma^*(j),a}^2 \widehat{M}_{ja} \right) \frac{s^2}{C_0^2}} - \kappa \sum_j \sum_a \Delta_{\sigma^*(j),a}^2 \widehat{M}_{ja} \\ &\leq \left(\frac{C}{C_0} - \kappa \right) \cdot \sum_j \sum_a \Delta_{\sigma^*(j),a}^2 \widehat{M}_{ja}, \end{aligned}$$

where the last step holds since $\tau^2 s^2 = \Delta^2$ and $\Delta^2 \widehat{\delta} = \Delta^2 \sum_{j,a} \widehat{M}_{ja} \leq \sum_{j,a} \Delta_{\sigma^*(j),a}^2 \widehat{M}_{ja}$ by definition. Since $C_0 > C/\kappa$ and $\sum_j \sum_a \Delta_{\sigma^*(j),a}^2 \widehat{M}_{ja} > 0$, the RHS above is negative, which is a contradiction. Therefore, the previous assumption is false and we must have

$$\widehat{\delta} \leq 3nke^{-s^2/C_0^2} \stackrel{(i)}{\leq} 3nk \cdot \frac{1}{3k} \cdot e^{-s^2/(2C_0^2)} = ne^{-s^2/(2C_0^2)},$$

where step (i) holds under the SNR condition $s^2 \gtrsim k$ assumed in Theorem 2. The theorem then follows from the fact that $\|\mathbf{F}^*\|_1 = n$.

Now that we have proved Theorem 2, the proof of Theorem 5 follows exactly the same argument, except that all instances of $\mathbf{g}_j, \beta_{ja}, \eta$ above are replaced by their semi-random versions $\tilde{\mathbf{g}}_j, \tilde{\beta}_{ja}, \tilde{\eta}$ (here $\tilde{\beta}_{ja}$ is equal to β_{ja} defined in (18) but with $\mathbf{g}_j$ replaced by $\tilde{\mathbf{g}}_j$), and accordingly each Z_{ja} is replaced by $\tilde{Z}_{ja} = \tilde{\beta}_{ja} + \kappa \Delta_{\sigma^*(j),a}^2 = \langle \boldsymbol{\mu}_a - \boldsymbol{\mu}_{\sigma^*(j)}, \tilde{\mathbf{g}}_j \rangle$. In the proof we use the distribution of the data only when invoking Lemma 4, whose conclusion (27) remains valid for $\{\tilde{Z}_{ja}\}$. To see this, recall that $\tilde{\mathbf{g}}_j = \alpha_j \mathbf{g}_j$ for some $\alpha_j \in [0, 1]$. It follows that $|\tilde{Z}_{ja}| = \alpha_j |Z_{ja}| \leq |Z_{ja}|$, so the LHS of (27) does not increase.

6.1. Proof of Lemma 4 We define the quantities

$$L(\mathbf{M}, \mathbf{b}) := \sum_{j,a} b_j Z_{ja} M_{ja}, \quad \text{for each } \mathbf{M} \in \{0, 1\}^{n \times k}, \mathbf{b} \in \{\pm 1, 0\}^n,$$

$$R(\mathbf{M}, t) := C \sqrt{t \left(\sum_{j,a} \rho_{ja}^2 M_{ja} \right) \log(3nk/t)}, \quad \text{for each } \mathbf{M} \in \{0, 1\}^{n \times k}, t \in [n]$$

and the sets

$$\mathcal{M}(t) := \left\{ \mathbf{M} \in \{0, 1\}^{n \times k} : \mathbf{M} \mathbf{1}_k \leq \mathbf{1}_n, \|\mathbf{M}\|_1 = t \right\}, \quad \text{for each } t \in [n],$$

$$\mathcal{B}(\mathbf{M}, t) := \left\{ \mathbf{b} \in \{\pm 1, 0\}^n : b_j = 0 \text{ if } M_{ja} = 0, \forall a \in [k] \right\}, \quad \text{for each } \mathbf{M} \in \mathcal{M}(t), t \in [n].$$

We begin by bounding the probability

$$\alpha_t := \mathbb{P} \left\{ \exists \mathbf{M} \in \mathcal{M}(t) : \sum_{j,a} |Z_{ja}| M_{ja} > R(\mathbf{M}, t) \right\}$$

for each integer $t \in [n]$. To this end, note that $|Z_{ja}| = \max_{b_j \in \{\pm 1\}} \{b_j Z_{ja}\}$, hence

$$\begin{aligned} \alpha_t &\leq \mathbb{P} \left\{ \exists \mathbf{M} \in \mathcal{M}(t), \exists \mathbf{b} \in \mathcal{B}(\mathbf{M}, t) : L(\mathbf{M}, \mathbf{b}) > R(\mathbf{M}, t) \right\} \\ &\leq \sum_{\mathbf{M} \in \mathcal{M}(t)} \sum_{\mathbf{b} \in \mathcal{B}(\mathbf{M}, t)} \mathbb{P} \left\{ L(\mathbf{M}, \mathbf{b}) > R(\mathbf{M}, t) \right\}. \end{aligned} \quad (28)$$

By assumption, the Z_{ja} 's are independent zero-mean sub-Gaussian random variables, so the squared sub-Gaussian norm of the sum $L(\mathbf{M}, \mathbf{b})$ is at most $C_{\psi_2} \sum_{j,a} \rho_{ja}^2 M_{ja}$ where $C_{\psi_2} > 0$ is a universal constant. We can therefore apply Hoeffding's inequality (Lemma 11) to bound each summand on the RHS of (28):

$$\begin{aligned} \mathbb{P} \{ L(\mathbf{M}, \mathbf{b}) > R(\mathbf{M}, t) \} &\leq \exp \left\{ - \frac{c_0 C^2 t \left(\sum_{j,a} \rho_{ja}^2 M_{ja} \right) \log(3nk/t)}{C_{\psi_2} \sum_{j,a} \rho_{ja}^2 M_{ja}} \right\} \\ &\leq \exp \{ -4t \log(3nk/t) \}, \end{aligned}$$

where $c_0 > 0$ is a universal constant. Plugging this back to (28), we have for each $t \in [n]$:

$$\begin{aligned} \alpha_t &\leq \sum_{\mathbf{M} \in \mathcal{M}(t)} \sum_{\mathbf{b} \in \mathcal{B}(\mathbf{M}, t)} \exp \{ -4t \log(3nk/t) \} \\ &= \binom{n}{t} k^t \cdot 2^t \cdot \exp \{ -4t \log(3nk/t) \} \\ &\leq \left(\frac{ne}{t} \right)^t k^t \cdot 2^t \cdot \exp \{ -4t \log(3nk/t) \} \\ &\stackrel{(i)}{\leq} \exp \{ 2t \log(3nk/t) - 4t \log(3nk/t) \} \\ &\leq \left(\frac{t}{3nk} \right)^t, \end{aligned}$$

where step (i) follows from $t \leq t \log(3nk/t)$ for $t \in [n]$ and $k \geq 2$.

With the above bound on α_t and a union bound over $t \in [n]$, we can control the probability that the conclusion of the lemma fails to hold:

$$\mathbb{P}\left\{\exists t \in [n], \exists \mathbf{M} \in \mathcal{M}(t) : \sum_{j,a} |Z_{ja}| M_{ja} > R(\mathbf{M}, t)\right\} \leq \sum_{t=1}^n \alpha_t \leq \sum_{t=1}^n \left(\frac{t}{3nk}\right)^t.$$

The proof is complete if we can show that the last RHS is at most $\frac{1.5}{n}$. Note that

$$\sum_{t=1}^n \left(\frac{t}{3nk}\right)^t \leq \frac{1}{3n} + \sum_{t=2}^n \left(\frac{t}{3n}\right)^t \leq \frac{1}{3n} + n \cdot \max_{t \in [2, n]} \left(\frac{t}{3n}\right)^t.$$

Hence it suffices to show that for all $t \in [2, n]$, there holds $\left(\frac{t}{3n}\right)^t \leq \frac{1}{n^2}$ or equivalently $f(t) := t(\log 3n - \log t) \geq 2 \log n$. For $t \in [2, n]$, the function f has derivative

$$f'(t) = \log 3n - \log t - 1 \geq \log 3n - \log(n) - 1 = \log 3 - 1 \geq 0.$$

Therefore, $f(t)$ is non-decreasing for $t \in [2, n]$ and thus $f(t) \geq f(2) = 2 \log 3n - 2 \log 2 \geq 2 \log n$ as desired. We conclude that $\sum_{t=1}^n \left(\frac{t}{3nk}\right)^t \leq \frac{1.5}{n}$, thereby completing the proof of Lemma 4.

7. Conclusion In this paper, we study the performance of SDP relaxation for clustering SGMM and its semi-random version. Our analysis proceeds in two steps: (a) bound the clustering error of the SDP by that of the idealized Oracle IP; (b) show that the error of the Oracle IP decays exponentially in the SNR. As mentioned, this two-step framework allows for a decoupling of the computational and statistical mechanisms that drive the performance of the SDP approach. The oracle bound in step (a) represents a fundamental performance limit of the SDP relaxation. We expect that further progress in understanding SDP relaxations is likely to come from improvements in this step. On the other hand, step (b) is problem-specific and makes use of the probabilistic structures of the model. By modifying and sharpening this step, one may generalize our results to other variants of SGMM.

Our work points to several interesting future directions. An immediate problem is to obtain tighter, less pessimistic bounds for mixtures with inhomogeneous components, which may have different pairwise separation and variance along different directions. It is also of interest to study other forms of robustness of SDP relaxations, e.g., in the presence of arbitrary outliers and model misspecification. Other directions worth exploring include obtaining better constants in error bounds, identifying sharp thresholds for different types of recovery, and developing scalable computational procedures for solving the SDP.

Last but not least, another interesting future direction is to better understand the relation between the clustering problem—which we have focused on—and the related parameter estimation problem (see Section 1). Note that a guarantee for one problem can be converted to a guarantee for the other. While this conversion may not be tight in general, it is nevertheless interesting to compare our SNR condition with those in existing work on both clustering and parameter estimation. The work of [56] shows that the SNR condition $s^2 \gtrsim \log k$ is sufficient and necessary to achieve a constant error in parameter estimation with polynomial sample complexity. Note that the algorithm in [56] requires an initialization procedure that runs in time exponential in k , the number of clusters. In the special case of Gaussian Mixture Model, the work [58] proves that spectral methods achieve exact cluster recovery under the SNR condition $s^2 \gtrsim \sqrt{k \log n}$; see Section 4.5.2 for

further discussion of the work [58]. These results seem to suggest that our SNR condition $s^2 \gtrsim k + \log n$ may have a suboptimal dependence on k . It is of interest to investigate whether this potential suboptimality is intrinsic to the SDP relaxation approach or can be avoided by a tighter analysis.

Appendix A: Proof of Theorem 3 We only need to prove the first part of the theorem. The second part follows immediately from the first part and Theorem 1.

The proof follows a similar strategy as those of [45, Theorem 17 and Lemma 18] and is divided into two lemmas. Recall that Algorithm 1 outputs a collection of sets $\{B_t\}_{t \geq 1}$ that are not necessarily equal-sized. The first lemma characterizes the quality of these sets. Define the shorthands $k' := |\{B_t\}_{t \geq 1}|$ (which satisfies $k' \geq k$) and $\epsilon := \|\hat{\mathbf{Y}} - \mathbf{Y}^*\|_1 / \|\mathbf{Y}^*\|_1$.

LEMMA 5. *There exists a partial matching π' (i.e., an injection from $[k]$ to $[k']$) and a universal constant $C > 0$ such that*

$$\left| \bigcup_{a \in [k]} C_a^* \cap B_{\pi'(a)} \right| \geq (1 - C\epsilon) n.$$

The proof is given in Section A.1. Building on the above lemma, the next lemma further characterizes the (equal-sized) clusters $\{U_t\}_{t \in [k]}$ obtained in Algorithm 2.

LEMMA 6. *There exists a permutation π on $[k]$ and a universal constant $C > 0$ such that*

$$\left| \bigcup_{a \in [k]} C_a^* \cap U_{\pi(a)} \right| \geq (1 - C\epsilon) n.$$

The proof is given in Section A.2. In light of the last lemma and the fact that

$$\text{err}(\hat{\boldsymbol{\sigma}}, \boldsymbol{\sigma}^*) = 1 - \frac{1}{n} \max_{\pi \in S_k} \left| \bigcup_{a \in [k]} C_a^* \cap U_{\pi(a)} \right|,$$

we obtain the bound $\text{err}(\hat{\boldsymbol{\sigma}}, \boldsymbol{\sigma}^*) \leq C\epsilon$ as stated in the first part of Theorem 3.

A.1. Proof of Lemma 5 Let $\mathbf{y}_a \in \mathbb{R}^n$ be one of the ℓ identical rows of $\mathbf{Y}^*$ whose row indices are in C_a^* . Define the sets

$$G_a := \left\{ i \in C_a^* : \|\hat{\mathbf{Y}}_{i\bullet} - \mathbf{y}_a\|_1 \leq \frac{\ell}{8} \right\}, \quad \forall a \in [k]$$

and let $G := \bigcup_{a \in [k]} G_a$ and $H := V \setminus G$, where $V := [n]$.

A partial matching between the sets $\{C_a^*\}_{a \in [k]}$ and $\{B_t\}_{t \in [k']}$ is given by an injective function π' from $[k]$ to $[k']$. We construct such a partial matching by matching each cluster C_a^* with the first B_t that intersects G_a ; i.e., we set $\pi'(a) = \min\{t \in [k'] : B_t \cap G_a \neq \emptyset\}$. Since each $i \in [n]$ belongs to some B_t , every C_a^* is matched with some B_t . Moreover, we show below that π' is indeed injective as it cannot match two distinct clusters C_a^* and C_b^* with the same B_t .

CLAIM 1. *For each $a \in [k]$ and $t \in [k']$ such that $t = \pi'(a)$, we have that $B_t \cap G_b = \emptyset$ for any $b \in [k] \setminus \{a\}$ and that $B_t \subset G_a \cup H$.*

Proof. Suppose that there exist B_t and $b \in [k] \setminus \{a\}$ such that $B_t \cap G_b \neq \emptyset$. Let $u \in B_t \cap G_a$ and $v \in B_t \cap G_b$. Since G_a and G_b are disjoint, we know that $u \neq v$. Let $w \in B_t$. Then we have

$$\|\hat{\mathbf{Y}}_{u\bullet} - \hat{\mathbf{Y}}_{w\bullet}\|_1 \leq \frac{\ell}{4} \quad \text{and} \quad \|\hat{\mathbf{Y}}_{v\bullet} - \hat{\mathbf{Y}}_{w\bullet}\|_1 \leq \frac{\ell}{4},$$

whence

$$\|\hat{\mathbf{Y}}_{u\bullet} - \hat{\mathbf{Y}}_{v\bullet}\|_1 \leq \|\hat{\mathbf{Y}}_{u\bullet} - \hat{\mathbf{Y}}_{w\bullet}\|_1 + \|\hat{\mathbf{Y}}_{v\bullet} - \hat{\mathbf{Y}}_{w\bullet}\|_1 \leq \frac{\ell}{2}.$$

This implies that

$$\begin{aligned} \|\mathbf{y}_a - \mathbf{y}_b\|_1 &\leq \|\mathbf{y}_a - \hat{\mathbf{Y}}_{u\bullet}\|_1 + \|\hat{\mathbf{Y}}_{u\bullet} - \hat{\mathbf{Y}}_{v\bullet}\|_1 + \|\mathbf{y}_b - \hat{\mathbf{Y}}_{v\bullet}\|_1 \\ &\leq \frac{\ell}{8} + \frac{\ell}{2} + \frac{\ell}{8} < \ell, \end{aligned}$$

which is a contradiction to the fact that $\|\mathbf{y}_a - \mathbf{y}_b\|_1 = 2\ell$. To complete the proof, we note that for any $i \in B_t$ we have either $i \in G_a$ or $i \in H$, hence $B_t \subset G_a \cup H$. $\square$

The rest of the proof proceeds by establishing the following three claims.

CLAIM 2. For each $a \in [k]$ and $t \in [k']$ such that $t = \pi'(a)$, we have

$$|B_t \cap C_a^*| \geq |G_a| - |B_t \cap H|. \quad (29)$$

Proof. Fix $i \in G_a$ for some $a \in [k]$. For any $j \in G_a$ we have $j \in B(i)$ since

$$\|\hat{\mathbf{Y}}_{i\bullet} - \hat{\mathbf{Y}}_{j\bullet}\|_1 \leq \|\mathbf{y}_a - \hat{\mathbf{Y}}_{i\bullet}\|_1 + \|\mathbf{y}_a - \hat{\mathbf{Y}}_{j\bullet}\|_1 \leq \frac{\ell}{4}.$$

Therefore, by definition we have $|B_t| \geq |B(i)| \geq |G_a|$. It follows that

$$\begin{aligned} |B_t \cap C_a^*| &\stackrel{(i)}{\geq} |B_t \cap G_a| \\ &= |B_t| - |B_t \setminus G_a| \\ &\stackrel{(ii)}{=} |B_t| - |B_t \cap H| \\ &\geq |G_a| - |B_t \cap H|, \end{aligned}$$

where step (i) holds since $G_a \subset C_a^*$ and step (ii) holds since $B_t \subset G_a \cup H$ by the previous claim. $\square$

CLAIM 3. We have

$$\sum_{(t,a): t=\pi'(a)} |B_t \cap C_a^*| \geq |V| - 2|H|.$$

Proof. Summing both sides of Equation (29) over $\{(t,a) : t = \pi'(a)\}$, we obtain

$$\begin{aligned} \sum_{(t,a): t=\pi'(a)} |B_t \cap C_a^*| &= \sum_{a \in [k]} |G_a| - \sum_{(t,a): t=\pi'(a)} |B_t \cap H| \\ &\geq \sum_{a \in [k]} |G_a| - \sum_{t \geq 1} |B_t \cap H| \\ &\stackrel{(i)}{=} |G| - |V \cap H| \\ &= |V| - 2|H|, \end{aligned}$$

where step (i) holds since the sets $\{B_t \cap H\}$ are disjoint and $\bigcup_{t \geq 1} B_t = V$. $\square$

CLAIM 4. There exists a universal constant $C > 0$ such that $|H| \leq C\epsilon n$.

Proof. We have

$$|H| \cdot \frac{\ell}{8} \leq \sum_{i \in H} \|\hat{\mathbf{Y}}_{i\bullet} - \mathbf{y}_{\sigma^*(i)}\|_1 \leq \|\hat{\mathbf{Y}} - \mathbf{Y}^*\|_1 \leq \epsilon \|\mathbf{Y}^*\|_1 = \epsilon \cdot n\ell$$

where the last step follows from the fact that $\|\mathbf{Y}^*\|_1 = n\ell$. $\square$

Combining the last two claims proves Lemma 5.

A.2. Proof of Lemma 6 Without loss of generality, assume that the output of Algorithm 1 is ordered as $|B_1| \geq |B_2| \geq \dots \geq |B_{k'}|$. Consequently, the sets $\{U_t\}_{t \in [k]}$ maintained in Algorithm 2 are such that each U_t consists of B_t and some elements from the sets B_u with $u > k$.

Let π' be the partial matching between $\{C_a^*\}_{a \in [k]}$ and $\{B_t\}_{t \in [k']}$ given by Lemma 5. Define $\pi(a) = \pi'(a)$ for each $a \in [k]$ with $\pi'(a) \leq k$, and extend π to a full permutation on $[k]$ in an arbitrary way. We have

$$\begin{aligned} \left| \bigcup_{(t,a): t=\pi(a)} C_a^* \cap U_t \right| &\geq \left| \bigcup_{(t,a): t=\pi'(a) \leq k} C_a^* \cap B_t \right| \\ &= \left| \bigcup_{(t,a): t=\pi'(a)} C_a^* \cap B_t \right| - \left| \bigcup_{(t,a): t=\pi'(a) > k} C_a^* \cap B_t \right| \\ &\geq (1 - C'\epsilon)n - \left| \bigcup_{(t,a): t=\pi'(a) > k} C_a^* \cap B_t \right|, \end{aligned}$$

where the last step follows from Lemma 5 and $C' > 0$ is a universal constant. Define the sets

$$\begin{aligned} T_1 &:= \{t > k : t = \pi'(a) \text{ for some } a \in [k]\}, \\ T_2 &:= \{t \in [k] : t \neq \pi'(a) \text{ for all } a \in [k]\}. \end{aligned}$$

It is easy to verify that $|T_1| = |T_2|$ and that $|B_{t_1}| \leq |B_{t_2}|$ for each $t_1 \in T_1$ and $t_2 \in T_2$. It follows that

$$\begin{aligned} \left| \bigcup_{(t,a): t=\pi'(a) > k} C_a^* \cap B_t \right| &\leq \left| \bigcup_{t \in T_1} B_t \right| \leq \left| \bigcup_{t \in T_2} B_t \right| \\ &\leq |V| - \left| \bigcup_{(t,a): t=\pi'(a)} C_a^* \cap B_t \right| \\ &\leq C'\epsilon n, \end{aligned}$$

where the last step follows again from Lemma 5. Combining pieces and setting $C := 2C'$, we obtain the desired bound

$$\left| \bigcup_{(t,a): t=\pi(a)} C_a^* \cap U_t \right| \geq (1 - C'\epsilon)n - C'\epsilon n = (1 - C\epsilon)n.$$

Appendix B: Proof of Theorem 4 Let $\tilde{\mathbf{G}} \in \mathbb{R}^{n \times d}$ be the matrix whose j -th row is $\tilde{\mathbf{g}}_j$, and recall that the j -th row of $\mathbb{E}\mathbf{H}$ is $\boldsymbol{\mu}_{\sigma^*(j)}$. Then using the same notations as in Section 5, we have the following analogue of Equation (16):

$$\begin{aligned} 0 \leq & \underbrace{\left\langle \hat{\mathbf{Y}} - \mathbf{Y}^*, \mathcal{P}_T \left(\tilde{\mathbf{G}} \tilde{\mathbf{G}}^\top \right) \right\rangle}_{S_1} + \underbrace{\left\langle \hat{\mathbf{Y}} - \mathbf{Y}^*, \mathcal{P}_{T^\perp} \left(\tilde{\mathbf{G}} \tilde{\mathbf{G}}^\top \right) \right\rangle}_{S_2} \\ & + 2 \underbrace{\left\langle \hat{\mathbf{Y}} - \mathbf{Y}^*, \tilde{\mathbf{G}} (\mathbb{E}\mathbf{H})^\top \right\rangle}_{S_3} + \underbrace{\left\langle \hat{\mathbf{Y}} - \mathbf{Y}^*, (\mathbb{E}\mathbf{H}) (\mathbb{E}\mathbf{H})^\top \right\rangle}_{S_4}. \end{aligned} \quad (30)$$

As will become clear shortly, the impact of the semi-random adversary is essentially limited to the term S_1 . With this effect accounted for, the proof of Theorem 4 is basically the same as that of Theorem 1 for the non-semi-random setting.

We begin by controlling S_1 . This is done in the following proposition, whose proof is given in Section B.1. Recall that $\tilde{O}(\cdot)$ hides a multiplicative factor of $\log(d + \log n)$.

PROPOSITION 4. *If $n \gtrsim k(d + \log n)$, then there exists some universal constant $C > 0$ such that*

$$S_1 \leq \frac{1}{100} \gamma \Delta^2 + \tau^2 \cdot \tilde{O} \left(\frac{nd}{s^2} + n^2 s^2 e^{-s^4/C} \right).$$

with probability at least $1 - 2n^{-1} - 2^{-d}$.

In comparison with its non-semi-random counterpart in Proposition 1, the bound in Proposition 4 has an additional error term that is due to the effect of the adversary.

Controlling the term S_2 is straightforward. In particular, the following proposition establishes a bound that is exactly the same as in the non-semi-random setting (cf. Proposition 2), since the adversary cannot make the bound worse.

PROPOSITION 5. *If $s^2 \geq Ck \left(\frac{d}{n} + 1 \right)$ for some universal constant $C > 0$, then $S_2 \leq \frac{1}{100} \Delta^2 \gamma$ with probability at least $1 - e^{-n/2}$.*

Proof. By the first part of Proposition 2, we have $S_2 \leq \frac{\gamma}{\ell} \|\tilde{\mathbf{G}}\|_{\text{op}}^2$. Since each row of $\tilde{\mathbf{G}}$ is a shrunk version of the corresponding row of $\mathbf{G}$, we have

$$\|\tilde{\mathbf{G}}\|_{\text{op}} = \max_{\mathbf{v} \in \mathbb{R}^d: \|\mathbf{v}\|_2=1} \|\tilde{\mathbf{G}}\mathbf{v}\|_2 \leq \max_{\mathbf{v} \in \mathbb{R}^d: \|\mathbf{v}\|_2=1} \|\mathbf{G}\mathbf{v}\|_2 = \|\mathbf{G}\|_{\text{op}},$$

whence $S_2 \leq \frac{\gamma}{\ell} \|\mathbf{G}\|_{\text{op}}^2$. Applying the second part of Proposition 2 proves the desired bound. $\square$

Note that the SNR condition in Proposition 5 is satisfied by our assumption $n \gtrsim kd$ and $s^2 \gtrsim k$.

The term S_4 is unaffected by the adversary and thus can be bounded as before using Proposition 3, which is re-stated below for readers' convenience.

PROPOSITION 3. *We have $S_4 = -\frac{1}{2} \sum_{a,b \in [k]: a \neq b} \Delta_{ab}^2 [\sum_{i \in C_a^*, j \in C_b^*} (\hat{\mathbf{Y}} - \mathbf{Y}^*)_{ij}]$. Furthermore, $S_4 \leq -\frac{1}{4} \Delta^2 \gamma$.*

We are now ready to prove Theorem 4. Propositions 5 and 3 imply that $S_2 \leq -\frac{1}{2} S_4$ with probability at least $1 - e^{-n/2} \geq 1 - 2n^{-1}$. Plugging into the basic inequality (30), we obtain that $0 \leq S_1 + 2S_3 + \frac{1}{2} S_4$, hence we must have either (a) $0 \leq 2S_3 + \frac{1}{4} S_4$ or (b) $0 \leq S_1 + \frac{1}{4} S_4$. Let us analyze each of these two cases.

Case (a): $0 \leq 2S_3 + \frac{1}{4} S_4$. In this case, we continue the proof in exactly the same way as in Section 5 following Equation (17), but replace $\{\mathbf{g}_j\}$ and η therein with $\{\tilde{\mathbf{g}}_j\}$ and $\tilde{\eta}$. Doing so establishes the desired inequality (15) (without the second RHS term).

Case (b): $0 \leq S_1 + \frac{1}{4} S_4$. Applying Proposition 4 to bound S_1 and Proposition 3 to bound S_4 , we obtain that with probability at least $1 - 2n^{-1} - 2^{-d}$,

$$\gamma \Delta^2 \leq \tau^2 \cdot \tilde{O} \left(\frac{nd}{s^2} + n^2 s^2 e^{-s^4/C} \right).$$

Dividing both sides of the above equation by $\Delta^2 \|\mathbf{Y}^*\|_1 = \Delta^2 \frac{n^2}{k}$ and recalling that $s := \frac{\Delta}{\tau}$, we get

$$\frac{\|\hat{\mathbf{Y}} - \mathbf{Y}^*\|_1}{\|\mathbf{Y}^*\|_1} \leq \tilde{O} \left(\frac{kd}{ns^4} + ke^{-s^4/C} \right) \leq \tilde{O} \left(\frac{kd}{ns^4} + e^{-s^4/C_0} \right),$$

where the last step holds since our assumption $s^2 \gtrsim k$ implies that $ke^{-s^4/C} \leq e^{-s^4/C_0}$ for some universal constant $C_0 \geq C$. We have therefore established the desired inequality (15) (without the first RHS term).

Combining the above two cases completes the proof of Theorem 4.

B.1. Proof of Proposition 4 We first state several technical lemmas concerning the vectors $\{\tilde{\mathbf{g}}_i\}$ produced by the semi-random adversary. Introduce the shorthand

$$d_0 := d + \log n, \quad (31)$$

which can be interpreted as the effective dimension and is motivated by the following two lemmas.

LEMMA 7. *With probability at least $1 - n^{-1}$, for some universal constant $c \geq 1$ we have*

$$\|\tilde{\mathbf{g}}_i\|_2 \leq c\tau\sqrt{d_0} \quad \text{uniformly for all } i \in [n].$$

Proof. By definition of $\tilde{\mathbf{g}}_i$ we have $\|\tilde{\mathbf{g}}_i\|_2 \leq \|\mathbf{g}_i\|_2$. Standard results on the norms of sub-Gaussian vectors (e.g., [59, Theorem 3.1.1]) ensure that $\|\mathbf{g}_i\|_2 \leq \frac{1}{2}c\tau(\sqrt{d} + \sqrt{\log n}) \leq c\tau\sqrt{d_0}$ with probability at least $1 - n^{-2}$. Taking a union bound over all $i \in [n]$ proves the lemma. $\square$

LEMMA 8. *If $\ell \geq 4d_0$, then with probability at least $1 - n^{-1}$, we have*

$$\left\| \frac{1}{\ell} \sum_{i \in C_a^*} \tilde{\mathbf{g}}_i \right\|_2 \leq 3\tau \quad \text{uniformly for all } a \in [k].$$

Proof. Fix $a \in [k]$. We have the expression $\frac{1}{\ell} \sum_{i \in C_a^*} \tilde{\mathbf{g}}_i = \frac{1}{\ell} \mathbf{B}_a \mathbf{D} \mathbf{1}$, where $\mathbf{B}_a \in \mathbb{R}^{d \times \ell}$ is the matrix with columns $\{\mathbf{g}_i, i \in C_a^*\}$, $\mathbf{D} \in \mathbb{R}^{\ell \times \ell}$ is a diagonal matrix with diagonal entries $\{\|\tilde{\mathbf{g}}_i\|_2 / \|\mathbf{g}_i\|_2, i \in C_a^*\}$, and $\mathbf{1} \in \mathbb{R}^\ell$ is the all-one vector. It follows that $\left\| \frac{1}{\ell} \sum_{i \in C_a^*} \tilde{\mathbf{g}}_i \right\|_2 \leq \frac{1}{\ell} \|\mathbf{B}_a\|_{\text{op}} \cdot \|\mathbf{D}\|_{\text{op}} \cdot \|\mathbf{1}\|_2$. But $\|\mathbf{D}\|_{\text{op}} \leq 1$, $\|\mathbf{1}\|_2 = \sqrt{\ell}$, and Lemma 3 guarantees that $\|\mathbf{B}_a\|_{\text{op}} = \sqrt{\|\mathbf{B}_a \mathbf{B}_a^\top\|_{\text{op}}} \leq \sqrt{6\tau^2(\ell + d)}$ with probability at least $1 - e^{-\ell/2} \geq 1 - n^{-2}$ (since $\ell \geq 4 \log n$ by assumption). Combining pieces, we have $\left\| \frac{1}{\ell} \sum_{i \in C_a^*} \tilde{\mathbf{g}}_i \right\|_2 \leq \frac{1}{\ell} \cdot \sqrt{6\tau^2(\ell + d)} \cdot \sqrt{\ell} \leq 3\tau$, where the last step holds since $\ell \geq 4d$ by assumption. Taking a union bound over all $a \in [k]$ proves the lemma. $\square$

To state the next lemma, we define the set $\mathcal{B}_\lambda(\mathbf{v}) := \{i \in [n] : |\langle \tilde{\mathbf{g}}_i, \mathbf{v} \rangle| > \lambda\tau\}$ for each unit vector $\mathbf{v}$ in $\mathbb{R}^d$ and each positive real number λ . Also recall that $O(\cdot)$ hides a multiplicative factor of $\log(d + \log n)$. The following key lemma is proved at the end of this section.

LEMMA 9. *Suppose that $n \geq d$. There exists a universal constant $C > 0$ such that with probability at least $1 - 2^{-d} - n^{-1}$, uniformly for all numbers $\lambda \geq 1$ and all unit vectors $\mathbf{v} \in \mathbb{R}^d$, we have*

$$|\mathcal{B}_\lambda(\mathbf{v})| = \tilde{O} \left(\frac{d}{\lambda^2} + n \exp \left(-\frac{\lambda^2}{C} \right) \right). \quad (32)$$

A similar result to Lemma 9 has appeared in [10, Lemma 2.11]. The main difference between Lemma 9 and the existing result is that we have the exponential term $ne^{-\lambda^2/C}$ in Equation (32). As it will turn out later, the quantity λ plays the role of the SNR s . Therefore, in the low-SNR regime, this exponential term dominates (32) and leads to the final exponential rate in Corollary 2.

We are now ready to prove Proposition 4. Note that the conditions in Lemma 8 and 9 are satisfied under the assumption of Proposition 4. Assume that the conclusions of Lemmas

7–9 hold simultaneously; this event has probability at least $1 - 2n^{-1} - 2^{-d}$, which is the probability claimed in Proposition 4.

We begin by using the definition of $\mathcal{P}_T$ to obtain the decomposition

$$\begin{aligned} \left| \left(\mathcal{P}_T \left(\tilde{\mathbf{G}} \tilde{\mathbf{G}}^\top \right) \right)_{ij} \right| &\leq \underbrace{\left| \left\langle \tilde{\mathbf{g}}_i, \frac{1}{\ell} \sum_{u \in C_{\sigma^*}^*(j)} \tilde{\mathbf{g}}_u \right\rangle \right|}_{T_{ij}} + \underbrace{\left| \left\langle \tilde{\mathbf{g}}_j, \frac{1}{\ell} \sum_{u \in C_{\sigma^*}^*(i)} \tilde{\mathbf{g}}_u \right\rangle \right|}_{T_{ji}} \\ &\quad + \underbrace{\left| \left\langle \frac{1}{\ell} \sum_{w \in C_{\sigma^*}^*(i)} \tilde{\mathbf{g}}_w, \frac{1}{\ell} \sum_{u \in C_{\sigma^*}^*(j)} \tilde{\mathbf{g}}_u \right\rangle \right|}_{S_{ij}} \end{aligned}$$

for each $i, j \in [n]$. It is easy to see that $S_{ij} \leq \frac{1}{\ell} \sum_{w \in C_{\sigma^*}^*(i)} T_{wj}$ by triangle inequality. With these inequalities and introducing the shorthand $D_{ij} := \left| (\hat{\mathbf{Y}} - \mathbf{Y}^*)_{ij} \right|$, we can bound the quantity of interest $S_1 := \langle \hat{\mathbf{Y}} - \mathbf{Y}^*, \mathcal{P}_T(\tilde{\mathbf{G}} \tilde{\mathbf{G}}^\top) \rangle$ as follows:

$$S_1 \leq \sum_{i,j} D_{ij} \cdot (T_{ij} + T_{ji} + S_{ij}) \leq 2 \sum_j \sum_i D_{ij} T_{ij} + \frac{1}{\ell} \sum_j \sum_i \sum_{w \in C_{\sigma^*}^*(i)} D_{wj} T_{ij}, \quad (33)$$

where the last step follows from the symmetry of the matrix $\hat{\mathbf{Y}} - \mathbf{Y}^*$ and a change of indexing.

To proceed, we fix an arbitrary $j \in [n]$. Let $\mathbf{v}_j := \frac{1}{\ell} \sum_{u \in C_{\sigma^*}^*(j)} \tilde{\mathbf{g}}_u / \left\| \frac{1}{\ell} \sum_{u \in C_{\sigma^*}^*(j)} \tilde{\mathbf{g}}_u \right\|_2$ and note that $\mathbf{v}_j$ is a unit vector. Set $\lambda_0 := \frac{\Delta^2}{C_0 \tau^2} = \frac{s^2}{C_0}$ for some sufficiently large constant $C_0 > 0$, and note that $\lambda_0 \geq k \geq 1$ by our assumption on s . Recalling the definition of the set $\mathcal{B}_\lambda(\cdot)$, we consider two cases:

- For each $i \in [n] \setminus \mathcal{B}_{\lambda_0}(\mathbf{v}_j)$, we have

$$T_{ij} = |\langle \tilde{\mathbf{g}}_i, \mathbf{v}_j \rangle| \cdot \left\| \frac{1}{\ell} \sum_{u \in C_{\sigma^*}^*(j)} \tilde{\mathbf{g}}_u \right\|_2 \leq \lambda_0 \tau \cdot 4\tau,$$

where in the last step we use the defining property of the set $[n] \setminus \mathcal{B}_j$ to bound the first term, and use Lemma 8 to bound the second term.

- For $i \in \mathcal{B}_{\lambda_0}(\mathbf{v}_j)$, we have

$$\begin{aligned} \sum_{i \in \mathcal{B}_{\lambda_0}(\mathbf{v}_j)} T_{ij} &\leq \sum_{i \in \mathcal{B}_{\lambda_0}(\mathbf{v}_j)} |\langle \tilde{\mathbf{g}}_i, \mathbf{v}_j \rangle| \cdot 4\tau \\ &= \sum_{i \in \mathcal{B}_{\lambda_0}(\mathbf{v}_j)} \int_0^\infty \mathbb{I}\{\lambda\tau < |\langle \tilde{\mathbf{g}}_i, \mathbf{v}_j \rangle|\} d(\lambda\tau) \cdot 4\tau \\ &= 4\tau^2 \int_0^\infty \sum_{i \in \mathcal{B}_{\lambda_0}(\mathbf{v}_j)} \mathbb{I}\{\lambda\tau < |\langle \tilde{\mathbf{g}}_i, \mathbf{v}_j \rangle|\} d\lambda \\ &= 4\tau^2 \int_0^\infty |\mathcal{B}_\lambda(\mathbf{v}_j) \cap \mathcal{B}_{\lambda_0}(\mathbf{v}_j)| d\lambda. \end{aligned}$$

Applying Lemma 9 to bound the sizes of $\mathcal{B}_\lambda(\mathbf{v}_j)$ and $\mathcal{B}_{\lambda_0}(\mathbf{v}_j)$, we obtain

$$\sum_{i \in \mathcal{B}_{\lambda_0}(\mathbf{v}_j)} T_{ij} = \tilde{O} \left(\tau^2 \int_0^\infty \min \left\{ \frac{d}{\lambda^2} + ne^{-\lambda^2/C}, \frac{d}{\lambda_0^2} + ne^{-\lambda_0^2/C} \right\} d\lambda \right)$$

$$\begin{aligned} &= \tau^2 \tilde{O} \left(\int_{\lambda_0}^{\infty} \left(\frac{d}{\lambda^2} + ne^{-\lambda^2/C} \right) d\lambda + \lambda_0 \cdot \left(\frac{d}{\lambda_0^2} + ne^{-\lambda_0^2/C} \right) \right) \\ &= \tau^2 \tilde{O} \left(\frac{d}{\lambda_0} + n\lambda_0 e^{-\lambda_0^2/C} \right), \end{aligned}$$

where C is the universal constant C given in Lemma 9.

With the above two bounds and the fact that $D_{ij} \leq 1, \forall i, j$ (implied by Fact 2(b)), we can bound the first RHS term in Equation (33) as

$$\begin{aligned} \sum_j \sum_i D_{ij} T_{ij} &\leq \sum_j \sum_{i \in [n] \setminus \mathcal{B}_{\lambda_0}(\mathbf{v}_j)} D_{ij} T_{ij} + \sum_j \sum_{i \in \mathcal{B}_{\lambda_0}(\mathbf{v}_j)} D_{ij} T_{ij} \\ &\leq \left(\sum_j \sum_{i \in [n] \setminus \mathcal{B}_{\lambda_0}(\mathbf{v}_j)} D_{ij} \right) \cdot 4\lambda_0 \tau^2 + \sum_j \left(\sum_{i \in \mathcal{B}_{\lambda_0}(\mathbf{v}_j)} T_{ij} \right) \\ &\lesssim \gamma \cdot \lambda_0 \tau^2 + n \cdot \tau^2 \tilde{O} \left(\frac{d}{\lambda_0} + n\lambda_0 e^{-\lambda_0^2/C} \right). \end{aligned}$$

The second RHS term in Equation (33) can be controlled in a similar fashion and obey the same bound. Combining these bounds and recalling that $\lambda_0 := \frac{\Delta^2}{C_0 \tau^2} = \frac{s^2}{C_0}$, we obtain

$$S_1 \lesssim \gamma \lambda_0 \tau^2 + n \tau^2 \tilde{O} \left(\frac{d}{\lambda_0} + n\lambda_0 e^{-\lambda_0^2/C} \right) \leq \frac{1}{100} \gamma \Delta^2 + \tau^2 \cdot \tilde{O} \left(\frac{nd}{s^2} + n^2 s^2 e^{-s^4/C_1} \right)$$

for some universal constant $C_1 > 0$, thereby proving Proposition 4.

B.1.1. Proof of Lemma 9 Our strategy involves two steps: (i) prove the desired inequality for fixed λ and $\mathbf{v}$, and (ii) establish a uniform bound using an ϵ -net argument.

Step (i): Controlling $|\mathcal{B}_\lambda(\mathbf{v})|$ for fixed λ and $\mathbf{v}$. Set $\epsilon := 1/(2c\sqrt{d_0})$, where c is as defined in Lemma 7 and d_0 is defined in Equation (31). We assume without loss of generality that $c\sqrt{d_0}$ is an integer; otherwise, we replace c in the definition of ϵ with a larger constant so as to fulfill this assumption. Let us first show that the inequality (32) holds for a fixed number $\lambda \geq 1$ and a unit vector $\mathbf{v}$ with probability at least $1 - (cd_0)^{-1}(3/\epsilon)^{-2d}$. Define the indicator random variable $X_i := \mathbb{I}\{|\langle \mathbf{g}_i, \mathbf{v} \rangle| > \lambda\tau\}$ for $i \in [n]$. Since each $\tilde{\mathbf{g}}_i$ is a shrunk version of $\mathbf{g}_i$, the quantity of interest satisfies the bound $|\mathcal{B}_\lambda(\mathbf{v})| \leq \sum_{i \in [n]} X_i$. The last RHS is the sum of independent Bernoulli RVs, where

$$\mu := \mathbb{E} \left[\sum_{i \in [n]} X_i \right] = \sum_{i \in [n]} \mathbb{P}\{|\langle \mathbf{g}_i, \mathbf{v} \rangle| > \lambda\tau\} \leq ne^{-\lambda^2/8}$$

as $\langle \mathbf{g}_i, \bar{\mathbf{v}} \rangle$ has sub-Gaussian norm at most τ . To bound the sum $\sum_i X_i$, we record the standard Chernoff bound:

LEMMA 10 (Chernoff bound; Theorem 4.4 in [46]). *Under the above setting, for each $\delta > 0$ we have $\mathbb{P}\{\sum_i X_i \geq (1 + \delta)\mu\} \leq e^{-\mu(1+\delta)\log(1+\delta) + \mu\delta} \leq e^{-\mu \min\{\delta^2, \delta\}/3}$*

We proceed by considering two cases:

Case 1: If $d < ne^{-\lambda^2/16}$, then we set $\delta = \left(\sqrt{\frac{9d}{\mu}} + \frac{9d}{\mu} \right) \log(\frac{3}{\epsilon})$. Applying the second inequality in the Chernoff bound, we obtain that with probability at least $1 - \exp(-3d \log(\frac{3}{\epsilon})) \geq 1 - (c\sqrt{d_0})^{-1}(3/\epsilon)^{-2d}$, there holds

$$\sum_i X_i \lesssim (\mu + d) \log(\frac{3}{\epsilon}) = \tilde{O} \left(ne^{-\lambda^2/8} + ne^{-\lambda^2/16} \right) = \tilde{O} \left(ne^{-\lambda^2/16} \right),$$

where we recall that note $\tilde{O}(\cdot)$ hides multiplicative factors of $\log(\frac{3}{\epsilon}) \asymp \log(d + \log n)$.

Case 2: If $d \geq ne^{-\lambda^2/16}$, then we set $\delta = \left(3 + \frac{18d}{\mu} / \log(\frac{18d}{\mu})\right) \log(\frac{3}{\epsilon})$. In this case, we have the inequalities

$$\delta \geq 3 \implies (1 + \delta) \log(1 + \delta) \geq \frac{3}{2} \delta \implies (1 + \delta) \log(1 + \delta) - \delta \geq \frac{1}{3} \delta \log \delta$$

and

$$\mu \leq ne^{-\lambda^2/8} \leq \frac{d^2}{n} \implies \frac{18d}{\mu} \geq \frac{18n}{d} \geq 18 \implies \log\left(\frac{18d}{\mu}\right) - \log \log\left(\frac{18d}{\mu}\right) \geq \frac{1}{2} \log\left(\frac{18d}{\mu}\right) \geq 0,$$

whence

$$\mu(1 + \delta) \log(1 + \delta) - \mu\delta \geq \frac{1}{3} \mu\delta \cdot \log \delta \geq \frac{6d \log(\frac{3}{\epsilon})}{\log(\frac{18d}{\mu})} \cdot \left(\log\left(\frac{18d}{\mu}\right) - \log \log\left(\frac{18d}{\mu}\right)\right) \geq 3d \log\left(\frac{3}{\epsilon}\right).$$

Applying the first inequality in the Chernoff bound, we obtain that with probability at least $1 - \exp(-3d \log(\frac{3}{\epsilon})) \geq 1 - (c\sqrt{d_0})^{-1}(3/\epsilon)^{-2d}$, there holds

$$\sum_i X_i \lesssim \mu + \frac{d \log(\frac{3}{\epsilon})}{\log(18d/\mu)} = \tilde{O}\left(\mu + \frac{d}{\log(18) + \lambda^2/8 - \log(n/d)}\right) = \tilde{O}\left(ne^{-\lambda^2/8} + \frac{d}{\lambda^2/16}\right),$$

where the last step holds because $d \geq ne^{-\lambda^2/16} \implies \log(n/d) \leq \frac{\lambda^2}{16}$.

Step (ii): Applying union bound on ϵ -net and for all λ . In both cases above, the inequality (32) holds with probability $\geq 1 - (c\sqrt{d_0})^{-1}(3/\epsilon)^{-2d}$, for each fixed number $\lambda \geq 1$ and unit vector $\mathbf{v}$. Now, let $\mathcal{N}$ be an ϵ -net of the set of d -dimensional unit vectors, where $|\mathcal{N}| \leq (3/\epsilon)^d$. Applying a union bound, we find that the inequality (32) holds simultaneously for all integers $\bar{\lambda} = 1, 2, \dots, c\sqrt{d_0}$ and vectors $\bar{\mathbf{v}} \in \mathcal{N}$ with probability at least $1 - (3/\epsilon)^{-d} \geq 1 - 6^{-d}$ (since $\epsilon := 1/(2c\sqrt{d_0}) \leq 1/2$). Also note that $|\mathcal{B}_{c\sqrt{d_0}}(\mathbf{v})| = 0$ with probability at least $1 - n^{-1}$ (cf. Lemma 7).

On the above event, for all *real numbers* $\lambda \in [1, c\sqrt{d_0}]$ and all vectors $\bar{\mathbf{v}} \in \mathcal{N}$, we have

$$|\mathcal{B}_\lambda(\bar{\mathbf{v}})| \leq |\mathcal{B}_{\lfloor \lambda \rfloor}(\bar{\mathbf{v}})| = \tilde{O}\left(\frac{d}{\lfloor \lambda \rfloor^2} + n \exp\left(-\frac{\lfloor \lambda \rfloor^2}{C}\right)\right) \leq \tilde{O}\left(\frac{4d}{\lambda} + n \exp\left(-\frac{\lambda^2}{4C}\right)\right),$$

where the last two steps hold since $\lfloor \lambda \rfloor$ is an integer and satisfy $\lfloor \lambda \rfloor \geq \lambda/2$. Moreover, for all $\lambda > c\sqrt{d_0}$ we have $|\mathcal{B}_\lambda(\bar{\mathbf{v}})| \leq |\mathcal{B}_{c\sqrt{d_0}}(\bar{\mathbf{v}})| = 0$. We hence see that the inequality (32) holds (with a change of absolute constants) for all $\lambda \geq 1$ and $\bar{\mathbf{v}} \in \mathcal{N}$. Finally, for each unit vector $\mathbf{v}$ in $\mathbb{R}^d$, let $\bar{\mathbf{v}}$ be the nearest vector in the ϵ -net $\mathcal{N}$. If $i \in \mathcal{B}_\lambda(\mathbf{v})$, then

$$|\langle \tilde{\mathbf{g}}_i, \bar{\mathbf{v}} \rangle| \geq |\langle \tilde{\mathbf{g}}_i, \mathbf{v} \rangle| - |\langle \tilde{\mathbf{g}}_i, \bar{\mathbf{v}} - \mathbf{v} \rangle| \stackrel{(i)}{\geq} \lambda\tau - c\tau\sqrt{d_0} \cdot \epsilon \stackrel{(ii)}{\geq} \frac{1}{2}\lambda\tau,$$

where step (i) follows from Lemma 7 and step (ii) follows from our choice $\epsilon := 1/(2c\sqrt{d_0})$ and $\lambda \geq 1$. This means that $\mathcal{B}_\lambda(\mathbf{v}) \subseteq \mathcal{B}_{\lambda/2}(\bar{\mathbf{v}})$, and thus inequality (32) holds (with a change of absolute constants) for all unit vectors $\mathbf{v}$ in $\mathbb{R}^d$ as well. We have completed the proof of Lemma 9.

Appendix C: Standard Results for Sub-Gaussian Random Variables We collect several standard tail bounds that are used in the proofs of our main theorems.

C.1. Sub-Gaussian Tail Bounds The first lemma is the standard Hoeffding’s inequality as given in [59, Theorem 2.6.2].

LEMMA 11 (Hoeffding’s inequality for Sub-Gaussians). *Let $X_1, \dots, X_N$ be independent, mean zero, sub-Gaussian random variables. Then, for every $t \geq 0$ we have*

$$\mathbb{P} \left[\left| \sum_{i=1}^N X_i \right| \geq t \right] \leq 2 \exp \left[- \frac{ct^2}{\sum_{i=1}^N \|X_i\|_{\psi_2}^2} \right],$$

where $c > 0$ is a universal constant.

The next lemma controls the inner product between sub-Gaussian random vectors.

LEMMA 12. *Let $\{\mathbf{x}_i\}_{i \in [n]}$ be independent sub-Gaussian random vectors such that $\|\mathbf{x}_i\|_{\psi_2} \leq \rho$ for each $i \in [n]$. For any fixed $i \in [n]$, $\mathcal{M} \subset [n]$, $t > 0$ and $\delta > 0$, there exists a universal constant $C > 0$ such that*

$$\left\langle \mathbf{x}_i, \frac{1}{|\mathcal{M}|} \sum_{j \in \mathcal{M}} \mathbf{x}_j \right\rangle \leq \frac{3\rho^2 \left(5C\sqrt{|\mathcal{M}|} (\sqrt{d \log n} + \log n) + d \right)}{|\mathcal{M}|}$$

with probability at least $1 - n^{-10}$.

Proof. We record the following lemma, which is Lemma A.3 in [42].

LEMMA 13. *Let $\{\mathbf{x}_i\}_{i \in [n]}$ be independent sub-Gaussian random vectors such that $\|\mathbf{x}_i\|_{\psi_2} \leq \rho$ for each $i \in [n]$. For any fixed $i \in [n]$, $\mathcal{M} \subset [n]$, $t > 0$ and $\delta > 0$, we have*

$$\mathbb{P} \left[\left\langle \mathbf{x}_i, \frac{1}{|\mathcal{M}|} \sum_{j \in \mathcal{M}} \mathbf{x}_j \right\rangle \geq \frac{3\rho^2 \left(t\sqrt{|\mathcal{M}|} + d + \log(1/\delta) \right)}{|\mathcal{M}|} \right] \leq \exp \left(- \min \left\{ \frac{t^2}{4d}, \frac{t}{4} \right\} \right) + \delta.$$

Taking $t = 4C(\sqrt{d \log n} + \log n)$ and $\delta = n^{-20}$ (where $C > 0$ is a sufficiently large universal constant), we prove the bound in Lemma 12. $\square$

C.2. Random Vectors on the Unit Ball In this section we prove Fact 1 concerning the sub-Gaussian norm of a rotationally invariant random vector on the unit ℓ_2 ball.

Proof of Fact 1. Let $\mathbf{g} = (g_1, \dots, g_d)^\top$ be a random vector drawn from a rotationally invariant distribution supported on the unit ℓ_2 ball in $\mathbb{R}^d$. Also let $C, C' > 0$ be universal constants whose values may change line by line. By [12, Proposition 4.10], $\mathbf{g}$ can be represented as $\mathbf{g} \stackrel{d}{=} r\mathbf{u}$, where $\stackrel{d}{=}$ means equality in distribution, $r \stackrel{d}{=} \|\mathbf{g}\|_2$, $\mathbf{u}$ is uniformly distributed on the unit sphere in $\mathbb{R}^d$, and r and $\mathbf{u}$ are independent. The random vector $\mathbf{u}$ is sub-Gaussian with norm $\|\mathbf{u}\|_{\psi_2} \leq C\sqrt{\frac{1}{d}}$ [59, Theorem 3.4.5], so its one-dimensional margin satisfies $\mathbb{P}\{|u_1| > t\} \leq 2 \exp\left(-\frac{t^2}{C'/d}\right)$. On the other hand, we have $r \stackrel{d}{=} \|\mathbf{g}\|_2 \in [0, 1]$ since $\mathbf{g}$ is supported on the unit ball. Putting together the above facts gives

$$\mathbb{P}\{|g_1| > t\} = \mathbb{P}\{r|u_1| > t\} \leq \mathbb{P}\{|u_1| > t\} \leq 2 \exp\left(-\frac{t^2}{C'/d}\right),$$

whence $\|g_1\|_{\psi_2} \leq C\sqrt{\frac{1}{d}}$. By rotation invariance, we know that $\langle \mathbf{a}, \mathbf{g} \rangle \stackrel{d}{=} g_1$ for all unit vector $\mathbf{a}$ [12, Proposition 4.8]. Therefore, all one-dimensional margins $\langle \mathbf{a}, \mathbf{g} \rangle$ of $\mathbf{g}$ is sub-Gaussian with norm at most $C\sqrt{\frac{1}{d}}$. This completes the proof. $\square$

Acknowledgments. Y. Fei and Y. Chen is partially supported by NSF grant CCF-1704828 and CAREER Award CCF-2047910.

References

- [1] Abbe E (2017) Community detection and the stochastic block model: recent developments. *Journal of Machine Learning Research*, to appear URL http://www.princeton.edu/~eabbe/publications/sbm_jmlr_4.pdf.
- [2] Abbe E, Bandeira AS, Hall G (2016) Exact recovery in the stochastic block model. *IEEE Transactions on Information Theory* 62(1):471–487.
- [3] Abbe E, Sandon C (2015) Community detection in general stochastic block models: Fundamental limits and efficient algorithms for recovery. *IEEE 56th Annual Symposium on Foundations of Computer Science (FOCS)*, 670–688 (IEEE).
- [4] Achlioptas D, McSherry F (2005) On spectral learning of mixtures of distributions. *International Conference on Computational Learning Theory*, 458–469 (Springer).
- [5] Aloise D, Deshpande A, Hansen P, Popat P (2009) NP-hardness of Euclidean sum-of-squares clustering. *Machine Learning* 75(2):245–248.
- [6] Ames BPW, Vavasis SA (2014) Convex optimization for the planted k-disjoint-clique problem. *Mathematical Programming* 143(1-2):299–337.
- [7] Amini AA, Levina E (2018) On semidefinite relaxations for the block model. *The Annals of Statistics* 46(1):149–179.
- [8] Awasthi P, Bandeira AS, Charikar M, Krishnaswamy R, Villar S, Ward R (2015) Relax, no need to round: Integrality of clustering formulations. *Proceedings of the 2015 Conference on Innovations in Theoretical Computer Science*, 191–200 (ACM).
- [9] Awasthi P, Sheffet O (2012) Improved spectral-norm bounds for clustering. *Approximation, Randomization, and Combinatorial Optimization. Algorithms and Techniques*, 37–49 (Springer).
- [10] Awasthi P, Vijayaraghavan A (2017) Clustering semi-random mixtures of Gaussians. *arXiv preprint arXiv:1711.08841*.
- [11] Balakrishnan S, Wainwright MJ, Yu B (2017) Statistical guarantees for the EM algorithm: From population to sample-based analysis. *The Annals of Statistics* 45(1):77–120.
- [12] Bilodeau M, Brenner D (2008) *Theory of multivariate statistics* (Springer Science & Business Media).
- [13] Charikar M, Guha S, Tardos É, Shmoys DB (1999) A constant-factor approximation algorithm for the k-median problem. *Proceedings of the 31st Annual ACM Symposium on Theory of Computing*, 1–10 (ACM).
- [14] Chen X, Yang Y (2018) Hanson-Wright inequality in Hilbert spaces with application to k-means clustering for non-Euclidean data. *arXiv preprint arXiv:1810.11180*.
- [15] Chen Y, Sanghavi S, Xu H (2014) Improved graph clustering. *IEEE Transactions on Information Theory* 60(10):6440–6455.
- [16] Coja-Oghlan A (2004) Coloring semirandom graphs optimally. *Automata, Languages and Programming* 383–395.
- [17] Dasgupta S (1999) Learning mixtures of gaussians. *40th Annual Symposium on Foundations of Computer Science*, 634–644 (IEEE).
- [18] Daskalakis C, Tzamos C, Zampetakis M (2016) Ten steps of EM suffice for mixtures of two Gaussians. *arXiv preprint arXiv:1609.00368*.
- [19] Dempster AP, Laird NM, Rubin DB (1977) Maximum likelihood from incomplete data via the EM algorithm. *Journal of the Royal Statistical Society. Series B (Methodological)* 1–38.
- [20] Diakonikolas I, Kane DM, Stewart A (2017) Statistical query lower bounds for robust estimation of high-dimensional gaussians and gaussian mixtures. *2017 IEEE 58th Annual Symposium on Foundations of Computer Science (FOCS)*, 73–84 (IEEE).

- [21] Fei Y, Chen Y (2018) Hidden integrality of SDP relaxations for sub-Gaussian mixture models. *Conference On Learning Theory*, 1931–1965.
- [22] Fei Y, Chen Y (2019) Achieving the Bayes error rate in stochastic block model by sdp, robustly. *Conference on Learning Theory*, 1235–1269.
- [23] Fei Y, Chen Y (2019) Achieving the Bayes error rate in synchronization and block models by sdp, robustly. *arXiv preprint arXiv:1904.09635* .
- [24] Fei Y, Chen Y (2019) Exponential error rates of SDP for block models: Beyond Grothendieck’s inequality. *IEEE Transactions on Information Theory* 65(1):551–571.
- [25] Feige U, Kilian J (2001) Heuristics for semirandom graph problems. *Journal of Computer and System Sciences* 63(4):639–671.
- [26] Giraud C, Verzelen N (2019) Partial recovery bounds for clustering with the relaxed k -means. *Mathematical Statistics and Learning* 1(3):317–374.
- [27] Guédon O, Vershynin R (2016) Community detection in sparse networks via Grothendieck’s inequality. *Probability Theory and Related Fields* 165(3-4):1025–1049.
- [28] Hopkins SB, Li J (2018) Mixture models, robustness, and sum of squares proofs. *Proceedings of the 50th Annual ACM SIGACT Symposium on Theory of Computing*, 1021–1034 (ACM).
- [29] Iguchi T, Mixon DG, Peterson J, Villar S (2017) Probably certifiably correct k -means clustering. *Mathematical Programming* 165(2):605–642.
- [30] Jain K, Mahdian M, Saberi A (2002) A new greedy approach for facility location problems. *Proceedings of the Thiry-Fourth Annual ACM Symposium on Theory of Computing*, 731–740 (ACM).
- [31] Jin C, Zhang Y, Balakrishnan S, Wainwright MJ, Jordan MI (2016) Local maxima in the likelihood of gaussian mixture models: Structural results and algorithmic consequences. *Advances in Neural Information Processing Systems*, 4116–4124.
- [32] Kalai AT, Moitra A, Valiant G (2010) Efficiently learning mixtures of two Gaussians. *Proceedings of the Forty-second ACM Symposium on Theory of Computing*, 553–562 (ACM).
- [33] Kanungo T, Mount DM, Netanyahu NS, Piatko CD, Silverman R, Wu AY (2004) A local search approximation algorithm for k -means clustering. *Computational Geometry* 28(2-3):89–112.
- [34] Klusowski JM, Brinda WD (2016) Statistical guarantees for estimating the centers of a two-component gaussian mixture by EM. *arXiv preprint arXiv:1608.02280* .
- [35] Kothari PK, Steinhart J (2017) Better agnostic clustering via relaxed tensor norms. *arXiv preprint arXiv:1711.07465* .
- [36] Krivelevich M, Vilenchik D (2006) Semirandom models as benchmarks for coloring algorithms. *Third Workshop on Analytic Algorithmics and Combinatorics (ANALCO)*, 211–221.
- [37] Kumar A, Kannan R (2010) Clustering with spectral norm and the k -means algorithm. *Proceedings of the 2010 IEEE 51st Annual Symposium on Foundations of Computer Science*, 299–308 (IEEE Computer Society).
- [38] Li S, Svensson O (2016) Approximating k -median via pseudo-approximation. *SIAM Journal on Computing* 45(2):530–547.
- [39] Li X, Chen Y, Xu J (2018) Convex relaxation methods for community detection. *arXiv preprint arXiv:1810.00315* .
- [40] Li X, Li Y, Ling S, Strohmer T, Wei K (2017) When do birds of a feather flock together? K -means, proximity, and conic programming. *arXiv preprint arXiv:1710.06008* .
- [41] Lloyd S (1982) Least squares quantization in PCM. *IEEE Transactions on Information Theory* 28(2):129–137.
- [42] Lu Y, Zhou HH (2016) Statistical and computational guarantees of Lloyd’s algorithm and its variants. *arXiv preprint arXiv:1612.02099* .

- [43] Mahajan M, Nimbhorkar P, Varadarajan K (2009) The planar k-means problem is NP-hard. *International Workshop on Algorithms and Computation*, 274–285 (Springer).
- [44] Makarychev K, Makarychev Y, Vijayaraghavan A (2012) Approximation algorithms for semi-random partitioning problems. *Proceedings of the forty-fourth annual ACM symposium on Theory of computing*, 367–384 (ACM).
- [45] Makarychev K, Makarychev Y, Vijayaraghavan A (2016) Learning communities in the presence of errors. *29th Annual Conference on Learning Theory*, 1258–1291.
- [46] Mitzenmacher M, Upfal E (2005) *Probability and computing: Randomized algorithms and probabilistic analysis* (Cambridge University Press).
- [47] Mixon DG, Villar S, Ward R (2017) Clustering subgaussian mixtures by semidefinite programming. *Information and Inference: A Journal of the IMA* 6(4):389–415.
- [48] Moitra A, Perry W, Wein AS (2016) How robust are reconstruction thresholds for community detection? *Proceedings of the 48th Annual ACM SIGACT Symposium on Theory of Computing*, 828–841 (ACM).
- [49] Moitra A, Valiant G (2010) Settling the polynomial learnability of mixtures of gaussians. *2010 51st Annual IEEE Symposium on Foundations of Computer Science (FOCS)*, 93–102 (IEEE).
- [50] Montanari A, Sen S (2016) Semidefinite programs on sparse random graphs and their application to community detection. *Proceedings of the 48th Annual ACM SIGACT Symposium on Theory of Computing (STOC)*, 814–827.
- [51] Ndaoud M (2018) Sharp optimal recovery in the two Gaussian mixture model. *arXiv preprint arXiv:1812.08078* .
- [52] Nellore A, Ward R (2015) Recovery guarantees for exemplar-based clustering. *Information and Computation* 245:165–180.
- [53] Pearson K (1936) Method of moments and method of maximum likelihood. *Biometrika* 28(1/2):34–59.
- [54] Peng J, Wei Y (2007) Approximating k-means-type clustering via semidefinite programming. *SIAM Journal on Optimization* 18(1):186–205.
- [55] Peng J, Xia Y (2005) A new theoretical framework for k-means-type clustering. *Foundations and Advances in Data Mining*, 79–96 (Springer).
- [56] Regev O, Vijayaraghavan A (2017) On learning mixtures of well-separated gaussians. *2017 IEEE 58th Annual Symposium on Foundations of Computer Science (FOCS)*, 85–96 (IEEE).
- [57] Royer M (2017) Adaptive clustering through semidefinite programming. *Advances in Neural Information Processing Systems*, 1795–1803.
- [58] Vempala S, Wang G (2004) A spectral algorithm for learning mixture models. *Journal of Computer and System Sciences* 68(4):841–860.
- [59] Vershynin R (2017) High-dimensional probability .
- [60] Xu J, Hsu DJ, Maleki A (2016) Global analysis of expectation maximization for mixtures of two gaussians. *Advances in Neural Information Processing Systems*, 2676–2684.
- [61] Yan B, Yin M, Sarkar P (2017) Convergence analysis of gradient EM for multi-component Gaussian mixture. *arXiv preprint arXiv:1705.08530* .