

Improved Approximation Algorithms for Orthogonally Constrained Problems Using Semidefinite Optimization

Ryan Cory-Wright¹ and Jean Pauphilet^{2*}

¹Imperial Business School, South Kensington Campus, London, SW7
2AZ, United Kingdom.

^{2*}London Business School, Regent's Park, London, NW1 4SA, United
Kingdom.

*Corresponding author(s). E-mail(s): jpauphilet@london.edu;
Contributing authors: r.cory-wright@imperial.ac.uk;

Abstract

Building on the blueprint from Goemans and Williamson (1995) for the Max-Cut problem, we construct a polynomial-time approximation algorithm for orthogonally constrained quadratic optimization problems. First, we derive a semidefinite relaxation and propose a randomized rounding algorithm to generate feasible solutions from the relaxation. Second, we derive constant-factor approximation guarantees for our algorithm. When optimizing for m orthonormal vectors in dimension n , we leverage strong duality and semidefinite complementary slackness to show that our algorithm achieves a $1/3$ -approximation ratio. For any m of the form 2^q for some integer q , we also construct an instance where the performance of our algorithm is exactly $(m+2)/(3m)$, which can be made arbitrarily close to $1/3$ by taking $m \rightarrow +\infty$, hence showing that our analysis is tight.

Keywords: Orthogonality constraints, Semidefinite relaxation, Randomized rounding, Approximation algorithm

1 Introduction

Many important optimization problems, such as quantum nonlocality (Briët et al. 2011), and control theory (Ben-Tal and Nemirovski 2002) problems feature semi-orthogonal matrices, i.e., matrices $U \in \mathbb{R}^{n \times m}$ where $U^\top U = I_m$. Orthogonality constraints are also related to the rank of a matrix, which models a matrix’s complexity in imputation (Bell and Koren 2007), factor analysis (Bertsimas et al. 2017), and multi-task regression (Negahban and Wainwright 2011) settings.

In combinatorial optimization, a major advance in the design of approximation algorithms occurred with Goemans and Williamson (1995), who proposed a 0.87856-approximation algorithm for Max-Cut. Their algorithm also provides a $2/\pi$ -approximation for general binary quadratic optimization (BQO) problems (Nesterov 1998), and can be extended to linearly-constrained BQO problems (Bertsimas and Ye 1998). Conceptually, Goemans and Williamson (1995) established semidefinite optimization and correlated rounding at the core of approximation algorithms (see Wolkowicz et al. 1998; Williamson and Shmoys 2011).

In this work, we extend the core ideas underpinning the Goemans–Williamson algorithm to quadratic semi-orthogonal optimization problems and provide analogous constant-factor guarantees on the quality of semidefinite relaxations in the semi-orthogonal setting.

1.1 The Original Goemans–Williamson Algorithm

BQO is a canonical optimization problem (see Luo et al. 2010, for a review of applications). It is also an important building block for logically constrained optimization problems with quadratic objectives (see, e.g., Dong et al. 2015). Formally, given a matrix $Q \succeq 0$, BQO selects a vector z in $\{-1, 1\}^m$ that solves

$$\max_{z \in \{-1, 1\}^m} \sum_{i,j} Q_{i,j} z_i z_j = \max_{z \in \{-1, 1\}^m} \langle Q, zz^\top \rangle, \quad (1)$$

where $\langle \cdot, \cdot \rangle$ denotes the Frobenius inner product between matrices. Problem (1) is NP-hard and often challenging to solve to optimality when $m \geq 100$ (Rehfeldt et al. 2023). Accordingly, a popular approach for obtaining near-optimal solutions is to sample from a distribution parameterized by the solution of (1)’s convex relaxation. Specifically, we can reformulate (1) in terms of the rank-one matrix $Z = zz^\top$. Then, a valid relaxation of Problem (1) is given by

$$\max_{Z \in S_+^m} \langle Q, Z \rangle \text{ s.t. } \text{diag}(Z) = e, \quad (2)$$

which would be an exact reformulation with the additional (non-convex) constraint $\text{rank}(Z) = 1$. Probabilistically speaking, (2) is a device for constructing a pseudodistribution over $z \in \{-1, 1\}^m$ (d’Aspremont and Boyd 2003). This suggests sampling vectors from a distribution with second moment Z^* and rounding to restore feasibility, as proposed by Goemans and Williamson (1995) for Max-Cut and described in Algorithm 1. The overall idea of Algorithm 1 is that the projection step (i.e., taking the

coordinate-wise sign of $\mathbf{y}$) partially preserves the second moment of the distribution of $\mathbf{y}$, $\mathbb{E}[\mathbf{y}\mathbf{y}^\top] = \mathbf{Z}^*$. Precisely, we have $\mathbb{E}[\hat{\mathbf{z}}\hat{\mathbf{z}}^\top] \succeq \frac{2}{\pi}\mathbf{Z}^*$ (see [Nesterov 1998](#); [Bertsimas and Ye 1998](#)).

Algorithm 1 The Goemans–Williamson rounding algorithm for Problem (1)

Compute $\mathbf{Z}^*$ an optimal solution of (2)
Sample $\mathbf{y} \sim \mathcal{N}(\mathbf{0}, \mathbf{Z}^*)$
Construct $\hat{\mathbf{z}} \in \{-1, 1\}^m : \hat{z}_i := \text{sign}(y_i)$
return $\hat{\mathbf{z}}$ a solution to Problem (1)

1.2 Orthogonally Constrained Quadratic Optimization

In this paper, we consider a family of orthogonally constrained quadratic problems that subsumes binary quadratic optimization. Formally, we search for m orthogonal vectors $\mathbf{u}_i \in \mathbb{R}^n$ which solve

$$\max_{\mathbf{u}_i \in \mathbb{R}^n, i \in [m]} \sum_{i,j=1}^m \mathbf{u}_i^\top \mathbf{A}^{(i,j)} \mathbf{u}_j \text{ s.t. } \mathbf{u}_i^\top \mathbf{u}_j = \delta_{i,j}, \quad \forall i, j \in [m], \quad (3)$$

where $\mathbf{A}$ is an $nm \times nm$ positive semidefinite matrix (in short, $\mathbf{A} \in \mathcal{S}_+^{nm}$) with block matrices $\mathbf{A}^{(i,j)} \in \mathbb{R}^{n \times n}$, and $\delta_{i,j} = 1$ if $i = j$ and 0 otherwise. We require $n \geq m$. By introducing the semi-orthogonal matrix $\mathbf{U} \in \mathbb{R}^{n \times m}$ whose columns are the vectors $\mathbf{u}_i \in \mathbb{R}^n$, we can write our problem as

$$\max_{\mathbf{U} \in \mathbb{R}^{n \times m}} \text{vec}(\mathbf{U})^\top \mathbf{A} \text{vec}(\mathbf{U}) \text{ s.t. } \mathbf{U}^\top \mathbf{U} = \mathbf{I}_m, \quad (4)$$

where the $\text{vec}(\cdot)$ operator stacks the columns of $\mathbf{U}$ together into a single vector.

The similarities between Problems (4) and (1) are striking: For example, we can reduce any BQO instance (1) to a special case of Problem (4). Our reduction (see [Appendix A](#)) not only shows that Problem (4) is NP-hard (as also proved in [Lai et al. 2026](#), Theorem 3.1) but is also approximation-preserving. Therefore, inheriting the inapproximability results of BQO, Problem (4) cannot be approximated in polynomial time within a factor of $2/\pi + \varepsilon$ unless $\text{P}=\text{NP}$ ([Br  t et al. 2015](#), Theorem I.3).

In addition to its applications in clustering, quantum nonlocality, or generalized trust-region problems (see [Burer and Park 2024](#), and references therein), Problem (4) appears as a relevant substructure for mixed-projection formulations of rank-constrained quadratic optimization problems ([Bertsimas et al. 2022](#)).

In this paper, inspired by the Goemans–Williamson algorithm for BQO, we develop a relax-then-round strategy with a $1/3$ -approximation ratio for Problem (4).

1.3 Related Work

Our work is most closely related to [Burer and Park \(2024\)](#), who develop a hierarchy of semidefinite relaxations for Problem (4). To numerically evaluate the tightness of their relaxations, they apply several ‘feasible rounding procedures’ but do not provide any theoretical approximation ratios. In contrast, we develop a randomized rounding procedure and show that it achieves a multiplicative factor guarantee, which is independent of both the ambient dimension n and the number of vectors m . As a non-convex quadratic optimization problem, Problem (4) can also be solved to optimality via global solvers such as [Gurobi](#) or [BARON](#). However, the scalability of these global solvers is currently limited.

Special Cases. A larger body of work considers a special case of Problem (4), where the matrix $\mathbf{A}$ is block-diagonal, namely $\mathbf{A}^{(i,j)} = \mathbf{0}$ if $i \neq j$. In this case, Problem (4) reduces to

$$\max_{\mathbf{U} \in \mathbb{R}^{n \times m}} \sum_{i \in [m]} \mathbf{u}_i^\top \mathbf{A}^{(i,i)} \mathbf{u}_i \quad \text{s.t.} \quad \mathbf{U}^\top \mathbf{U} = \mathbf{I}_m, \quad (5)$$

which is referred to as the sum of heterogeneous quadratic forms or the heterogeneous PCA problem. Indeed, when all the matrices $\mathbf{A}^{(i,i)}$ are equal, we recover the Principal Component Analysis (PCA) problem. [Bolla et al. \(1998, Section 5\)](#) solve Problem (5) in polynomial time via linear algebra techniques when the matrices $\mathbf{A}^{(i,i)}$ are diagonal or commute with each other. For general matrices, [Gilman et al. \(2025\)](#) further tailor the semidefinite relaxations of [Burer and Park \(2024\)](#). Although tighter, their relaxations are not always tight. For some instances, they even obtain optimality gaps exceeding 100%. We are not aware of any approximation algorithms with guarantees specifically for general (non-diagonal) instances of Problem (5).

Approximation Algorithms. To our knowledge, many of the existing approximation algorithms apply to optimization problems with different orthogonality structures. [Briët et al. \(2010\)](#) propose an approximation algorithm for problems of the form

$$\max_{\mathbf{U} \in \mathbb{R}^{n \times m}} \sum_{i,j \in [m]} A_{i,j} \mathbf{u}_i^\top \mathbf{u}_j \quad \text{s.t.} \quad \mathbf{u}_i^\top \mathbf{u}_i = 1 \quad \forall i \in [m], \quad (6)$$

which also subsumes BQO (for $n = 1$), but does not enforce orthogonality between the columns of $\mathbf{U}$. They devise a relax-and-round strategy analogous to Goemans–Williamson that achieves an approximation ratio of $2/\pi + \Theta(1/n)$.

A second line of work ([Nemirovski 2007](#); [So 2011](#)) proposes $\Omega(1/\log(n+m))$ -approximation algorithms for quadratic optimization problems over matrices $\mathbf{U}$ that satisfy $\mathbf{U}^\top \mathbf{U} \preceq \mathbf{I}_m$, i.e., whose largest singular value is at most one. This constraint does not ensure that the columns of $\mathbf{U}$ are orthogonal. Nonetheless, [Nemirovski \(2007\)](#) shows that, in several special cases such as the orthogonal Procrustes or quadratic assignment problems (but not in the case of Problem (4)), orthogonality constraints can be relaxed into $\mathbf{U}^\top \mathbf{U} \preceq \mathbf{I}_m$ without loss of optimality. Our algorithm differs in that it generates matrices $\mathbf{U}$ with orthogonal columns, while they only guarantee

$\sigma_{\max}(\mathbf{U}) \leq 1$. In addition, their approximation ratio vanishes both with n and m , while ours is dimension-independent.

Finally, [Bandeira et al. \(2016\)](#) study approximation algorithms for problems of the form

$$\max_{\mathbf{U}_i \in \mathbb{R}^{n \times m}, i=1, \dots, k} \sum_{i,j \in [k]} \langle \mathbf{A}^{(i,j)}, \mathbf{U}_i^\top \mathbf{U}_j \rangle \text{ s.t. } \mathbf{U}_i^\top \mathbf{U}_i = \mathbf{I}_m \ \forall i \in [k]. \quad (7)$$

Unfortunately, Problem (7) is not equivalent to (4). In particular, heterogeneous PCA is a special case of our Problem (4) but cannot be cast in the form (7). There are two key differences in the objective function of (7): it involves the *inner* products between columns of *different* semi-orthogonal matrices $\mathbf{U}_i, \mathbf{U}_{i'}$ for $i \neq i'$. On the other hand, the objective in (4) depends on *outer* products between columns of the *same* matrix $\mathbf{U}$. In particular, [Bandeira et al. \(2016\)](#) can restore feasibility for each $\mathbf{U}_i$, $i = 1, \dots, k$ in (7) separately, while the columns of $\mathbf{U}$ in (4) need to be orthogonalized together. Finally, [Bandeira et al. \(2016\)](#)'s proof technique involves taking a singular value decomposition of a random matrix and arguing that the singular vectors are distributed according to the Haar measure and independent of the singular values. Unfortunately, our setting (4) is more general than (7), and the singular vectors in the random matrices studied here are not distributed according to the Haar measure.

However, both our Problem (4) with $\mathbf{A} \succeq \mathbf{0}$ and Problem (7) are special cases of the non-commutative Grothendieck problem. To clarify this connection, let us reformulate (4) as a bilinear optimization problem, using a trick analogous to that of [Naor et al. \(2014, Section 5.3\)](#).

Lemma 1 *Assume that $\mathbf{A} \succeq \mathbf{0}$. Then, Problem (4) achieves the same objective value as*

$$\max_{\mathbf{U}, \mathbf{V} \in \mathbb{R}^{n \times m}} \langle \mathbf{A}, \text{vec}(\mathbf{U}) \text{vec}(\mathbf{V})^\top \rangle \text{ s.t. } \mathbf{U}^\top \mathbf{U} = \mathbf{V}^\top \mathbf{V} = \mathbf{I}_m. \quad (8)$$

Proof Any feasible solution $\mathbf{U}$ to Problem (4) defines a feasible solution $(\mathbf{U}, \mathbf{U})$ to Problem (8) with the same objective value, so (4) $\leq$ (8). Conversely, given a feasible solution to (8), we can apply Cauchy-Schwarz because $\mathbf{A} \succeq \mathbf{0}$ and get

$$\langle \mathbf{A}, \text{vec}(\mathbf{U}) \text{vec}(\mathbf{V})^\top \rangle \leq \|\mathbf{A}^{1/2} \text{vec}(\mathbf{U})\|_2 \|\mathbf{A}^{1/2} \text{vec}(\mathbf{V})\|_2 \leq (4),$$

which concludes the proof. $\square$

By embedding $\mathbf{U}$ and $\mathbf{V}$ as the first m columns of $n \times n$ orthogonal matrices, $\tilde{\mathbf{U}}$ and $\tilde{\mathbf{V}}$, and appropriately padding the matrix $\mathbf{A}$ with zeros, Problem (8) can be expressed as an instance of the non-commutative Grothendieck problem studied, for example, in [Naor et al. \(2014\)](#) and [Briët et al. \(2015\)](#). [Naor et al. \(2014\)](#) propose a $1/(2\sqrt{2})$ -factor approximation algorithm for this problem ([Naor et al. 2014, Figure 3](#)). Although our approximation ratio is marginally weaker, our algorithm relies on a more compact semidefinite relaxation (involving one $nm \times nm$ semidefinite matrix vs. $2n^2 \times 2n^2$), and is therefore more tractable, especially when $n \gg m$. It also applies directly to real-valued variables (rather than their complex-valued rounding step) and recovers the original Goemans–Williamson algorithm for BQO instances.

1.4 Contributions and Structure

Our main contribution is the development of a Goemans–Williamson sampling algorithm for the class of semi-orthogonal problems (4).

In Section 2, we derive a semidefinite relaxation and propose a sampling procedure. Our algorithm achieves a $1/3$ -approximation guarantee (Theorem 3) for Problem (4). Our analysis is tight in the sense that for any $\epsilon > 0$, we can construct an instance such that the approximation ratio of our algorithm is at most $1/3 + \epsilon$ (Proposition 4). While our approximation ratio does not depend on n, m , we can also derive tighter dimension-dependent approximation ratios (Proposition 5 and Theorem D.1). The proofs of our main results are presented in Section 3. In particular, our analysis relies heavily on properties of the dual of the semidefinite relaxation. To better judge the quality of our approximations, we develop two simpler algorithms in Section 4: uniform sampling and a stronger benchmark inspired by PCA. We show that they achieve $1/nm$ and $1/m^2$ approximation ratios, respectively, which are dominated by our semidefinite relax-and-project procedure. Then, we evaluate the performance of our approach numerically in Section 5.

1.5 Notations

We let lowercase boldfaced characters such as $\mathbf{x}$ denote vectors and uppercase boldfaced characters such as $\mathbf{X}$ denote matrices. We denote by $\mathbf{e}$ the vector of all ones. We let $[n]$ denote the set of running indices $\{1, \dots, n\}$. The cone of $n \times n$ symmetric (resp. positive semidefinite) matrices is denoted by $\mathcal{S}^n$ (resp. $\mathcal{S}_+^n$). For a matrix $\mathbf{X} \in \mathbb{R}^{n \times m}$, we let $\mathbf{x}_i$ denote its i th column. We let $\text{vec}(\mathbf{X}) : \mathbb{R}^{n \times m} \rightarrow \mathbb{R}^{nm}$ denote the vectorization operator which maps matrices to vectors by stacking columns. For a matrix $\mathbf{W}$, we may describe it as a block matrix composed of equally sized blocks and denote the (i, i') block by $\mathbf{W}^{(i, i')}$. The dimension of each block will be clear from the context. For any matrices $\mathbf{A}, \mathbf{B}$, we denote their Kronecker product $\mathbf{A} \otimes \mathbf{B}$. We will extensively use the identity $\text{vec}(\mathbf{ABC}) = (\mathbf{C}^\top \otimes \mathbf{A}) \text{vec}(\mathbf{B})$ (see, e.g., [Henderson and Searle 1981](#)). Finally, we let $\mathcal{N}(\mathbf{0}, \Sigma)$ denote a centered multivariate normal distribution with covariance matrix Σ .

2 A Goemans–Williamson Approach

In this section, we propose a new Goemans–Williamson-type approach for semi-orthogonal quadratic optimization problems. First, we review a semidefinite relaxation for semi-orthogonal quadratic optimization in Section 2.1. Second, in Section 2.2, we propose a randomized rounding scheme to generate feasible solutions. Finally, we present a multiplicative approximation ratio of $1/3$ and a dimension-dependent approximation ratio for our algorithm in Section 2.3.

2.1 A Shor Relaxation

As reviewed in Section 1.3, [Burer and Park \(2024, Section 2.2\)](#) derive the following semidefinite relaxation for Problem (4):

$$\max_{\mathbf{W} \in \mathcal{S}_+^{mn}} \langle \mathbf{A}, \mathbf{W} \rangle \quad \text{s.t.} \quad \sum_{i \in [m]} \mathbf{W}^{(i,i)} \preceq \mathbf{I}_n, \quad \text{tr}(\mathbf{W}^{(j,j')}) = \delta_{j,j'}, \forall j, j' \in [m], \quad (9)$$

where the matrix $\mathbf{W}$ encodes the outer-product of $\text{vec}(\mathbf{U})$ with itself, and the trace constraints on the blocks of $\mathbf{W}$ stem from the columns of $\mathbf{U}$ having unit norm and being pairwise orthogonal.

Similarly to the semidefinite relaxation of (1), imposing the constraint that $\mathbf{W}$ is rank-one in (9) would result in an exact reformulation of (4). Accordingly, the relaxation (9) is tight whenever some optimal solution is rank-one. However, the optimal solutions to (9) are often high-rank (the case $m = 1$ is one of the special cases where this semidefinite relaxation is tight). Actually, it follows from manipulating the Barvinok–Pataki bound ([Barvinok 2001](#); [Pataki 1998](#)) that there exists some optimal solution to Problem (9) with rank at most $n + m$.¹ However, not all optimal solutions obey this bound; thus, we do not use this observation in our analysis. An interesting question is how to generate a high-quality feasible solution to (4), with a provable approximation ratio, by leveraging a solution of (9), which is the focus of the rest of the section.

Note that Problem (9), which is essentially a Shor relaxation ([Shor 1987](#)), corresponds to the ‘DiagSum’ relaxation of [Burer and Park \(2024\)](#). They also derive an even stronger relaxation, which they call a ‘Kronecker’ relaxation. We do not explicitly analyze their Kronecker relaxation here because it is significantly less tractable (as reported in [Burer and Park 2024, Table 1](#)).

2.2 A Relax-and-Project Procedure

We propose a randomized rounding scheme to generate high-quality feasible solutions to (4) from an optimal solution to (9).

Our algorithm involves three main steps: First, we solve (9) and obtain a semidefinite matrix $\mathbf{W}^*$. Second, using $\mathbf{W}^*$, we sample an $n \times m$ matrix $\mathbf{G}$ such that $\text{vec}(\mathbf{G})$ follows a normal distribution with mean $\mathbf{0}_{nm}$ and covariance matrix $\mathbf{W}^*$. Third, from the matrix $\mathbf{G}$, we generate a feasible solution to (4) by projecting $\mathbf{G}$ onto the set of semi-orthogonal matrices. Specifically, we compute a singular value decomposition of $\mathbf{G}$, $\mathbf{G} = \mathbf{U}\mathbf{\Sigma}\mathbf{V}^\top$ and define $\mathbf{Q} := \mathbf{U}\mathbf{V}^\top$. We summarize our procedure in Algorithm 2.

In Algorithm 2, we can sample $\text{vec}(\mathbf{G}) \sim \mathcal{N}(\mathbf{0}_{nm}, \mathbf{W}^*)$ even when $\mathbf{W}^*$ is rank-deficient via the following construction—which will also be relevant for our theoretical analysis. Denoting $r = \text{rank}(\mathbf{W}^*)$, we first construct a Cholesky decomposition of $\mathbf{W}^*$: $\mathbf{W}^* = \sum_{k \in [r]} \text{vec}(\mathbf{B}_k) \text{vec}(\mathbf{B}_k)^\top$ with $\mathbf{B}_k \in \mathbb{R}^{n \times m}$. Then, we

¹After introducing a slack matrix $\mathbf{S}$ to write the semidefinite inequality constraint $\mathbf{S} = \mathbf{I}_n - \sum_{i \in [m]} \mathbf{W}^{(i,i)}$, $\mathbf{S} \succeq \mathbf{0}$, we have $m(m+1)/2 + n(n+1)/2$ scalar inequalities. Thus, [Pataki \(1998, Theorem 2.2\)](#) states that there exists some optimal solution $(\mathbf{W}, \mathbf{S})$ with $\text{rank}(\mathbf{W})(\text{rank}(\mathbf{W}) + 1)/2 + \text{rank}(\mathbf{S})(\text{rank}(\mathbf{S}) + 1)/2 \leq m(m+1)/2 + n(n+1)/2$. Since $\text{rank}(\mathbf{S}) \geq 0$, it implies $\text{rank}(\mathbf{W}) \leq n + m$.

sample $\text{vec}(\mathbf{G}) = \sum_{k \in [r]} \text{vec}(\mathbf{B}_k) z_k$ with $\mathbf{z} \sim \mathcal{N}(\mathbf{0}_r, \mathbf{I}_r)$. This procedure ensures that $\text{vec}(\mathbf{G}) \in \text{span}(\mathbf{W}^*)$ almost surely. In particular, if our semidefinite relaxation is tight (e.g., when $m = 1$) and $\mathbf{W}^*$ has rank 1, then $\mathbf{Q}$ is optimal almost surely (our rounding is exact).

Algorithm 2 A relax-and-project algorithm for Problem (4)

Compute $\mathbf{W}^*$ an optimal solution of (9)
Sample $\mathbf{G}$ according to $\text{vec}(\mathbf{G}) \sim \mathcal{N}(\mathbf{0}_{nm}, \mathbf{W}^*)$
Compute the SVD of $\mathbf{G}$, $\mathbf{G} = \mathbf{U}\mathbf{\Sigma}\mathbf{V}^\top$
Construct $\mathbf{Q} = \mathbf{U}\mathbf{V}^\top$
return the semi-orthogonal matrix $\mathbf{Q}$

By properties of the Gaussian distribution, we can explicitly control some second and fourth moments of the random matrix $\mathbf{G}$, which will be crucial in deriving our constant-factor approximation guarantees.

Lemma 2 Consider $\mathbf{W} \in \mathcal{S}_+^{nm}$ a feasible solution to (9). The random matrix $\mathbf{G}$ generated by Algorithm 2 satisfies

$$\mathbb{E}[\mathbf{G}^\top \mathbf{G}] = \mathbf{I}_m, \quad \text{and} \quad \mathbb{E}[\mathbf{G}\mathbf{G}^\top] \preceq \mathbf{I}_n.$$

Moreover it satisfies

$$\mathbb{E}[(\mathbf{G}^\top \mathbf{G})^2] \preceq 3\mathbf{I}_m, \quad \text{and} \quad \mathbb{E}[(\mathbf{G}\mathbf{G}^\top)^2] \preceq 3\mathbf{I}_n.$$

The equations on second moments in Lemma 2 stem from the constraints imposed in our semidefinite relaxation. Indeed, denoting $\mathbf{g}_j, j = 1, \dots, m$ the columns of $\mathbf{G}$, we have that the (j, j') block of $\mathbf{W}$, $\mathbf{W}^{(j, j')}$, corresponds to $\mathbb{E}[\mathbf{g}_j \mathbf{g}_{j'}^\top]$. In addition, we can express the matrices $\mathbf{G}^\top \mathbf{G}$ and $\mathbf{G}\mathbf{G}^\top$ as $[\mathbf{g}_j^\top \mathbf{g}_{j'}]_{j, j'}$ and $\sum_{j \in [m]} \mathbf{g}_j \mathbf{g}_j^\top$ respectively. The bounds on the fourth moments are a generalization of the known equality $\mathbb{E}[z^4] = 3$ for $z \sim \mathcal{N}(0, 1)$ to matrix Gaussian series. We defer a formal proof of Lemma 2 to Section C.1.

Similar to the original algorithm of Goemans and Williamson (1995), the intuition behind Algorithm 2 is that the sampled matrix $\mathbf{G}$ achieves an average performance equal to the relaxation value ($\mathbb{E}[\langle \mathbf{A}, \text{vec}(\mathbf{G}) \text{vec}(\mathbf{G})^\top \rangle] = \langle \mathbf{A}, \mathbf{W}^* \rangle$) and is feasible on average ($\mathbb{E}[\mathbf{G}^\top \mathbf{G}] = \mathbf{I}_m$). Therefore, the average objective value of the feasible solution $\mathbf{Q}$ should not be too different from that of $\mathbf{G}$, as we now theoretically study.

2.3 A Constant-Factor Guarantee

In this section, we study the theoretical performance of Algorithm 2 in the case where the objective matrix $\mathbf{A}$ in (4) is positive semidefinite.

Our main result is a constant-factor guarantee for Algorithm 2:

Theorem 3 Assume that $\mathbf{A} \succeq \mathbf{0}$. Let $\mathbf{W}^*$ denote an optimal solution to the semidefinite relaxation (9). The random matrix $\mathbf{Q} \in \mathbb{R}^{n \times m}$ generated by Algorithm 2 satisfies the inequality

$$\mathbb{E} \left[\text{vec}(\mathbf{Q})^\top \mathbf{A} \text{vec}(\mathbf{Q}) \right] \geq \frac{1}{3} \langle \mathbf{A}, \mathbf{W}^* \rangle.$$

We defer the proof of Theorem 3 to Section 3. Theorem 3 provides a worst-case approximation ratio of $1/3$ for Algorithm 2, which is independent of the ambient dimension n as well as the number of vectors m . Although the actual performance of our algorithm can be much stronger (as we demonstrate numerically in Section 5), we can show that our worst-case analysis is tight.

Proposition 4 Consider an integer m of the form $m = 2^q$ for some $q \geq 1$. There exists an integer $n \geq m$, a positive semidefinite matrix $\mathbf{A}$, and an optimal solution $\mathbf{W}$ to the semidefinite relaxation (9), such that the solutions generated by Algorithm 2 satisfy

$$\text{vec}(\mathbf{Q})^\top \mathbf{A} \text{vec}(\mathbf{Q}) = \frac{m+2}{3m} \langle \mathbf{A}, \mathbf{W} \rangle \quad \text{a.s.}$$

For any $\epsilon > 0$, applying Proposition 4 (proof in Section B) for a sufficiently large m shows the existence of an instance on which the approximation ratio of Algorithm 2 is at most $1/3 + \epsilon$, hence showing that, from a worst-case perspective, the $1/3$ -approximation ratio of Theorem 3 is essentially tight.

However, it does not rule out the existence of algorithms with tighter guarantees or the possibility to derive tighter constants for fixed values of n and m . In particular, we can show the following (proof deferred to Section 3.3):

Proposition 5 Assume that $\mathbf{A} \succeq \mathbf{0}$. Let $\mathbf{W}^*$ denote an optimal solution to the semidefinite relaxation (9). The random matrix $\mathbf{Q} \in \mathbb{R}^{n \times m}$ generated by Algorithm 2 satisfies the inequality

$$\mathbb{E} \left[\text{vec}(\mathbf{Q})^\top \mathbf{A} \text{vec}(\mathbf{Q}) \right] \geq \frac{2}{\pi m} \langle \mathbf{A}, \mathbf{W}^* \rangle.$$

For $m = 1$, Proposition 5 provides a $2/\pi \approx 0.636$ approximation factor, which is strictly better than $1/3$. However, for $m \geq 2$, the guarantee of Proposition 5 is weaker than that of Theorem 3.

Remark 1 In Theorem D.1, we derive an even tighter approximation guarantee, which decreases monotonically with n and m . In particular, for $m = 1$, our improved constant is at least equal to 0.6804 (in the limit $n \rightarrow \infty$). For $m = 2$, it is at least equal to 0.3402. However, it fails to improve upon the $1/3$ guarantee of Theorem 3 for $m > 2$.

3 Proof of the Main Results

The proof of Theorem 3 and Proposition 5 rely on strong duality and complementarity slackness results for the semidefinite relaxation (9). We present these intermediate technical results in Section 3.1. We then present the proof of Theorem 3 and Proposition 5 in Sections 3.2 and 3.3 respectively.

3.1 Semidefinite relaxation and duality

We first derive the dual of Problem (9).

Lemma 6 Assume that $\mathbf{A} \in \mathcal{S}_+^{mn}$. The dual of (9) is equal to

$$\begin{aligned} \min_{\mathbf{X} \in \mathcal{S}_+^n, \mathbf{\Lambda} \in \mathcal{S}^m} \quad & \text{tr}(\mathbf{X}) + \text{tr}(\mathbf{\Lambda}) \\ \text{s.t.} \quad & \mathbf{I}_m \otimes \mathbf{X} + \mathbf{\Lambda} \otimes \mathbf{I}_n \succeq \mathbf{A}. \end{aligned} \tag{10}$$

and strong duality holds. Furthermore, we can take $\mathbf{\Lambda} \succeq \mathbf{0}$ in (10) without loss of optimality.

The fact that we can take $\mathbf{\Lambda} \succeq \mathbf{0}$ will be crucial in our proof. Indeed, with this additional assumption, for any positive semidefinite matrices $\mathbf{M}_1, \mathbf{M}_2$, the relationship $\mathbf{M}_1 \succeq \mathbf{M}_2$ implies $\text{tr}(\mathbf{\Lambda} \mathbf{M}_1) \geq \text{tr}(\mathbf{\Lambda} \mathbf{M}_2)$, which would not be true if $\mathbf{\Lambda} \not\succeq \mathbf{0}$.

Proof We introduce a semidefinite dual variable $\mathbf{X} \in \mathcal{S}_+^n$ for the primal semidefinite constraint $\sum_{i \in [m]} \mathbf{W}^{(i,i)} \preceq \mathbf{I}_n$, and scalars $\lambda_{j,j'}$ for the trace equalities $\text{tr}(\mathbf{W}^{(j,j')}) = \delta_{j,j'}$, for all $(j, j') \in \{1, \dots, m\}^2$. Collectively, the variables $(\lambda_{j,j'})$ form a symmetric matrix $\mathbf{\Lambda} \in \mathcal{S}^m$. The Lagrangian is defined

$$\mathcal{L}(\mathbf{W}, \mathbf{X}, \mathbf{\Lambda}) = \langle \mathbf{A}, \mathbf{W} \rangle + \left\langle \mathbf{X}, \mathbf{I}_n - \sum_{i \in [m]} \mathbf{W}^{(i,i)} \right\rangle + \sum_{j, j' \in [m]} \lambda_{j,j'} (\delta_{j,j'} - \text{tr}(\mathbf{W}^{(j,j')})).$$

Denoting $\{\mathbf{E}_{j,j'}\}_{j,j'=1,\dots,m}$ the canonical basis of $\mathbb{R}^{m \times m}$, we can concisely rewrite

$$\sum_{i \in [m]} \langle \mathbf{X}, \mathbf{W}^{(i,i)} \rangle = \sum_{i \in [m]} \langle \mathbf{E}_{i,i} \otimes \mathbf{X}, \mathbf{W} \rangle = \langle \mathbf{I}_m \otimes \mathbf{X}, \mathbf{W} \rangle,$$

and

$$\sum_{j, j' \in [m]} \lambda_{j,j'} \text{tr}(\mathbf{W}^{(j,j')}) = \sum_{j, j' \in [m]} \lambda_{j,j'} \langle \mathbf{E}_{j,j'} \otimes \mathbf{I}_n, \mathbf{W} \rangle = \langle \mathbf{\Lambda} \otimes \mathbf{I}_n, \mathbf{W} \rangle$$

Taking the supremum of the Lagrangian over $\mathbf{W} \in \mathcal{S}_+^{nm}$ then yields the dual of the form (10).

Observe that $(\mathbf{X}, \mathbf{\Lambda}) = (\mathbf{I}_n, \lambda_{\max}(\mathbf{A}) \mathbf{I}_m)$ is a strictly feasible solution to Problem (10). Therefore, Slater's conditions are satisfied and strong duality holds (Vandenberghe and Boyd 1996, Theorem 3.1).

For the second part of the proof, consider a feasible dual solution $\mathbf{X} \in \mathcal{S}_+^n, \mathbf{\Lambda} \in \mathcal{S}^m$. The semidefinite constraint in (10) implies

$$\mathbf{I}_m \otimes \mathbf{X} + \mathbf{\Lambda} \otimes \mathbf{I}_n \succeq \mathbf{A} \succeq \mathbf{0}.$$

Take a unit vector $\mathbf{u} \in \mathbb{R}^m$ and $\mathbf{x} \in \mathbb{R}^n$. We have

$$(\mathbf{u} \otimes \mathbf{x})^\top (\mathbf{I}_m \otimes \mathbf{X} + \mathbf{\Lambda} \otimes \mathbf{I}_n) (\mathbf{u} \otimes \mathbf{x}) = \mathbf{x}^\top \mathbf{X} \mathbf{x} + (\mathbf{u}^\top \mathbf{\Lambda} \mathbf{u}) \|\mathbf{x}\|^2 \geq 0$$

which implies $\mathbf{X} \succeq -\lambda_{\min}(\mathbf{\Lambda})\mathbf{I}_n$. Set $\gamma := \max(0, -\lambda_{\min}(\mathbf{\Lambda})) \geq 0$. We have $\mathbf{X} \succeq \gamma\mathbf{I}_n$.

Define $\mathbf{X}^* = \mathbf{X} - \gamma\mathbf{I}_n$, $\mathbf{\Lambda}^* = \mathbf{\Lambda} + \gamma\mathbf{I}_m$. Then, $\mathbf{X}^* \succeq \mathbf{0}$ and $\mathbf{I}_m \otimes \mathbf{X}^* + \mathbf{\Lambda}^* \otimes \mathbf{I}_n = \mathbf{I}_m \otimes \mathbf{X} + \mathbf{\Lambda} \otimes \mathbf{I}_n$. So, $(\mathbf{X}^*, \mathbf{\Lambda}^*)$ is feasible for (10). By construction, $\mathbf{\Lambda}^* \succeq \mathbf{0}$. In terms of objective value, $\text{tr}(\mathbf{X}^*) + \text{tr}(\mathbf{\Lambda}^*) = \text{tr}(\mathbf{X}) + \text{tr}(\mathbf{\Lambda}) - \gamma(n-m) \leq \text{tr}(\mathbf{X}) + \text{tr}(\mathbf{\Lambda})$ because $n \geq m$. So, $(\mathbf{X}^*, \mathbf{\Lambda}^*)$ achieves a lower objective value. $\square$

The second technical lemma uses complementarity slackness of a primal-dual pair to simplify bilinear expressions of the form $\text{vec}(\mathbf{Q})^\top \mathbf{A} \text{vec}(\mathbf{G})$.

Lemma 7 Consider $\mathbf{W}^*$ an optimal solution to the semidefinite relaxation (9) and $(\mathbf{X}^*, \mathbf{\Lambda}^*) \in \mathcal{S}_+^n \times \mathcal{S}_+^m$ an optimal solution to the dual problem (10).

Consider a random matrix $\mathbf{G} \in \mathbb{R}^{n \times m}$ generated as in Algorithm 2: $\text{vec}(\mathbf{G}) \sim \mathcal{N}(\mathbf{0}, \mathbf{W}^*)$ and denote $\mathbf{Q}$ its projection onto the set of semi-orthogonal matrices. We have

$$\text{tr}(\mathbf{\Lambda}^* \mathbb{E}[\mathbf{G}^\top \mathbf{G}]) = \text{tr}(\mathbf{\Lambda}^*) \quad \text{and} \quad \text{tr}(\mathbf{X}^* \mathbb{E}[\mathbf{G} \mathbf{G}^\top]) = \text{tr}(\mathbf{X}^*). \quad (11)$$

Furthermore, the following relationships hold almost surely:

$$\text{vec}(\mathbf{Q})^\top \mathbf{A} \text{vec}(\mathbf{G}) = \text{tr}(\mathbf{X}^* (\mathbf{G} \mathbf{G}^\top)^{1/2}) + \text{tr}(\mathbf{\Lambda}^* (\mathbf{G}^\top \mathbf{G})^{1/2}) \quad (12)$$

$$\geq \frac{1}{\|\mathbf{G}\|_F} \text{vec}(\mathbf{G})^\top \mathbf{A} \text{vec}(\mathbf{G}). \quad (13)$$

Proof The first equality in (11) follows directly from the fact that $\mathbb{E}[\mathbf{G}^\top \mathbf{G}] = \mathbf{I}_m$ (Lemma 2).

For the second equality, $\mathbb{E}[\mathbf{G} \mathbf{G}^\top] \preceq \mathbf{I}_n$ (Lemma 2) only provides an inequality. To obtain an equality, we leverage the fact that $(\mathbf{W}^*, \mathbf{X}^*, \mathbf{\Lambda}^*)$ are primal-dual optimal. In particular, they satisfy $\langle \mathbf{X}^*, \mathbf{I}_n - \sum_{i \in [m]} \mathbf{W}^{*(i,i)} \rangle = 0$ by complementarity slackness. Given that $\mathbb{E}[\mathbf{G} \mathbf{G}^\top] = \sum_{i \in [m]} \mathbf{W}^{*(i,i)}$, we obtain

$$\text{tr}(\mathbf{X}^* \mathbb{E}[\mathbf{G} \mathbf{G}^\top]) = \text{tr}(\mathbf{X}^*).$$

Equations (12)-(13) also follow from complementarity slackness. From $\langle \mathbf{W}^*, \mathbf{I}_m \otimes \mathbf{X}^* + \mathbf{\Lambda}^* \otimes \mathbf{I}_n - \mathbf{A} \rangle = 0$, we get that

$$(\mathbf{I}_m \otimes \mathbf{X}^* + \mathbf{\Lambda}^* \otimes \mathbf{I}_n - \mathbf{A}) \mathbf{W}^* = \mathbf{0},$$

because both matrices are positive semidefinite. In other words, the range of $\mathbf{W}^*$ is included in the kernel of the matrix $(\mathbf{I}_m \otimes \mathbf{X}^* + \mathbf{\Lambda}^* \otimes \mathbf{I}_n - \mathbf{A})$. Since, by construction, $\text{vec}(\mathbf{G}) \in \text{span}(\mathbf{W}^*)$ almost surely (as discussed in Section 2.2), we must have

$$(\mathbf{I}_m \otimes \mathbf{X}^* + \mathbf{\Lambda}^* \otimes \mathbf{I}_n - \mathbf{A}) \text{vec}(\mathbf{G}) = \mathbf{0},$$

i.e.,

$$\mathbf{A} \text{vec}(\mathbf{G}) = (\mathbf{I}_m \otimes \mathbf{X}^* + \mathbf{\Lambda}^* \otimes \mathbf{I}_n) \text{vec}(\mathbf{G}). \quad (14)$$

Hence,

$$\text{vec}(\mathbf{Q})^\top \mathbf{A} \text{vec}(\mathbf{G}) = \text{vec}(\mathbf{Q})^\top (\mathbf{I}_m \otimes \mathbf{X}^* + \mathbf{\Lambda}^* \otimes \mathbf{I}_n) \text{vec}(\mathbf{G}) = \text{tr}(\mathbf{X}^* \mathbf{Q} \mathbf{G}^\top) + \text{tr}(\mathbf{\Lambda}^* \mathbf{Q}^\top \mathbf{G}).$$

Equation (12) follows immediately after observing that $\mathbf{Q} = (\mathbf{G} \mathbf{G}^\top)^{-1/2} \mathbf{G} = \mathbf{G} (\mathbf{G}^\top \mathbf{G})^{-1/2}$, so $\mathbf{Q} \mathbf{G}^\top = (\mathbf{G} \mathbf{G}^\top)^{1/2}$ and $\mathbf{Q}^\top \mathbf{G} = (\mathbf{G}^\top \mathbf{G})^{1/2}$.

Finally, to obtain (13), we apply the bound $\mathbf{\Sigma} \succeq \mathbf{\Sigma}^2 / \|\mathbf{\Sigma}\|_F$ and get $\mathbf{Q} \mathbf{G}^\top = (\mathbf{G} \mathbf{G}^\top)^{1/2} \succeq (\mathbf{G} \mathbf{G}^\top) / \|\mathbf{G}\|_F$ and $\mathbf{Q}^\top \mathbf{G} = (\mathbf{G}^\top \mathbf{G})^{1/2} \succeq (\mathbf{G}^\top \mathbf{G}) / \|\mathbf{G}\|_F$. Because $\mathbf{X}^* \succeq \mathbf{0}$ and $\mathbf{\Lambda}^* \succeq \mathbf{0}$, these semidefinite bounds lead to

$$\text{vec}(\mathbf{Q})^\top \mathbf{A} \text{vec}(\mathbf{G}) \geq \frac{1}{\|\mathbf{G}\|_F} \text{tr}(\mathbf{X}^* \mathbf{G} \mathbf{G}^\top) + \frac{1}{\|\mathbf{G}\|_F} \text{tr}(\mathbf{\Lambda}^* \mathbf{G}^\top \mathbf{G}) = \frac{1}{\|\mathbf{G}\|_F} \text{vec}(\mathbf{G})^\top \mathbf{A} \text{vec}(\mathbf{G}),$$

where the equality follows from applying (14) again. $\square$

3.2 Proof of Theorem 3

The proof of Theorem 3 leverages the bounds on the second and fourth moments of $\mathbf{G}^\top \mathbf{G}$ and $\mathbf{G}\mathbf{G}^\top$ in Lemma 2. However, the average performance of Algorithm 2 involves the second moment of $\text{vec}(\mathbf{Q})$, $\mathbb{E}[\text{vec}(\mathbf{Q}) \text{vec}(\mathbf{Q})^\top]$. Instead of trying to control this matrix directly, the proof follows three steps: (i) We apply Cauchy-Schwarz to lower bound $\mathbb{E}[\text{vec}(\mathbf{Q})^\top \mathbf{A} \text{vec}(\mathbf{Q})]$ by $\mathbb{E}[\text{vec}(\mathbf{Q})^\top \mathbf{A} \text{vec}(\mathbf{G})]$; (ii) We use duality results from Section 3.1 (in particular, Equation (12) from Lemma 7) to show that $\mathbb{E}[\text{vec}(\mathbf{Q})^\top \mathbf{A} \text{vec}(\mathbf{G})]$ only depends on $\mathbb{E}[\mathbf{Q}^\top \mathbf{G}]$ and $\mathbb{E}[\mathbf{Q}\mathbf{G}^\top]$; (iii) We leverage concentration inequalities for matrix Gaussian series from Lemma 2 combined with the following technical lemma (proof deferred to Section C.2).

Lemma 8 *Consider a random positive semidefinite matrix $\mathbf{S} \in \mathcal{S}_+^d$ and a fixed semidefinite matrix $\mathbf{M} \in \mathcal{S}_+^d$ such that $\text{tr}(\mathbf{M}\mathbb{E}[\mathbf{S}]) = \text{tr}(\mathbf{M})$ and $\text{tr}(\mathbf{M}\mathbb{E}[\mathbf{S}^2]) \leq c \text{tr}(\mathbf{M})$ for some constant $c > 0$. Then, $\text{tr}(\mathbf{M}\mathbb{E}[\mathbf{S}^{1/2}]) \geq (1/\sqrt{c}) \text{tr}(\mathbf{M})$.*

Proof of Theorem 3 Given that $\mathbf{A} \succeq \mathbf{0}$, by Cauchy-Schwarz, we have

$$\mathbb{E}[\text{vec}(\mathbf{Q})^\top \mathbf{A} \text{vec}(\mathbf{Q})] \geq \frac{\mathbb{E}[\text{vec}(\mathbf{Q})^\top \mathbf{A} \text{vec}(\mathbf{G})]^2}{\mathbb{E}[\text{vec}(\mathbf{G})^\top \mathbf{A} \text{vec}(\mathbf{G})]}.$$

Let us introduce optimal dual solutions to the semidefinite relaxation, $(\mathbf{X}^*, \mathbf{\Lambda}^*) \in \mathcal{S}_+^n \times \mathcal{S}_+^m$, whose existence is guaranteed by Lemma 6. By strong duality, we have $\langle \mathbf{A}, \mathbf{W}^* \rangle = \text{tr}(\mathbf{X}^*) + \text{tr}(\mathbf{\Lambda}^*)$. Furthermore, Equation (12) in Lemma 7 yields

$$\mathbb{E}[\text{vec}(\mathbf{Q})^\top \mathbf{A} \text{vec}(\mathbf{G})] = \text{tr}(\mathbf{X}^* \mathbb{E}[(\mathbf{G}\mathbf{G}^\top)^{1/2}]) + \text{tr}(\mathbf{\Lambda}^* \mathbb{E}[(\mathbf{G}^\top \mathbf{G})^{1/2}]).$$

The random matrix $\mathbf{G}\mathbf{G}^\top$ satisfies $\text{tr}(\mathbf{X}^* \mathbb{E}[\mathbf{G}\mathbf{G}^\top]) = \text{tr}(\mathbf{X}^*)$ (Equation (11) in Lemma 7) and $\text{tr}(\mathbf{X}^* \mathbb{E}[(\mathbf{G}\mathbf{G}^\top)^2]) \leq 3 \text{tr}(\mathbf{X}^*)$ (see Lemma 2), so, Lemma 8 implies that

$$\text{tr}(\mathbf{X}^* \mathbb{E}[(\mathbf{G}\mathbf{G}^\top)^{1/2}]) \geq \text{tr}(\mathbf{X}^*)/\sqrt{3}.$$

Similarly, with $\mathbf{G}^\top \mathbf{G}$, we get $\text{tr}(\mathbf{\Lambda}^* \mathbb{E}[(\mathbf{G}^\top \mathbf{G})^{1/2}]) \geq \text{tr}(\mathbf{\Lambda}^*)/\sqrt{3}$.

As a result, we have

$$\begin{aligned} \mathbb{E}[\text{vec}(\mathbf{Q})^\top \mathbf{A} \text{vec}(\mathbf{G})] &\geq \frac{1}{\sqrt{3}} (\text{tr}(\mathbf{X}^*) + \text{tr}(\mathbf{\Lambda}^*)) = \frac{1}{\sqrt{3}} \langle \mathbf{A}, \mathbf{W}^* \rangle, \\ \text{and } \mathbb{E}[\text{vec}(\mathbf{Q})^\top \mathbf{A} \text{vec}(\mathbf{Q})] &\geq \frac{1}{3} \langle \mathbf{A}, \mathbf{W}^* \rangle. \end{aligned}$$

□

3.3 Proof of Proposition 5

The proof of Proposition 5 adopts a similar strategy as that of Theorem 3 but uses Equation (13) from Lemma 7 instead of Equation (12) to lower bound $\text{vec}(\mathbf{Q})^\top \mathbf{A} \text{vec}(\mathbf{Q})$.

Proof of Proposition 5 Given that $\mathbf{A} \succeq \mathbf{0}$, by Cauchy-Schwarz, we have

$$\text{vec}(\mathbf{Q})^\top \mathbf{A} \text{vec}(\mathbf{Q}) \geq \frac{(\text{vec}(\mathbf{Q})^\top \mathbf{A} \text{vec}(\mathbf{G}))^2}{\text{vec}(\mathbf{G})^\top \mathbf{A} \text{vec}(\mathbf{G})}.$$

From Equation (13) in Lemma 7, we have

$$\text{vec}(\mathbf{Q})^\top \mathbf{A} \text{vec}(\mathbf{G}) \geq \frac{1}{\|\mathbf{G}\|_F} \text{vec}(\mathbf{G})^\top \mathbf{A} \text{vec}(\mathbf{G}),$$

so

$$\text{vec}(\mathbf{Q})^\top \mathbf{A} \text{vec}(\mathbf{Q}) \geq \frac{1}{\|\mathbf{G}\|_F^2} \text{vec}(\mathbf{G})^\top \mathbf{A} \text{vec}(\mathbf{G}).$$

Taking expectation yields

$$\mathbb{E}[\text{vec}(\mathbf{Q})^\top \mathbf{A} \text{vec}(\mathbf{Q})] \geq \mathbb{E} \left[\frac{\text{vec}(\mathbf{G})^\top \mathbf{A} \text{vec}(\mathbf{G})}{\|\mathbf{G}\|_F^2} \right] = \langle \mathbf{A}, \mathbb{E} \left[\frac{\text{vec}(\mathbf{G}) \text{vec}(\mathbf{G})^\top}{\|\mathbf{G}\|_F^2} \right] \rangle$$

To conclude, we show that

$$\mathbb{E} \left[\frac{\text{vec}(\mathbf{G}) \text{vec}(\mathbf{G})^\top}{\|\mathbf{G}\|_F^2} \right] \succeq \frac{2}{\pi m} \mathbf{W}^\star.$$

Consider a fixed vector $\mathbf{v} \in \mathbb{R}^{nm}$. From Cauchy-Schwarz, we have that for, any random variables $A \geq 0, B > 0$, $\mathbb{E}[A/B] \geq \mathbb{E}[\sqrt{A}]^2 / \mathbb{E}[B]$. Applied to $A = (\mathbf{v}^\top \text{vec}(\mathbf{G}))^2$ and $B = \|\mathbf{G}\|_F^2$, we get

$$\mathbb{E} \left[\frac{(\mathbf{v}^\top \text{vec}(\mathbf{G}))^2}{\|\mathbf{G}\|_F^2} \right] \geq \frac{\mathbb{E} [|\mathbf{v}^\top \text{vec}(\mathbf{G})|]^2}{\mathbb{E} [\|\mathbf{G}\|_F^2]}$$

For the numerator, $\mathbf{v}^\top \text{vec}(\mathbf{G}) \sim \mathcal{N}(0, \mathbf{v}^\top \mathbf{W}^\star \mathbf{v})$ so $\mathbb{E} [|\mathbf{v}^\top \text{vec}(\mathbf{G})|] = \sqrt{\frac{2}{\pi} \mathbf{v}^\top \mathbf{W}^\star \mathbf{v}}$. For the denominator, $\mathbb{E} [\|\mathbf{G}\|_F^2] = \mathbb{E} [\text{tr}(\mathbf{G}^\top \mathbf{G})] = m$, which concludes the proof. $\square$

Remark 2 The proof of Proposition 5 shows that

$$\mathbb{E}[\text{vec}(\mathbf{Q})^\top \mathbf{A} \text{vec}(\mathbf{Q})] \geq \left\langle \mathbf{A}, \mathbb{E} \left[\frac{\text{vec}(\mathbf{G}) \text{vec}(\mathbf{G})^\top}{\|\mathbf{G}\|_F^2} \right] \right\rangle.$$

However, we would like to emphasize that, while this scalar inequality holds, we do not necessarily have the semidefinite relationship

$$\mathbb{E}[\text{vec}(\mathbf{Q}) \text{vec}(\mathbf{Q})^\top] \succeq \mathbb{E} \left[\frac{\text{vec}(\mathbf{G}) \text{vec}(\mathbf{G})^\top}{\|\mathbf{G}\|_F^2} \right].$$

Indeed, consider the counterexample with $n = 4, m = 2$ and

$$\mathbf{W}^\star = \begin{pmatrix} \mathbf{A} & \mathbf{B} \\ \mathbf{B}^\top & \mathbf{C} \end{pmatrix},$$

with

$$\mathbf{A} = \text{Diag} \begin{pmatrix} 0.019 \\ 0.132 \\ 0.268 \\ 0.581 \end{pmatrix}, \quad \mathbf{B} = \text{Diag} \begin{pmatrix} 0.016 \\ 0.193 \\ 0.192 \\ -0.401 \end{pmatrix}, \quad \mathbf{C} = \text{Diag} \begin{pmatrix} 0.035 \\ 0.294 \\ 0.154 \\ 0.517 \end{pmatrix}.$$

For each seed s , we sample $N = 25,000$ matrices $(\mathbf{G}, \mathbf{Q})$ as in Algorithm 2 and compute $\lambda_s = \lambda_{\min} \left(\widehat{\mathbb{E}} \left[\text{vec}(\mathbf{Q}) \text{vec}(\mathbf{Q})^\top - \frac{\text{vec}(\mathbf{G}) \text{vec}(\mathbf{G})^\top}{\|\mathbf{G}\|_F^2} \right] \right)$. We generate $B = 100$ bootstrapped samples from the same N pairs of matrices $(\mathbf{G}, \mathbf{Q})$ and compute the corresponding minimum

eigenvalue $\lambda_{s,b}$. The average variance $\sigma_s^2 = \frac{1}{B} \sum_{b=1}^B (\lambda_{s,b} - \lambda_s)^2$ captures the within-seed error in estimating $\lambda_{\min}$ due to using a finite number of observations N to approximate the expectation. We repeat this process over $s = 50$ seeds to evaluate the estimation error due to using these particular N samples. An unbiased estimate of the smallest eigenvalue is given by $\hat{\lambda} = \frac{1}{S} \sum_{s=1}^S \lambda_s$, with standard error $\frac{1}{S^2} \sum_{s=1}^S (\lambda_s - \hat{\lambda})^2 + \frac{1}{S^2} \sum_{s=1}^S \sigma_s^2$, and we construct a 95% confidence interval. We find

$$\lambda_{\min} \left(\mathbb{E} \left[\text{vec}(\mathbf{Q}) \text{vec}(\mathbf{Q})^\top - \frac{\text{vec}(\mathbf{G}) \text{vec}(\mathbf{G})^\top}{\|\mathbf{G}\|_F^2} \right] \right) = -0.0136 \pm 0.0003 < 0.$$

4 Benchmark: Uniform Sampling and Deflation

To appreciate the strength of our approximation ratios for Algorithm 2, we analyze the performance of two baselines: uniform sampling (Section 4.1) and a deflation heuristic inspired by PCA (Section 4.2). The results of this section are summarized in Table 1.

Table 1 Summary of our guarantees for Algorithm 2 and two benchmarks.

Name	Algorithm 2	Uniform	Deflation
Ratio	$\frac{1}{3}$	$\frac{1}{nm}$	$\frac{1}{m^2}$
Source	Theorem 3	Proposition 9	Proposition 10

4.1 Uniform Sampling

We consider a naive baseline where we draw $\mathbf{Q}$ uniformly from the set of semi-orthogonal matrices. Note that this is analogous to generating i.i.d. Bernoulli vectors in BQO, which achieves a $1/2$ approximation ratio in the Max-Cut case (e.g., [Williamson and Shmoys 2011](#), Theorem 5.3).

In practice, to sample $\mathbf{U}$ uniformly from the set of semi-orthogonal matrices, we construct $\mathbf{M} \in \mathbb{R}^{n \times m}$ with entries $M_{ij} \sim \mathcal{N}(0, 1)$ i.i.d. and define $\mathbf{Q}$ from the singular value decomposition of $\mathbf{G}$, $\mathbf{Q} = \mathbf{U}$ where $\mathbf{G} = \mathbf{U} \mathbf{\Sigma} \mathbf{V}^\top$ (Lemma 2.6 [Tulino et al. 2004](#)).

We now show that it provides a $1/nm$ -approximation guarantee for Problem (4).

Proposition 9 *Let $\mathbf{Q} \in \mathbb{R}^{n \times m}$ be distributed uniformly over the set of semi-orthogonal matrices, $\{\mathbf{U} \in \mathbb{R}^{n \times m} : \mathbf{U}^\top \mathbf{U} = \mathbf{I}_m\}$. Then, for any $\mathbf{A} \in \mathcal{S}_+^{nm}$, we have that:*

$$\mathbb{E}[\text{vec}(\mathbf{Q})^\top \mathbf{A} \text{vec}(\mathbf{Q})] \leq \max_{\mathbf{U} \in \mathbb{R}^{n \times m} : \mathbf{U}^\top \mathbf{U} = \mathbf{I}_m} \text{vec}(\mathbf{U})^\top \mathbf{A} \text{vec}(\mathbf{U}) \leq nm \mathbb{E}[\text{vec}(\mathbf{Q})^\top \mathbf{A} \text{vec}(\mathbf{Q})].$$

Proof Because $\mathbf{Q}$ is feasible for (4), we have

$$\text{vec}(\mathbf{Q})^\top \mathbf{A} \text{vec}(\mathbf{Q}) \leq \max_{\mathbf{U} \in \mathbb{R}^{n \times m}: \mathbf{U}^\top \mathbf{U} = \mathbf{I}_m} \text{vec}(\mathbf{U})^\top \mathbf{A} \text{vec}(\mathbf{U}),$$

which leads to the first inequality.

Furthermore,

$$\max_{\mathbf{U} \in \mathbb{R}^{n \times m}: \mathbf{U}^\top \mathbf{U} = \mathbf{I}_m} \text{vec}(\mathbf{U})^\top \mathbf{A} \text{vec}(\mathbf{U}) \leq \max_{\mathbf{u} \in \mathbb{R}^{nm}: \|\mathbf{u}\|^2 = m} \mathbf{u}^\top \mathbf{A} \mathbf{u} = m \lambda_{\max}(\mathbf{A}) \leq m \text{tr}(\mathbf{A}).$$

To conclude, observe that since $\mathbf{Q}$ is distributed according to the Haar measure, we have $\mathbb{E}[\mathbf{q}_i \mathbf{q}_i^\top] = \frac{1}{n} \mathbf{I}_n$ and $\mathbb{E}[\mathbf{q}_i \mathbf{q}_j^\top] = \mathbf{0}$ for $i \neq j$ (cf. Meckes 2019). Therefore, we have $\mathbb{E}[\text{vec}(\mathbf{Q}) \text{vec}(\mathbf{Q})^\top] = \frac{1}{n} \mathbf{I}_{nm}$ and $\mathbb{E}[\text{vec}(\mathbf{Q})^\top \mathbf{A} \text{vec}(\mathbf{Q})] = \frac{1}{n} \text{tr}(\mathbf{A})$. $\square$

Remark 3 Proposition 9's upper bound is tight for uniform sampling. Indeed, if $\mathbf{A}$ is an identity matrix, then any uniformly sampled $\mathbf{Q}$ is optimal and the left inequality is tight. Moreover, if $\mathbf{A}$ is a matrix such that $A_{i,j}^{(i,j)} = 1$ for $i, j \in [m]$ and $A_{\ell_1, \ell_2}^{(i,j)} = 0$ for $\ell_1 \neq i$ or $\ell_2 \neq j$ otherwise, then $\text{tr}(\mathbf{A}) = m$ and an optimal choice of $\mathbf{U}$ is $\mathbf{U}_i = \mathbf{e}_i$, giving both $\max_{\mathbf{U} \in \mathbb{R}^{n \times m}: \mathbf{U}^\top \mathbf{U} = \mathbf{I}_m} \text{vec}(\mathbf{U})^\top \mathbf{A} \text{vec}(\mathbf{U}) = m^2$ and $nm \mathbb{E}[\text{vec}(\mathbf{Q})^\top \mathbf{A} \text{vec}(\mathbf{Q})] = m^2$.

4.2 A Deflation Heuristic

We now propose a second baseline inspired by the deflation approach for Principal Component Analysis (e.g., Mackey 2008). Namely, we consider the first diagonal block of $\mathbf{A}$, $\mathbf{A}^{(1,1)}$, and compute its leading eigenvector. This defines the column $\mathbf{q}_1$. We then update (or deflate) the other diagonal blocks $\mathbf{A}^{(j,j)} \leftarrow (\mathbf{I}_n - \mathbf{q}_1 \mathbf{q}_1^\top) \mathbf{A}^{(j,j)} (\mathbf{I}_n - \mathbf{q}_1 \mathbf{q}_1^\top)$ and proceed with another block. We describe the procedure in Algorithm 3, where we assume the diagonal blocks are processed in the natural order, $1, 2, \dots, m$.

Algorithm 3 A Deflation-Inspired Benchmark

Require: Positive semidefinite matrix $\mathbf{A} \in \mathcal{S}_+^{nm}$

Initialize $\mathbf{Q} = \mathbf{0}$

for $i = 1, \dots, m$ **do**

 Deflate the block $\mathbf{A}^{(i,i)}$ according to $\mathbf{q}_1, \dots, \mathbf{q}_{i-1}$, i.e., compute

$$\mathbf{B} = (\mathbf{I}_n - \mathbf{q}_{i-1} \mathbf{q}_{i-1}^\top) \cdots (\mathbf{I}_n - \mathbf{q}_1 \mathbf{q}_1^\top) \mathbf{A}^{(i,i)} (\mathbf{I}_n - \mathbf{q}_1 \mathbf{q}_1^\top) \cdots (\mathbf{I}_n - \mathbf{q}_{i-1} \mathbf{q}_{i-1}^\top)$$

 Compute $\mathbf{v}_i \in \arg \max_{\mathbf{x} \|\mathbf{x}\|_2=1} \mathbf{x}^\top \mathbf{B} \mathbf{x}$

 Define $\mathbf{q}_i = z_i \mathbf{v}_i$ with $\mathbb{P}(z_i = 1) = 1 - \mathbb{P}(z_i = -1) = 1/2$.

end for

return Semi-orthogonal matrix $\mathbf{Q}$

Remark 4 Observe that if $\mathbf{A}$ is a block diagonal matrix with identical diagonal blocks $\mathbf{A}^{(i,i)} = \mathbf{\Sigma}$ and zero off diagonal blocks $\mathbf{A}^{(i,j)} = \mathbf{0}$ for $i \neq j$, as in principal component analysis, then the proposed algorithm corresponds to deflation in PCA and is thus exact.

In practice, we can process the diagonal blocks in any order, with each ordering leading to a different candidate solution. In our implementation, we consider N random permutations of $\{1, \dots, m\}$ and generate N feasible solutions to allow for a fair comparison with our sampling-based approach. We show (Proposition 10) that it provides a $1/m^2$ -factor approximation. We show that Algorithm 3, with a random ordering of the blocks, leads to a $1/m^2$ -approximation guarantee.

Proposition 10 *Let $\mathbf{Q}$ be generated according to Algorithm 3 with a random ordering of the blocks. Then, for any $\mathbf{A} \in \mathcal{S}_+^{nm}$, we have that:*

$$\mathbb{E} [\text{vec}(\mathbf{Q})^\top \mathbf{A} \text{vec}(\mathbf{Q})] \leq \max_{\mathbf{U} \in \mathbb{R}^{n \times m} : \mathbf{U}^\top \mathbf{U} = \mathbf{I}_m} [\text{vec}(\mathbf{U})^\top \mathbf{A} \text{vec}(\mathbf{U})] \leq m^2 \mathbb{E} [\text{vec}(\mathbf{Q})^\top \mathbf{A} \text{vec}(\mathbf{Q})].$$

Proof First, at iteration i , since $\mathbf{q}_i$ is collinear to the leading eigenvector of the matrix obtained by projecting $\mathbf{A}^{(i,i)}$ onto a space orthogonal to $\mathbf{q}_1, \dots, \mathbf{q}_{i-1}$, we have that $\mathbf{q}_i^\top \mathbf{q}_j = 0$ for each $j < i$ and thus $\mathbf{Q}$ is semi-orthogonal, so the left inequality holds.

Second, since z_i, z_j are i.i.d. with mean 0, in expectation, we have that $\mathbb{E} [\mathbf{q}_i^\top \mathbf{A}^{(i,j)} \mathbf{q}_j] = 0$ for $i \neq j$, and thus the expected objective value attained by $\mathbf{Q}$ is

$$\mathbb{E} [\text{vec}(\mathbf{Q})^\top \mathbf{A} \text{vec}(\mathbf{Q})] = \sum_{i \in [m]} \mathbb{E} [\mathbf{q}_i^\top \mathbf{A}^{(i,i)} \mathbf{q}_i] = \sum_{i \in [m]} \mathbf{v}_i^\top \mathbf{A}^{(i,i)} \mathbf{v}_i.$$

When treating the blocks in the order $1, \dots, m$, we have $\mathbb{E} [\text{vec}(\mathbf{Q})^\top \mathbf{A} \text{vec}(\mathbf{Q})] \geq \mathbf{v}_1^\top \mathbf{A}^{(1,1)} \mathbf{v}_1 = \lambda_{\max}(\mathbf{A}^{(1,1)})$. By taking the average over random permutations of $\{1, \dots, m\}$, we get $\mathbb{E} [\text{vec}(\mathbf{Q})^\top \mathbf{A} \text{vec}(\mathbf{Q})] \geq \frac{1}{m} \sum_{i \in [m]} \lambda_{\max}(\mathbf{A}^{(i,i)})$.

On the other hand, for any 2×2 block of $\mathbf{A}$ indexed by (i, j) with $i \neq j$, we have

$$\begin{pmatrix} \mathbf{A}^{(i,i)} & \mathbf{A}^{(i,j)} \\ \mathbf{A}^{(j,i)} & \mathbf{A}^{(j,j)} \end{pmatrix} \succeq \mathbf{0},$$

so for any orthogonal matrix $\mathbf{U}$,

$$\mathbf{u}_i^\top \mathbf{A}^{(i,j)} \mathbf{u}_j + \mathbf{u}_j^\top \mathbf{A}^{(j,i)} \mathbf{u}_i \leq \mathbf{u}_i^\top \mathbf{A}^{(i,i)} \mathbf{u}_i + \mathbf{u}_j^\top \mathbf{A}^{(j,j)} \mathbf{u}_j,$$

and

$$\sum_{i,j \in [m]} \mathbf{u}_i^\top \mathbf{A}^{(i,j)} \mathbf{u}_j \leq m \sum_{i \in [m]} \mathbf{u}_i^\top \mathbf{A}^{(i,i)} \mathbf{u}_i \leq m \sum_{i \in [m]} \lambda_{\max}(\mathbf{A}^{(i,i)}) \leq m^2 \mathbb{E} [\text{vec}(\mathbf{Q})^\top \mathbf{A} \text{vec}(\mathbf{Q})],$$

where the last inequality follows from $\mathbb{E} [\text{vec}(\mathbf{Q})^\top \mathbf{A} \text{vec}(\mathbf{Q})] \geq \frac{1}{m} \sum_{i \in [m]} \lambda_{\max}(\mathbf{A}^{(i,i)})$ and concludes the proof. $\square$

Remark 5 Our proof technique shows that $\mathbb{E} [\text{vec}(\mathbf{Q})^\top \mathbf{A} \text{vec}(\mathbf{Q})] \geq \frac{1}{m} \sum_{i \in [m]} \lambda_{\max}(\mathbf{A}^{(i,i)})$. Thus, Algorithm 3 yields a $1/m$ -factor approximation when $\mathbf{A}$ is a block diagonal matrix and a $1/(2m)$ -factor approximation when $\mathbf{A}$ is block diagonally dominant. In general, however, the block diagonal objective bounds the full objective within a factor of m , hence the overall $1/m^2$ guarantee. It is also worth noting that the semidefinite inequality

$$\begin{pmatrix} \mathbf{A}^{(1,1)} & \mathbf{A}^{(1,2)} & \dots & \mathbf{A}^{(1,m)} \\ \mathbf{A}^{(2,1)} & \mathbf{A}^{(2,2)} & \dots & \mathbf{A}^{(2,m)} \\ \vdots & \vdots & \ddots & \vdots \\ \mathbf{A}^{(m,1)} & \mathbf{A}^{(m,2)} & \dots & \mathbf{A}^{(m,m)} \end{pmatrix} \preceq m \begin{pmatrix} \mathbf{A}^{(1,1)} & \mathbf{0} & \dots & \mathbf{0} \\ \mathbf{0} & \mathbf{A}^{(2,2)} & \dots & \mathbf{0} \\ \mathbf{0} & \vdots & \ddots & \mathbf{0} \\ \mathbf{0} & \mathbf{0} & \dots & \mathbf{A}^{(m,m)} \end{pmatrix}$$

which we implicitly prove as part of our approximation guarantee, is actually a special case of the pinching inequality from quantum information theory (see [Mosonyi and Ogawa 2015](#), Lemma II.2).

5 Numerical Results

We numerically evaluate the performance of Algorithm 2 for semi-orthogonal quadratic optimization problems (4). All experiments are conducted on one Intel(R) Xeon(R) Platinum 8370C 2.80GHz CPU with 128GB of RAM, using the Julia programming language and `Mosek` version 11.0. For fixed (n, m) , we generate a random semidefinite matrix $\mathbf{A} = \mathbf{B}\mathbf{B}^\top \in \mathcal{S}_+^{nm}$ where the entries of $\mathbf{B} \in \mathbb{R}^{nm \times 10}$ are i.i.d. standard normal random variables. We solve the Shor relaxation (9) and sample $N = 100$ feasible solutions from Algorithm 2. For comparison, we also implement the following benchmarks:

- We sample N solutions uniformly at random (Uniform).
- We sample N solutions by applying the earlier deflation heuristic (Deflation).
- We follow the heuristic in [Burer and Park \(2024\)](#), namely, we project the reshaped leading eigenvector of $\mathbf{W}^*$ (Burer and Park).
- We implement the $1/(2\sqrt{2})$ -algorithm described in [Naor et al. \(2014, Figure 3\)](#).

We consider $n = 100$ and $m \in \{1, 2, 5, 10, 15, 20, 25, 30, 40, 60, 80, 100\}$. We generate five instances for each (n, m) .

The left panel of Figure 1 compares the average approximation ratio for these five algorithms. Confirming our theoretical analysis, we observe that our Algorithm 2 strongly outperforms Deflation and Uniform—theoretically, Uniform and Deflation achieve a $1/nm$ - and $1/m^2$ -approximation ratio respectively (Proposition 9 and 10). Despite its mildly stronger worst-case theoretical approximation ratio ($1/(2\sqrt{2})$ vs $1/3$), the algorithm of [Naor et al. \(2014\)](#) achieves a worst average empirical performance. We also find that the heuristic in [Burer and Park \(2024\)](#) achieves strong performance, comparable to that of Algorithm 2.

A more relevant performance metric in practice is the performance of the best solution found (out of N), rather than the average performance. Regarding the best solution found, the right panel of Figure 1 shows that the relative ordering of the methods remains unchanged, although the gap between methods shrinks.

To support the above discussion, Table 2 reports the time required to solve our semidefinite relaxation (9) for different values of m , using `Mosek` as the semidefinite optimization solver. Table 3 reports the time required by each feasibility heuristic (excluding time to solve the relaxation when needed). We observe that using the `Mosek` solver, our relaxation requires 1–2 orders of magnitude more computational time to solve than to run benchmark heuristics like deflation or uniform rounding, but facilitates the selection of significantly higher quality feasible solutions.

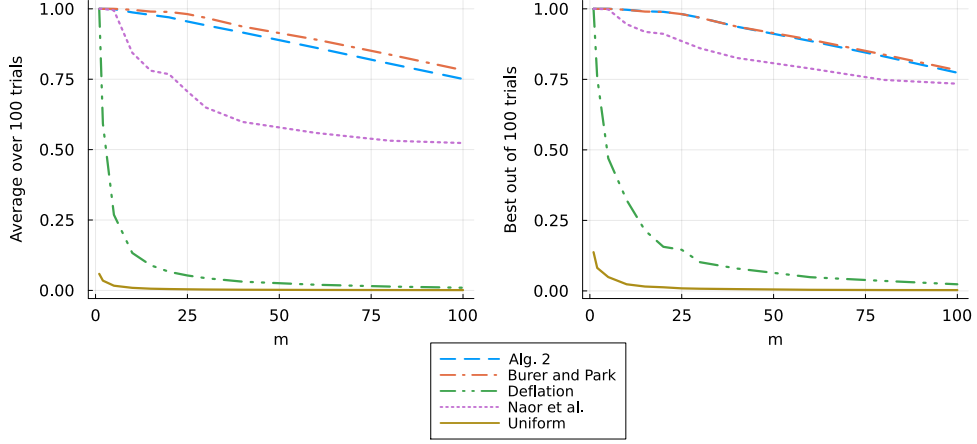


Fig. 1 Average approximation ratio $\text{vec}(\mathbf{Q})^\top \mathbf{A} \text{vec}(\mathbf{Q}) / \langle \mathbf{A}, \mathbf{W}^* \rangle$ (left panel) and approximation ratio $\text{vec}(\mathbf{Q})^\top \mathbf{A} \text{vec}(\mathbf{Q}) / \langle \mathbf{A}, \mathbf{W}^* \rangle$ of the best solution (right panel) over $N = 100$ samples, for different feasibility heuristics. Note that the method of [Burer and Park \(2024\)](#) only returns one solution. For each value of m , results are averaged over 5 instances.

Table 2 Computational time (average and standard deviation) for solving (9) for $n = 100$ and various values of m . Results are aggregated over 5 instances.

m	Average Time (s)	Std Dev (s)
1	2.01	0.15
2	3.71	0.34
5	7.35	0.63
10	19.06	1.69
15	38.07	2.94
20	80.33	7.34
25	140.67	21.99
30	244.0	37.19
40	510.24	82.15
60	1789.24	321.35
80	3240.3	552.06
100	7233.63	69.91

6 Conclusion

This paper proposes a new technique for relaxing and rounding quadratic optimization problems over semi-orthogonal matrices, which generalizes the blueprint of the Goemans–Williamson algorithm for BQO. Our algorithm provides a constant-factor approximation for the orthogonally constrained quadratic optimization problem (4), which subsumes the heterogeneous PCA problem (5) among others. Future work could

Table 3 Computational time (average and standard deviation) for different feasibility heuristics for $n = 100$ and various values of m . For methods that require solving the relaxation (9) (Alg. 2, [Burer and Park \(2024\)](#)), time for solving the relaxation is not included but reported in Table 2. Results are aggregated over 5 instances.

m	Alg. 2	Burer and Park	Deflation	Naor et al.	Uniform
2	0.01 (0.0)	0.01 (0.0)	0.23 (0.09)	49.25 (0.77)	0.01 (0.0)
5	0.07 (0.01)	0.06 (0.01)	0.8 (0.45)	48.27 (0.76)	0.07 (0.01)
10	0.34 (0.03)	0.22 (0.02)	2.7 (1.76)	47.43 (0.35)	0.33 (0.03)
15	0.6 (0.35)	0.5 (0.02)	5.81 (3.98)	48.13 (0.79)	0.53 (0.19)
20	0.99 (0.31)	0.92 (0.01)	9.17 (6.44)	47.34 (0.77)	0.92 (0.15)
25	1.44 (0.24)	1.49 (0.02)	9.32 (0.86)	47.91 (0.54)	1.44 (0.24)
30	2.59 (0.35)	2.36 (0.03)	12.1 (0.53)	48.31 (1.43)	2.59 (0.36)
40	7.0 (0.71)	5.51 (0.1)	20.37 (1.03)	50.92 (1.04)	7.0 (0.71)
60	26.98 (1.93)	17.01 (0.24)	43.14 (2.38)	73.5 (3.46)	26.97 (1.93)
80	79.83 (23.52)	37.32 (0.3)	77.44 (6.37)	111.68 (3.59)	79.82 (23.52)
100	124.65 (3.67)	66.53 (0.38)	116.52 (9.99)	172.1 (7.01)	124.34 (3.54)

investigate the theoretical analysis of other sampling schemes or extend the approach (namely, solving a Shor relaxation followed by a sampling algorithm) to a broader class of problems, such as low-rank optimization problems.

References

- Afonso S Bandeira, Christopher Kennedy, and Amit Singer. Approximating the little Grothendieck problem over the orthogonal and unitary groups. *Mathematical Programming*, 160:433–475, 2016.
- Alexander Barvinok. A remark on the rank of positive semidefinite matrices subject to affine constraints. *Discrete & Computational Geometry*, 25(1):23–31, 2001.
- Robert M Bell and Yehuda Koren. Lessons from the Netflix prize challenge. *ACM SIGKDD Explorations Newsletter*, 9(2):75–79, 2007.
- Aharon Ben-Tal and Arkadi Nemirovski. On tractable approximations of uncertain linear matrix inequalities affected by interval uncertainty. *SIAM Journal on Optimization*, 12(3):811–833, 2002.
- Dimitris Bertsimas and Yinyu Ye. Semidefinite relaxations, multivariate normal distributions, and order statistics. In *Handbook of Combinatorial Optimization*, pages 1473–1491. Springer, 1998.
- Dimitris Bertsimas, Martin S Copenhaver, and Rahul Mazumder. Certifiably optimal low rank factor analysis. *Journal of Machine Learning Research*, 18(29):1–53, 2017.
- Dimitris Bertsimas, Ryan Cory-Wright, and Jean Pauphilet. Mixed-projection conic optimization: A new paradigm for modeling rank constraints. *Operations Research*, 70(6):3321–3344, 2022.

- Marianna Bolla, György Michaletzky, Gábor Tusnády, and Margit Ziermann. Extrema of sums of heterogeneous quadratic forms. *Linear Algebra and its Applications*, 269(1-3):331–365, 1998.
- Jop Briët, Fernando Mário de Oliveira Filho, and Frank Vallentin. The positive semidefinite Grothendieck problem with rank constraint. In *International Colloquium on Automata, Languages, and Programming*, pages 31–42. Springer, 2010.
- Jop Briët, Harry Buhrman, and Ben Toner. A generalized Grothendieck inequality and nonlocal correlations that require high entanglement. *Communications in Mathematical Physics*, 305(3):827–843, 2011.
- Jop Briët, Oded Regev, and Rishi Saket. Tight hardness of the non-commutative Grothendieck problem. In *2015 IEEE 56th Annual Symposium on Foundations of Computer Science*, pages 1108–1122. IEEE, 2015.
- Samuel Burer and Kyungchan Park. A strengthened SDP relaxation for quadratic optimization over the Stiefel manifold. *Journal of Optimization Theory and Applications*, 202(1):320–339, 2024.
- Hongbo Dong, Kun Chen, and Jeff Linderoth. Regularization vs. relaxation: A conic optimization perspective of statistical variable selection. *arXiv preprint arXiv:1510.06083*, 2015.
- Alexandre d’Aspremont and Stephen Boyd. Relaxations and randomized methods for nonconvex QCQPs. *EE392o Class Notes, Stanford University*, 1:1–16, 2003.
- Kyle Gilman, Sam Burer, and Laura Balzano. A semidefinite relaxation for sums of heterogeneous quadratics on the Stiefel manifold. *SIAM Journal on Matrix Analysis and Applications*, 46(2):1091–1116, 2025.
- Michel X Goemans and David P Williamson. Improved approximation algorithms for maximum cut and satisfiability problems using semidefinite programming. *Journal of the ACM (JACM)*, 42(6):1115–1145, 1995.
- Harold V Henderson and Shayle R Searle. The vec-permutation matrix, the vec operator and Kronecker products: A review. *Linear and Multilinear Algebra*, 9(4):271–288, 1981.
- Zehua Lai, Lek-Heng Lim, and Tianyun Tang. Stiefel optimization is NP-hard. *Optimization Letters*, 2026. doi: 10.1007/s11590-025-02271-9.
- Zhi-Quan Luo, Wing-Kin Ma, Anthony Man-Cho So, Yinyu Ye, and Shuzhong Zhang. Semidefinite relaxation of quadratic optimization problems. *IEEE Signal Processing Magazine*, 27(3):20–34, 2010.

- Lester Mackey. Deflation methods for sparse PCA. *Advances in Neural Information Processing Systems*, 21, 2008.
- Elizabeth S Meckes. *The Random Matrix Theory of the Classical Compact Groups*, volume 218. Cambridge University Press, 2019.
- Milán Mosonyi and Tomohiro Ogawa. Quantum hypothesis testing and the operational interpretation of the quantum rényi relative entropies. *Communications in Mathematical Physics*, 334(3):1617–1648, 2015.
- Assaf Naor, Oded Regev, and Thomas Vidick. Efficient rounding for the noncommutative Grothendieck inequality. *Theory of Computing*, 10(1):257–295, 2014.
- Sahand Negahban and Martin J Wainwright. Estimation of (near) low-rank matrices with noise and high-dimensional scaling. *Annals of Statistics*, 39:1069–1097, 2011.
- Arkadi Nemirovski. Sums of random symmetric matrices and quadratic optimization under orthogonality constraints. *Mathematical Programming*, 109(2-3):283–317, 2007.
- Yu Nesterov. Semidefinite relaxation and nonconvex quadratic optimization. *Optimization Methods and Software*, 9(1-3):141–160, 1998.
- Gábor Pataki. On the rank of extreme matrices in semidefinite programs and the multiplicity of optimal eigenvalues. *Mathematics of Operations Research*, 23(2):339–358, 1998.
- Daniel Rehfeldt, Thorsten Koch, and Yuji Shinano. Faster exact solution of sparse MaxCut and QUBO problems. *Mathematical Programming Computation*, 15(3):445–470, 2023.
- Naum Z Shor. Quadratic optimization problems. *Soviet Journal of Computer and Systems Sciences*, 25:1–11, 1987.
- Anthony Man-Cho So. Moment inequalities for sums of random matrices and their applications in optimization. *Mathematical Programming*, 130(1):125–151, 2011.
- Antonia M Tulino, Sergio Verdú, et al. *Random Matrix Theory and Wireless Communications*, volume 1 of *Foundations and Trends® in Communications and Information Theory*. Now Publishers, Inc., 2004.
- Lieven Vandenbergh and Stephen Boyd. Semidefinite programming. *SIAM Review*, 38(1):49–95, 1996.
- David P Williamson and David B Shmoys. *The Design of Approximation Algorithms*. Cambridge University Press, 2011.

Henry Wolkowicz, Romesh Saigal, and Lieven Vandenbergh. *Handbook of Semidefinite Programming: Theory, Algorithms, and Applications*, volume 27. Springer Science & Business Media, 1998.

Improved Approximation Algorithms for Orthogonally Constrained Problems Using Semidefinite Optimization (Appendix)

Appendix A Connection with Binary Quadratic Optimization and the Original Goemans–Williamson Algorithm

We connect our quadratic semi-orthogonal optimization problem (4), its semidefinite relaxation (9), and our rounding algorithm (Algorithm 2) to the canonical binary quadratic optimization problem.

First, we show the following reduction result between Problems (1) and (4):

Proposition A.1 *Consider an instance of the binary quadratic optimization problem (1) with $\mathbf{Q} \succeq \mathbf{0}$. We can construct an optimization problem of the form (4) with $n \geq m$ and such that:*

- *any feasible solution to (1) can be converted (in polynomial time) into a feasible solution to (4) with the same objective value;*
- *any feasible solution to (4) can be converted (in polynomial time) into a feasible solution to (1) with equal or higher objective value.*

Proof Consider a binary quadratic optimization problem (1):

$$\max_{\mathbf{z} \in \{-1,1\}^m} \mathbf{z}^\top \mathbf{Q} \mathbf{z}.$$

Fix an integer $n \geq m$ and denote $\{\mathbf{e}_i\}_{i=1,\dots,n}$ the canonical basis of $\mathbb{R}^n$. Define the (i, j) block of the matrix $\mathbf{A}$ as $\mathbf{A}^{(i,j)} := Q_{i,j} \mathbf{e}_i \mathbf{e}_j^\top$. In particular, we have that $\mathbf{Q} \succeq \mathbf{0} \iff \mathbf{A} \succeq \mathbf{0}$. Consider the corresponding instance of Problem (4).

For any feasible solution to (1), we can construct a feasible solution to (4) with equal cost. Indeed, for each $i \in [m]$, we can define $\mathbf{u}_i := z_i \mathbf{e}_i$. By construction, we have $\mathbf{u}_i^\top \mathbf{u}_j = 0$ if $i \neq j$ and $\mathbf{u}_i^\top \mathbf{u}_i = z_i^2 = 1$. With this construction,

$$\mathbf{z}^\top \mathbf{Q} \mathbf{z} = \sum_{i,j \in [m]} Q_{i,j} z_i z_j = \sum_{i,j \in [m]} (\mathbf{u}_i^\top \mathbf{e}_i) Q_{i,j} (\mathbf{e}_j^\top \mathbf{u}_j) = \sum_{i,j \in [m]} \mathbf{u}_i^\top \mathbf{A}^{(i,j)} \mathbf{u}_j.$$

Alternatively, consider a feasible solution to this instance of (4), with objective value

$$\sum_{i,j \in [m]} \mathbf{u}_i^\top \mathbf{A}^{(i,j)} \mathbf{u}_j = \sum_{i,j \in [m]} (\mathbf{u}_i^\top \mathbf{e}_i) Q_{i,j} (\mathbf{e}_j^\top \mathbf{u}_j) = \sum_{i,j \in [m]} u_{i,i} Q_{i,j} u_{j,j}.$$

We show how to construct a feasible solution $\mathbf{z} \in \{-1,1\}^m$ to (1) with objective greater than or equal to $\sum_{i,j} u_{i,i} Q_{i,j} u_{j,j}$. Start from $z_i := u_{i,i} \in [-1,1]$. Fix z_i for all $i > 1$. The function $z_1 \in [-1,1] \mapsto \sum_{i,j} Q_{i,j} z_i z_j$ is convex because $\mathbf{Q} \succeq \mathbf{0}$. Thus we can shift z_1 to one of the endpoints, -1 or 1 , without decreasing the value. Proceeding in this way with the other coordinates $z_2, \dots, z_m$, we construct a solution $\mathbf{z} \in \{-1,1\}^m$ such that $\mathbf{z}^\top \mathbf{Q} \mathbf{z} \geq \sum_{i,j} u_{i,i} Q_{i,j} u_{j,j}$. $\square$

First of all, Proposition A.1 reduces any BQO problem with a positive semidefinite objective matrix $\mathbf{Q} \succeq \mathbf{0}$ to a problem of the form (4) with equal objective value. We note that our procedure to construct a feasible solution to BQO from a solution to (4) is analogous to that of Naor et al. (2014) and leverages the convexity of the objective. In particular, the standard (nonnegative weight) Max-Cut can be formulated as a BQO problem of this class, with $\mathbf{Q}$ being the (weighted) Laplacian of the graph. So Proposition A.1 shows that Problem (4) is NP-hard. Lai et al. (2026, Theorem 3.1) provide an alternative proof of this result. They also use a reduction from Max-Cut. However, their BQO formulation of Max-Cut is different, and the corresponding matrix $\mathbf{Q}$ is not PSD so they need to use a different decoding scheme.

Second, the reduction in Proposition A.1 is approximation-preserving: Any polynomial-time α -approximation (in objective value) for Problem (4) would immediately yield an α -approximation for Max-Cut. Therefore, the inapproximability threshold of $2/\pi + \varepsilon$ for BQO (Br  t et al. 2015, Theorem I.3) transfers to Problem (4).

In the rest of this section, we show how our semidefinite relaxation and our relax-and-project algorithm would work in the special case where Problem (4) encodes an instance of (1).

First, in this case, let us observe that optimal solutions to (4) can be found among solutions of the form $\mathbf{u}_i = z_i \mathbf{e}_i$ with $z_i \in \{-1, 1\}$. In other words, although we do not impose the constraint that each vector $\mathbf{u}_i$ should be colinear to $\mathbf{e}_i$, optimality naturally enforces this constraint.

Proof For any feasible solution to (4), its objective value is $\sum_{i,j} \mathbf{u}_i^\top \mathbf{A}^{(i,j)} \mathbf{u}_j = \sum_{i,j} u_{i,i} Q_{i,j} u_{j,j}$. Denoting $z_i := u_{i,i} \in [-1, 1]$, we have that

$$(4) \leq \max_{\mathbf{z} \in [-1,1]^m} \mathbf{z}^\top \mathbf{Q} \mathbf{z} = \max_{\mathbf{z} \in \{-1,1\}^m} \mathbf{z}^\top \mathbf{Q} \mathbf{z},$$

where the last equality follows from the fact that $\mathbf{Q} \succeq \mathbf{0}$. Conversely, for each $\mathbf{z} \in \{-1, 1\}^m$ the matrix defined as $\mathbf{u}_i = z_i \mathbf{e}_i$ is feasible for (4) and achieves an objective value of $\mathbf{z}^\top \mathbf{Q} \mathbf{z}$. $\square$

We now generalize this observation to the semidefinite relaxation of (4) and show that, without loss of optimality, the semidefinite variable $\mathbf{W}$ is of the form $\mathbf{W}^{(i,j)} = Z_{i,j} \mathbf{e}_i \mathbf{e}_j^\top$ for some matrix $\mathbf{Z} \in \mathcal{S}_+^m$.

Proof For any feasible solution to the semidefinite relaxation (9), its objective value is

$$\langle \mathbf{W}, \mathbf{A} \rangle = \sum_{i,j \in [m]} \langle \mathbf{W}^{(i,j)}, \mathbf{A}^{(i,j)} \rangle = \sum_{i,j \in [m]} Q_{i,j} \mathbf{e}_i^\top \mathbf{W}^{(i,j)} \mathbf{e}_j = \langle \mathbf{Z}, \mathbf{Q} \rangle,$$

with $Z_{i,j} := \mathbf{e}_i^\top \mathbf{W}^{(i,j)} \mathbf{e}_j$. The constraint $\mathbf{W} \succeq \mathbf{0}$ implies $\mathbf{Z} \succeq \mathbf{0}$ and the constraints $\text{tr}(\mathbf{W}^{(j,j)}) = 1$, $j \in [m]$ imply $Z_{j,j} \leq 1$, $j \in [m]$. For any $j \in [m]$, the function $Z_{j,j} \in [0, 1] \mapsto \langle \mathbf{Z}, \mathbf{Q} \rangle$ is linear with slope $Q_{j,j} \geq 0$ so, for any feasible $\mathbf{Z}$, setting $Z_{j,j}$ to 1 cannot decrease the objective value (and does not break feasibility). Consequently, at optimality, we can assume that $Z_{j,j} = W_{j,j}^{(j,j)} = 1$. However, $\text{tr}(\mathbf{W}^{(j,j)}) = 1$. Thus, all other diagonal coefficients of the block $\mathbf{W}^{(j,j)}$ are equal to 0. Recall that, for a positive semidefinite matrix, if a diagonal coefficient is equal to 0, then its entire column/row has to be equal to 0. Because $\mathbf{W}^{(j,j)} \succeq \mathbf{0}$, it means that $\mathbf{W}^{(j,j)} = Z_{j,j} \mathbf{e}_j \mathbf{e}_j^\top$. Regarding the off-diagonal blocks of $\mathbf{W}$, for (i, j) with $i \neq j$, because $\mathbf{W} \succeq \mathbf{0}$, the columns j' of $\mathbf{W}^{(i,j)}$ with $j' \neq j$

are equal to 0, i.e., only the j th column of $\mathbf{W}^{(i,j)}$ can be nonzero. Similarly, all the rows i' with $i' \neq i$ of $\mathbf{W}^{(i,j)}$ are equal to 0. All in all, we have that the blocks of $\mathbf{W}$ are of the form $\mathbf{W}^{(i,j)} = Z_{i,j} \mathbf{e}_i \mathbf{e}_j^\top$. $\square$

Considering a solution to the semidefinite relaxation of the form $\mathbf{W}^{(i,j)} = Z_{i,j} \mathbf{e}_i \mathbf{e}_j^\top$. The objective of (9) can thus be written as

$$\langle \mathbf{A}, \mathbf{W} \rangle = \sum_{i,j} \langle \mathbf{A}^{(i,j)}, \mathbf{W}^{(i,j)} \rangle = \sum_{i,j} Q_{i,j} Z_{i,j},$$

and the constraints on the matrix $\mathbf{W}$ are equivalent to:

$$\begin{aligned} \mathbf{W} \succeq \mathbf{0} : & & \mathbf{Z} \succeq \mathbf{0}, \\ \text{tr}(\mathbf{W}^{(j,j')}) = \delta_{j,j'} : & & Z_{j,j} = 1, \\ \sum_{i \in [m]} \mathbf{W}^{(i,i)} \preceq \mathbf{I}_n : & & Z_{i,i} \leq 1, \forall i \in [m]. \end{aligned}$$

So, we recover the semidefinite relaxation of BQO, (2), exactly.

Consider a solution to the semidefinite relaxation (9), $\mathbf{W}^*$, and $\mathbf{Z}^*$ such that $\mathbf{W}^{*(i,j)} = Z_{i,j}^* \mathbf{e}_i \mathbf{e}_j^\top$. By sampling $\text{vec}(\mathbf{G}) \sim \mathcal{N}(\mathbf{0}, \mathbf{W}^*)$, the sparsity pattern of $\mathbf{W}^*$ implies that each column of $\mathbf{G}$, $\mathbf{g}_i$, is of the form $\mathbf{g}_i = y_i \mathbf{e}_i$, with $\mathbf{y} \sim \mathcal{N}(\mathbf{0}, \mathbf{Z}^*)$. In this case, the matrix $\mathbf{G}$ is diagonal and its SVD can be written

$$\mathbf{G} = \mathbf{U} \mathbf{\Sigma} \mathbf{V}^\top := \begin{pmatrix} \mathbf{I}_m \\ \mathbf{0}_{(n-m) \times m} \end{pmatrix} \begin{pmatrix} |y_1| & & \\ & \ddots & \\ & & |y_m| \end{pmatrix} \begin{pmatrix} \text{sign}(y_1) & & \\ & \ddots & \\ & & \text{sign}(y_m) \end{pmatrix}.$$

We then generate $\mathbf{Q} = \mathbf{U} \mathbf{V}^\top$ so each column of $\mathbf{Q}$, $\mathbf{q}_i$, can be expressed as $\mathbf{q}_i = \text{sign}(y_i) \mathbf{e}_i$, i.e., $\hat{z}_i = \text{sign}(y_i)$, which is precisely the original Goemans–Williamson algorithm (Algorithm 1).

Appendix B Proof of Proposition 4

The construction of our instance relies on a group of orthogonal matrices known as the real Clifford group (see, e.g., Hashagen et al. 2018). We first show an intermediate result.

Lemma B.1 *Consider an integer $m = 2^q$ for some $q \geq 1$. Denoting $n := 1 + 2^{q^2+q+2}(2^q - 1) \prod_{j \in [q-1]} (4^j - 1)$, there exist $n - 1$ matrices $\mathbf{R}_i \in \mathcal{S}^m$ of the form $\mathbf{R}_i = \mathbf{U}_i \mathbf{D} \mathbf{U}_i^\top$ with $\mathbf{U}_i \in \mathbb{R}^{m \times m}$ orthogonal and $\mathbf{D} = \text{Diag}(\mathbf{I}_{m/2}, -\mathbf{I}_{m/2})$ such that*

$$\lambda_{\max} \left(\left[\frac{1}{m} \text{tr}(\mathbf{R}_i^\top \mathbf{R}_j) \right]_{i,j} \right) < \frac{n-1}{m-1}, \quad (\text{B.1})$$

and, for any unit-norm vector $\mathbf{u} \in \mathbb{R}^m$,

$$\frac{1}{n-1} \sum_{i \in [n-1]} \mathbf{R}_i \mathbf{u} \mathbf{u}^\top \mathbf{R}_i^\top = \frac{m}{(m-1)(m+2)} \mathbf{I}_m + \frac{m-2}{(m-1)(m+2)} \mathbf{u} \mathbf{u}^\top. \quad (\text{B.2})$$

Proof The proof relies on the real Clifford group, $\mathcal{H}$, a finite subset of the set of $m \times m$ orthogonal matrices (see, e.g., Hashagen et al. 2018). Denote $n-1 := |\mathcal{H}| = 2^{q^2+q+2}(2^q-1) \prod_{j \in [q-1]} (4^j-1)$ (see Nebe et al. 2006, Equation 6.3.4) and enumerate $\mathcal{H} = \{\mathbf{H}_1, \dots, \mathbf{H}_{n-1}\}$. The real Clifford group is a finite set of orthogonal matrices satisfying (among others) a property referred to as the exact twirling identity (see, e.g., Equation (50) in Hashagen et al. 2018): For any matrix $\mathbf{X} \in \mathbb{R}^{m^2 \times m^2}$,

$$\frac{1}{|\mathcal{H}|} \sum_{\mathbf{H} \in \mathcal{H}} (\mathbf{H} \otimes \mathbf{H}) \mathbf{X} (\mathbf{H} \otimes \mathbf{H})^\top = \mathbb{E}_{\mathbf{U}}[(\mathbf{U} \otimes \mathbf{U}) \mathbf{X} (\mathbf{U} \otimes \mathbf{U})^\top],$$

where $\mathbf{U}$ is distributed according to the uniform (Haar) measure over the set of orthogonal matrices. In other words, the average of $(\mathbf{H} \otimes \mathbf{H}) \mathbf{X} (\mathbf{H} \otimes \mathbf{H})^\top$ over the finite set of matrices $\mathcal{H}$ equals the average over $\mathbf{H}$ distributed according to the Haar measure. Audenaert et al. (2002, Section IV.B) give an explicit expression for the Haar-average, namely

$$\begin{aligned} \mathbf{T}(\mathbf{X}) &:= \mathbb{E}_{\mathbf{U}}[(\mathbf{U} \otimes \mathbf{U}) \mathbf{X} (\mathbf{U} \otimes \mathbf{U})^\top] \\ &= \text{tr}(\mathbf{P}_\Omega \mathbf{X}) \mathbf{P}_\Omega + \frac{2}{m(m-1)} \text{tr}(\mathbf{P}_A \mathbf{X}) \mathbf{P}_A + \frac{2}{(m+2)(m-1)} [\text{tr}((\mathbf{P}_S - \mathbf{P}_\Omega) \mathbf{X})] (\mathbf{P}_S - \mathbf{P}_\Omega) \end{aligned}$$

where

- $\mathbf{P}_\Omega$ is the orthogonal projection onto the span of $\text{vec}(\mathbf{I}_m)$, i.e., $\mathbf{P}_\Omega = \text{vec}(\mathbf{I}_m) \text{vec}(\mathbf{I}_m)^\top / m$
- $\mathbf{P}_A$ is the orthogonal projection onto the span of $\{\text{vec}(\mathbf{A}) : \mathbf{A}^\top = -\mathbf{A}\}$,
- $\mathbf{P}_S$ is the orthogonal projection onto the span of $\{\text{vec}(\mathbf{S}) : \mathbf{S}^\top = \mathbf{S}\}$.

In particular, if $\mathbf{K} \in \mathbb{R}^{m^2 \times m^2}$ denotes the commutation matrix (see, e.g., Magnus and Neudecker 1979), i.e., the matrix $\mathbf{K}$ such that, for any $\mathbf{X} \in \mathbb{R}^{m \times m}$, $\text{vec}(\mathbf{X}^\top) = \mathbf{K} \text{vec}(\mathbf{X})$, then $\mathbf{P}_A = (\mathbf{I}_{m^2} - \mathbf{K})/2$ and $\mathbf{P}_S = (\mathbf{I}_{m^2} + \mathbf{K})/2$.

For $i = 1, \dots, n-1$, we define $\mathbf{R}_i = \mathbf{H}_i \mathbf{D} \mathbf{H}_i^\top \in \mathcal{S}^m$. Let us show that the matrices $\mathbf{R}_i$ constructed in this manner satisfy the desired properties.

To prove the first property, let us denote $\mathbf{v}_i = \text{vec}(\mathbf{R}_i) = (\mathbf{H}_i \otimes \mathbf{H}_i) \text{vec}(\mathbf{D})$ and $C_{i,j} := \mathbf{v}_i^\top \mathbf{v}_j / m$. With this notation,

$$\begin{aligned} \lambda_{\max}(\mathbf{C}) &= \frac{1}{m} \lambda_{\max}(\mathbf{V}^\top \mathbf{V}) = \frac{1}{m} \lambda_{\max}(\mathbf{V} \mathbf{V}^\top) \\ &= \frac{n-1}{m} \lambda_{\max} \left(\frac{1}{|\mathcal{H}|} \sum_{\mathbf{H} \in \mathcal{H}} (\mathbf{H} \otimes \mathbf{H}) \text{vec}(\mathbf{D}) \text{vec}(\mathbf{D})^\top (\mathbf{H} \otimes \mathbf{H}) \right). \end{aligned}$$

The matrix on the right-hand side is exactly an average over the real Clifford group, hence it is equal to $\mathbf{T}(\text{vec}(\mathbf{D}) \text{vec}(\mathbf{D})^\top)$. We compute

$$\begin{aligned} \text{tr}(\mathbf{P}_\Omega \text{vec}(\mathbf{D}) \text{vec}(\mathbf{D})^\top) &= \frac{1}{m} (\text{vec}(\mathbf{D})^\top \text{vec}(\mathbf{I}_m))^2 = \frac{1}{m} (\text{tr}(\mathbf{D}))^2 = 0, \\ \text{tr}(\mathbf{P}_A \text{vec}(\mathbf{D}) \text{vec}(\mathbf{D})^\top) &= \text{tr}(\mathbf{0}) = 0, \\ \text{tr}(\mathbf{P}_S \text{vec}(\mathbf{D}) \text{vec}(\mathbf{D})^\top) &= \text{vec}(\mathbf{D})^\top \text{vec}(\mathbf{D}) = m, \end{aligned}$$

so

$$\mathbf{T}(\text{vec}(\mathbf{D}) \text{vec}(\mathbf{D})^\top) = \frac{2m}{(m+2)(m-1)}(\mathbf{P}_S - \mathbf{P}_\Omega).$$

Consequently,

$$\lambda_{\max}(\mathbf{C}) \leq \frac{n-1}{m} \frac{2m}{(m+2)(m-1)} = \frac{n-1}{m-1} \frac{2}{m+2} < \frac{n-1}{m-1},$$

for $m \geq 2$ (or, equivalently, $q \geq 1$).

For the second property, let us consider the vectorized version of the matrix identity

$$\begin{aligned} \frac{1}{n-1} \sum_{i \in [n-1]} \text{vec}(\mathbf{R}_i \mathbf{u} \mathbf{u}^\top \mathbf{R}_i^\top) &= \frac{1}{n-1} \sum_{i \in [n-1]} (\mathbf{R}_i \otimes \mathbf{R}_i) \text{vec}(\mathbf{u} \mathbf{u}^\top) \\ &= \left(\frac{1}{|\mathcal{H}|} \sum_{\mathbf{H} \in \mathcal{H}} (\mathbf{H} \otimes \mathbf{H})(\mathbf{D} \otimes \mathbf{D})(\mathbf{H} \otimes \mathbf{H}) \right) \text{vec}(\mathbf{u} \mathbf{u}^\top) \\ &= \mathbf{T}(\mathbf{D} \otimes \mathbf{D}) \text{vec}(\mathbf{u} \mathbf{u}^\top). \end{aligned}$$

We have

$$\text{tr}(\mathbf{P}_\Omega(\mathbf{D} \otimes \mathbf{D})) = \frac{1}{m} \text{vec}(\mathbf{D} \mathbf{I}_m \mathbf{D})^\top \text{vec}(\mathbf{I}_m) = 1.$$

For $\mathbf{P}_A(\mathbf{D} \otimes \mathbf{D})$ and $\mathbf{P}_S(\mathbf{D} \otimes \mathbf{D})$, we use the fact that the commutation matrix satisfies $\text{tr}(\mathbf{K}(\mathbf{A} \otimes \mathbf{B})) = \text{tr}(\mathbf{A}\mathbf{B})$ for any matrices $\mathbf{A}, \mathbf{B} \in \mathbb{R}^{m \times m}$ (Theorem 3.1(xiii) in [Magnus and Neudecker 1979](#)) to write

$$\begin{aligned} \text{tr}(\mathbf{P}_A(\mathbf{D} \otimes \mathbf{D})) &= \frac{1}{2} (\text{tr}(\mathbf{D} \otimes \mathbf{D}) - \text{tr}(\mathbf{K}(\mathbf{D} \otimes \mathbf{D}))) = -\frac{1}{2}m, \\ \text{tr}(\mathbf{P}_S(\mathbf{D} \otimes \mathbf{D})) &= \frac{1}{2} (\text{tr}(\mathbf{D} \otimes \mathbf{D}) + \text{tr}(\mathbf{K}(\mathbf{D} \otimes \mathbf{D}))) = \frac{1}{2}m. \end{aligned}$$

As a result,

$$\mathbf{T}(\mathbf{D} \otimes \mathbf{D}) = \mathbf{P}_\Omega - \frac{1}{m-1} \mathbf{P}_A + \frac{m-2}{(m+2)(m-1)}(\mathbf{P}_S - \mathbf{P}_\Omega).$$

The matrix $\mathbf{u} \mathbf{u}^\top$ is symmetric, so $\mathbf{P}_A \text{vec}(\mathbf{u} \mathbf{u}^\top) = \mathbf{0}$ and $\mathbf{P}_S \text{vec}(\mathbf{u} \mathbf{u}^\top) = \text{vec}(\mathbf{u} \mathbf{u}^\top)$. Furthermore,

$$\mathbf{P}_\Omega \text{vec}(\mathbf{u} \mathbf{u}^\top) = \frac{1}{m} \text{vec}(\mathbf{I}_m) \text{vec}(\mathbf{I}_m)^\top \text{vec}(\mathbf{u} \mathbf{u}^\top) = \frac{1}{m} \text{tr}(\mathbf{u} \mathbf{u}^\top) \text{vec}(\mathbf{I}_m) = \frac{1}{m} \text{vec}(\mathbf{I}_m).$$

Therefore, we proved that

$$\begin{aligned} \text{vec} \left(\frac{1}{n-1} \sum_{i \in [n-1]} \mathbf{R}_i \mathbf{u} \mathbf{u}^\top \mathbf{R}_i^\top \right) &= \mathbf{T}(\mathbf{D} \otimes \mathbf{D}) \text{vec}(\mathbf{u} \mathbf{u}^\top) \\ &= \frac{1}{m} \left(1 - \frac{m-2}{(m+2)(m-1)} \right) \text{vec}(\mathbf{I}_m) + \frac{m-2}{(m+2)(m-1)} \text{vec}(\mathbf{u} \mathbf{u}^\top) \\ &= \frac{m}{(m+2)(m-1)} \text{vec}(\mathbf{I}_m) + \frac{m-2}{(m+2)(m-1)} \text{vec}(\mathbf{u} \mathbf{u}^\top) \end{aligned}$$

which concludes the proof. $\square$

With this intermediate result, we can turn to the proof of Proposition 4.

Proof of Proposition 4 Fix $m = 2^q$ with $q \geq 1$. By Lemma B.1, there exists $n \geq m$ and symmetric matrices $\mathbf{R}_2, \dots, \mathbf{R}_n$ of the form $\mathbf{R}_i = \mathbf{U}_i \mathbf{D} \mathbf{U}_i^\top$ with $\mathbf{U}_i \in \mathbb{R}^{m \times m}$ orthogonal and $\mathbf{D} = \text{Diag}(\mathbf{I}_{m/2}, -\mathbf{I}_{m/2})$ satisfying (B.1) and (B.2).

For this n , let $\mathbf{e}_1 \in \mathbb{R}^n$ denote the first standard basis vector in $\mathbb{R}^n$ and define $\mathbf{A} := \mathbf{I}_m \otimes (\mathbf{e}_1 \mathbf{e}_1^\top) = \text{Diag}(\mathbf{e}_1 \mathbf{e}_1^\top, \dots, \mathbf{e}_1 \mathbf{e}_1^\top) \in \mathcal{S}_+^{nm}$.

Given that $\mathbf{A}$ is block-diagonal with identical diagonal blocks, Problem (4) is equivalent to finding the top- m eigenvectors of $\mathbf{e}_1 \mathbf{e}_1^\top$. The optimal objective value is 1 (e.g., achieved by the semi-orthogonal matrix $\mathbf{U} = [\mathbf{e}_1, \dots, \mathbf{e}_m]$). Similarly, the objective value of the semidefinite relaxation (9) is 1: for any feasible $\mathbf{W} \in \mathcal{S}_+^{nm}$, we have

$$\langle \mathbf{A}, \mathbf{W} \rangle = \sum_{j=1}^m \langle \mathbf{e}_1 \mathbf{e}_1^\top, \mathbf{W}^{(j,j)} \rangle = \langle \mathbf{e}_1 \mathbf{e}_1^\top, \sum_{j=1}^m \mathbf{W}^{(j,j)} \rangle \leq \langle \mathbf{e}_1 \mathbf{e}_1^\top, \mathbf{I}_n \rangle = 1,$$

and the reverse inequality holds because it is a relaxation. For any matrix $\mathbf{Q} \in \mathbb{R}^{n \times m}$, its objective value is $\text{vec}(\mathbf{Q})^\top \mathbf{A} \text{vec}(\mathbf{Q}) = \sum_{j \in [m]} (\mathbf{q}_j^\top \mathbf{e}_1)^2 = \|\mathbf{e}_1^\top \mathbf{Q}\|_2^2$, i.e., it is equal to the squared norm of the first row of $\mathbf{Q}$.

We construct the matrix $\mathbf{G}$ for Algorithm 2 row by row: we sample $\mathbf{x} \sim \mathcal{N}(\mathbf{0}, \mathbf{I}_m)$ and define

$$\mathbf{G} = \frac{1}{\sqrt{m}} \begin{pmatrix} \mathbf{x}^\top \\ \rho \mathbf{x}^\top \mathbf{R}_2 \\ \vdots \\ \rho \mathbf{x}^\top \mathbf{R}_n \end{pmatrix} \quad \text{or equivalently} \quad \mathbf{G}^\top = \frac{1}{\sqrt{m}} [\mathbf{x} \mid \rho \mathbf{R}_2^\top \mathbf{x} \mid \dots \mid \rho \mathbf{R}_n^\top \mathbf{x}],$$

with $\rho = \sqrt{(m-1)/(n-1)} \leq 1$. The vector $\text{vec}(\mathbf{G})$ is normally distributed, $\mathbb{E}[\text{vec}(\mathbf{G})] = \mathbf{0}$, so we have implicitly constructed $\mathbf{W} := \mathbb{E}[\text{vec}(\mathbf{G}) \text{vec}(\mathbf{G})^\top] \succeq \mathbf{0}$. In the remainder of this proof, we show that this matrix $\mathbf{W}$ is feasible and optimal for the semidefinite relaxation, and compute $\text{vec}(\mathbf{Q})^\top \mathbf{A} \text{vec}(\mathbf{Q})$ analytically.

First, let us show that $\mathbf{W}$ is feasible for the SDP relaxation. The matrix $\mathbf{W}$ is feasible if and only if the corresponding random matrix $\mathbf{G}$ satisfies $\mathbb{E}[\mathbf{G}^\top \mathbf{G}] = \mathbf{I}_m$ and $\mathbb{E}[\mathbf{G} \mathbf{G}^\top] \preceq \mathbf{I}_n$. In our case, we have

$$\mathbf{G}^\top \mathbf{G} = \sum_{i \in [n]} \mathbf{G}^\top \mathbf{e}_i \mathbf{e}_i^\top \mathbf{G} = \frac{1}{m} \mathbf{x} \mathbf{x}^\top + \frac{\rho^2}{m} \sum_{i=2}^n \mathbf{R}_i^\top \mathbf{x} \mathbf{x}^\top \mathbf{R}_i.$$

Taking expectations and using the facts that $\mathbb{E}[\mathbf{x} \mathbf{x}^\top] = \mathbf{I}_m$ and $\mathbf{R}_i^\top \mathbf{R}_i = \mathbf{U}_i^\top \mathbf{D}^2 \mathbf{U}_i = \mathbf{I}_m$, we have

$$\mathbb{E}[\mathbf{G}^\top \mathbf{G}] = \frac{1}{m} (1 + (n-1)\rho^2) \mathbf{I}_m = \mathbf{I}_m.$$

In addition,

$$\mathbb{E}[\mathbf{G} \mathbf{G}^\top] = \begin{pmatrix} 1 & \mathbf{0} \\ \mathbf{0} & \rho^2 \mathbf{C} \end{pmatrix} \quad \text{with } C_{i,j} := \frac{1}{m} \text{tr}(\mathbf{R}_i^\top \mathbf{R}_j).$$

By construction (Equation (B.1) in Lemma B.1), $\lambda_{\max}(\mathbf{C}) < \frac{n-1}{m-1} = \rho^{-2}$ so $\mathbb{E}[\mathbf{G} \mathbf{G}^\top] \preceq \mathbf{I}_n$. Hence, $\mathbf{W}$ is a feasible solution to the semidefinite relaxation (9).

Second, we have that $\langle \mathbf{A}, \mathbf{W} \rangle = \langle \mathbf{e}_1 \mathbf{e}_1^\top, \sum_{j=1}^m \mathbf{W}^{(j,j)} \rangle = \langle \mathbf{e}_1 \mathbf{e}_1^\top, \mathbb{E}[\mathbf{G} \mathbf{G}^\top] \rangle = 1$. So $\mathbf{W}$ is optimal.

Finally, let us evaluate the performance of $\mathbf{Q}$. By definition, $\mathbf{Q} = \mathbf{G}(\mathbf{G}^\top \mathbf{G})^{-1/2}$ so

$$\|\mathbf{e}_1^\top \mathbf{Q}\|_2^2 = \mathbf{e}_1^\top \mathbf{Q} \mathbf{Q}^\top \mathbf{e}_1 = \mathbf{e}_1^\top \mathbf{G} (\mathbf{G}^\top \mathbf{G})^{-1} \mathbf{G}^\top \mathbf{e}_1 = \frac{1}{m} \mathbf{x}^\top (\mathbf{G}^\top \mathbf{G})^{-1} \mathbf{x}.$$

Furthermore,

$$\mathbf{G}^\top \mathbf{G} = \frac{1}{m} \mathbf{x} \mathbf{x}^\top + \frac{\rho^2}{m} \sum_{i=2}^n \mathbf{R}_i^\top \mathbf{x} \mathbf{x}^\top \mathbf{R}_i.$$

Denoting $\mathbf{u} = \mathbf{x}/\|\mathbf{x}\|$, we have

$$\begin{aligned} \mathbf{G}^\top \mathbf{G} &= \|\mathbf{x}\|^2 \left(\frac{1}{m} \mathbf{u} \mathbf{u}^\top + \frac{\rho^2}{m} \sum_{i=2}^n \mathbf{R}_i^\top \mathbf{u} \mathbf{u}^\top \mathbf{R}_i \right), \\ \|\mathbf{e}_1^\top \mathbf{Q}\|_2^2 &= \mathbf{u}^\top \left(\mathbf{u} \mathbf{u}^\top + \rho^2 \sum_{i=2}^n \mathbf{R}_i^\top \mathbf{u} \mathbf{u}^\top \mathbf{R}_i \right)^{-1} \mathbf{u}. \end{aligned}$$

By construction (Equation (B.2) in Lemma B.1), we have, for any unit-norm vector $\mathbf{u}$,

$$\frac{1}{n-1} \sum_{i=2}^n \mathbf{R}_i^\top \mathbf{u} \mathbf{u}^\top \mathbf{R}_i = \frac{m}{(m-1)(m+2)} \mathbf{I}_m + \frac{m-2}{(m-1)(m+2)} \mathbf{u} \mathbf{u}^\top.$$

So,

$$\mathbf{u} \mathbf{u}^\top + \rho^2 \sum_{i=2}^n \mathbf{R}_i^\top \mathbf{u} \mathbf{u}^\top \mathbf{R}_i = \frac{m}{m+2} \mathbf{I}_m + \frac{2m}{m+2} \mathbf{u} \mathbf{u}^\top.$$

In particular, we get that $\mathbf{u}$ is an eigenvector of the matrix above with eigenvalue $3m/(m+2)$ so we have

$$\mathbf{u}^\top \left(\mathbf{u} \mathbf{u}^\top + \rho^2 \sum_{i=2}^n \mathbf{R}_i^\top \mathbf{u} \mathbf{u}^\top \mathbf{R}_i \right)^{-1} \mathbf{u} = \frac{m+2}{3m}.$$

In other words, $\text{vec}(\mathbf{Q})^\top \mathbf{A} \text{vec}(\mathbf{Q}) = \frac{m+2}{3m}$ almost surely. $\square$

Appendix C Technical Appendix to Section 3

C.1 Proof of Lemma 2

Proof As described in Section 2.2, in our implementation of Algorithm 2, we sample $\text{vec}(\mathbf{G}) \sim \mathcal{N}(\mathbf{0}_{nm}, \mathbf{W}^*)$ as $\text{vec}(\mathbf{G}) = \sum_{k \in [r]} \text{vec}(\mathbf{B}_k) z_k$ with $\mathbf{z} \sim \mathcal{N}(\mathbf{0}_r, \mathbf{I}_r)$ and $\mathbf{W}^* = \sum_{k \in [r]} \text{vec}(\mathbf{B}_k) \text{vec}(\mathbf{B}_k)^\top$ a Cholesky decomposition of $\mathbf{W}^*$. This construction interprets $\mathbf{G}$ as a matrix series, $\mathbf{G} = \sum_{k \in [r]} \mathbf{B}_k z_k$, as studied in the statistics literature (see, e.g., Tropp 2015).

Let us denote

$$\mathbf{V}_1 := \sum_{k \in [r]} \mathbf{B}_k \mathbf{B}_k^\top \text{ and } \mathbf{V}_2 := \sum_{k \in [r]} \mathbf{B}_k^\top \mathbf{B}_k,$$

so $\mathbb{E}[\mathbf{G} \mathbf{G}^\top] = \mathbf{V}_1$ and $\mathbb{E}[\mathbf{G}^\top \mathbf{G}] = \mathbf{V}_2$. Noting that $\mathbf{W}^{(i,j)} = \sum_{k \in [r]} \mathbf{B}_k \mathbf{e}_i \mathbf{e}_j^\top \mathbf{B}_k^\top$, we have

$$\begin{aligned} \mathbb{E} \left[\mathbf{G}^\top \mathbf{G} \right]_{i,j} &= \left(\sum_{k \in [r]} \mathbf{B}_k^\top \mathbf{B}_k \right)_{i,j} = \sum_{k \in [r]} \mathbf{e}_i^\top \mathbf{B}_k^\top \mathbf{B}_k \mathbf{e}_j = \text{tr}(\mathbf{W}^{(i,j)}), \\ \text{and } \mathbb{E} \left[\mathbf{G} \mathbf{G}^\top \right] &= \sum_{k \in [r]} \mathbf{B}_k \mathbf{B}_k^\top = \sum_{k \in [r]} \sum_{i \in [m]} \mathbf{B}_k \mathbf{e}_i \mathbf{e}_i^\top \mathbf{B}_k^\top = \sum_{i \in [m]} \mathbf{W}^{(i,i)}. \end{aligned}$$

The fact that $\mathbf{W}$ satisfies the constraints in (9) yields $\mathbf{V}_1 \preceq \mathbf{I}_n$ and $\mathbf{V}_2 = \mathbf{I}_m$, as claimed.

For the fourth moments, we define the following *Hermitian* Gaussian series

$$\mathbf{Y} := \sum_{k \in [r]} z_k \mathbf{A}_k \quad \text{with} \quad \mathbf{A}_k := \begin{pmatrix} \mathbf{0} & \mathbf{B}_k \\ \mathbf{B}_k^\top & \mathbf{0} \end{pmatrix}.$$

By construction,

$$\sum_{k \in [r]} \mathbf{A}_k^2 = \begin{pmatrix} \sum_k \mathbf{B}_k \mathbf{B}_k^\top & \mathbf{0} \\ \mathbf{0} & \sum_k \mathbf{B}_k^\top \mathbf{B}_k \end{pmatrix} = \begin{pmatrix} \mathbf{V}_1 & \mathbf{0} \\ \mathbf{0} & \mathbf{V}_2 \end{pmatrix} \quad \text{and} \quad \mathbf{Y}^4 = \begin{pmatrix} (\mathbf{G} \mathbf{G}^\top)^2 & \mathbf{0} \\ \mathbf{0} & (\mathbf{G}^\top \mathbf{G})^2 \end{pmatrix}.$$

Hence it suffices to upper bound $\mathbb{E}[\mathbf{Y}^4]$ and then read the diagonal blocks.

Since $\mathbf{z} \sim \mathcal{N}(\mathbf{0}, \mathbf{I}_r)$, Isserlis' theorem gives $\mathbb{E}[z_a z_b z_c z_d] = \delta_{ab} \delta_{cd} + \delta_{ac} \delta_{bd} + \delta_{ad} \delta_{bc}$ (Isserlis 1918). Expanding $\mathbf{Y}^4$ and taking expectation yields

$$\begin{aligned} \mathbb{E}[\mathbf{Y}^4] &= \sum_{a,b,c,d \in [r]} \mathbb{E}[z_a z_b z_c z_d] \mathbf{A}_a \mathbf{A}_b \mathbf{A}_c \mathbf{A}_d \\ &= \sum_{a,c \in [r]} \mathbf{A}_a^2 \mathbf{A}_c^2 + \sum_{a,b \in [r]} \mathbf{A}_a \mathbf{A}_b \mathbf{A}_a \mathbf{A}_b + \sum_{a,b \in [r]} \mathbf{A}_a \mathbf{A}_b^2 \mathbf{A}_a \\ &= \left(\sum_{k \in [r]} \mathbf{A}_k^2 \right)^2 + \sum_{a,b \in [r]} \mathbf{A}_a \mathbf{A}_b \mathbf{A}_a \mathbf{A}_b + \sum_{a,b \in [r]} \mathbf{A}_a \mathbf{A}_b^2 \mathbf{A}_a. \end{aligned}$$

Fix two indices a, b . By expanding the inequality $(\mathbf{A}_a \mathbf{A}_b - \mathbf{A}_b \mathbf{A}_a)^\top (\mathbf{A}_a \mathbf{A}_b - \mathbf{A}_b \mathbf{A}_a) \succeq \mathbf{0}$ and using the fact that $\mathbf{A}_a, \mathbf{A}_b$ are symmetric, we have

$$\mathbf{A}_a \mathbf{A}_b \mathbf{A}_a \mathbf{A}_b + \mathbf{A}_b \mathbf{A}_a \mathbf{A}_b \mathbf{A}_a \preceq \mathbf{A}_a \mathbf{A}_b^2 \mathbf{A}_a + \mathbf{A}_b \mathbf{A}_a^2 \mathbf{A}_b.$$

Summing over all (a, b) yields

$$\sum_{a,b \in [r]} \mathbf{A}_a \mathbf{A}_b \mathbf{A}_a \mathbf{A}_b \preceq \sum_{a,b \in [r]} \mathbf{A}_a \mathbf{A}_b^2 \mathbf{A}_a.$$

Therefore,

$$\mathbb{E}[\mathbf{Y}^4] \preceq \left(\sum_{k \in [r]} \mathbf{A}_k^2 \right)^2 + 2 \sum_{a,b \in [r]} \mathbf{A}_a \mathbf{A}_b^2 \mathbf{A}_a = \left(\sum_{a \in [r]} \mathbf{A}_a^2 \right)^2 + 2 \sum_{a \in [r]} \mathbf{A}_a \left(\sum_{b=1}^r \mathbf{A}_b^2 \right) \mathbf{A}_a.$$

Reading the first diagonal block, we have

$$\mathbb{E}[(\mathbf{G} \mathbf{G}^\top)^2] \preceq \mathbf{V}_1^2 + 2 \sum_{k \in [r]} \mathbf{B}_k \mathbf{V}_2 \mathbf{B}_k^\top \preceq \mathbf{V}_1^2 + 2 \lambda_{\max}(\mathbf{V}_2) \mathbf{V}_1 \preceq 3 \mathbf{I}_n.$$

Similarly for the second diagonal block,

$$\mathbb{E}[(\mathbf{G}^\top \mathbf{G})^2] \preceq \mathbf{V}_2^2 + 2 \sum_{k \in [r]} \mathbf{B}_k^\top \mathbf{V}_1 \mathbf{B}_k \preceq \mathbf{V}_2^2 + 2 \lambda_{\max}(\mathbf{V}_1) \mathbf{V}_2 \preceq 3 \mathbf{I}_m,$$

which concludes the proof. $\square$

C.2 Proof of Lemma 8

Proof Denote $\mathbf{N}$ a square root of $\mathbf{M}$. By Cauchy-Schwarz,

$$\text{tr}(\mathbf{M}) = \text{tr}(\mathbf{M} \mathbb{E}[\mathbf{S}]) = \mathbb{E}[\text{tr}(\mathbf{N} \mathbf{S}^{1/4} \mathbf{S}^{3/4} \mathbf{N})] \leq \sqrt{\mathbb{E}[\text{tr}(\mathbf{N} \mathbf{S}^{1/2} \mathbf{N})]} \sqrt{\mathbb{E}[\text{tr}(\mathbf{N} \mathbf{S}^{3/2} \mathbf{N})]}.$$

Applying Cauchy-Schwarz again, we have

$$\mathbb{E}[\text{tr}(\mathbf{N} \mathbf{S}^{3/2} \mathbf{N})] = \mathbb{E}[\text{tr}(\mathbf{N} \mathbf{S} \mathbf{S}^{1/2} \mathbf{N})] \leq \sqrt{\mathbb{E}[\text{tr}(\mathbf{N} \mathbf{S}^2 \mathbf{N})]} \sqrt{\mathbb{E}[\text{tr}(\mathbf{N} \mathbf{S} \mathbf{N})]} \leq \sqrt{c} \text{tr}(\mathbf{M}).$$

Combining these inequalities together, we have

$$\text{tr}(\mathbf{M})^2 \leq \mathbb{E}[\text{tr}(\mathbf{N} \mathbf{S}^{1/2} \mathbf{N})] \sqrt{c} \text{tr}(\mathbf{M}).$$

Rearranging the terms concludes the proof. $\square$

Appendix D A tighter dimension-dependent approximation ratio

Theorem D.1 Assume that $\mathbf{A} \succeq \mathbf{0}$. Let $\mathbf{W}^*$ denote an optimal solution to the semidefinite relaxation (9). The random matrix $\mathbf{Q} \in \mathbb{R}^{n \times m}$ generated by Algorithm 2 satisfies the inequality

$$\mathbb{E} [\text{vec}(\mathbf{Q})^\top \mathbf{A} \text{vec}(\mathbf{Q})] \geq \beta_{n,m} \langle \mathbf{A}, \mathbf{W}^* \rangle.$$

with

$$\beta_{n,m} := \min_{\lambda \in [0,1]} \int_0^\infty \left(1 + 2tm \frac{1-\lambda}{nm-1}\right)^{-(nm-1)/2} (1+2tm\lambda)^{-3/2} dt.$$

In particular, the constant $\beta_{n,m}$ satisfies the following properties:

- (a) For any integer m , $\beta_{n,m}$ is non-increasing in n . For any integer n , $\beta_{n,m}$ is non-increasing in m .
- (b) For any integer m , we have $\beta_{n,m} \rightarrow \beta_{\infty,m}$ as $n \rightarrow \infty$ with

$$\beta_{\infty,m} := \min_{\lambda \in [0,1]} \int_0^\infty e^{-tm(1-\lambda)} (1+2tm\lambda)^{-3/2} dt$$

- (c) For $m = 1$, $\beta_{n,1}$ is optimal, i.e., there exists a covariance matrix $\mathbf{W}^*$ satisfying $\mathbb{E} [\text{vec}(\mathbf{Q})^\top \mathbf{A} \text{vec}(\mathbf{Q})] = \beta_{n,1} \langle \mathbf{A}, \mathbf{W}^* \rangle$.

Proof From the proof of Proposition 5, we have

$$\mathbb{E} [\text{vec}(\mathbf{Q})^\top \mathbf{A} \text{vec}(\mathbf{Q})] \geq \langle \mathbf{A}, \mathbb{E} \left[\frac{\text{vec}(\mathbf{G}) \text{vec}(\mathbf{G})^\top}{\|\mathbf{G}\|_F^2} \right] \rangle.$$

We use the integral representation $1/x^2 = \int_{t \geq 0} e^{-tx^2} dt$ to obtain

$$\mathbb{E} \left[\frac{\text{vec}(\mathbf{G}) \text{vec}(\mathbf{G})^\top}{\|\mathbf{G}\|_F^2} \right] = \int_0^\infty \mathbb{E} [\text{vec}(\mathbf{G}) \text{vec}(\mathbf{G})^\top e^{-t\|\mathbf{G}\|_F^2}] dt.$$

Denote $r = \text{rank}(\mathbf{W}^*) \leq nm$ and consider an eigenvalue decomposition of $\mathbf{W}^*$, $\mathbf{W}^* = \mathbf{H} \mathbf{\Lambda} \mathbf{H}^\top$. We have $\text{vec}(\mathbf{G}) = \mathbf{H} \mathbf{\Lambda}^{1/2} \mathbf{z}$, with $\mathbf{z} \sim \mathcal{N}(\mathbf{0}_r, \mathbf{I}_r)$ and thus

$$\mathbb{E} [\text{vec}(\mathbf{G}) \text{vec}(\mathbf{G})^\top \exp(-t\|\mathbf{G}\|_F^2)] = \mathbf{H} \mathbf{\Lambda}^{1/2} \mathbb{E} [\mathbf{z} \mathbf{z}^\top \exp(-t\mathbf{z}^\top \mathbf{\Lambda} \mathbf{z})] \mathbf{\Lambda}^{1/2} \mathbf{H}^\top.$$

Furthermore, for $\mathbf{z} \sim \mathcal{N}(\mathbf{0}_r, \mathbf{I}_r)$, we can compute the expectation analytically:

$$\mathbb{E} [\mathbf{z} \mathbf{z}^\top \exp(-t\mathbf{z}^\top \mathbf{\Lambda} \mathbf{z})] = \frac{1}{\sqrt{\det(\mathbf{I}_r + 2t\mathbf{\Lambda})}} (\mathbf{I}_r + 2t\mathbf{\Lambda})^{-1}.$$

So, the integral is lower bounded by

$$\mathbf{H} \mathbf{\Lambda}^{1/2} \mathbf{B} \mathbf{\Lambda}^{1/2} \mathbf{H}^\top \text{ with } \mathbf{B} := \int_0^\infty \frac{1}{\sqrt{\det(\mathbf{I}_r + 2t\mathbf{\Lambda})}} (\mathbf{I}_r + 2t\mathbf{\Lambda})^{-1} dt.$$

To conclude, we show that there exists a scalar $\beta > 0$ such that $\mathbf{B} \succeq \beta \mathbf{I}_r$.

To find such β , observe that $\mathbf{B}$ is a diagonal matrix. Hence, it is sufficient to find a lower bound on its diagonal entries. Given the constraints on $\mathbf{W}^*$, the eigenvalues $\mathbf{\Lambda}$ must satisfy:

$\Lambda_i \geq 0$ (from $\mathbf{W}^* \succeq 0$), $\sum_{i=1}^r \Lambda_i = m$ (from $\text{tr}(\mathbf{W}^*) = \sum_{i \in [m]} \text{tr}(\mathbf{W}^{*(i,i)}) = m$). Hence, we can take

$$\beta = \min_{\mathbf{\Lambda} \in [0, m]^r : \sum_{i=1}^r \Lambda_i = m} \int_0^\infty \prod_{i=1}^r (1 + 2t\Lambda_{i'})^{-1/2} (1 + 2t\Lambda_1)^{-1} dt.$$

The function $\mathbf{\Lambda} \mapsto \int_0^\infty \prod_{i=1}^r (1 + 2t\Lambda_{i'})^{-1/2} (1 + 2t\Lambda_1)^{-1} dt$ is convex, and invariant under any permutation of the Λ_i , $i > 1$. So, by Jensen's inequality, we can restrict our attention to minimizers of the form $\Lambda_1 = \lambda$, $\Lambda_i = \frac{m-\lambda}{r-1}$, $i > 1$:

$$\beta = \min_{\lambda \in [0, m]} \int_0^\infty \left(1 + 2t \frac{m-\lambda}{r-1}\right)^{-(r-1)/2} (1 + 2t\lambda)^{-3/2} dt.$$

For a fixed value of (t, λ) , the integrand is decreasing in $r \leq nm$, so

$$\beta \geq \beta_{n,m} := \min_{\lambda \in [0, m]} \int_0^\infty \left(1 + 2t \frac{m-\lambda}{nm-1}\right)^{-(nm-1)/2} (1 + 2t\lambda)^{-3/2} dt.$$

The change of variable $\lambda \leftarrow \lambda/m$ leads to the desired expression for $\beta_{n,m}$.

Claim (a) follows from the fact that for any $\lambda \in [0, 1]$, the integrand is decreasing in n for m fixed, and decreasing in m for n fixed.

Claim (b) follows from the fact that the sequence $(1 + x/k)^{-k}$ converges monotonically to e^{-x} combined with the dominated convergence theorem.

Claim (c) follows from the fact that, for $m = 1$, $\mathbf{Q} = \mathbf{G}/\|\mathbf{G}\|_F$ so our analysis is exact. $\square$

Remark D.1 We observe that the bound $\beta_{n,m}$ is obtained by looking at the worst-case instance over all covariance matrices $\mathbf{W}^*$. In particular, we can obtain tighter values of $\beta_{n,m}$ by allowing for dependency on the rank of $\mathbf{W}^*$, r , instead of the ambient dimension n . For instance, if $r = 1$, we can get

$$\beta = \min_{\lambda \in [0, m]} \int_0^\infty (1 + 2t\lambda)^{-3/2} dt = \min_{\lambda \in [0, m]} \frac{1}{\lambda} = \frac{1}{m}.$$

Alternatively, by the Barvinok-Pataki bound, we know there exists some optimal solution $\mathbf{W}^*$ with rank at most $n + m$ and we could use this bound to refine our constant. Furthermore, if $\mathbf{W}^*$ has additional structure (e.g., $\mathbf{W}^*$ is block diagonal), we can derive additional constraints on the eigenvalues $\mathbf{\Lambda}$, hence tighter constants $\beta_{n,m}$.

Compared with Proposition 5, the value of Theorem D.1 is primarily computational. By solving numerically the one-dimensional minimization problem in λ , it provides tighter estimates of the performance of our algorithm, especially for small values of m , as reported in Table D.1. While the guarantee from Proposition 5 is independent of n , the constant $\beta_{n,m}$ in Theorem D.1 is monotonically decreasing with n , obtaining stronger approximation ratios for finite values of n . However, we should acknowledge that $\beta_{n,m}$ is worse than $1/3$ for $m > 2$ (actually, Remark D.1 identifies a class of matrices $\mathbf{W}^*$ for which $\beta_{n,m} \leq 1/m$).

Table D.1 Values of the approximation factor from Proposition 5 and Theorem D.1 for some values of n and m , with $m \leq n$.

m	Proposition 5	$\beta_{n,m}$ (Theorem D.1)			
		$n = 5$	$n = 10$	$n = 15$	$n = \infty$
1	0.636620	0.735264	0.706972	0.697920	0.680415
2	0.318310	0.353486	0.346734	0.344533	0.340208
3	0.212207	0.232640	0.229689	0.228720	0.226805
4	0.159155	0.173367	0.171721	0.171179	0.170104
5	0.127324	0.138164	0.137116	0.136770	0.136083
10	0.063662	—	0.068299	0.068213	0.068042
15	0.042441	—	—	0.045437	0.045361

References

- Koenraad Audenaert, Bart De Moor, Karl Gerd H Vollbrecht, and Reinhard F Werner. Asymptotic relative entropy of entanglement for orthogonally invariant states. *Physical Review A*, 66(3):032310, 2002.
- Jop Briët, Oded Regev, and Rishi Saket. Tight hardness of the non-commutative Grothendieck problem. In *2015 IEEE 56th Annual Symposium on Foundations of Computer Science*, pages 1108–1122. IEEE, 2015.
- AK Hashagen, Steven T Flammia, David Gross, and Joel J Wallman. Real randomized benchmarking. *Quantum*, 2:85, 2018.
- Leon Isserlis. On a formula for the product-moment coefficient of any order of a normal frequency distribution in any number of variables. *Biometrika*, 12(1/2):134–139, 1918.
- Zehua Lai, Lek-Heng Lim, and Tianyun Tang. Stiefel optimization is NP-hard. *Optimization Letters*, 2026. doi: 10.1007/s11590-025-02271-9.
- Jan R Magnus and Heinz Neudecker. The commutation matrix: some properties and applications. *The Annals of Statistics*, 7(2):381–394, 1979.
- Assaf Naor, Oded Regev, and Thomas Vidick. Efficient rounding for the noncommutative Grothendieck inequality. *Theory of Computing*, 10(1):257–295, 2014.
- Gabriele Nebe, Eric M Rains, and Neil JA Sloane. *Self-Dual Codes and Invariant Theory*. Springer Science & Business Media, 2006.
- Joel A Tropp. An introduction to matrix concentration inequalities. *Foundations and Trends® in Machine Learning*, 8(1-2):1–230, 2015.