



DISSERTATION | DOCTORAL THESIS

Titel | Title

Conic approaches to mixed-integer QPs with complementarity constraints and applications to sparse machine learning models

verfasst von | submitted by

Bo Peng

angestrebter akademischer Grad | in partial fulfilment of the requirements for the degree of
Doctor of Philosophy (PhD)

Wien | Vienna, 2024

Studienkennzahl lt. Studienblatt | Degree
programme code as it appears on the
student record sheet:

UA 794 370 403

Dissertationsgebiet lt. Studienblatt | Field of
study as it appears on the student record
sheet:

Business Analytics, Logistics and Operations
Research

Betreut von | Supervisor:

Univ.-Prof. i.R. Mag. Dr. Immanuel Bomze

Acknowledgements

My PhD journey started in October 2020 at the University of Vienna where I was given an opportunity to continue my passion for doing research. These past years have been a time of learning, personal growth, and professional development, and I am grateful to the many people who have supported and guided me along the way.

First and foremost, I would like to express my deepest gratitude to my supervisor, Professor Immanuel Bomze, whose expertise, patience, and guidance have been invaluable throughout my research. When I encountered difficulties in my journey, he always encouraged me and provided me with professional and thoughtful suggestions, which I found out later, were the best solutions. He never lost his patience in teaching and guiding me to be a qualified PhD candidate from every perspective, and I could not have asked for a more dedicated and encouraging supervisor.

I would also like to extend my thanks to my colleagues and collaborators. Working and studying alongside them has been a source of inspiration and motivation. In particular, I would like to deliver my thanks to the speaker, the invited lecturers, and colleagues from the Vienna Doctoral School on Computational Optimization (VGSCO), funded by FWF (Austrian Science Fund), Project W1260-N35: Professor Radu Ioan Bot, Professor Francesco Rinaldi, Professor Yurii Malitskyi, Professor Gabriele Eichfelder, Dr. Markus Gabl, Mikhail Karapetyants, Bettina Hiebl. Moreover, I would also like to thank Professor E. Alper Yildirim and Professor Laura Palagi, who were the hosts during my research stays at the University of Edinburgh and the University of Rome for their attentive assistance with my personal life and thoughtful discussions on the joint projects. Finally, I would like to extend my special thanks to Dr. Yuzhou Qiu, a collaborator and one of my closest friends now. The experience of working with him during my stay in Edinburgh made me more determined to continue on the academic path, as I admire his passion and dedication to academia.

To my dear friends, I extend my gratitude for your encouragement and belief in me, even in moments when I doubted myself, which helped me maintain a balanced perspective. I would like to offer special thanks to Yang Zhang, Yiwen Yang, Guojun Lai, Shiyang Zhou, and Zhijin Xue, whose presence brought color and joy to the monotonous PhD journey. Finally, I am especially grateful to Claudya (Yun) Hu, who accompanied me during the thesis writing process. The time spent together not only lifted my spirits but also provided me with inspiration and courage.

Last but not least, I would like to express my deepest appreciation to my mother Shu-ying Shu, and my father Cheng-long Peng. Without your unconditional love and support, I would not have been able to complete this PhD journey. As an only child, I am deeply aware of the emotional hardship my parents have suffered, particularly during holidays when families traditionally come together. I hope to make up one day for the moments we have missed.

Thank you all for being part of this journey with me.

Abstract

The thesis addresses mixed-integer (nonconvex) quadratic optimization problems (QPs) with complementarity constraints (CC), a crucial subclass of quadratically constrained quadratic optimization problems (QCQPs). These problems serve as powerful modeling tools in various fields, including interpretable machine learning, mathematical finance, and selection in biosciences. Given the challenge of solving these problems globally, rather than solving the problems directly, the focus is placed on convexification techniques, such as linear conic relaxations. As demonstrated in the thesis, a tight yet tractable convex relaxation not only provides a robust lower bound but also yields a high-quality feasible point through heuristics based on the optimal solution of the relaxation. These relaxations can be further integrated into the branch-and-bound framework, offering a comprehensive approach for globally solving the original problems.

This cumulative thesis is primarily organized as a collection of papers already published or submitted to internationally renowned journals [Pen22, BP23, BDPP24, BPQY24, BPQY25] and a nearly complete work, and the main results can be divided into two parts. The first part consists of two chapters: In Chapter 2, I establish several completely positive (CP) reformulations for the QPCCs under new conditions that are significantly less restrictive than those in the current literature. Essentially, CP reformulations package the nonconvexity into the completely positive cone, a matrix cone that has received increasing attention in the optimization domain recently. Additionally, novel and compact CP reformulations are proposed for the (nonconvex) QPs involving a cardinality term in Chapter 3, which is based on a work nearing completion to be ready for submission.

The second part (Chapters 4 and 5) focuses on tractable convex relaxations based on the CP reformulations, which are of particular interest in real-world applications. For sparse standard quadratic optimization problems (StQPs), an analytical comparison is made among these tractable relaxations, resulting in more compact hybrid relaxations. Furthermore, novel and computationally scalable relaxations are proposed for feature selection in support vector machines (SVMs), a widely used tool in supervised machine learning. In this context, the cardinality constraint enhances the interpretability of the resulting AI models. Extensive numerical experiments demonstrate the effectiveness and efficiency of these innovative relaxations.

Kurzfassung

Die Dissertation befasst sich mit gemischt-ganzzahligen (nichtkonvexen) quadratischen Optimierungsproblemen (QPs) mit Komplementaritätsbedingungen (CC), einer wichtigen Unterklasse von quadratischen Optimierungsproblemen unter (linearen und) quadratischen Nebenbedingungen (QCQPs). Diese Probleme dienen als leistungsstarke Modellierungswerkzeuge in verschiedenen Bereichen, darunter interpretierbares maschinelles Lernen, mathematische Finanzen und die Selektion in den Biowissenschaften. Angesichts der Herausforderung, diese Probleme global zu lösen, wird der Fokus nicht auf die direkte Lösung der Probleme gelegt, sondern auf Konvexifizierungstechniken wie lineare konische Relaxationen. Wie in der Dissertation gezeigt wird, liefert eine starke, aber dennoch handhabbare konvexe Relaxation nicht nur eine robuste untere Schranke, sondern ermöglicht auch Heuristiken, die auf der optimalen Lösung der Relaxation basieren, und die einen qualitativ hochwertigen zulässigen Punkt liefern. Diese Relaxationen können zudem in den Branch-and-Bound-Ansatz integriert werden, was einen umfassenden Ansatz zur globalen Lösung der ursprünglichen Probleme bietet.

Diese kumulative Dissertation ist in erster Linie als eine Sammlung von (bereits publizierten oder bei international angesehenen Zeitschriften eingereichten) Aufsätzen [Pen22, BP23, BDPP24, BPQY24, BPQY25] und einem fast fertiggestellten Artikel organisiert. Die Hauptresultate lassen sich in zwei Teile gliedern. Der erste Teil besteht aus zwei Kapiteln: In Kapitel 2 stelle ich mehrere completely positive (CP) Reformulierungen für die QPCCs unter neuen Bedingungen vor, die deutlich weniger restriktiv sind als die in der aktuellen Literatur. Im Wesentlichen codieren die CP-Reformulierungen die Nichtkonvexität mit Hilfe des completely positive cone, eines Matrixkegels, der in der Optimierungsforschung in letzter Zeit zunehmend Beachtung findet. Darüber hinaus werden im Kapitel 3 neuartige und kompakte CP-Reformulierungen für die (nichtkonvexen) QPs mit Kardinalitätsbedingungen vorgeschlagen, basierend auf einer Arbeit, die beinahe zur Einreichung bereit ist.

Der zweite Teil (Kapitel 4 und Kapitel 5) konzentriert sich auf handhabbare konvexe Relaxationen, die auf den CP-Reformulierungen basieren, was besonders relevant in realen Anwendungen ist. Für Standard-Quadratische-Optimierungsprobleme (StQPs) unter Kardinalitätsbedingungen wird ein analytischer Vergleich dieser handhabbaren Relaxationen vorgenommen, was zu kompakteren hybriden Relaxationen führt. Des Weiteren werden neuartige und rechnerisch skalierbare Relaxationen für die Merkmalsauswahl in Support-Vektor-Maschinen (SVMs) vorgeschlagen, ein weit verbreitetes Werkzeug im überwachten maschinellen Lernen. In diesem Zusammenhang verbessert die Kardinalitätsbedingung die Interpretierbarkeit der resultierenden KI-Modelle. Umfangreiche numerische Experimente belegen die Effektivität und Effizienz dieser innovativen Relaxationen.

Contents

Acknowledgements	i
Abstract	iii
Kurzfassung	v
List of Tables	ix
List of Figures	xi
1. Introduction	1
1.1. Sparsity in machine learning	1
1.2. Complementarity constraint	3
1.3. Copositive optimization	6
1.4. Copositivity tests: comparison of two recent proposals	7
1.4.1. Copositivity test based on almost copositivity	8
1.4.2. Copositivity test based on KKT conditions	8
1.4.3. Numerical comparison on two copositivity tests	9
1.4.4. Instance generation	10
1.4.5. Computational results	11
1.5. An improved variant of one copositivity test	12
2. CP reformulation for QPCCs	15
2.1. The problem	15
2.2. CP reformulation for QPCCs	16
2.2.1. Conic formulation: Burer's key condition	16
2.2.2. Relaxing the key condition: vertex covers	17
2.2.3. A simplification of (C) and its dual	22
2.3. QPCC and CP formulations for quadratic problems under sparsity constraints . .	23
2.3.1. Soft sparsity: regularization by zero-norm	24
2.3.2. Hard sparsity constraints	27
2.3.3. Sparse linear regression: conic reformulation	28
2.4. An iterative algorithm for dual problems (CD)	30
2.5. Conclusion	31
3. CP reformulation for sparse QPs	33
3.1. The problem	33
3.2. CP reformulation for soft sparse QPs	34
3.2.1. Burer's representation theorem and its extension	34
3.2.2. New and more compact CP reformulation for soft sparse QPs	36
3.2.3. The decomposed CP reformulation	40

3.3.	CP reformulation for hard sparse QPs	42
3.3.1.	New and more compact CP reformulation for hard sparse QPs	42
3.4.	Conclusion	47
4.	Conic relaxations for sparse StQPs	49
4.1.	The problem	49
4.2.	Two exact completely positive formulations	50
4.2.1.	Conic formulation of direct MIQP	50
4.2.2.	Conic formulation of MIQP with complementarity constraint	51
4.3.	Two convex relaxations and their smaller equivalents	52
4.3.1.	DNN relaxations of the direct MIQP formulation	53
4.3.2.	DNN relaxations of the complementarity constraint formulation	57
4.3.3.	Comparing the two DNN relaxations	60
4.4.	Generating nontrivial instances of sparse StQPs	61
4.4.1.	Generating StQP instances with a unique optimal solution	61
4.4.2.	Two algorithms for generating nontrivial sparse StQP instances	63
4.4.3.	Set of instances	65
4.5.	Computational results	66
4.5.1.	Experimental setup	66
4.5.2.	Solution time	67
4.5.3.	Quality of lower bounds	75
4.6.	Conclusion	76
5.	A scalable conic approach for feature selection in linear SVMs	79
5.1.	The problem	79
5.2.	Exact Mixed-Integer approaches	80
5.3.	Tractable Relaxations	82
5.3.1.	Box relaxation	82
5.3.2.	The Shor relaxation	83
5.3.3.	Decomposed SDP-based relaxations	85
5.4.	The algorithms	86
5.4.1.	Two upper bounding strategies	87
5.4.2.	A strategy to tighten the big-M parameter	89
5.4.3.	The heuristic algorithm	90
5.4.4.	The exact algorithm	90
5.5.	Numerical experiments	93
5.5.1.	Computational analysis on the MIQP models	94
5.5.2.	Results on the heuristics	96
5.5.3.	Results on the exact approach	97
5.6.	Conclusion	98
	Bibliography	101
A.	Appendix	109
A.1.	Notation	109

List of Tables

1.1. Comparison between (COP1) and (COP1 ₊) on the test sets with small c	13
4.1. Choices of the parameters n, ρ_0 , and ρ	67
4.2. Number of instances (out of 25) terminated due to the time limit (excluding (D1A(ρ)) that was terminated due to the time limit on <i>all</i> instances with $n = 50$)	75
5.1. First datasets	94
5.2. Second datasets	94
5.3. Comparison of BigMP and CoP on the large datasets	94
5.4. Results on large datasets of the Heuristic algorithm implementing either the Local search (H-LS) or the Kernel Search (H-KS) as the upper bound strategies	96
5.5. Comparison of the objective functions of CoP returned by Gurobi with and without using the Heuristic Kernel Search algorithm as a warm start for large datasets	97
5.6. Lower bounds comparison between the SDP one and the starting bound found by Gurobi for the BigMP and CoP models for different datasets	99
5.7. Comparison of the Gurobi solution to the CoP model using the Heuristic Kernel Search solution as a warm start heuristic, and the Exact algorithm using the Kernel Search procedure as the upper bound strategy.	100

List of Figures

1.1. Results on the copositive instances in $\mathfrak{A}_{0.02i,50}^+$	12
1.2. Results on the non-copositive instances in $\mathfrak{A}_{0.02i,50}^-$	13
1.3. Results of the second experiment	14
4.1. Average solution times (in seconds) of (P1(ρ)), (P2(ρ)), (D1A(ρ)), (D1B(ρ)), (D2A(ρ)), and (D2B(ρ)) for PSD, SPN, and COP instances	68
4.2. Empirical cumulative distribution functions of solution times of (P1(ρ)), (P2(ρ)), (D1A(ρ)), (D1B(ρ)), (D2A(ρ)), and (D2B(ρ))	70
4.3. Empirical cumulative distribution functions of solution times of (P1(ρ)) and (P2(ρ)) for all instances with $n = 25$	72
4.4. Empirical cumulative distribution functions of solution times of (P1(ρ)) and (P2(ρ)) for all instances with $n = 50$	73
4.5. Absolute gaps of the best of (P1(ρ)) and (P2(ρ)) versus the best of (D1B(ρ)) and (D2B(ρ))	77
5.1. Comparison of solving BigMP and CoP with Gurobi for the first set of datasets: mean computational times (in seconds)	95
5.2. Comparison of the mean MipGap values of the solution found by Gurobi solving CoP after 1 hour, and the Relaxed Gap of the Heuristic Kernel Search solution using the solution of DSCoP as a lower bound	98

1. Introduction

In this section, we provide a brief introduction to the background and motivation of this thesis. Additionally, we present an overview of the primary technique employed: copositive optimization.

1.1. Sparsity in machine learning

Sparsity is becoming an increasingly important concept in machine learning, especially in the era of high-dimensional and high-frequency data. According to Albert Einstein, "The supreme goal of all theory is to make the irreducible basic elements as simple and as few as possible", this principle of simplicity is highly relevant to machine learning as well. A term that represents this idea is data sparsity, which describes a situation where a substantial proportion of data points are either missing or set to zero. Mathematically, we use the notation $\|\cdot\|_0$ to capture the sparsity of a given data (vector). For a vector $\mathbf{x} \in \mathbb{R}^n$, by

$$\|\mathbf{x}\|_0 := |\{i \in [1:n] : x_i \neq 0\}|,$$

where by $[a:b]$, for any integers a, b , we denote the set of all integers k satisfying $a \leq k \leq b$.

we count the number of non-zero elements of $\mathbf{x}$. Consider a minimization optimization problem over the variable $\mathbf{x}$. By adding the term $\|\mathbf{x}\|_0$ to the objective function with a positive coefficient λ , it is evident that the solution tends to exhibit sparsity. In an extreme case where λ approaches infinity, the optimization problem effectively seeks the sparsest feasible point of the original problem. Alternatively, we can directly limit the number of nonzero elements in the variable $\mathbf{x}$ using a cardinality constraint

$$\|\mathbf{x}\|_0 \leq k,$$

where k is a non-negative integer. In this case, $\mathbf{x}$ is said to be k -sparse if the constraint is satisfied. Incorporating the sparse term $\|\cdot\|_0$ into many machine learning tasks promotes sparsity in the solutions, leading to significant advantages, as demonstrated in the following subsections.

Sparse principal component analysis

Sparse principal component analysis (sparse PCA) is a widely used data processing and reduction technique that seeks to find a low-dimensional representation of the data while enforcing sparsity on the principal components. For a given data $\mathbf{X} \in \mathbb{R}^{n \times p}$ of n observations on p variables, PCA is a process of reconstructing the empirical covariance matrix of data by linear combinations of the original p variables. This leads to the following optimization problem [Jol02]:

$$\begin{aligned} \max_{\mathbf{v} \in \mathbb{R}^p} \quad & \mathbf{v}^\top \Sigma \mathbf{v} \\ \text{s.t.} \quad & \mathbf{v}^\top \mathbf{v} = 1, \end{aligned}$$

where Σ is the empirical covariance matrix of data $\mathbf{X}$ and we call the optimal solution $\mathbf{v}^*$ principal component. The principal component is typically dense, e.g., a linear combination of all original variables. As a result, interpreting the principal component can be difficult, especially when

you have a large number of variables. In other words, the individual component can be too abstract to be related back to the original data. To overcome the weakness of traditional PCA, the following sparse PCA is proposed

$$\begin{aligned} \max_{\mathbf{v} \in \mathbb{R}^p} \quad & \mathbf{v}^\top \Sigma \mathbf{v} \\ \text{s.t.} \quad & \mathbf{v}^\top \mathbf{v} = 1 \\ & \|\mathbf{v}\|_0 \leq k, \end{aligned}$$

where we restrict the number of variables involved in each principal component. Since in each principal component, only a small subset of variables are involved, this gives more direct insight into which specific variables (features) are responsible for the most significant patterns. Therefore it makes the results more interpretable and meaningful, especially in domains like biology, finance, or social sciences, where understanding the influence of individual variables is crucial.

Compressed sensing

Compressed sensing is one of the significant breakthroughs in signal processing, which allows the recovery of signals via much fewer samples than required by the Nyquist–Shannon sampling theorem [Sha48] under the condition:

- The signal is sparse in some domains (e.g., Fourier, wavelet);
- The measurement matrix $\mathbf{A}$ must satisfy the restricted isometry property (RIP) [CRT06, CT06]. RIP ensures that the lengths of sparse signals are approximately preserved when projected onto the lower-dimensional measurement space:

$$(1 - \delta_k) \|\mathbf{x}\|_2^2 \leq \|\mathbf{A}\mathbf{x}\|_2^2 \leq (1 + \delta_k) \|\mathbf{x}\|_2^2$$

Mathematically, compressed sensing turns out to find a sparse solution from an underdetermined linear system:

$$\mathbf{A}\mathbf{x} = \mathbf{b},$$

where $\mathbf{A} \in \mathbb{R}^{m \times n}$ ($m \ll n$) is the measurement matrix and $\mathbf{b} \in \mathbb{R}^m$ is m observations. The related optimization problem reads as follows:

$$\begin{aligned} \max_{\mathbf{x} \in \mathbb{R}^n} \quad & \|\mathbf{x}\|_0 \\ \text{s.t.} \quad & \mathbf{A}\mathbf{x} = \mathbf{b}. \end{aligned}$$

Sparse dictionary learning

Sparse Dictionary Learning is a technique in machine learning and signal processing that seeks to represent data as a linear combination of a few elements from a larger set, known as the dictionary. The dictionary $\mathbf{D} \in \mathbb{R}^{n \times m}$ is a set of vectors that forms the building blocks for reconstructing data $\mathbf{x}_i \in \mathbb{R}^n$ for $i \in I$. Our goal is to find a set of sparse vector $\alpha_i \in \mathbb{R}^m$ and the dictionary $\mathbf{D}$, such that

$$\mathbf{x}_i \approx \mathbf{D}\alpha_i.$$

The corresponding optimization problem is usually of the forms:

$$\begin{aligned} \min_{\mathbf{D}, \alpha_i} \quad & \sum_{i \in I} \|\mathbf{x}_i - \mathbf{D}\alpha_i\|_2^2 \\ \text{s.t.} \quad & \|\alpha_i\|_0 \leq k, \end{aligned}$$

or

$$\min_{\mathbf{D}, \alpha_i} \quad \sum_{i \in I} \|\mathbf{x} - \mathbf{D}\alpha_i\|_2^2 + \lambda \|\alpha_i\|_0,$$

where the first problem directly limits the variables α_i to be k -sparse, and the second problem price the sparsity of the variables α_i in the objective by multiplying a user-defined parameter λ .

Sparse support vector machine (controlled feature selection)

Support vector machine (SVM) is a supervised learning technique mainly used for classification tasks. The key idea behind SVM is to find the optimal hyperplane that separates data points of different classes with the maximum margin. Sparse SVM extends traditional SVM by introducing the sparsity constraint in the feature space, which is particularly useful for feature selection, improving interpretability, reducing overfitting, and enhancing computational efficiency, especially in high-dimensional datasets. Here we take linear SVM with the soft margin as an example:

$$\begin{aligned} \min_{\mathbf{w}, b, \xi} \quad & \frac{1}{2} \|\mathbf{w}\|_2^2 + C \sum_{i=1}^m \xi_i \\ \text{s.t.} \quad & y^i (\mathbf{w}^\top \mathbf{x}^i + b) \geq 1 - \xi_i, \quad i \in [1:m] \\ & \|\mathbf{w}\|_0 \leq B \\ & \xi \geq \mathbf{0}, \end{aligned}$$

where the user-defined parameter $B \in \mathbb{Z}_+$ constitutes a budget on the number of features used, indicating that the number of features involved in the classifier is at most B .

After showing a few examples where we employ $\|\cdot\|_0$ for data sparsity, We close the subsection by listing the benefits of enforcing sparsity in machine learning:

- Sparsity allows models to be more interpretable because fewer variables are involved in the final model. Irrelevant or redundant features are assigned a coefficient of zero, making it clear which features are driving the model's predictions.
- In high-dimensional datasets, the number of features can be much larger than the number of observations. In these cases, enforcing sparsity in machine learning models automatically performs feature selection by excluding irrelevant variables, which reduces the risk of overfitting and simplifies the model.
- Sparse models require fewer computations because many features are excluded from the model. This reduces memory usage and computational time, especially important in large-scale applications such as natural language processing or computer vision.

1.2. Complementarity constraint

Complementarity constraints are widely used to model various combinatorial structures, such as binary constraints, cardinality constraints, logical relations, and disjunctive conditions. As a

result, they are attracting increased attention in interpretable machine learning where we focus on explainable and understandable machine learning models by using logical building blocks rather than relying on black-box approaches.

A complementarity constraint typically has the following form:

$$x \geq 0, \quad y \geq 0, \quad xy = 0.$$

In the following, we will demonstrate the strong modeling ability of the complementarity constraint with some examples.

Binary constraint

The binary constraint simply means a variable only takes two values. e.g., $x \in \{0, 1\}$. This can be modeled with complementarity constraint by introducing an auxiliary variable:

$$x + y = 1, \quad xy = 0.$$

It is important to note that mixed-integer optimization problems represent a special case of optimization problems with complementarity constraints. As a result, the latter problem class contains at least the former, which includes a wide range of problems in combinatorial optimization.

Cardinality constraint

The cardinality constraint refers to limiting the number of non-zero elements of a variable $\mathbf{x} \in \mathbb{R}_+^n$. For instance:

$$\|\mathbf{x}\|_0 \leq k.$$

The importance of this type of constraint has been fully addressed in the last subsection. In fact, after introducing auxiliary variables, it can be modeled with complementarity constraints as:

$$\mathbf{0} \leq \mathbf{u} \leq \mathbf{e}, \quad \mathbf{e}^\top \mathbf{u} \geq n - k, \quad \mathbf{x}^\top \mathbf{u} = 0$$

where $\mathbf{e}$ is the all-ones vector. Indeed, the first two constraints imply that at least $n - k$ entries of $\mathbf{u}$ are positive. Due to the complementarity constraint between $\mathbf{x}$ and $\mathbf{u}$, it follows immediately that $\mathbf{x}$ has at least $n - k$ zero entries.

Logical If-Then constraint

Logical If-Then constraints are used to model relationships where the satisfaction of one condition (the "if" part) implies the satisfaction of another condition (the "then" part). Given two constraints of the form:

$$f(x) > 0, \quad g(x) \geq 0,$$

consider the logical condition:

$$\text{If } f(x) > 0, \quad \text{then } g(x) \geq 0.$$

By introducing auxiliary variables, we arrive at the following formulation with complementarity constraint:

$$\begin{aligned} f(x) + s_1 - s_2 &= 0, & g(x) + s_3 - s_4 &= 0, \\ s_1 s_2 + s_3 s_4 + s_2 s_3 &= 0, & s_1, s_2, s_3, s_4 &\geq 0. \end{aligned}$$

Note that s_1 and s_3 are negative parts of $f(x)$ and $g(x)$ respectively. Similarly, s_2 and s_4 are positive parts of $f(x)$ and $g(x)$ respectively. Indeed, $f(x) > 0$ implies $s_2 > 0$. Together with the constraint $s_2 s_3 = 0$, we have $s_3 = 0$ and therefore $g(x) = s_4 \geq 0$. As we may know, the If-Then constraint can be modeled by the so-called big-M formulation, where we introduce a sufficiently large constant M and binary variables. However, the existence and the choice of such M are unclear in many cases, which may introduce numerical instability, especially when it is fed to a branch and bound method. Large values of M can cause the optimization problem to become ill-conditioned, leading to slower convergence or inaccurate solutions. Instead, the formulation with complementarity constraint does not involve any big M and binary variables, and the difficulty is only packaged in the complementarity constraint.

Logical Either-Or constraint

Logical Either-Or constraint refers to the situation where exactly one of two mutually exclusive conditions must hold, simplifying the optimization problem to choosing between two constraints. For example, consider a production planning problem in which a machine can either produce Product A or Product B, but not both at the same time. Let

- $f(x)$ represents the production level of Product A;
- $g(x)$ represents the production level of Product B.

The Either-Or constraint is:

$$\text{Either } f(x) \geq 0 \quad \text{or} \quad g(x) \geq 0.$$

This can be formulated as:

$$\begin{aligned} f(x) + s_1 &\geq 0, & g(x) + s_2 &\geq 0, \\ s_1 s_2 &= 0, & s_1 + s_2 &\geq \epsilon, \\ s_1, s_2 &\geq 0, \end{aligned}$$

where $\epsilon > 0$ is a small enough positive parameter. Due to the constraints on s_1 and s_2 , we have either $s_1 = 0$ and $s_2 > 0$, or $s_1 > 0$ and $s_2 = 0$. Together with the constraints $f(x) + s_1 \geq 0$ and $g(x) + s_2 \geq 0$, it follows that either $f(x) \geq 0$ or $g(x) \geq 0$, but not both simultaneously.

Disjunctive constraint

The disjunctive constraint can be viewed as an extension of the logical Either-Or constraint and refers to the feasible region described as a union of a finite collection of disjoint sets S_i , $i \in [1:k]$. A disjunctive constraint has the form:

$$x \in \cup_{i=1}^k S_i,$$

A piecewise defined function consists of different functional expressions on different parts, therefore it can be naturally represented using disjunctive constraints. We take the well-known rectified linear unit (ReLU) function as an example:

$$\text{ReLU}(x) = \max(0, x).$$

The function returns 0 if $x < 0$ and returns x if $x \geq 0$. Again, employing a complementarity constraint and auxiliary variables $s_i \geq 0$, the ReLU can be equivalently expressed as:

$$\begin{aligned} x + s_1 - s_2 &= 0, & s_1 s_2 + \text{ReLU}(x) s_1 &= 0, \\ \text{ReLU}(x) &\leq s_2, & \text{ReLU}(x) &\geq s_2 - s_1, \\ s_1, s_2, \text{ReLU}(x) &\geq 0. \end{aligned}$$

If $x < 0$, due to the nonnegativity of s_1 and s_2 , and the complementarity between them, we have $s_1 > 0$ and $s_2 = 0$. Together with the constraint $\text{ReLU}(x)s_1 = 0$, it follows immediately that $\text{ReLU}(x) = 0$; if $x \geq 0$, due to the nonnegativity of s_1 and s_2 , and the complementarity between them, we have $s_1 = 0$ and $s_2 \geq 0$. Since $s_2 - s_1 \leq \text{ReLU}(x) \leq s_2$, we have $\text{ReLU}(x) = x$.

As shown above, the complementarity constraint is a versatile and efficient tool for a wide range of optimization and decision-making problems. It has a strong modeling ability arising from its flexibility in handling logical relationships, piecewise functions, non-convexities, binarity, and more compact formulations for many mixed-integer optimization problems.

However, an optimization problem involving complementarity constraints is generally difficult to solve due to its non-smoothness and non-convexity. Even the simplest one, a linear optimization problem with complementarity constraints, is NP-hard [CPS92]. Nevertheless, these constraints have two distinguished properties: they are homogeneous, and the variables involved are nonnegative. Based on these two key features, we will introduce a tool that has the potential to handle this type of problem in the next section.

1.3. Copositive optimization

Copositive optimization is a branch of mathematical optimization that focuses on problems involving copositive matrices, where a matrix is said to be copositive if its quadratic form takes non-negative values for all non-negative vectors. The field has gained attention due to its ability to reformulate many hard combinatorial optimization problems, and the interested readers are referred to [Bom12, BSU12, Bur15, DR21] and references therein.

Copositive optimization is closely related to semidefinite programming (SDP), but it goes beyond SDP by encompassing problems that are generally more difficult to solve due to the structure of the copositive cone. Despite this, it has the advantage of transforming certain non-convex problems into a convex framework, offering theoretical and computational benefits.

The concept of copositive matrices was introduced in the 1950s in the context of quadratic forms and matrix theory [Mot52]. A matrix $\mathbf{Q} \in \mathcal{S}^n$ is said to be copositive if:

$$\mathbf{x}^\top \mathbf{Q} \mathbf{x} \geq 0, \quad \forall \mathbf{x} \in \mathbb{R}_+^n.$$

The set of copositive matrices is known as the copositive cone, denoted by $\mathcal{COP}^n$, which is a convex, closed, pointed, and full-dimensional cone. For a given matrix $\mathbf{Q} \in \mathcal{S}^n$, the membership problem to decide whether or not $\mathbf{Q} \in \mathcal{COP}^n$ holds is co-NP-complete [MK87]. It is equivalent to determining whether or not $v(\mathbf{Q}) \geq 0$ holds, for

$$v(\mathbf{Q}) := \min_{\mathbf{x} \in \Delta_n} \mathbf{x}^\top \mathbf{Q} \mathbf{x},$$

where $\Delta_n := \{\mathbf{x} \in \mathbb{R}_+^n : \|\mathbf{x}\|_1 = 1\}$ denotes the standard simplex. The above standard quadratic problem is NP-hard, since the maximum clique problem can be reduced to it [MS65].

For an arbitrary cone $\mathcal{K} \in \mathcal{S}$, the dual cone $\mathcal{K}^*$ is defined as:

$$\mathcal{K}^* := \{\mathbf{A} \in \mathcal{S} : \langle \mathbf{A}, \mathbf{B} \rangle \geq 0, \quad \text{for all } \mathbf{B} \in \mathcal{K}\}.$$

Therefore, the dual cone of the copositive cone equals

$$\mathcal{COP}^n = \left\{ \sum_{i \in I} a_i a_i^\top : a_i \in \mathbb{R}_+^n \right\},$$

which is known as the completely positive (CP) cone.

The idea of using copositive matrices in optimization started gaining attraction in the 1990s, particularly in the field of quadratic programming and combinatorial optimization. Researchers discovered that certain non-convex quadratic optimization problems could be reformulated as convex optimization problems over the copositive cone [QdKRT98].

A copositive optimization problem refers to a linear optimization problem over the completely positive cone and has the following form:

$$\begin{aligned} \min_{\mathbf{X} \in \mathcal{S}^n} \quad & \langle \mathbf{C}, \mathbf{X} \rangle \\ \text{s.t.} \quad & \mathcal{A}\mathbf{X} = \mathbf{b} \\ & \mathbf{X} \in \mathcal{CP}^n, \end{aligned} \tag{1.1}$$

where $\mathcal{A} : \mathcal{S}^n \rightarrow \mathbb{R}^m$ is a linear operator.

The Lagrangian dual of (1.1) is a linear conic optimization problem over the copositive cone:

$$\begin{aligned} \max_{\mathbf{y} \in \mathbb{R}_+^m} \quad & \mathbf{b}^\top \mathbf{y} \\ \text{s.t.} \quad & \mathcal{A}^* \mathbf{y} + \mathbf{Z} = \mathbf{C} \\ & \mathbf{Z} \in \mathcal{COP}^n, \end{aligned} \tag{1.2}$$

where $\mathcal{A}^* : \mathbb{R}^m \rightarrow \mathcal{S}^n$ is the adjoint operator of $\mathcal{A}$. For historical reasons and simplicity, both problems (sometimes the primal-dual pair of conic problems) are frequently addressed as copositive optimization problems.

In the early 2000s, researchers [BDDK⁺00] recognized the significant potential of copositive optimization when it was discovered that many NP-hard combinatorial problems could be reformulated as linear optimization problems over the copositive cone. A seminal contribution in this field was made by Burer [Bur09], who demonstrated that the quadratic optimization problem with binary variables could be exactly reformulated as a copositive optimization problem. Another important work [PVZ15] extends the scope of applying copositive optimization to polynomial optimization problems, presenting a general characterization of the class of problems that can be formulated as a conic program over the cone of completely positive tensors. Essentially, the reformulation effectively packages the complexity of the problem within the CP cone, enabling tighter bounds and novel approximation techniques based on the CP cone.

Despite the compact form of CP reformulation, their computational intractability arises from the presence of the CP cone. In practice, outer approximations of the CP reformulation are typically considered, yielding lower bounds that are significantly tighter than those obtained through SDP relaxations. It is important to note that any feasible solution to (1.2) gives, by weak duality, a rigorous lower bound to the optimal value of (1.1) which in turn is, by relaxation, a lower bound to many NP-hard problems (sometimes the same, for CP reformulations). Moreover, (1.2)-feasibility requires copositivity testing for $\mathbf{Z}$, which underlines the importance of the next section.

1.4. Copositivity tests: comparison of two recent proposals

As demonstrated in the previous section, the primary challenge of a copositive optimization problem lies in the cone constraint. In this section, we focus on methods for testing the copositivity

of a given matrix and provide a numerical comparison of these methods. Additionally, a modified variant is introduced, which offers improved efficiency in certain scenarios. This section is mainly based on a peer-reviewed and published short communication [Pen22].

1.4.1. Copositivity test based on almost copositivity

A matrix $Q \in \mathcal{S}^n$ is said almost copositive if either $n = 1$ or every $n - 1 \times n - 1$ principal submatrix of Q is copositive. Based upon the theory of almost copositivity (see, e.g. [Dic19]), the approach of [Ans21] formulates an MILP as follows:

$$\begin{aligned} l_1(Q) := \max \quad & \lambda \\ \text{s.t.} \quad & \mathbf{q}_i^\top \mathbf{x} \leq -\lambda + m_i(1 - y_i), \quad i \in [1:n], \\ & \lambda \geq 0, \mathbf{o} \leq \mathbf{x} \leq \mathbf{y}, \\ & \mathbf{y} \in \{0, 1\}^n, \mathbf{e}^\top \mathbf{y} \geq \theta, \end{aligned} \tag{COP1}$$

where $m_i = 1 + \sum_{j \neq i} \max\{q_{ij}, 0\}$ for all $i \in [1:n]$ and $\theta > 0$ is a parameter.

Theorem 1. [Ans21] *Let $\theta = 1$. Then Q is copositive if and only if the optimal value of (COP1) is zero.*

Essentially, the optimization problem (COP1) identifies the existence of a copositive matrix among the principal submatrices of Q . According to [Dic19], if an almost copositive submatrix of Q exists, the objective value of (COP1) must be strictly greater than zero; otherwise, Q is copositive.

1.4.2. Copositivity test based on KKT conditions

Recall that a matrix $Q \in \mathcal{S}^n$ is copositive if and only if

$$\mathbf{x}^\top Q \mathbf{x} \geq 0, \quad \forall \mathbf{x} \in \mathbb{R}_+^n,$$

which is equivalent to $v(Q) \geq 0$ and it is defined as the optimal value of the optimization problem:

$$\begin{aligned} v(Q) := \min_{\mathbf{x} \in \mathbb{R}_+^n} \quad & \mathbf{x}^\top Q \mathbf{x} \\ \text{s.t.} \quad & \mathbf{e}^\top \mathbf{x} = 1. \end{aligned} \tag{1.3}$$

Due to the constraints of (1.3) satisfies linear independence constraint qualification at every feasible point, for any (1.3)-optimal solution, there exists $\lambda \in \mathbb{R}$ and $\mathbf{s} \in \mathbb{R}^n$ such that:

$$\begin{aligned} \mathbf{e}^\top \mathbf{x} &= 1 \\ Q\mathbf{x} - (\lambda \mathbf{e} + \mathbf{s}) &= \mathbf{o} \\ \mathbf{x}^\top \mathbf{s} &= 0 \\ \mathbf{x} &\geq \mathbf{o}, \quad \mathbf{s} \geq \mathbf{o}. \end{aligned}$$

Considering the first two constraints, we have $\lambda = \mathbf{x}^\top \mathbf{Q} \mathbf{x}$. Therefore, the problem (1.3) is equivalent to the following problem:

$$\begin{aligned} \min_{\mathbf{x}, \mathbf{s}, \lambda} \quad & \lambda \\ \text{s.t.} \quad & \mathbf{e}^\top \mathbf{x} = 1 \\ & \mathbf{Q} \mathbf{x} - (\lambda \mathbf{e} + \mathbf{s}) = \mathbf{o} \\ & \mathbf{x}^\top \mathbf{s} = 0 \\ & \mathbf{x} \geq \mathbf{o}, \quad \mathbf{s} \geq \mathbf{o}. \end{aligned}$$

By a standard linearization procedure for the complementarity constraint, [GY21] proposed the following equivalent MILP:

$$\begin{aligned} l_2(\mathbf{Q}) := \min \quad & \lambda \\ \text{s.t.} \quad & \mathbf{Q} \mathbf{x} - \lambda \mathbf{e} - \mathbf{s} = \mathbf{o} \\ & \mathbf{e}^\top \mathbf{x} = 1 \\ & x_i \leq y_i, \quad i \in [1:n] \\ & s_i \leq M_i(1 - y_i), \quad i \in [1:n] \\ & \mathbf{s}, \mathbf{x} \geq \mathbf{o} \\ & \mathbf{y} \in \{0, 1\}^n. \end{aligned} \tag{COP2}$$

We choose the parameter M_i according to the following considerations: by the first two constraints in (COP2), it follows that

$$s_i = \mathbf{q}_i^\top \mathbf{x} - \lambda \leq \max_{j \in [1:n]} q_{ij} - \lambda, \tag{1.4}$$

where $\mathbf{q}_i$ and q_{ij} are the i -th column and ij -th entry of $\mathbf{Q}$, respectively. Thus, we only need to estimate the lower bound of λ . Likewise, from the complementary slackness conditions $\mathbf{s}^\top \mathbf{x} = 0$ and (1.4) it follows that $\lambda = \mathbf{x}^\top \mathbf{Q} \mathbf{x} \geq v(\mathbf{Q})$. Therefore, lower bounds on $v(\mathbf{Q})$ can be used to bound λ from below. We tried two bounds on $v(\mathbf{Q})$ from [BLT08]:

$$\lambda \geq v(\mathbf{Q}) \geq v_1(\mathbf{Q}) := \min_{1 \leq i \leq j \leq n} q_{ij} + \frac{1}{\sum_{k=1}^n (1/(q_{kk} - \min_{1 \leq i \leq j \leq n} q_{ij}))}$$

or a tighter conic bound:

$$\lambda \geq v(\mathbf{Q}) \geq v_2(\mathbf{Q}) := \min \{ \langle \mathbf{Q}, \mathbf{X} \rangle : \langle \mathbf{E}, \mathbf{X} \rangle = 1, \mathbf{X} \in \mathcal{S}_+^n \cap \mathcal{N}^n \}.$$

Despite much higher computational cost for $v_2(\mathbf{Q})$, we found in a preliminary limited study no significant improvement of $v_2(\mathbf{Q})$ over $v_1(\mathbf{Q})$, so we decided to use $v_1(\mathbf{Q})$ throughout our experiments. Summarizing, we chose $M_i = \max_{j \in [1:n]} q_{ij} - v_1(\mathbf{Q})$ for use in the comparison described in the following.

1.4.3. Numerical comparison on two copositivity tests

In this section, we conduct a numerical comparison of the two copositivity tests using constructed instances. The experiment aims to demonstrate the strengths and weaknesses of each test when applied to matrices with different characteristics.

1.4.4. Instance generation

Based on the preliminary numerical study on (COP1) and (COP2), it is noteworthy that both methods perform poorly on matrices perturbed by the positive semidefinite cone. Motivated by this observation, we introduce a matrix parameter to quantify the distance between a matrix and the positive semidefinite cone, which can also be used to predict the difficulty of the copositivity test. For a given matrix $Q \in \mathcal{S}^n$, by $\{\lambda_i : i \in [1:n]\}$ we denote the set of its eigenvalues. By I_+ and I_- we denote the index set of positive eigenvalues and nonpositive eigenvalues, respectively. By $\mathcal{S}_1^n$ we denote a subset of $\mathcal{S}^n$, which contains matrices with at least one positive eigenvalue. We then define a function $h : \mathcal{S}_1^n \rightarrow \mathbb{R}_+$ by

$$h(Q) := \frac{\sum_{i \in I_-} |\lambda_i|}{\sum_{i \in I_+} \lambda_i}.$$

By $\mathfrak{A}_{c,n}$ we denote the test set

$$\mathfrak{A}_{c,n} := \{Q \in \mathcal{S}_1^n : h(Q) = c\},$$

where c is a nonnegative parameter. Note that the extreme cases $c \searrow 0$ correspond to the easy case $Q \in \mathcal{S}_+^n$ and likewise $c \nearrow \infty$ to the easy case $-Q \in \mathcal{S}_+^n$, if $\sum_{i=1}^n |\lambda_i|$ remains bounded.

To generate $Q \in \mathfrak{A}_{c,n}$, we use the eigendecomposition $Q = P^\top \Lambda P$ where P is a randomly generated orthogonal matrix and Λ is a constructed diagonal matrix. Note that $\mathfrak{A}_{c,n}$ includes not only copositive matrices but also non-copositive ones. We separate the problem set $\mathfrak{A}_{c,n}$ according to copositivity: let $\mathfrak{A}_{c,n}^+ := \mathfrak{A}_{c,n} \cap \mathcal{COP}^n$ and $\mathfrak{A}_{c,n}^- := \mathfrak{A}_{c,n} \setminus \mathcal{COP}^n$. In fact, we cannot generate $\mathfrak{A}_{c,n}^+$ and $\mathfrak{A}_{c,n}^-$ directly. However, we are able to analyse the numerical results separately to avoid the bias towards hardness of copositive instances, as compared to non-copositive ones.

We next adopt the method [HP73] in which a procedure of constructing an extreme copositive matrix with entries from $\{-1, 0, 1\}$ was proposed. We call a copositive matrix $Q \in \mathcal{COP}^n$ exceptional if $Q \notin \mathcal{SPN}^n$. Exceptional copositive matrices are always of great interest from the aspect of copositivity test since there are polynomial-time approaches for testing membership of $\mathcal{SPN}^n$, unlike the membership of $\mathcal{COP}^n \setminus \mathcal{SPN}^n$. We know that the extreme copositive matrices are either exceptional or trivial (e.g. lying in $\mathcal{S}_+^n \cup \mathcal{N}^n$) and therefore filter the non-trivial matrices as instances.

Theorem 2. [HP73] *A given symmetric matrix B with entries from $\{-1, 0, 1\}$ is on an extreme ray of $\mathcal{COP}^n$ if and only if the following conditions hold*

- the diagonal entries of B are all equal to 1,
- $G_{-1}(B)$ is connected and contains no triangle,
- $G_1(B)$ contains precisely those edges (i, j) where i and j are at distance 2 in $G_{-1}(B)$,
- $L^*(B)$ has no isolated vertices and no bipartite graph for a connected component,

where $G_{-1}(B)$, $G_0(B)$, $G_1(B)$ and $L^*(B)$ are four undirected graphs associated with B . In particular, $G_{-1}(B)$ ($G_0(B)$, $G_1(B)$) is a graph on n vertices with edges (i, j) if $B_{ij} = -1$ ($B_{ij} = 0$, $B_{ij} = 1$) and $L^*(B)$ is a graph whose vertices are the edges of $G_0(B)$ with two vertices adjacent if the corresponding edges in $G_0(B)$ have a common vertex and the third edge of the associated triangle is in $G_{-1}(B)$.

By $\mathfrak{B}_n$ we denote the test set of all $n \times n$ symmetric matrices satisfying the conditions of Theorem 2. Note that, by construction, $\mathfrak{B}_n$ contains copositive matrices only, and that the famous Horn and Hoffman-Pereira matrices [HP73] are all included in $\mathfrak{B}_n$.

1.4.5. Computational results

In the first experiment, we test (COP1) and (COP2) on the test sets $\mathfrak{A}_{0.02i,50}^+$ and $\mathfrak{A}_{0.02i,50}^-$. For each $i \in [1:10]$, we generate 500 instances in each test set. The box plots, bar charts and curves in Figure 1.1 below present the results on copositive part $\mathfrak{A}_{0.02i,50}^+$ for $i \in [1:10]$ while Figure 1.2 summarizes results from the non-copositive part.

In the copositive case, a result immediately from Figures 1.1a and 1.1c is that in general, the performance of both methods worsens when c decreases. From Figure 1.1a we notice that several instances cannot be solved within 200 seconds by (COP1) in sets 1-4. We count solved instances taking more than 50 seconds as hard as well for the percentages shown in Figure 1.1b. Approximately 16% of the cases in the set 1 are hard for method (COP1), meaning they need even more time by (COP1) than the worst case for method (COP2).

The situation improves quickly with increasing c . A similar pattern can also be seen for method (COP2) in Figure 1.1c, however, among the hardest test sets 1-4, all of the instances are resolved within 40 seconds, which is much better than (COP1). This effect is also confirmed by the average CPU time reported in Figure 1.1d. As c increases, the times for each instance become similar for both methods (COP1) and (COP2), almost coinciding when $c \geq 0.14$.

For $\mathfrak{A}_{0.02i,50}^-$, the behavior of (COP1) and (COP2) follows a pattern similar to that for $\mathfrak{A}_{0.02i,50}^+$, computational cost for both methods decreases with increasing c . However, it seems that (COP1) is more efficient to solve non-copositive instances as with copositive cases, taking less time across all sets. Moreover, we note that (COP1) takes less time than (COP2) if $i \geq 2$, as shown in Figure 1.2d. This effect is not seen in the copositive case.

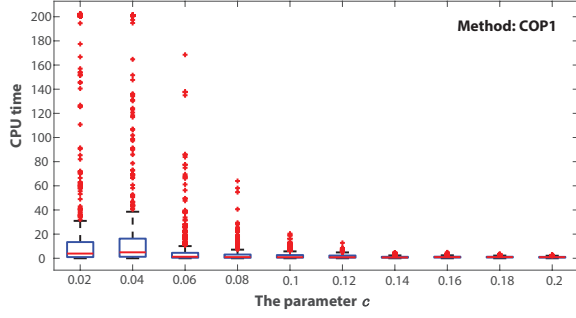
We conclude from the first experiment:

- Both methods perform worse for smaller values of c .
- For the copositive case, the performance of the method (COP2) is better than the method (COP1), while the gap reduces with increasing c .
- For the non-copositive case, performance of (COP1) seems better than that of (COP2) if $c \geq 0.04$.

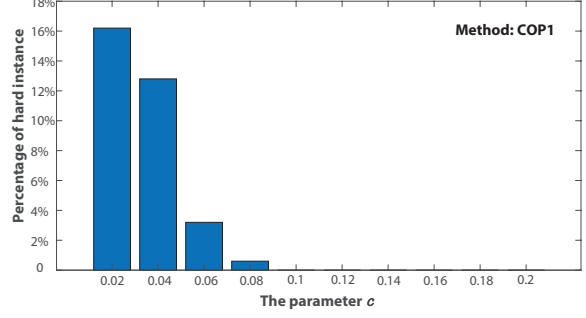
In the second experiment, we generate 500 test instances from $\mathcal{SPN}^n + \mathfrak{B}_n$ for $n \in \{50, 100, 150, 200\}$. For each generated instance $\mathbf{Q}$ we verified that $h(\mathbf{Q}) \geq 0.14$, which indicates that the hardness of the copositivity test from a close distance to the positive semidefinite cone is removed. Note that $\mathcal{SPN}^n + \mathfrak{B}_n \subseteq \mathcal{COP}^n$. In both Figure 1.3a and Figure 1.3b, the horizontal axis represents the dimension of different problem sets while the vertical axis shows the distribution of CPU time in each set. Overall, a clearly discernible trend is that the CPU time taken by (COP2) grows nearly exponentially with the problem size. In contrast, (COP1) is not sensitive to problem dimension, taking slightly more time solving high dimension (200) instances compared to lower dimension instances. Average CPU times reported in Figure 1.3c seem to support this observation.

We conclude from the second experiment:

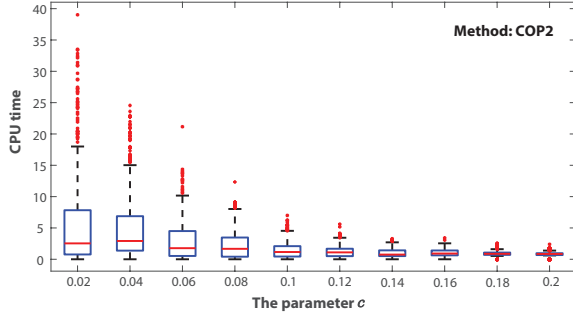
- Apart from hard instances, (COP1) performs much better than (COP2) when dealing with high dimension problems.
- Computational complexity of (COP1) grows slowly when the dimension of problems increases.
- Computational complexity of (COP2) is highly related to the dimension of problems.



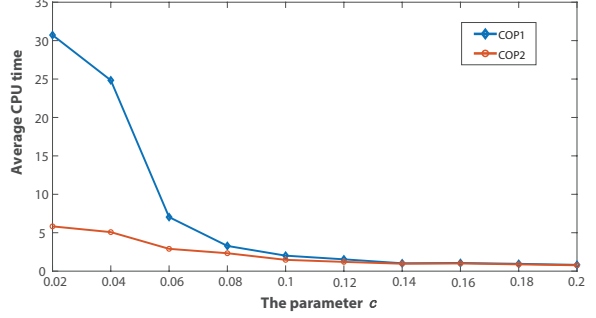
(a) CPU time by COP1



(b) Percentage of hard instances for COP1



(c) CPU time by COP2



(d) CPU time taken by COP1

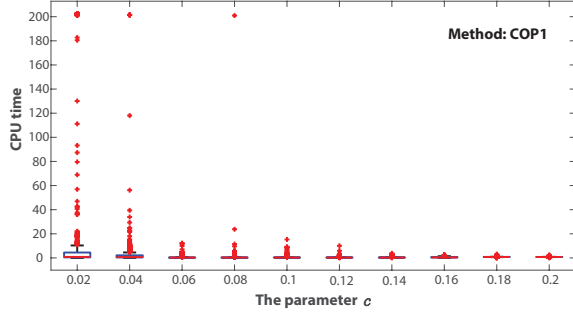
Figure 1.1.: Results on the copositive instances in $\mathfrak{A}_{0.02i,50}^+$

1.5. An improved variant of one copositivity test

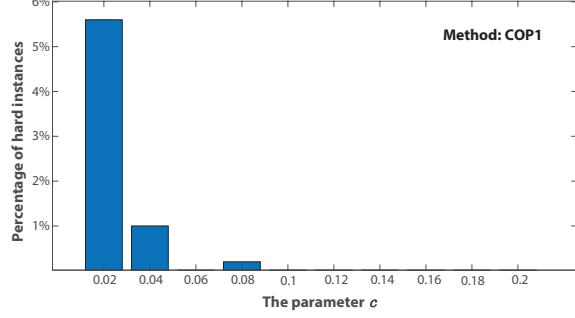
This section also contains material published in [Pen22]. As presented in the previous section, (COP1) can perform poorly in the cases when $h(Q) \leq 0.04$, and this can be explained by the nature of the method (COP1). To check copositivity, (COP1) iterates over all principal submatrices of the targeted matrix Q (done under the binary variables y) until an almost copositive matrix [Dic19] is found. Given the extreme case that $Q \in \mathcal{S}_+^n$, all of the principal submatrices of Q need to be checked, which is of exponential complexity. When the matrix parameter c is small, the distance between Q and $\mathcal{S}_+^n$ is small in the sense of the Frobenius norm. Therefore, by continuity of the condition for almost copositivity, we can infer that (COP1) needs to undergo a large number of principal submatrix checks to determine copositivity. However, compared with (COP2), (COP1) is much more competitive when tackling high-dimensional cases. To retain this nice property and to improve the poor performance of (COP1) when $h(Q)$ is small, we propose the following model:

$$\begin{aligned}
 l_3(Q) := \min \quad & \lambda \\
 \text{s.t.} \quad & \mathbf{q}_i^\top \mathbf{x} \leq \lambda + m_i(1 - y_i), \quad i \in [1:n], \\
 & \mathbf{q}_i^\top \mathbf{x} \geq \lambda, \quad i \in [1:n], \\
 & y \in \{0, 1\}^n, \mathbf{o} \leq \mathbf{x} \leq y, \\
 & \mathbf{e}^\top y \geq 1.
 \end{aligned} \tag{COP1+}$$

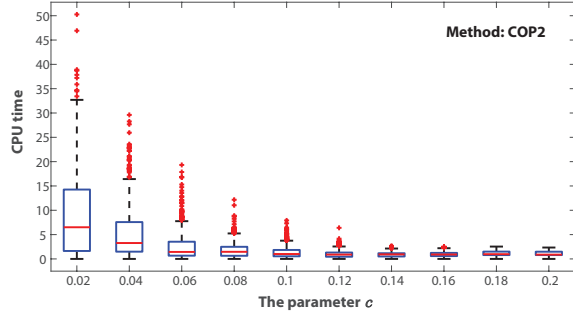
Note that the only difference between (COP1₊) and (COP1) lies in the additional inequalities $\mathbf{q}_i^\top \mathbf{x} \geq \lambda, i \in [1:n]$. The model (COP1₊) now is actually a MILP formulation for the Lemma 3.2(4)



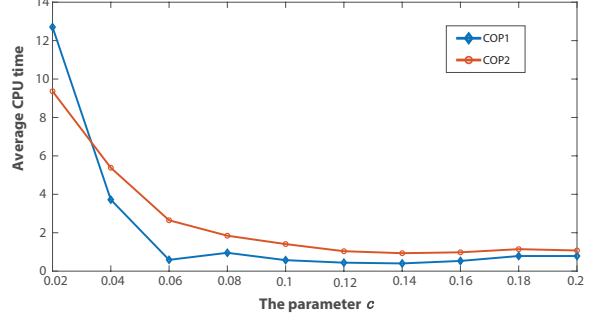
(a) CPU time by COP1



(b) Percentage of hard instances for COP1



(c) CPU time by COP2



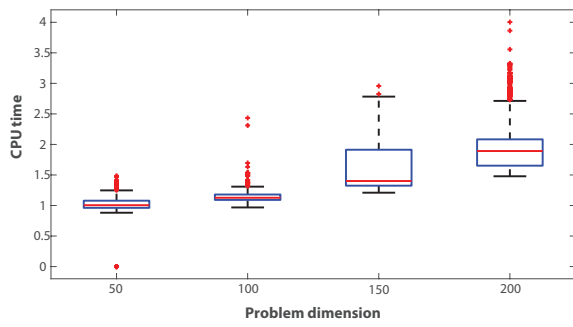
(d) Comparison on average CPU time

Figure 1.2.: Results on the non-copositive instances in $\mathfrak{A}_{0.02i,50}^-$

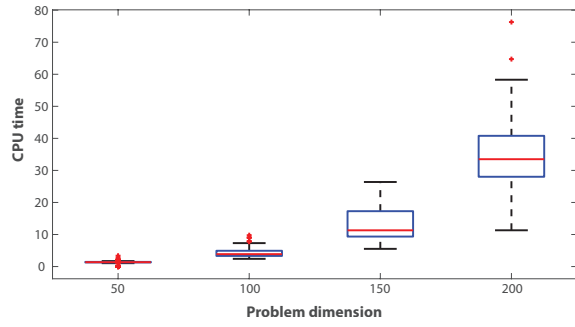
in [Dic19]. With these additional constraints, fewer principal submatrices satisfy the stronger condition for almost copositivity in the process of searching almost copositive matrices, which results in a reduction in the nodes used in solving MILP. The following Table 1.1 also confirms our speculation with roughly 80% reduction in CPU time for the test set when $c = 0.01$. Here $\mathfrak{A}_{0.01,50}$ and $\mathfrak{A}_{0.02,50}$ are two randomly generated test sets containing 500 instances respectively. It is worth noting that there are many instances that cannot be resolved within 200 seconds by (COP1) while (COP1₊) solved all of the instances within the time limit. One may also notice that the gap between average times comes closer as the parameter c increases from 0.01 to 0.02. The reason for it is that the test set contains more easy instances(noncopositive cases) when we increase the parameter c . However, as long as Q is copositive, the method (COP1) still needs much computational effort to determine its copositivity.

Table 1.1.: Comparison between (COP1) and (COP1₊) on the test sets with small c

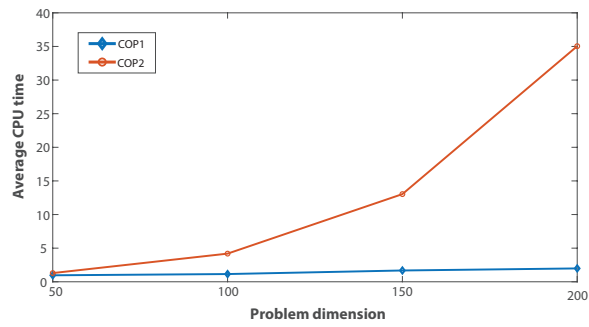
test sets	methods	average time	# hard instances	# unsolved instances
$\mathfrak{A}_{0.01,50}$	COP1	101.2	376	212
	COP1 ₊	21.2	83	0
$\mathfrak{A}_{0.02,50}$	COP1	30.3	89	62
	COP1 ₊	15.3	72	0



(a) CPU time by COP1



(b) CPU time by COP2



(c) Comparison on average CPU time

Figure 1.3.: Results of the second experiment

2. CP reformulation for QPCCs

In this chapter, we study (nonconvex) quadratic optimization problems (QPs) with complementarity constraints, establishing an exact completely positive reformulation under — apparently new — mild conditions involving only the constraints, not the objective. Moreover, we also give the conditions for strong conic duality between the obtained completely positive problem and its dual. The content of this chapter is mainly based on the published paper [BP23].

2.1. The problem

Many real-world applications can be modeled by (nonconvex) quadratic optimization problems with complementarity constraints (QPCCs) [LPR96]. Due to nonconvexity of the objective function and the complementarity constraints, QPCCs form an NP-hard class of problems. We study in this chapter the following problem where $\mathbf{Q} \in \mathbb{R}^{n \times n}$, $\mathbf{c} \in \mathbb{R}^n$, $\mathbf{x} \in \mathbb{R}^n$, $\mathbf{A} \in \mathbb{R}^{m \times n}$, $\mathbf{b} \in \mathbb{R}^m$ and $\mathbf{F} \in \mathbb{R}_+^{n \times n}$:

$$\begin{aligned} \min_{\mathbf{x} \in \mathbb{R}_+^n} \quad & \mathbf{x}^\top \mathbf{Q} \mathbf{x} + 2\mathbf{c}^\top \mathbf{x} \\ \text{s.t.} \quad & \mathbf{A} \mathbf{x} = \mathbf{b}, \mathbf{x}^\top \mathbf{F} \mathbf{x} = 0. \end{aligned} \tag{2.1}$$

Since all $F_{ij} \geq 0$, observe that $F_{ii} > 0$ would entail $x_i^2 = 0$ for all feasible $\mathbf{x}$, so that this variable can be dropped. Hence we may and do assume without loss of generality that all $F_{ii} = 0$; likewise we may and do assume symmetry: $\mathbf{F}^\top = \mathbf{F}$ and $\mathbf{Q}^\top = \mathbf{Q}$ as well. Finally, we assume $\mathbf{F} \neq \mathbf{O}$.

Problem (2.1) actually covers a wide range of relevant problems, including, for instance, all mixed-binary quadratic optimization problems which are hard as well [DDM17, AOHE19]:

$$\begin{aligned} \min_{\mathbf{x} \in \mathbb{R}_+^n} \quad & \mathbf{x}^\top \mathbf{Q} \mathbf{x} + 2\mathbf{c}^\top \mathbf{x} \\ \text{s.t.} \quad & \mathbf{A} \mathbf{x} = \mathbf{b}, x_i \in \{0, 1\}, \text{ all } i \in I, \end{aligned}$$

with I a suitable subset of all coordinate indices.

One of the possible approaches to tackle problem (2.1) is conic relaxation. In particular, [Bur09] proposed a conic relaxation model, which links QPCCs to conic linear optimization problems over the completely positive cone, and which is exact (i.e., a conic reformulation) under some conditions. However, these conditions are not necessarily satisfied in many real-world applications including quadratic optimization problems with a true sparsity term $\|\mathbf{x}\|_0$ (see below). The purpose of this work is to generalize the results of [Bur09] such that the exact relaxation technique is applicable, for instance, to all sparse QPs.

2.2. CP reformulation for QPCCs

2.2.1. Conic formulation: Burer's key condition

As announced, we focus on QPCCs of the form

$$\begin{aligned} z_P^* &:= \inf_{\mathbf{x} \in \mathbb{R}_+^n} f(\mathbf{x}) := \mathbf{x}^\top \mathbf{Q} \mathbf{x} + 2\mathbf{c}^\top \mathbf{x} \\ \text{s.t. } & \mathbf{a}_i^\top \mathbf{x} = b_i, \quad i \in [1:m], \\ & \mathbf{x}^\top \mathbf{F} \mathbf{x} = 0, \end{aligned} \tag{P}$$

where $\mathbf{Q} \in \mathcal{S}^n$, $\mathbf{F} \in \mathcal{N}^n \setminus \{\mathbf{O}\}$ with zero diagonal, $\{\mathbf{c}, \mathbf{a}_i\} \subset \mathbb{R}^n$ and $\mathbf{b} = [b_1, \dots, b_m]^\top \in \mathbb{R}^m$. We collect $\mathbf{A}^\top = [\mathbf{a}_1, \dots, \mathbf{a}_m]$ to get an $m \times n$ matrix $\mathbf{A}$ and represent the linear constraints in the form

$$\mathbf{x} \in L := \{\mathbf{x} \in \mathbb{R}_+^n : \mathbf{A}\mathbf{x} = \mathbf{b}\}.$$

Problem (P) is an all-quadratic problem of special structure; remark that the case where $\{\mathbf{Q}, \mathbf{F}\} \subset \mathcal{S}_+^n$ reduces to a traditional convex QP of the form

$$\inf_{\mathbf{x} \in \mathbb{R}_+^n} \{f(\mathbf{x}) : \mathbf{A}\mathbf{x} = \mathbf{b}, \mathbf{F}\mathbf{x} = \mathbf{o}\},$$

which can be solved in polynomial time to arbitrary accuracy. In all other cases, this problem class includes NP-hard nonconvex subclasses. Problems of form (P) have been addressed in the seminal paper by Burer [Bur09] who related (P) to the following conic optimization problem:

$$\begin{aligned} z_C^* &:= \min g(\mathbf{Y}) := \langle \mathbf{Q}, \mathbf{X} \rangle + 2\mathbf{c}^\top \mathbf{x} \\ \text{s.t. } & \mathbf{a}_i^\top \mathbf{x} = b_i, \quad i \in [1:m], \\ & \mathbf{a}_i^\top \mathbf{X} \mathbf{a}_i = b_i^2, \quad i \in [1:m], \\ & \langle \mathbf{F}, \mathbf{X} \rangle = 0, \\ & \mathbf{Y} := \begin{pmatrix} 1 & \mathbf{x}^\top \\ \mathbf{x} & \mathbf{X} \end{pmatrix} \in \mathcal{CP}^{n+1}. \end{aligned} \tag{C}$$

The following horizon cone of L plays a crucial role in this analysis:

$$L_\infty := \{\mathbf{x} \in \mathbb{R}_+^n : \mathbf{A}\mathbf{x} = \mathbf{o}\}.$$

The set of indices belonging to variables that remain bounded over L ,

$$B := \{i \in [1:n] : \text{there is an } M_i > 0 \text{ such that } x_i \leq M_i \text{ if } \mathbf{x} \in L\}, \tag{2.2}$$

is related to L_∞ by the implication $\mathbf{z} \in L_\infty \implies \mathbf{z}_B = \mathbf{o}$. We need one more index set, namely

$$S_F := \{i \in [1:n] : F_{ij} > 0 \text{ for some } j \in [1:n]\},$$

the projection of the support of $\mathbf{F}$,

$$\mathcal{E} := \{\{i, j\} : F_{ij} > 0\}.$$

We will view the sparsity pattern of $\mathbf{F}$ as an undirected graph $\mathcal{G}_F := (V, \mathcal{E})$ on the vertex set $V = [1:n]$ and edge set $\mathcal{E}$. Since by assumption $\mathbf{F} \in \mathcal{N}^n$, we have $\mathbf{x}^\top \mathbf{F} \mathbf{x} = 0$ if and only if $x_i x_j = 0$

for all $\{i, j\} \in \mathcal{E}$, so $\mathcal{G}_F$ comprises all information on F relevant to (P). Thus we can always replace F by $E := \sum_{\{i, j\} \in \mathcal{E}} e_i e_j^\top$ without changing the problem.

With these notational preliminaries, we easily can formulate the so-called key condition put forward by Burer [Bur09] in the analysis of QPCCs:

$$S_F \subseteq B, \quad (\text{K})$$

in other words, any variable x_j actually involved in the complementarity constraints, must remain bounded over the linear portion L of the feasible set. Under (K), Burer proved that the conic relaxation (C) of (P) is exact.

2.2.2. Relaxing the key condition: vertex covers

Unfortunately, condition (K) is not necessarily met in some relevant applications, and in other applications with a convex objective, the key condition is superfluous (even if the problem is not convex as may happen for $F \notin \mathcal{S}_+^n$). In this section, we will propose less restrictive conditions ensuring exactness which can be successfully applied to several optimization problems, including sparse least-squares regression with soft or hard sparsity constraints, a hot topic in data science.

Given a graph $\mathcal{G} = (V, \mathcal{E})$, recall that a subset $C \subseteq V$ is said to be a vertex cover for $\mathcal{G}$ if $\{i, j\} \cap C \neq \emptyset$ for all $\{i, j\} \in \mathcal{E}$ (this holds in particular if $\{C, V \setminus C\}$ is a bipartition of $\mathcal{G}$ in the sense that $(i, j) \in C \times (V \setminus C) \cup (V \setminus C) \times C$ for any $\{i, j\} \in \mathcal{E}$, but more general vertex covers C may also have insider edges, $C \times C \cap \mathcal{E} \neq \emptyset$). Obviously, V itself is a vertex cover but it may be too large for our purposes, as we aim to replace Burer's S_F with a (typically smaller) vertex cover C for $\mathcal{G}_F$.

The *relaxed key condition* reads as follows, using notation (2.2):

$$\text{There is a vertex cover } C \subseteq B \text{ for } \mathcal{G}_F \text{ and a point } \hat{x} \in L \text{ with } \hat{x}_C = \mathbf{o}, \quad (\text{R})$$

or equivalently, there is a vertex cover C for $\mathcal{G}_F$ and some (P)-feasible $\hat{x}$ such that for all $j \in C$ there exists an $M_j > 0$ such that

$$\hat{x}_j = 0 \quad \text{and} \quad x_j \leq M_j \text{ for all } x \in L.$$

Obviously, the conditions in (R) are met more easily if C is smaller, so the system $\mathfrak{C}(F)$ in which we collect all minimal vertex covers¹ for $\mathcal{G}_F$ is of particular relevance. But to keep our approach more transparent, we will focus on one vertex cover C first (which is sufficient for application to sparse regression) and later on, extend the following conditions for future applications to more general QPCCs. First let us note that indeed (R) relaxes (K):

Proposition 3. *Assume that (P) is feasible, so that there exists a point $\hat{x} \in L$ such that $\hat{x}^\top F \hat{x} = 0$. Then (K) implies (R).*

Proof. We show that $C := \{j \in S_F : \hat{x}_j = 0\}$ and $\hat{x} \in L$ satisfy (R) under condition (K). Indeed, by definition $\hat{x}_C = \mathbf{o}$ and $C \subseteq S_F$. Furthermore, $\hat{x}^\top F \hat{x} = 0$ implies $\hat{x}_i \hat{x}_j = 0$ for all $\{i, j\} \in \mathcal{E}$, so that $\{i, j\} \cap C \neq \emptyset$, hence C is a vertex cover for $\mathcal{G}_F$. $\square$

¹note that we do not suggest to determine $\mathfrak{C}(F)$ which may be exponential in size (as with the maximum-clique problem) but rather propose to exploit structural knowledge of the problem at hand.

In the case of a convex objective, condition (K) and even condition (R) are unnecessary. Actually, we only need the conclusion that $\mathbf{Q}$ is copositive on the cone

$$L_\infty \cap F := \left\{ \mathbf{z} \in \mathbb{R}_+^n : \mathbf{A}\mathbf{z} = \mathbf{o} \text{ and } \mathbf{z}^\top \mathbf{F}\mathbf{z} = 0 \right\},$$

the intersection of the horizon cone of L with the complementarity set

$$F := \left\{ \mathbf{x} \in \mathbb{R}_+^n : \mathbf{x}^\top \mathbf{F}\mathbf{x} = 0 \right\}.$$

Condition (R) guarantees that $L_\infty \cap F = L_\infty$ and the mentioned copositivity property of $\mathbf{Q}$, which is also implied by the convexity of the objective. Essential arguments are summarized in the following.

Lemma 4. *Assume $z_P^* > -\infty$, so that problem (P) is bounded, and that (R) holds. Then $L_\infty \subseteq F$, and for any $\mathbf{z} \in L_\infty$ we have $\mathbf{z}^\top \mathbf{Q}\mathbf{z} \geq 0$. The latter property holds if $\mathbf{Q} \in \mathcal{COP}^n$ (e.g., because $\mathbf{Q}$ is positive-semidefinite).*

Proof. If condition (R) holds for some vertex cover $C \subseteq V$ of $\mathcal{G}_F$, pick an arbitrary $\mathbf{z} \in L_\infty$ and select $\hat{\mathbf{x}} \in L$ with $\hat{\mathbf{x}}_C = \mathbf{o}$. Then for all $t \geq 0$, we have $\mathbf{x}(t) = \hat{\mathbf{x}} + t\mathbf{z} \in L$ which implies, by $C \subseteq B$ from (R), that $\mathbf{z}_C = \mathbf{o}$, which in turn implies $\mathbf{x}(t)_C = \mathbf{o}$ for all $t \geq 0$. So $x_i(t)x_j(t) = 0$ for all $j \in C$ (and all $i \in V$) which yields $\mathbf{x}(t)^\top \mathbf{F}\mathbf{x}(t) = 0$ since C is a vertex cover for $\mathcal{G}_F$. Hence $\mathbf{x}(t) \in F$ and in a similar fashion, we derive $\mathbf{z} \in F$. We conclude $L_\infty \subseteq F$ and furthermore that $\mathbf{x}(t)$ is (P)-feasible, hence

$$-\infty < z_P^* \leq f(\mathbf{x}(t)) = f(\hat{\mathbf{x}}) + 2t\mathbf{z}^\top (\mathbf{Q}\hat{\mathbf{x}} + \mathbf{c}) + t^2 \mathbf{z}^\top \mathbf{Q}\mathbf{z} \quad \text{for all } t \geq 0,$$

which implies $\mathbf{z}^\top \mathbf{Q}\mathbf{z} \geq 0$. This implication is trivial if $\mathbf{Q} \in \mathcal{COP}^n$, since $L_\infty \subseteq \mathbb{R}_+^n$ by definition. $\square$

Note: if $\mathbf{z}^\top \mathbf{Q}\mathbf{z} = 0$, we could infer, by the same reasoning, that $\mathbf{z}^\top \nabla f(\hat{\mathbf{x}}) = 2\mathbf{z}^\top (\mathbf{Q}\hat{\mathbf{x}} + \mathbf{c}) \geq 0$, i.e., that $\mathbf{z}$ could not be a strict descent direction for f at $\hat{\mathbf{x}}$. But we will not explore this first-order information further here.

Now we formulate a variant applicable to other QPCCs:

Lemma 5. *Assume $z_P^* > -\infty$, so that problem (P) is bounded. Suppose that for any $C \in \mathfrak{C}(F)$ with nontrivial cone $\{\mathbf{x} \in L_\infty : \mathbf{x}_C = \mathbf{o}\}$, there exists a point $\hat{\mathbf{x}} \in L$ such that $\hat{\mathbf{x}}_C = \mathbf{o}$. Then $\mathbf{Q}$ is $(L_\infty \cap F)$ -copositive.*

Proof. First observe that

$$F \subseteq \bigcup_{C \in \mathfrak{C}(F)} \{\mathbf{z} \in \mathbb{R}_+^n : \mathbf{z}_C = \mathbf{o}\}.$$

Indeed, for any $\mathbf{z} \in F$ consider $C_z = \{i \in V : z_i = 0\} \neq \emptyset$ (due to $\mathbf{F} \neq \mathbf{O}$), which obviously is a vertex cover for $\mathcal{G}_F$. Clearly there is an $C \in \mathfrak{C}(F)$ with $C \subseteq C_z$, so $\mathbf{z}_C = \mathbf{o}$. Then the result follows by arguments similar to those employed in the previous proof. $\square$

The following examples shall demonstrate wider applicability of our proposed conditions. But for the readers' convenience let us start with a counterexample already presented in [BMP16] which shows that some conditions are needed to avoid a positive, or even infinite, conic relaxation gap:

Example 1.

$$\begin{aligned}
& \min -x_2^2 \\
& \text{s.t. } x_1 + x_2 - x_3 = 3, \\
& \quad x_1 - x_4 = 2, \\
& \quad x_1 x_2 = 0, \\
& \quad \mathbf{x} \in \mathbb{R}_+^4.
\end{aligned} \tag{2.3}$$

As argued in [BMP16, Ex.1], the relaxation (C) is unbounded below while the optimal value of (P) is zero. As can be seen easily, conditions (K) and (R) are violated, and as well the one of Lemma 5: all $\mathbf{x} \in L$ must have $x_1 \geq 2 > 0$, so $\mathbf{x}_C \neq \mathbf{o}$ for $C = \{1\} \in \mathfrak{C}(\mathbf{F})$ while L_∞ contains, e.g., $(0, 1, 1, 0)^\top$.

Example 2.

$$\begin{aligned}
& x_1 - x_2 - x_3 + x_4 = 1, \\
& \quad x_1 + x_2 = 1, \\
& \quad x_1 x_2 + x_1 x_3 = 0, \\
& \quad \mathbf{x} \in \mathbb{R}_+^4.
\end{aligned} \tag{2.4}$$

Note that the second constraint implies boundedness of x_1 and x_2 , so we have $B = \{1, 2\}$. However, since the set S of variables involved in the complementary constraint is $S = \{1, 2, 3\} \supset \{1, 2\}$, condition (K) is violated. On the other hand, the class of minimal vertex covers of (2.4) is $\mathfrak{C}(\mathbf{F}) = \{\{1\}, \{2, 3\}\}$, with a vertex cover $C = \{1\} \subset B$. Moreover, $(0, 1, 1, 3)^\top$ is (2.4)-feasible and therefore condition (R) is met. Also, the condition of Lemma 5 holds, as $C = \{2, 3\}$ yields a trivial cone $\{\mathbf{z} \in L_\infty : \mathbf{z}_C = \mathbf{o}\}$.

Example 3.

$$\begin{aligned}
& x_1 - x_2 - x_3 + x_4 - x_5 = 1, \\
& \quad x_1 + x_2 = 1, \\
& \quad x_2 x_3 + x_3 x_4 = 0, \\
& \quad \mathbf{x} \in \mathbb{R}_+^5.
\end{aligned} \tag{2.5}$$

As in (2.4), we have $B = \{1, 2\}$. Again, condition (K) is violated since $S = \{2, 3, 4\}$ is not contained in B . The class of minimal vertex covers of (2.5) is $\mathfrak{C}(\mathbf{F}) = \{\{3\}, \{2, 4\}\}$. None of them is contained in B , so condition (R) is violated. For $C = \{3\}$, the cone

$$\{\mathbf{x} \in L_\infty : \mathbf{x}_C = \mathbf{o}\} = \{\mathbf{x} \in \mathbb{R}_+^5 : x_4 = x_5, x_1 = x_2 = x_3 = 0\}$$

is nontrivial, and the point $\hat{\mathbf{x}} = (1, 0, 0, 0, 0)^\top \in L$ satisfies $\hat{\mathbf{x}}_C = \mathbf{o}$. But for $C = \{2, 4\}$, the cone

$$\{\mathbf{x} \in L_\infty : \mathbf{x}_C = \mathbf{o}\} = \{\mathbf{x} \in \mathbb{R}_+^5 : x_3 + x_5 = 0, x_1 = x_2 = x_4 = 0\} = \{\mathbf{o}\}.$$

Therefore the system satisfies the condition in Lemma 5.

Example 4.

$$\begin{aligned}
& x_1 + x_4 - x_5 + x_6 = 5, \\
& \quad x_2 + x_3 - x_5 + x_6 = 1, \\
& \quad x_1 + x_2 = 1, \\
& \quad x_1 x_3 + x_1 x_4 = 0, \\
& \quad \mathbf{x} \in \mathbb{R}_+^6.
\end{aligned} \tag{2.6}$$

As in (2.4), we have $B = \{1, 2\}$. Again, condition (K) is violated since $S = \{1, 3, 4\}$ is not contained in B . On the other hand, the class of minimal vertex covers of (2.6) is $\mathfrak{C}(F) = \{\{1\}, \{3, 4\}\}$, with a vertex cover $C = \{1\} \subset B$. Moreover, $(0, 1, 0, 5, 2, 2)^\top$ is (2.6)-feasible and therefore condition (R) is met. For $C = \{3, 4\}$, the cone

$$\{\mathbf{x} \in L_\infty : \mathbf{x}_C = \mathbf{o}\} = \{\mathbf{x} \in \mathbb{R}_+^6 : x_5 = x_6, x_1 = x_2 = x_3 = x_4 = 0\}$$

is nontrivial. However, the set $\{\mathbf{x} \in L : \mathbf{x}_C = \mathbf{o}\}$ is empty now, because $\mathbf{x}_C = \mathbf{o}$ implies a contradiction, looking at the sum of the first two constraints and the third one, in conjunction with the sign constraints on x_1 and x_2 . Therefore, the condition of Lemma 5 is violated.

While (K) implies (R), it does not imply the condition of Lemma 5; we give the following example:

Example 5.

$$\begin{aligned} x_3 - x_4 + x_5 &= 2, \\ x_2 - x_4 + x_5 &= 1, \\ x_1 + x_2 + x_3 &= 2, \\ x_1 x_2 + x_1 x_3 &= 0, \\ \mathbf{x} &\in \mathbb{R}_+^5. \end{aligned} \tag{2.7}$$

Since we have $S = \{1, 2, 3\} = B$, condition (K) holds. On the other hand, the class of minimal vertex covers of (2.7) is $\mathfrak{C}(F) = \{\{1\}, \{2, 3\}\}$. For $C = \{2, 3\}$, the cone

$$\{\mathbf{x} \in L_\infty : \mathbf{x}_C = \mathbf{o}\} = \{\mathbf{x} \in \mathbb{R}_+^5 : x_4 = x_5, x_1 = x_2 = x_3 = 0\}$$

is nontrivial. However, the set $\{\mathbf{x} \in L : \mathbf{x}_C = \mathbf{o}\}$ is empty now because of the contradiction between the first two constraints. Therefore, the condition of Lemma 5 is violated.

As a final note, it is easy to construct examples for which both (K) and the condition of Lemma 5 are satisfied [BMP16].

Now we are in a position to state and prove the following result which generalizes the one in [Bur09]. Although the main arguments are well known by now, we repeat them here for the readers' convenience. Note that even without any assumptions, the first part of the following proof also shows that (P) is feasible if (C) is feasible.

Theorem 6. *Suppose that $\mathbf{Q}$ is $(L_\infty \cap F)$ -copositive in the sense of Lemmas 4, 5. Then the QPCC problem (P) is equivalent to the conic problem (C).*

Proof. It is easy to see that any (P)-feasible $\mathbf{x} \in \mathbb{R}_+^n$ gives rise to a (C)-feasible point $\mathbf{Y} = \mathbf{u}\mathbf{u}^\top \in \mathcal{CP}^{n+1}$ of rank one with $\mathbf{u}^\top = [1, \mathbf{x}^\top]$ and same objective value $f(\mathbf{x}) = \langle \mathbf{Q}, \mathbf{x}\mathbf{x}^\top \rangle + 2\mathbf{c}^\top \mathbf{x} = g(\mathbf{Y})$. Therefore (C) is a conic relaxation of (P) and $z_C^* \leq z_P^*$. We proceed to show equivalence by establishing the reverse inequality. To this end, take a (C)-feasible and optimal $\mathbf{Y}$ (or approximately optimal, in case of non-attainability for either of the problems (P) or (C); we will address primal attainability later). Then consider the following decomposition for this completely positive matrix:

$$\mathbf{Y} = \begin{pmatrix} 1 & \mathbf{x}^\top \\ \mathbf{x} & \mathbf{X} \end{pmatrix} = \sum_{k \in K} \begin{pmatrix} \lambda_k \\ \mathbf{z}_k \end{pmatrix} \begin{pmatrix} \lambda_k \\ \mathbf{z}_k \end{pmatrix}^\top, \quad \lambda_k \in \mathbb{R}_+, \mathbf{z}_k \in \mathbb{R}_+^n, k \in K, \tag{2.8}$$

with $g(\mathbf{Y}) = \langle \mathbf{Q}, \mathbf{X} \rangle + 2\mathbf{c}^\top \mathbf{x} \approx z_C^*$. Denote by $K_0 := \{k \in K : \lambda_k = 0\}$ and $K_+ := \{k \in K : \lambda_k > 0\}$ respectively. The following equalities hold directly by the decomposition (2.8) and the constraints in (C):

$$1 = \sum_{k \in K} \lambda_k^2, \quad (\text{so } K_+ \neq \emptyset), \quad (2.9)$$

$$b_i = \mathbf{a}_i^\top \mathbf{x} = \sum_{k \in K} \lambda_k \mathbf{a}_i^\top \mathbf{z}_k, \quad i \in [1:m], \quad (2.10)$$

$$b_i^2 = \mathbf{a}_i^\top \mathbf{X} \mathbf{a}_i = \sum_{k \in K} (\mathbf{a}_i^\top \mathbf{z}_k)^2, \quad i \in [1:m]. \quad (2.11)$$

Combining the equalities (2.9), (2.10) and (2.11), we have for any fixed $i \in [1:m]$,

$$\sum_{k \in K} \lambda_k^2 \cdot \sum_{k \in K} (\mathbf{a}_i^\top \mathbf{z}_k)^2 = 1 \cdot b_i^2 = b_i^2 = \left(\sum_{k \in K} \lambda_k \mathbf{a}_i^\top \mathbf{z}_k \right)^2.$$

By the condition of equality case of the Cauchy-Schwarz inequality, there exists $\beta_i \in \mathbb{R}$ such that

$$\mathbf{a}_i^\top \mathbf{z}_k = \beta_i \lambda_k \quad \text{for all } k \in K. \quad (2.12)$$

Multiplying by λ_k and summing relations (2.12) across k and by (2.9) and (2.10), we obtain

$$\beta_i = b_i \quad \text{for all } i \in [1:m]. \quad (2.13)$$

Moreover, by (C)-feasibility of $\mathbf{Y}$ we have

$$0 = \langle \mathbf{F}, \mathbf{X} \rangle = \sum_{k \in K} \mathbf{z}_k^\top \mathbf{F} \mathbf{z}_k \implies \mathbf{z}_k^\top \mathbf{F} \mathbf{z}_k = 0, \quad \text{for all } k \in K.$$

Therefore, by taking $\mathbf{x}_k := \frac{1}{\lambda_k} \mathbf{z}_k$ if $\lambda_k > 0$, we get $\mathbf{x}_k \in L \cap F$ for all $k \in K_+$ while for $\lambda_k = 0$ we obtain $\mathbf{A} \mathbf{z}_k = \mathbf{o}$, i.e., $\mathbf{z}_k \in L_\infty \cap F$ and therefore $\mathbf{z}_k^\top \mathbf{Q} \mathbf{z}_k \geq 0$ for all $k \in K_0$. We obtain that all $\mathbf{x}_k$ are (P)-feasible, so that

$$z_P^* \leq z_+^* := \min \{f(\mathbf{x}_k) : k \in K_+\} \leq z_C^*,$$

where the last inequality follows from

$$\begin{aligned} z_+^* &= \sum_{k \in K_+} (\lambda_k)^2 z_+^* \leq \sum_{k \in K_+} (\lambda_k)^2 f(\mathbf{x}_k) \\ &= \sum_{k \in K_+} [\mathbf{z}_k^\top \mathbf{Q} \mathbf{z}_k + \lambda_k \mathbf{c}^\top \mathbf{z}_k] \leq \sum_{k \in K_+} [\mathbf{z}_k^\top \mathbf{Q} \mathbf{z}_k + \lambda_k \mathbf{c}^\top \mathbf{z}_k] + \sum_{k \in K_0} \mathbf{z}_k^\top \mathbf{Q} \mathbf{z}_k \\ &= \sum_{k \in K} [\mathbf{z}_k^\top \mathbf{Q} \mathbf{z}_k + \lambda_k \mathbf{c}^\top \mathbf{z}_k] = g(\mathbf{Y}) \approx z_C^*, \end{aligned}$$

which establishes the desired inequality $z_P^* \leq z_C^*$ and moreover that any $\mathbf{x}^k$ realizing the minimum in z_+^* provides an (approximate) optimal solution to (P). $\square$

Note: since the feasible set of (P) is a finite union of polyhedra (if not empty), the Frank-Wolfe theorem guarantees attainability of z_P^* if this value is finite. Slater's condition (strict feasibility of the conic dual below) would ensure attainability of z_C^* , but above equivalence ensures this automatically: indeed, $\mathbf{Y}^* = \mathbf{u}^*(\mathbf{u}^*)^\top$ solves (C) if $(\mathbf{u}^*)^\top = [1, (\mathbf{x}^*)^\top]$ and $\mathbf{x}^*$ solves (P).

2.2.3. A simplification of (C) and its dual

We show a possible simplification of (C) which is initially proposed in [Bur09], reducing the cone in the feasible set from $\mathcal{CP}^{n+1}$ to $\mathcal{CP}^n$, and which works under

Condition 1. *There exists $\mathbf{y} \in \mathbb{R}_+^m$ such that*

$$\mathbf{d} := \mathbf{A}^\top \mathbf{y} = \sum_{i=1}^m y_i \mathbf{a}_i \in \mathbb{R}_+^n \quad \text{and} \quad \mathbf{b}^\top \mathbf{y} = 1.$$

Condition 1 is certainly met if $\mathbf{a}_i \in \mathbb{R}_+^n$ and $b_i > 0$ for at least one $i \in [1:m]$, taking $\mathbf{y} = \frac{1}{b_i} \mathbf{e}_i \in \mathbb{R}_+^m$.

It is clear that the constraints $\mathbf{Ax} = \mathbf{b}$ imply $\mathbf{d}^\top \mathbf{x} = \mathbf{y}^\top \mathbf{b} = 1$. From the decomposition (2.8) and as in the proof of Theorem 6 we have $\mathbf{d}^\top \mathbf{z}_k = \mathbf{y}^\top \mathbf{Az}_k = \mathbf{y}^\top \mathbf{o} = 0$ for all $k \in K_0$ and

$$\mathbf{x} = \sum_{k \in K_+} \lambda_k \mathbf{z}_k = \sum_{k \in K_+} \lambda_k^2 \mathbf{x}_k$$

with $\mathbf{Ax}_k = \mathbf{b}$, entailing $\mathbf{d}^\top \mathbf{x}_k = 1$ for all $k \in K_+$, yielding

$$\mathbf{Xd} = \sum_{k \in K} \mathbf{z}_k \mathbf{z}_k^\top \mathbf{d} = \sum_{k \in K_+} \lambda_k^2 \mathbf{x}_k \mathbf{x}_k^\top \mathbf{d} + \sum_{k \in K_0} \mathbf{z}_k \mathbf{z}_k^\top \mathbf{d} = \sum_{k \in K_+} \lambda_k^2 \mathbf{x}_k + \mathbf{o} = \mathbf{x},$$

which allows us to rewrite (C) as:

$$\begin{aligned} \inf \quad & \langle \mathbf{Q} + (\mathbf{dc}^\top + \mathbf{cd}^\top), \mathbf{X} \rangle = \langle \mathbf{Q}, \mathbf{X} \rangle + 2\mathbf{c}^\top \mathbf{Xd} \\ \text{s.t.} \quad & \langle \mathbf{da}_i^\top + \mathbf{a}_i \mathbf{d}^\top, \mathbf{X} \rangle = 2\mathbf{a}_i^\top \mathbf{Xd} = 2b_i, \quad i \in [1:m], \\ & \langle \mathbf{a}_i \mathbf{a}_i^\top, \mathbf{X} \rangle = \mathbf{a}_i^\top \mathbf{Xa}_i = b_i^2, \quad i \in [1:m], \\ & \langle \mathbf{F}, \mathbf{X} \rangle = 0, \\ & \mathbf{X} \in \mathcal{CP}^n. \end{aligned} \tag{C_{sim}}$$

We abbreviate $\mathbf{G} := \mathbf{Q} + (\mathbf{dc}^\top + \mathbf{cd}^\top)$ and $\mathbf{A}_i := \frac{1}{2}(\mathbf{da}_i^\top + \mathbf{a}_i \mathbf{d}^\top)$ as well as $\bar{\mathbf{A}}_i := \mathbf{a}_i \mathbf{a}_i^\top$ for all $i \in [1:m]$ to rewrite (C_{sim}) in a simpler form, recalling that $\langle \mathbf{F}, \mathbf{X} \rangle = 0 \Leftrightarrow \langle \mathbf{E}, \mathbf{X} \rangle = 0$:

$$\begin{aligned} \inf \quad & \langle \mathbf{G}, \mathbf{X} \rangle \\ \text{s.t.} \quad & \langle \mathbf{A}_i, \mathbf{X} \rangle = b_i, \quad i \in [1:m], \\ & \langle \bar{\mathbf{A}}_i, \mathbf{X} \rangle = b_i^2, \quad i \in [1:m], \\ & \langle \mathbf{E}, \mathbf{X} \rangle = 0, \\ & \mathbf{X} \in \mathcal{CP}^n. \end{aligned} \tag{CP}$$

By introducing dual variables $\mathbf{r} \in \mathbb{R}^m$ for constraints $\langle \mathbf{A}_i, \mathbf{X} \rangle = b_i$, $\mathbf{t} \in \mathbb{R}^m$ for constraints $\langle \bar{\mathbf{A}}_i, \mathbf{X} \rangle = b_i^2$ and $h \in \mathbb{R}$ for the constraint $\langle \mathbf{E}, \mathbf{X} \rangle = 0$, the dual problem of (CP) can be obtained:

$$\begin{aligned} \sup \quad & \sum_{i=1}^m (r_i b_i + t_i b_i^2) \\ \text{s.t.} \quad & \mathbf{S}(\mathbf{r}, \mathbf{t}, h) := \mathbf{G} - \sum_{i=1}^m (r_i \mathbf{A}_i + t_i \bar{\mathbf{A}}_i) - h \mathbf{E} \in \mathcal{COP}^n \\ & (\mathbf{r}, \mathbf{t}, h) \in \mathbb{R}^m \times \mathbb{R}^m \times \mathbb{R}. \end{aligned} \tag{CD}$$

While we have, by preceding arguments under the assumption of $(L_\infty \cap F)$ -copositivity of $\mathbf{Q}$ and $z_P^* \in \mathbb{R}$,

$$z_P^* = z_C^* = \inf \{ \langle \mathbf{G}, \mathbf{X} \rangle : \mathbf{X} \text{ is (CP)-feasible} \}$$

as the optimal value for (CP), we denote by

$$z_D^* := \sup \left\{ \sum_{i=1}^m (r_i b_i + t_i b_i^2) : \mathbf{S}(\mathbf{r}, \mathbf{t}, h) \in \mathcal{COP}^n \right\}$$

the optimal value for (CD).

We will now establish sufficient conditions for strict feasibility of (CD), which means that Slater's condition holds for (CD). The condition seems quite natural as it is a sharpening of the conclusion of Lemma 4. It is met if $\mathbf{Q}$ is positive-definite, and for general $\mathbf{Q}$ if L is bounded (implying $L_\infty = \{\mathbf{o}\}$), by default:

Theorem 7. *Suppose that $z_P^* \in \mathbb{R}$ and that $\mathbf{Q}$ is strictly L_∞ -copositive: $\mathbf{z}^\top \mathbf{Q} \mathbf{z} > 0$ if $\mathbf{z} \in L_\infty \setminus \{\mathbf{o}\}$. Then*

$$z_C^* = z_D^* \quad \text{and the optimal value } z_C^* \text{ is attained.}$$

Proof. First observe that by definition of $\mathbf{d} = \mathbf{A}^\top \mathbf{y}$, we have $\mathbf{z}^\top \mathbf{d} = (\mathbf{A} \mathbf{z})^\top \mathbf{d} = 0$ for all $\mathbf{z} \in L_\infty$, hence $\mathbf{z}^\top \mathbf{G} \mathbf{z} = \mathbf{z}^\top \mathbf{Q} \mathbf{z}$ on L_∞ . Now, under the assumptions, there is a constant $\sigma > 0$ such that

$$\mathbf{z}^\top \mathbf{G} \mathbf{z} = \mathbf{z}^\top \mathbf{Q} \mathbf{z} \geq 2\sigma \quad \text{for all } \mathbf{z} \in \Delta_\infty := L_\infty \cap \Delta^n.$$

By continuity, there is a $\delta > 0$ such that also in a small relative neighbourhood

$$U_\delta := \{\mathbf{x} \in \Delta^n : \text{dist}(\mathbf{x}, \Delta_\infty) < \delta\}$$

we have $\mathbf{z}^\top \mathbf{G} \mathbf{z} \geq \sigma > 0$. The continuous function $\|\mathbf{A} \mathbf{x}\|^2$ takes its positive minimum $\mu > 0$ on the compact set $\Delta^n \setminus U_\delta$ which is disjoint from L_∞ ; furthermore, we have that $\gamma := \min_{\mathbf{x} \in \Delta^n} \mathbf{z}^\top \mathbf{G} \mathbf{z} \in \mathbb{R}$.

Choosing $\alpha := \max \left\{ 0, \frac{\sigma - \gamma}{\mu} \right\}$ gives therefore a strictly copositive matrix $\mathbf{S}(\mathbf{o}, -\alpha \mathbf{e}, 0)$ with a quadratic form $\mathbf{z}^\top \mathbf{G} \mathbf{z} + \alpha \|\mathbf{A} \mathbf{z}\|^2$ taking values not smaller than σ over U_δ but also, by choice of α , not smaller than σ over $\Delta^n \setminus U_\delta$. Hence Slater's condition holds for (CD) and the claim follows. $\square$

2.3. QPCC and CP formulations for quadratic problems under sparsity constraints

In this section, we discuss two models for pursuing sparsity in quadratic optimization problems with the direct use of $\|\cdot\|_0$ put forward by [BKS16, FMP⁺18]. Both are based on the following idea: negating the binary variables $|\text{sign}(y_i)| \in \{0, 1\}$, we arrive again at binary variables $u_i := 1 - |\text{sign}(y_i)| \in \{0, 1\}$ which satisfy $u_i |y_i| = 0$ for all i . As no negative variables appear in this complementarity relation, we can proceed towards a QPCC with a single homogeneous bilinear equality constraint. An additional interesting feature of this approach is that binarity of u_i is not needed here explicitly as a constraint, only the upper bounds $u_i \leq 1$ are already sufficient, as will be argued below. As already noted in the introduction, a naive approach via Burer's result [Bur09] would result in additional constraints in the conic formulation which may be detrimental to algorithmic performance. Anyhow, some of them may be added in the implementation of specially structured problems.

2.3.1. Soft sparsity: regularization by zero-norm

No sign constraints on the variables

We will first treat sparse quadratic problems under linear equality constraints, but no sign constraints on the variables. At first sight, this may seem a restriction of models but actually any sign constraint on some variables would reduce dimensionality rather than increase complexity, as will be clear in the next subsection.

So consider the following soft sparsity model:

$$\min_{\mathbf{y} \in \mathbb{R}^n} \left\{ \mathbf{y}^\top \mathbf{Q}_0 \mathbf{y} + 2\mathbf{c}_0^\top \mathbf{y} + \rho \|\mathbf{y}\|_0 : \mathbf{c}_i^\top \mathbf{y} = b_i, i \in [1:m] \right\}, \quad (\text{P}_{S_0})$$

where $\rho > 0$ is a regularization parameter for balancing accuracy against sparsity: the larger ρ , the more sparse the solution obtained from (P_{S_0}) will be.

By introduction of new variables $\mathbf{u} \in \mathbb{R}_+^n$ and decomposition of $\mathbf{y} = \mathbf{v} - \mathbf{w}$, where $\{\mathbf{v}, \mathbf{w}\} \subset \mathbb{R}_+^n$, problem (P_{S_0}) can be exactly expressed by a quadratic problem:

$$\begin{aligned} \min_{\{\mathbf{v}, \mathbf{w}, \mathbf{u}\} \subset \mathbb{R}_+^n} \quad & (\mathbf{v} - \mathbf{w})^\top \mathbf{Q}_0 (\mathbf{v} - \mathbf{w}) + 2\mathbf{c}_0^\top (\mathbf{v} - \mathbf{w}) + \rho n - \rho \mathbf{e}^\top \mathbf{u} \\ \text{s.t.} \quad & \mathbf{c}_i^\top (\mathbf{v} - \mathbf{w}) = b_i, \quad i \in [1:m], \\ & \mathbf{u}^\top (\mathbf{v} + \mathbf{w}) = 0, \\ & \mathbf{u} \leq \mathbf{e}. \end{aligned} \quad (\text{P}_{S'})$$

Because of nonnegativity of $(\mathbf{u}, \mathbf{v}, \mathbf{w})$, the homogeneous bilinear equality constraint $\mathbf{u}^\top (\mathbf{v} + \mathbf{w}) = 0$ is equivalent to the equality system $u_i(v_i + w_i) = 0, i \in [1:n]$.

Theorem 8. *Problems (P_{S_0}) and $(\text{P}_{S'})$ are equivalent.*

Proof. This was proved in [FMP⁺18], but we provide a short argument for the readers' convenience here. For any $\alpha \in \mathbb{R}$, let $\alpha_+ := \max\{\alpha, 0\}$ and $\alpha_- = \alpha_+ - \alpha$; then $\alpha_\pm \geq 0$ and $|\alpha| = \alpha_+ + \alpha_-$ as well as $\alpha = \alpha_+ - \alpha_-$. If $\mathbf{y}$ is feasible to (P_{S_0}) , then putting $v_i = (y_i)_+$ and $w_i = (y_i)_-$ as well as $u_i = 1 - |\text{sign } y_i| \in \{0, 1\}$ one gets a feasible solution to $(\text{P}_{S'})$ with same objective value, so $\text{opt}(\text{P}_{S'}) \leq \text{opt}(\text{P}_{S_0})$. Conversely, let

$$F' := \{(\mathbf{v}, \mathbf{w}, \mathbf{u}) \in \mathbb{R}_+^{3n} : \mathbf{u}^\top (\mathbf{v} + \mathbf{w}) = 0, \mathbf{u} \leq \mathbf{e}\}.$$

For $(\mathbf{v}, \mathbf{w}, \mathbf{u}) \in F'$ feasible to $(\text{P}_{S'})$, put $\mathbf{y} := \mathbf{v} - \mathbf{w} \in \mathbb{R}^n$ and consider $v'_i := (y_i)_+, w'_i := (y_i)_-$ and $u'_i := \text{sign } u_i$. Then also $(\mathbf{v}', \mathbf{w}', \mathbf{u}') \in F'$ is $(\text{P}_{S'})$ -feasible with lower or equal objective value due to $\mathbf{e}^\top \mathbf{u} \leq \mathbf{e}^\top \mathbf{u}'$ and $\mathbf{v}' - \mathbf{w}' = \mathbf{v} - \mathbf{w}$, and again, this value coincides with the objective value of $\mathbf{y}$ in (P_{S_0}) . Thus $\text{opt}(\text{P}_{S_0}) \leq \text{opt}(\text{P}_{S'})$. $\square$

Given $(\mathbf{v}, \mathbf{w}, \mathbf{u}) \in F'$, we introduce slack variables $\mathbf{s} := \mathbf{e} - \mathbf{u} \in \mathbb{R}_+^n$ if $\mathbf{u} \leq \mathbf{e}$ and stack all variables to $\mathbf{x} := [\mathbf{v}^\top, \mathbf{w}^\top, \mathbf{u}^\top, \mathbf{s}^\top]^\top \in \mathbb{R}_+^{4n}$. Then problem $(\text{P}_{S'})$ can be rewritten, dropping the additive constant ρn :

$$\begin{aligned} \min_{\mathbf{x} \in \mathbb{R}_+^{4n}} \quad & \mathbf{x}^\top \mathbf{Q} \mathbf{x} + 2\mathbf{c}^\top \mathbf{x} \\ \text{s.t.} \quad & \mathbf{a}_i^\top \mathbf{x} = b_i, \quad i \in [1:m+n], \\ & \sum_{(i,j) \in \mathcal{E}} x_i x_j = 0, \end{aligned} \quad (\text{P}_S)$$

where

$$\mathbf{Q} = \begin{pmatrix} \mathbf{Q}_0 & -\mathbf{Q}_0 & \mathbf{O} \\ -\mathbf{Q}_0 & \mathbf{Q}_0 & \mathbf{O} \\ \mathbf{O} & \mathbf{O} & \mathbf{O} \end{pmatrix} \in \mathcal{S}^{4n} \text{ and } \mathbf{c} = [\mathbf{c}_0^\top, -\mathbf{c}_0^\top, -\frac{\rho}{2}\mathbf{e}^\top, \mathbf{o}^\top]^\top \in \mathbb{R}^{4n},$$

as well as the linear constraints, putting $b_i = 1$ for all $i \in [m+1:m+n]$ while keeping all b_i unchanged for $i \in [1:m]$:

$$\mathbf{a}_i = \begin{cases} [\mathbf{c}_i^\top, -\mathbf{c}_i^\top, \mathbf{o}^\top]^\top \in \mathbb{R}^{4n}, & \text{if } i \in [1:m], \\ \mathbf{e}_{i-m+2n} + \mathbf{e}_{i-m+3n} \in \mathbb{R}^{4n}, & \text{if } i \in [m+1:m+n]. \end{cases}$$

Furthermore, we use the graph $\mathcal{G}_F = (V, \mathcal{E})$ with

$$V = [1:4n] \text{ and edge set } \mathcal{E} = \{\{i+2n, j\} : i \in [1:n], j \in \{i, i+n\}\}$$

to express the bilinear constraints $\mathbf{u}^\top \mathbf{v} + \mathbf{u}^\top \mathbf{w} = 0$ modelling $u_i |y_i| = 0$.

The first $2n$ variables $(\mathbf{v}, \mathbf{w})$ are not necessarily bounded but nevertheless are involved in the complementarity constraint $\mathbf{u}^\top (\mathbf{v} + \mathbf{w}) = 0$, so the classical key condition (K) fails. Fortunately, by construction we find a vertex cover $C := [2n+1:3n]$ of $\mathcal{G}_F$ such that any $x_j = u_{j-2n} \leq 1$ for all $j \in C$. Moreover, for any feasible solution $\hat{\mathbf{y}} \in \mathbb{R}^n$ to the original QP (i.e., such that $\mathbf{C}\hat{\mathbf{y}} = \mathbf{b}$),

$$\hat{\mathbf{x}}^\top = [(\hat{y}_1)^+, \dots, (\hat{y}_n)^+, (\hat{y}_1)^-, \dots, (\hat{y}_n)^-, \mathbf{o}^\top, \mathbf{e}^\top]$$

satisfies $\hat{\mathbf{x}}_C = \mathbf{o}$, so that condition (R) is met if the original QP is feasible. Lemma 4 can be applied, and yields $L_\infty \subseteq F$ for this class of problems. Note that [BMP16, Thm.4] requires a priori knowledge of copositivity of the Hessian of the objective while our approach yields it automatically only by property of the constraints alone, and the fact that the optimal objective value is finite.

Therefore the conic relaxation is exact whenever the QP is feasible and bounded, which automatically would be the case for positive-semidefinite $\mathbf{Q}_0$, given that the linear system $\mathbf{c}_i^\top \mathbf{y} = b_i$, $i \in [1:m]$, has a solution. Furthermore, the slackness constraints $\mathbf{u} + \mathbf{s} = \mathbf{e}$ provide some rows $\mathbf{a}_i \in \mathbb{R}_+^{4n}$ with $b_i = 1 > 0$, so that also the reduction step as in Section 2.2.3 can be performed.

Unfortunately, the constructed $\mathbf{Q}$ is never strictly L_∞ -copositive even if $\mathbf{Q}_0$ is positive-definite, hence we cannot establish strong duality along the general lines in Section 2.2.3. Note that in the present case, by construction of $\mathbf{A}$,

$$L_\infty = \left\{ [\mathbf{v}^\top, \mathbf{w}^\top, \mathbf{o}^\top]^\top \in \mathbb{R}_+^{4n} : \mathbf{c}_i^\top (\mathbf{v} - \mathbf{w}) = 0, i \in [1:m] \right\}, \quad (2.14)$$

which includes all vectors of the form $[\mathbf{v}^\top, \mathbf{v}^\top, \mathbf{o}^\top]^\top \in \mathbb{R}_+^{4n}$ for which any slack matrix $\mathbf{S}(\mathbf{r}, \mathbf{t}, h)$ gives a zero quadratic form.

However, even if $\mathbf{Q}_0$ is merely positive-semidefinite, adding an ℓ^2 term $\tau \sum_{i=1}^n (v_i^2 + w_i^2)$ with small $\tau > 0$ to the objective provides a regularization which ensures strong conic duality. See Section 2.3.3 for more details.

Sparsity in QPs with sign constraints on variables

Additional linear inequality constraints can also be incorporated by introducing additional slack variables. We leave the obvious adaptations to the interested readers and concentrate for ease of presentation on the following model:

$$\min_{y \in \mathbb{R}_+^n} \left\{ y^\top Q_0 y + 2c_0^\top y + \rho \|y\|_0 : c_i^\top y = b_i, i \in [1:m] \right\}. \quad (P_S^+)$$

Here we need not split $y = v - w$ into two nonnegative vectors any more, but we will keep u and the slack variables as in the previous subsection. So, here $x^\top = [y^\top, u^\top, s^\top]$ with $x \in \mathbb{R}_+^{3n}$ now, and we end up with a QPCC formulation very similar to (P_S) , but with reduced dimensionality $3n$ and smaller primitive data: $x^\top Qx + 2c^\top x = y^\top Q_0 y + 2c_0^\top y$ for

$$Q = \begin{pmatrix} Q_0 & 0 \\ 0 & 0 \end{pmatrix} \in \mathcal{S}^{3n} \quad \text{and} \quad c = [c_0^\top, -\frac{\rho}{2}e^\top, o^\top]^\top \in \mathbb{R}^{3n},$$

The linear constraints read now, putting $b_i = 1$ for all $i \in [m+1:m+n]$ while keeping b_i unchanged for all $i \in [1:m]$:

$$a_i = \begin{cases} [c_i^\top, o^\top]^\top \in \mathbb{R}^{3n}, & \text{if } i \in [1:m], \\ e_{i-m+n} + e_{i-m+2n} \in \mathbb{R}^{3n}, & \text{if } i \in [m+1:m+n]. \end{cases}$$

Furthermore, we use the edge set $\mathcal{E} = \{\{i, i+n\} : i \in [1:n]\}$ to express the bilinear constraints $u^\top y = 0$ modelling $u_i y_i = 0$. Here $C = [n+1:2n]$ is a vertex cover for $\mathcal{G}_F$. Again, similar to the case with no sign constraints, Lemma 4 can be applied and yields $L_\infty \subseteq F$ also for this problem class. Indeed, for any feasible $\hat{y}$ here $\hat{x} := [\hat{y}, o^\top, e^\top]^\top \in L$ satisfies $\hat{x}_C = o$.

The two resulting, already simplified conic reformulation problems read exactly as in (CP) and (CD) for their duals, with n replaced by $3n$ and m replaced by $m+n$.

Apart from employing a conic representation with reduced dimensionality, we also can infer strong duality for the conic primal-dual pair by imposing conditions which in a realistic setting are met more frequently, compared to the setting with no sign constraints on the primitive variables. Indeed, as explained in the remark after the following proof, if the primitive polyhedron is bounded, then the strict copositivity condition holds by default.

Theorem 9. *In model (P_S^+) , assume that Q_0 is strictly L_∞^C -copositive with*

$$L_\infty^C := \{y \in \mathbb{R}_+^n : Cy = o\} \quad \text{where } C^\top = [c_1, \dots, c_m],$$

i.e., $y^\top Q_0 y > 0$ for all $y \in \mathbb{R}_+^n \setminus \{o\}$ with $Cy = o$. Then

- *any feasible problem instance is bounded below;*
- *the conic representation is exact;*
- *the primal-dual conic pair has zero duality gap.*

Proof. Observe that in our case $L_\infty = L_\infty^C \times \{o\} \subset \mathbb{R}_+^{3n}$, so that $z \in L_\infty \setminus \{o\}$ if and only if $z^\top = [y^\top, o^\top]$ with $y \in L_\infty^C \setminus \{o\}$, and then $z^\top Qz = y^\top Q_0 y > 0$ by assumption. The claims now follows by application of Theorem 7. $\square$

Note that the assumption that Q_0 be strictly L_∞^C -copositive is trivially satisfied if the primitive feasible polyhedron $M = \{y \in \mathbb{R}_+^n : Cy = b\}$ is bounded, because then $L_\infty^C = \{o\}$.

2.3.2. Hard sparsity constraints

Sometimes we want to force the solution of a quadratic problem to be k -sparse, where the integer $k > 0$ is a hard parameter for sparsity. To this end, we consider another sparsity model for constrained quadratic problems:

$$\min_{y \in \mathbb{R}^n} \left\{ y^\top Q_0 y + 2c_0^\top y : c_i^\top y = b_i, i \in [1:m], \quad \|y\|_0 \leq k \right\}. \quad (P_{H_0})$$

Instead of using the soft parameter ρ to pursuit sparsity, alternatively, we employ $\|\cdot\|_0$ directly in a hard constraint. In analogy to the formulation of the soft sparsity problem (P_{S_0}) , we can also exactly express the hard sparsity problem (P_{H_0}) as follows:

$$\begin{aligned} \min_{\{v, w, u\} \in \mathbb{R}_+^n} \quad & (v - w)^\top Q_0 (v - w) + 2c_0^\top (v - w) \\ \text{s.t.} \quad & c_i^\top (v - w) = b_i, \quad i \in [1:m], \\ & u^\top (v + w) = 0, \\ & -e^\top u \leq k - n, \\ & u \leq e. \end{aligned} \quad (P_{H'})$$

Theorem 10. *Problems (P_{H_0}) and $(P_{H'})$ are equivalent.*

Proof. The proof was presented in [BKS16] and is very similar to that of Theorem 8. $\square$

Similarly, with the help of slack variables $s := e - u \in \mathbb{R}_+^n$ and $\sigma := e^\top u + k - n \geq 0$ and after stacking all variables to $x := [v^\top, w^\top, u^\top, s^\top, \sigma]^\top \in \mathbb{R}_+^{4n+1}$, the problem $(P_{H'})$ can be rewritten:

$$\begin{aligned} \min_{x \in \mathbb{R}_+^{4n+1}} \quad & x^\top Q x + 2c^\top x \\ \text{s.t.} \quad & a_i^\top x = b_i, \quad i \in [1:m+n+1], \\ & \sum_{(i,j) \in \mathcal{E}} x_i x_j = 0, \end{aligned} \quad (P_H)$$

where

$$Q = \begin{pmatrix} Q_0 & -Q_0 & 0 \\ -Q_0 & Q_0 & 0 \\ 0 & 0 & 0 \end{pmatrix} \in \mathcal{S}^{4n+1} \quad \text{and} \quad c = [c_0^\top, -c_0^\top, o^\top]^\top \in \mathbb{R}^{4n+1},$$

as well as the linear constraints, putting $b_i = 1$ for all $i \in [m+1:m+n]$ and $b_i = k - n$ for $i = m+n+1$:

$$a_i = \begin{cases} [c_i^\top, -c_i^\top, o^\top]^\top \in \mathbb{R}^{4n+1}, & \text{if } i \in [1:m], \\ e_{i-m+2n} + e_{i-m+3n} \in \mathbb{R}^{4n+1}, & \text{if } i \in [m+1:m+n], \\ [o^\top, o^\top, -e^\top, o^\top, 1]^\top \in \mathbb{R}^{4n+1}, & \text{if } i = m+n+1. \end{cases}$$

Similarly to the soft sparsity case, the model with sign constraints on the variables can also be treated, leading to

$$\min_{y \in \mathbb{R}_+^n} \left\{ y^\top Q_0 y + 2c_0^\top y : c_i^\top y = b_i, i \in [1:m], \quad \|y\|_0 \leq k \right\}. \quad (P_H^+)$$

Again, we need not split $y = v - w$ into two nonnegative vectors anymore, but we will keep u and the slack variables as in the previous subsection. So, here $x^\top = [y^\top, u^\top, s^\top, \sigma]$ with $x \in \mathbb{R}_+^{3n+1}$ now, and again $x^\top Qx + 2c^\top x = y^\top Q_0 y + 2c_0^\top y$. The QPCC formulation is similar to (P_H) , using the same

$$Q = \begin{pmatrix} Q_0 & O \\ O & O \end{pmatrix} \in \mathcal{S}^{3n+1} \quad \text{and} \quad c = [c_0^\top, o^\top]^\top \in \mathbb{R}^{3n+1},$$

adding a linear constraint with $b_{m+n+1} = k - n$ while keeping b_i unchanged for all $i \in [1:m+n]$:

$$a_i = \begin{cases} [c_i^\top, o^\top]^\top \in \mathbb{R}^{3n+1}, & \text{if } i \in [1:m], \\ e_{i-m+n} + e_{i-m+2n} \in \mathbb{R}^{3n+1}, & \text{if } i \in [m+1:m+n] \\ [o^\top, -e^\top, o^\top, 1]^\top \in \mathbb{R}^{3n+1}, & \text{if } i = m+n+1. \end{cases}$$

The conic reformulation problems read exactly as in (CP) and (CD) for their duals, with n replaced by $3n+1$ and m replaced by $m+n+1$.

Theorem 11. *In model (P_H^+) , assume that Q_0 is strictly L_∞^C -copositive with*

$$L_\infty^C := \{y \in \mathbb{R}_+^n : Cy = o\} \quad \text{where } C^\top = [c_1, \dots, c_m],$$

i.e., $y^\top Q_0 y > 0$ for all $y \in \mathbb{R}_+^n \setminus \{o\}$ with $Cy = o$. Then

- *any feasible problem instance is bounded below;*
- *the conic representation is exact;*
- *the primal-dual conic pair has zero duality gap.*

Proof. Employ the same arguments as for Theorem 9. □

Again, the assumption that Q_0 be strictly L_∞^C -copositive is trivially satisfied if the primitive feasible polyhedron $M = \{y \in \mathbb{R}_+^n : Cy = b\}$ is bounded, because then $L_\infty^C = \{o\}$. It is also satisfied if Q_0 is positive-definite as in the following section which presents an application that motivated this study.

2.3.3. Sparse linear regression: conic reformulation

In a linear least-squares/regression context, $Q_0 = A_0^\top A_0$ with A_0 an $m \times n$ matrix with $m < n$, so Q_0 is always positive-semidefinite but will have some zero eigenvalues. Hence strict copositivity of Q_0 cannot be ensured by spectral properties of Q_0 alone; one possibility is the condition that A_0 contains some strictly positive rows which renders $Q_0 = P + N$ with $P \in \mathcal{S}_+^n$ and $N \in \mathcal{N}^n$. This is enough to ensure strict $\mathbb{R}_+^n$ -copositivity of Q_0 which establishes all favorable properties of our conic reformulation by application of Theorems 9 and 11. Another class of instances is that with bounded primitive feasible polyhedron $M = \{y \in \mathbb{R}_+^n : Cy = b\}$, as explained in the note after these results.

In absence of sign constraints, another option already hinted at in Section 2.3.1 would consist in adding a regularization term $\tau \sum_{i=1}^n (v_i^2 + w_i^2)$, with small $\tau > 0$:

Proposition 12. *Consider a regularized sparse linear regression problem with a small $\tau > 0$ where, in the notation of (P_{S_0}) or (P_{H_0}) ,*

$$y^\top Q_0 y = \|A_0 y\|^2 + \tau \|y\|^2.$$

Then

- *any feasible problem instance is bounded below;*
- *the conic representation is exact;*
- *the primal-dual conic pair has zero duality gap.*

Proof. By reformulation as before, we only need to add $\tau (\|\mathbf{v}\|^2 + \|\mathbf{w}\|^2)$ to the original objective of (P_{S_0}) or (P_{H_0}) instead of $\tau \|\mathbf{y}\|^2 = \tau \|\mathbf{v} - \mathbf{w}\|^2$; indeed, in the optimum we have $\|\mathbf{v}\|^2 + \|\mathbf{w}\|^2 = \|\mathbf{v} - \mathbf{w}\|^2$ as optimality already implies $v_i = y_i^+$ and $w_i = y_i^-$ and therefore $\mathbf{v}^\top \mathbf{w} = 0$, as argued before. Thus we can regularize by a positive-definite term in $(\mathbf{v}, \mathbf{w})$ which renders $\mathbf{Q}$ strictly L_∞ -copositive where L_∞ is given as in (2.14) (replace $4n$ with $4n + 1$ when dealing with hard sparsity constraints). The assertions then follow by application of Theorems 6 and 7. $\square$

Thus, for a large class of sparse regression problems, be they constrained or not, our theory applies in full strength. Apart from that, even if the last line of arguments cannot be carried through for all instances with no sign constraints on the variables, we still can use weak conic duality to obtain strong dual bounds, as our experiments show.

2.4. An iterative algorithm for dual problems (CD)

For the convenience of expression, we write (CD) in the following form:

$$\max_{\mathbf{x} \in \mathbb{R}^m} \left\{ \mathbf{a}^\top \mathbf{x} : g(\mathbf{x}) \in \mathcal{COP}^n \right\}, \quad (2.15)$$

where $g : \mathbb{R}^m \rightarrow \mathcal{S}^n$ is a linear-affine function in $\mathbf{x}$ and $\mathbf{a} \in \mathbb{R}^m$. By convexity of $\mathcal{COP}^n$, it is obvious that the set $\{\mathbf{x} : g(\mathbf{x}) \in \mathcal{COP}^n\}$ is convex too. Therefore it is possible to solve the convex problem (2.15) via a directional search scheme, provided a feasible initial guess. Recall that any feasible solution to (2.15) already provides a rigorous lower bound $\mathbf{a}^\top \mathbf{x}$ for (P).

Now we introduce a local search algorithm for (2.15) where a subscript i denotes the i -th entry of a vector and a superscript k stands for the iteration counter.

Data: a problem of the form (2.15), a feasible initial guess $\mathbf{x}^0 \in \mathbb{R}^m$, adaptive parameter $0 < \eta < 1$, stopping tolerance $\varepsilon > 0$, initial step size $\lambda^0 \geq \varepsilon \mathbf{e}$

Result: a solution $\mathbf{x}^*$ to (2.15)

```

while  $\|\lambda^k\|_\infty \geq \varepsilon$  do
    Generate a permutation  $\mathbf{p}^k \in \mathbb{R}^m$  of set  $[1:d]$  such that the following inequalities hold:
     $\lambda_{p_i}^k \text{sign}(a_{p_i}) \leq \lambda_{p_j}^k \text{sign}(a_{p_j})$  if  $i < j$ 
     $\mathbf{y}^0 \leftarrow \mathbf{x}^k, \quad \beta^0 \leftarrow \lambda^k, \quad \mathbf{p} \leftarrow \mathbf{p}^k$ 
    while  $i \in [1:m]$  do
         $\bar{\mathbf{y}}^i \leftarrow \mathbf{y}^{i-1}, \quad \beta^i \leftarrow \beta^{i-1}$ 
         $\bar{y}_{p_i}^i \leftarrow y_{p_i}^{i-1} + \beta_{p_i}^{i-1} a_{p_i}$ 
        if  $g(\bar{\mathbf{y}}^i) \in \mathcal{COP}^n$  then
             $\mathbf{y}^i \leftarrow \bar{\mathbf{y}}^i, \quad \beta_{p_i}^i \leftarrow \frac{1}{\eta} \beta_{p_i}^{i-1}$ 
        else
             $\mathbf{y}^i \leftarrow \mathbf{y}^{i-1}, \quad \beta_{p_i}^i \leftarrow \eta \beta_{p_i}^{i-1}$ 
        end
         $i \leftarrow i + 1$ 
    end
     $\mathbf{x}^{k+1} \leftarrow \mathbf{y}^m, \quad \lambda^{k+1} \leftarrow \beta^m, \quad k \leftarrow k + 1$ 
end

```

Algorithm 1: Greedy coordinate descent algorithm

As for the copositivity detection involved in the Algorithm 1, one of possible methods is [Ans21], which amounts to solve the following MILP:

$$\begin{aligned} \max_{\mathbf{x}, \mathbf{y}} \quad & \gamma \\ \text{s.t.} \quad & \mathbf{s}_i^\top \mathbf{x} \leq -\gamma + m_i(1 - y_i), \quad i \in [1:n], \\ & \gamma \geq 0, \mathbf{0} \leq \mathbf{x} \leq \mathbf{y}, \\ & \mathbf{y} \in \{0, 1\}^n, \mathbf{e}^\top \mathbf{y} \geq 1, \end{aligned} \quad (2.16)$$

where $\mathbf{s}_i$ is the i -th column of targeted matrix $\mathbf{S}$ and the parameter m_i is typically chosen $m_i = 1 + \sum_{i \neq j} \{s_{ij} : s_{ij} > 0\}$ as in [Ans21].

Theorem 13. [Ans21] Let $\mathbf{S} \in \mathcal{S}^n$. Then $\mathbf{S} \in \mathcal{COP}^n$ if and only if the optimal value of (2.16) is zero.

2.5. Conclusion

In this chapter, we propose an exact completely copositive reformulation of QPCCs under mild conditions on the constraints, which extend previously known ones. An application on pursuing sparse solutions of (nonconvex) quadratic optimization is shown to satisfy our proposed conditions and therefore an equivalent completely positive optimization problem is obtained. Moreover, the condition for strong conic duality is also studied. To solve the conic optimization problem from the dual side, we introduce a direct search method with embedded MILP-based copositivity detection.

3. CP reformulation for sparse QPs

In this chapter, we focus on the sparse quadratic optimization problems and propose novel and more compact completely positive reformulations for them. This part is based on a nearly complete joint work with Immanuel Bomze and Markus Gabl, which hopefully will be ready for submission soon.

3.1. The problem

In this chapter, we aim to obtain sparse solutions to quadratic optimization problems (QPs). We consider the QP in standard form:

$$\begin{aligned} z^* &:= \min_{\mathbf{x} \in \mathbb{R}_+^n} \mathbf{x}^\top \mathbf{Q} \mathbf{x} + 2\mathbf{c}^\top \mathbf{x} \\ \text{s.t.} \quad & \mathbf{A} \mathbf{x} = \mathbf{b}, \end{aligned} \tag{QP}$$

where $\mathbf{Q}$ is a (possibly indefinite) symmetry real matrix and the linear operator $\mathbf{A} : \mathbb{R}^n \rightarrow \mathbb{R}^m$, represented by an $m \times n$ matrix. Throughout the chapter, we assume the (non-sparse) problem (QP) is feasible and bounded below,

$$-\infty < \min \left\{ \mathbf{x}^\top \mathbf{Q} \mathbf{x} + 2\mathbf{c}^\top \mathbf{x} : \mathbf{A} \mathbf{x} = \mathbf{b}, \mathbf{x} \in \mathbb{R}_+^n \right\} < +\infty. \tag{3.1}$$

Incorporating the ℓ_0 -pseudonorm

$$\|\mathbf{x}\|_0 := |\{i \in [1:n] : x_i \neq 0\}|,$$

term into an optimization problem, one can obtain sparse solutions to this problem.

Generally, there are two ways to exploit the ℓ_0 -norm to boost the sparsity of the optimal solution. By adding the term to the objective, we consider the problem of the form, which is referred to as soft sparse QP:

$$\begin{aligned} \min_{\mathbf{x} \in \mathbb{R}_+^n} \quad & \mathbf{x}^\top \mathbf{Q} \mathbf{x} + 2\mathbf{c}^\top \mathbf{x} + \gamma \|\mathbf{x}\|_0 \\ \text{s.t.} \quad & \mathbf{A} \mathbf{x} = \mathbf{b}, \end{aligned} \tag{Soft-QP}$$

where $\gamma > 0$ is a trade-off parameter to balance between the sparsity of the optimal solution and the optimality of the original QP. It is worth noticing that the optimal solution of (Soft-QP) is in fact always attained. Indeed, considering all possible supports of $\mathbf{x}^*$, the problem (Soft-QP) reduces to a finite collection of 2^n QPs (some of them may be infeasible). To this end, the optimal solution to the problem (Soft-QP) can be interpreted as the best choice from a candidate pool consisting of optimal solutions to a finite collection of QPs and therefore the attainability of the optimal solution follows from the celebrated Frank/Wolfe Theorem [FW56].

Alternatively, we can also limit the number of nonzero elements of variables directly by a so-called cardinality constraint and arrive at the problem, which is referred to as hard sparse QP:

$$\begin{aligned} \min_{\mathbf{x} \in \mathbb{R}_+^n} \quad & \mathbf{x}^\top \mathbf{Q} \mathbf{x} + 2\mathbf{c}^\top \mathbf{x} \\ \text{s.t.} \quad & \mathbf{A} \mathbf{x} = \mathbf{b} \\ & \|\mathbf{x}\|_0 \leq \rho, \end{aligned} \tag{Hard-QP}$$

where $\rho \in [0, n]$ is a user-defined parameter.¹ Note that this problem is feasible if and only if $\rho \geq \underline{\rho} := \min_{\mathbf{x} \in \mathbb{R}_+^n} \{\|\mathbf{x}\|_0 : \mathbf{A} \mathbf{x} = \mathbf{b}\}$, see the discussion before Proposition 21 below. Then the optimal solution of (Hard-QP) is attainable for the same reason as in the soft sparse QP case, for a finite collection of classical QPs.

It is known that both (Soft-QP) and (Hard-QP) are in the NP-hard problem class even if the term $\|\cdot\|_0$ is absent. In this paper, we develop the completely positive (CP) technique following important works [Bur09] and propose a novel, hopefully more efficient and effective, CP reformulation of these two problems. CP reformulation is a convexification technique for some particular subsets of quadratic-constrained quadratic optimization problems (QCQPs), packaging the nonconvexity of the original problem into the completely positive cone (or its dual cone: copositive cone), and ends up with a linear conic optimization problem. For more excellent works and literature in this field, we refer interested readers to [Bom12, BDDK⁺00, Bur15, DR21, PVZ15] and references therein.

3.2. CP reformulation for soft sparse QPs

In this section, we propose two compact CP relaxations of (Soft-QP) and prove the exactness of the relaxations.

3.2.1. Burer's representation theorem and its extension

In this subsection, we review Burer's CP reformulation for quadratic optimization problems with complementarity constraints (QPCCs), which is mainly based on the seminal paper [Bur09]. For two vectors $\mathbf{a} \in \mathbb{R}^n$ and $\mathbf{b} \in \mathbb{R}^n$, we denote their Hadamard product by $\mathbf{a} \circ \mathbf{b} := [a_i b_i]_{i \in [1:n]} \in \mathbb{R}^n$.

First of all, we consider an optimization problem of the form (QP), and linearize the objective function by replacing $\mathbf{x} \mathbf{x}^\top$ by $\mathbf{X}$. To this end, we introduced a new variable $\mathbf{X}$ and the objective function becomes a linear function in $\mathbf{X}$ and $\mathbf{x}$, e.g., $\langle \mathbf{X}, \mathbf{Q} \rangle + 2\mathbf{c}^\top \mathbf{x}$, which is known as the lifting technique. It is worth noticing that the lifted objective is always a lower bound of the original objective, as for any feasible point $\mathbf{x}$, the point $(\mathbf{x}, \mathbf{x} \mathbf{x}^\top)$ is also feasible to the lifted problem and gives the same objective value. Hence, in order to have the lifted problem as tight as the original problem, some extra constraints are usually posed to the variable $(\mathbf{x}, \mathbf{X})$. Now consider the following lifted set, to shed more light on the effect of squaring a linear constraint:

Lemma 14. *If the point $(\mathbf{x}, \mathbf{X})$ is an element of the set*

$$\left\{ (\mathbf{x}, \mathbf{X}) \in \mathbb{R}_+^n \times \mathcal{S}^n : \begin{array}{l} \mathbf{A} \mathbf{x} = \mathbf{b} \\ \text{diag}(\mathbf{X} \mathbf{A} \mathbf{X}^\top) = \mathbf{b} \circ \mathbf{b} \end{array}, \begin{pmatrix} 1 & \mathbf{x}^\top \\ \mathbf{x} & \mathbf{X} \end{pmatrix} \in \mathcal{CP}^{n+1} \right\}, \tag{3.2}$$

¹ not necessarily integer.

it has the following decomposition:

$$\begin{pmatrix} 1 & \mathbf{x}^\top \\ \mathbf{x} & \mathbf{X} \end{pmatrix} = \sum_{k \in K_+} \beta_k \begin{pmatrix} 1 \\ \frac{\mathbf{x}^k}{\lambda_k} \end{pmatrix} \begin{pmatrix} 1 \\ \frac{\mathbf{x}^k}{\lambda_k} \end{pmatrix}^\top + \sum_{k \in K_0} \begin{pmatrix} 0 \\ \mathbf{x}^k \end{pmatrix} \begin{pmatrix} 0 \\ \mathbf{x}^k \end{pmatrix}^\top,$$

where $\beta_k = \lambda_k^2 > 0$ for all $k \in K_+$ with $\sum_{k \in K_+} \beta_k = 1$, $\mathbf{x}^k \in \mathbb{R}_+^n$ for all $k \in K_+ \cup K_0$, and the following equations hold:

$$\frac{\mathbf{x}^k}{\lambda_k} \in L := \{\mathbf{x} \in \mathbb{R}_+^n : \mathbf{A}\mathbf{x} = \mathbf{b}\}, \quad \forall k \in K_+,$$

and

$$\mathbf{x}^k \in L_0 := \{\mathbf{x} \in \mathbb{R}_+^n : \mathbf{A}\mathbf{x} = \mathbf{o}\}, \quad \forall k \in K_0.$$

Proof. See in the paper [Bur09]. □

Note that the lifted variable $\mathbf{X}$ consists of two parts, namely,

$$\mathbf{X} = \sum_{k \in K_+} \beta_k \frac{\mathbf{x}^k}{\lambda_k} \left(\frac{\mathbf{x}^k}{\lambda_k} \right)^\top + \sum_{k \in K_0} \mathbf{x}^k (\mathbf{x}^k)^\top,$$

where $\mathbf{x}^k$, β_k , and λ_k satisfy the property presented in Lemma 14. Now we consider the objective in the lifted problem:

$$\begin{aligned} & \langle \mathbf{Q}, \mathbf{X} \rangle + 2\mathbf{c}^\top \mathbf{x} \\ &= \sum_{k \in K_+} \beta_k \left(\frac{\mathbf{x}^k}{\lambda_k} \right)^\top \mathbf{Q} \frac{\mathbf{x}^k}{\lambda_k} + \sum_{k \in K_0} (\mathbf{x}^k)^\top \mathbf{Q} \mathbf{x}^k + 2\mathbf{c}^\top \mathbf{x} \\ &\geq z^* + \sum_{k \in K_0} (\mathbf{x}^k)^\top \mathbf{Q} \mathbf{x}^k, \end{aligned}$$

where z^* is the optimal value to the (QP).

We notice that if the quadratic form $\mathbf{x}^\top \mathbf{Q} \mathbf{x}$ takes no negative values over the set $\mathbf{x} \in L_0$, then at least one optimal solution in the above lifted set (3.2) can be expressed as the convex hull of the set of dyadic matrices $\mathbf{x}^k (\mathbf{x}^k)^\top$ based upon elements $\mathbf{x}^k$ from the feasible region of (QP), which leads to an exact CP relaxation.

Following a similar idea, Burer proposed a so-called key condition² ensuring the exactness of CP relaxation for QPCC. Here is the form of QPCC considered in [Bur09]:

$$\begin{aligned} & \min_{\mathbf{x} \in \mathbb{R}_+^n} \quad \mathbf{x}^\top \mathbf{Q} \mathbf{x} + 2\mathbf{c}^\top \mathbf{x} \\ & \text{s.t.} \quad \mathbf{A}\mathbf{x} = \mathbf{b}, \\ & \quad \quad x_i x_j = 0, \quad \forall (i, j) \in B_1 \times B_2 \subseteq [1:n] \times [1:n]. \end{aligned} \tag{QPCC}$$

Under the key condition:

$$\mathbf{x} \in L \implies \exists u \geq 0 \text{ such that } 0 \leq x_j \leq u, \quad \forall j \in B_1 \cup B_2,$$

²this condition ensuring exactness of the CP relaxation for QPCCs can be significantly relaxed [BP23].

(QPCC) is equivalent to the following CP relaxation:

$$\begin{aligned}
\min_{\mathbf{x} \in \mathbb{R}_+^n} \quad & \langle \mathbf{Q}, \mathbf{X} \rangle + 2\mathbf{c}^\top \mathbf{x} \\
\text{s.t.} \quad & \mathbf{A}\mathbf{x} = \mathbf{b}, \\
& \text{diag}(\mathbf{A}\mathbf{X}\mathbf{A}^\top) = \mathbf{b} \circ \mathbf{b}, \\
& X_{ij} = 0, \quad \forall (i, j) \in B_1 \times B_2, \\
& \begin{pmatrix} 1 & \mathbf{x}^\top \\ \mathbf{x} & \mathbf{X} \end{pmatrix} \in \mathcal{CP}^{n+1}.
\end{aligned}$$

In the next sections, we will see that both soft sparse QP and hard sparse QP can be expressed in the form of (QPCC) and therefore Burer's CP relaxation results are applicable. However, the key condition is not necessarily satisfied if the feasible region is unbounded. Even for the bounded (binary) variables, we need to introduce the slack variables to indicate the boundedness from the perspective of set L , which leads to an increase in the dimension of problem. To overcome this difficulty, we propose more compact CP relaxations and prove their exactness for both soft and hard sparse QP cases in a direct way.

3.2.2. New and more compact CP reformulation for soft sparse QPs

In this subsection, we propose new and compact CP reformulations for soft sparse QPs.

Recall that the soft sparse QP is of the form:

$$\begin{aligned}
\min_{\mathbf{x} \in \mathbb{R}_+^n} \quad & \mathbf{x}^\top \mathbf{Q}\mathbf{x} + 2\mathbf{c}^\top \mathbf{x} + \gamma \|\mathbf{x}\|_0 \\
\text{s.t.} \quad & \mathbf{A}\mathbf{x} = \mathbf{b}.
\end{aligned} \tag{Soft-QP}$$

To express the term $\|\mathbf{x}\|_0$ in a more tractable way, we introduce auxiliary binary variables and employ a complementarity constraint, arriving at a mixed-integer reformulation of (Soft-QP):

$$\begin{aligned}
\zeta_\gamma^* := \min_{\mathbf{x} \in \mathbb{R}_+^n} \quad & \mathbf{x}^\top \mathbf{Q}\mathbf{x} + 2\mathbf{c}^\top \mathbf{x} - \gamma \mathbf{e}^\top \mathbf{u} \\
\text{s.t.} \quad & \mathbf{A}\mathbf{x} = \mathbf{b}, \\
& \mathbf{x}^\top \mathbf{u} = 0, \\
& \mathbf{u} \in \{0, 1\}^n,
\end{aligned} \tag{MIQP-soft}(\gamma)$$

where $\mathbf{u}$ serves as an indicator variable. It is clear that, if x_i is positive for some i , the corresponding u_i has to be 0 due to the complementarity constraint, whereas, u_i has to be 1 if x_i equals 0 since we maximize u_i in the objective. To this end, the sparse term $\|\mathbf{x}\|_0$ is captured by $n - \mathbf{e}^\top \mathbf{u}$. After eliminating the constant term in the objective, we arrive at (MIQP-soft(γ)).

Note that the resulting problem is seemingly difficult due to the presence of two types of non-convexity, coming from the binarity and the complementarity of constraints. The most common way in the current literature to proceed with the problem is to linearize the complementarity constraint, involving bounds (big-M) on all variables. However, the consequence of the above procedure can be highly dependent on the estimation of big-M, which is impractical in many real-world applications.

Inspired by the previous work [Bur09, BKS16, BP23], we propose the following CP relaxation of the problem (MIQP-soft(γ)):

$$\begin{aligned}
z_\gamma^* = \inf \quad & z_\gamma(\mathbf{Y}) := \langle \mathbf{Q}, \mathbf{Y}^{\mathbf{x}\mathbf{x}} \rangle + 2\mathbf{c}^\top \mathbf{x} - \gamma \mathbf{e}^\top \mathbf{u} \\
\text{s.t.} \quad & \mathbf{A}\mathbf{x} = \mathbf{b}, \\
& \text{diag}(\mathbf{A}\mathbf{Y}^{\mathbf{x}\mathbf{x}}\mathbf{A}^\top) = \mathbf{b} \circ \mathbf{b}, \\
& \mathbf{e}^\top \text{diag}(\mathbf{Y}^{\mathbf{x}\mathbf{u}}) = 0, \\
& \text{diag}(\mathbf{Y}^{\mathbf{u}\mathbf{u}}) = \mathbf{u}, \\
& \mathbf{Y} := \begin{pmatrix} 1 & \mathbf{x}^\top & \mathbf{u}^\top \\ \mathbf{x} & \mathbf{Y}^{\mathbf{x}\mathbf{x}} & \mathbf{Y}^{\mathbf{x}\mathbf{u}} \\ \mathbf{u} & (\mathbf{Y}^{\mathbf{x}\mathbf{u}})^\top & \mathbf{Y}^{\mathbf{u}\mathbf{u}} \end{pmatrix} \in \mathcal{CP}^{2n+1}.
\end{aligned} \tag{CP-soft(γ)}$$

In the following, we will show that the problems (MIQP-soft(γ)) and (CP-soft(γ)) are equivalent. For the sake of compact expression, we abbreviate the above matrix relation by $\mathbf{Y} = \mathcal{Y}(\mathbf{x}, \mathbf{Y}^{\mathbf{x}\mathbf{x}}, \mathbf{u}, \mathbf{Y}^{\mathbf{x}\mathbf{u}}, \mathbf{Y}^{\mathbf{u}\mathbf{u}})$.

Lemma 15. *Given any $\mathbf{p} \in \mathbb{R}_+^d$ with finite d , consider the following maximization problem:*

$$\sup_{\mathbf{v}, s} \left\{ \mathbf{p}^\top \mathbf{v} : \mathbf{p}^\top \mathbf{v} = \|\mathbf{v}\|^2 + s, (s, \mathbf{v}) \in \mathbb{R}_+^{d+1} \right\}.$$

This problem has the unique optimal solution $[\mathbf{v}^{\top}, s^*] = [\mathbf{p}^\top, 0]$.*

Proof. It is clear that over the feasible set we have $\|\mathbf{v}\|^2 = \mathbf{p}^\top \mathbf{v} - s \leq \|\mathbf{p}\| \|\mathbf{v}\| - s$, which implies that $\|\mathbf{v}\| \leq \|\mathbf{p}\|$ and the objective function $\mathbf{p}^\top \mathbf{v} \leq \|\mathbf{p}\| \|\mathbf{v}\| \leq \|\mathbf{p}\|^2$. On the other hand, it is clear that $\mathbf{v}^* = \mathbf{p}$ is the only feasible point that attains the upper bound, namely, $\mathbf{p}^\top \mathbf{v}^* = \|\mathbf{p}\|^2$. In this case, $s^* = 0$. $\square$

Lemma 16. *For a given (CP-soft(γ))-feasible point $\mathbf{Y}_0$, there exists an improving (CP-soft(γ))-feasible point $\mathbf{Y}$ (i.e., $z_\gamma(\mathbf{Y}) \leq z_\gamma(\mathbf{Y}_0)$ holds) which has the decomposition*

$$\mathbf{Y} = \begin{pmatrix} 1 & \mathbf{x}^\top & \mathbf{u}^\top \\ \mathbf{x} & \mathbf{Y}^{\mathbf{x}\mathbf{x}} & \mathbf{Y}^{\mathbf{x}\mathbf{u}} \\ \mathbf{u} & (\mathbf{Y}^{\mathbf{x}\mathbf{u}})^\top & \mathbf{Y}^{\mathbf{u}\mathbf{u}} \end{pmatrix} = \sum_{k \in K_+} \beta_k \begin{pmatrix} 1 \\ \frac{\mathbf{x}^k}{\lambda_k} \\ \frac{\mathbf{u}^k}{\lambda_k} \end{pmatrix} \begin{pmatrix} 1 \\ \frac{\mathbf{x}^k}{\lambda_k} \\ \frac{\mathbf{u}^k}{\lambda_k} \end{pmatrix}^\top + \sum_{k \in K_0} \begin{pmatrix} 0 \\ \mathbf{x}^k \\ \mathbf{u}^k \end{pmatrix} \begin{pmatrix} 0 \\ \mathbf{x}^k \\ \mathbf{u}^k \end{pmatrix}^\top, \tag{3.3}$$

where $\beta_k = \lambda_k^2 > 0$ for all $k \in K_+$ with $\sum_{k \in K_+} \beta_k = 1$ and $\{\mathbf{x}^k, \mathbf{u}^k\} \subset \mathbb{R}_+^n$ for all $k \in K := K_+ \cup K_0$. Moreover, $\frac{\mathbf{u}^k}{\lambda_k} \in \{0, 1\}^n$ and $\mathbf{A}\mathbf{x}^k = \lambda_k \mathbf{b}$ for all $k \in K_+$, while $\mathbf{u}^k = \mathbf{o}$ and $\mathbf{A}\mathbf{x}^k = \mathbf{o}$ for all $k \in K_0$.

Proof. Let $\mathbf{Y}_0 = \mathcal{Y}(\mathbf{x}_0, \mathbf{Y}_0^{\mathbf{x}\mathbf{x}}, \dots)$ be (CP-soft(γ))-feasible. We construct another (CP-soft(γ))-feasible point $\mathbf{Y} = \mathcal{Y}(\mathbf{x}, \mathbf{Y}^{\mathbf{x}\mathbf{x}}, \mathbf{u}, \mathbf{Y}^{\mathbf{x}\mathbf{u}}, \mathbf{Y}^{\mathbf{u}\mathbf{u}})$ with $\mathbf{x} = \mathbf{x}_0$ and $\mathbf{Y}^{\mathbf{x}\mathbf{x}} = \mathbf{Y}_0^{\mathbf{x}\mathbf{x}}$. Hence $\mathbf{A}\mathbf{x} = \mathbf{b}$ and $\text{diag}(\mathbf{A}\mathbf{Y}^{\mathbf{x}\mathbf{x}}\mathbf{A}^\top) = \mathbf{b} \circ \mathbf{b}$ are granted. It remains to show $z_\gamma(\mathbf{Y}) \leq z_\gamma(\mathbf{Y}_0)$ and the property: $\frac{\mathbf{u}^k}{\lambda_k} \in \{0, 1\}^n$ for all $k \in K_+$ and $\mathbf{u}^k = \mathbf{o}$ for all $k \in K_0$, which entails (CP-soft(γ))-feasibility of $\mathbf{Y}$.

We start with the CP decomposition of the north-west principal submatrix of $\mathbf{Y}$ (which coincides with that of $\mathbf{Y}_0$) as follows, resulting from Lemma 14:

$$\begin{pmatrix} 1 & \mathbf{x}^\top \\ \mathbf{x} & \mathbf{Y}^{\mathbf{x}\mathbf{x}} \end{pmatrix} = \begin{pmatrix} 1 & \mathbf{x}_0^\top \\ \mathbf{x}_0 & \mathbf{Y}_0^{\mathbf{x}\mathbf{x}} \end{pmatrix} = \sum_{k \in K_+} \beta_k \begin{pmatrix} 1 \\ \frac{\mathbf{x}^k}{\lambda_k} \end{pmatrix} \begin{pmatrix} 1 \\ \frac{\mathbf{x}^k}{\lambda_k} \end{pmatrix}^\top + \sum_{k \in K_0} \begin{pmatrix} 0 \\ \mathbf{x}^k \end{pmatrix} \begin{pmatrix} 0 \\ \mathbf{x}^k \end{pmatrix}^\top,$$

where $\mathbf{A}\mathbf{x}^k = \lambda_k \mathbf{b}$ and $\beta_k = \lambda_k^2 > 0$ for all $k \in K_+$ with $\sum_{k \in K_+} \beta_k = 1$ while $\mathbf{A}\mathbf{x}^k = \mathbf{o}$ for all $k \in K_0$ and $\mathbf{x}^k \in \mathbb{R}_+^n$ for all $k \in K := K_+ \cup K_0$. With $(\mathbf{x}, \mathbf{Y}^{\mathbf{x}\mathbf{x}})$ and the above CP decomposition fixed, we aim to find a set of $\mathbf{u}^k \in \mathbb{R}_+^n$ for all $k \in K$ such that the extended point $\mathbf{Y} = \mathcal{Y}(\mathbf{x}, \mathbf{Y}^{\mathbf{x}\mathbf{x}}, \mathbf{u}, \mathbf{Y}^{\mathbf{x}\mathbf{u}}, \mathbf{Y}^{\mathbf{u}\mathbf{u}})$ defined by the CP decomposition (3.3) satisfies the stated property.

Regarding the constraint $\mathbf{e}^\top \text{diag}(\mathbf{Y}^{\mathbf{x}\mathbf{u}}) = 0$, we notice that this is equivalent to stating

$$x_i^k u_i^k = 0 \quad \text{for all } k \in K \text{ and } i \in [1:n], \quad (3.4)$$

where the subscript index i denotes the i -th coordinate of the underlying vector.

Focusing on the i -th coordinate of $\mathbf{x}^k$ across all $k \in K_+$, we split the index set $K_+ = P_i \cup N_i$ into two parts according to positivity. i.e.,

$$P_i = \{k \in K_+ : x_i^k > 0\} \quad \text{and} \quad N_i = K_+ \setminus P_i = \{k \in K_+ : x_i^k = 0\}.$$

To this end, to obtain (3.4), we let

$$u_i^k := 0 \quad \text{for all } i \in [1:n] \quad \text{and all } k \in P_i. \quad (3.5)$$

Moreover, we let the remaining part of $\mathbf{u}^k$ be the optimal solution to the following problem:

$$\begin{aligned} \max_{\mathbf{u}^k \in \mathbb{R}_+^n, k \in K} \quad & \sum_{k \in K_+} \lambda_k \mathbf{e}^\top \mathbf{u}^k \\ \text{s.t.} \quad & \sum_{k \in N_i} (u_i^k)^2 + \sum_{k \in P_i} (u_i^k)^2 + \sum_{k \in K_0} (u_i^k)^2 = \sum_{k \in K_+} \lambda_k u_i^k, \quad \forall i \in [1:n], \\ & x_i^k u_i^k = 0, \quad \forall k \in K, \forall i \in [1:n]. \end{aligned} \quad (3.6)$$

It follows immediately from the objective of $(\text{CP-soft}(\gamma))$ that then $z_\gamma(\mathbf{Y}) \leq z_\gamma(\mathbf{Y}_0)$.

Notice that, in the optimization problem (3.6), there is no constraint involving two variables from different coordinates and the objective function is separable across the coordinates. Hence, the problem (3.6) can be equivalently represented by a set of mini-problems indexed by $i \in [1:n]$:

$$\begin{aligned} \max_{u_i^k \in \mathbb{R}_+, k \in K} \quad & \sum_{k \in K_+} \lambda_k u_i^k \\ \text{s.t.} \quad & \sum_{k \in N_i} (u_i^k)^2 + \sum_{k \in P_i} (u_i^k)^2 + \sum_{k \in K_0} (u_i^k)^2 = \sum_{k \in K_+} \lambda_k u_i^k, \\ & x_i^k u_i^k = 0, \quad \forall k \in K. \end{aligned}$$

Taking the fact (3.5) into account, we arrive at the following problems across all $i \in [1:n]$:

$$\begin{aligned} \max_{s_i, u_i^k \in \mathbb{R}_+, k \in N_i} \quad & \sum_{k \in N_i} \lambda_k u_i^k \\ \text{s.t.} \quad & \sum_{k \in N_i} (u_i^k)^2 + s_i = \sum_{k \in N_i} \lambda_k u_i^k, \end{aligned} \quad (3.7)$$

where $s_i = \sum_{k \in K_0} (u_i^k)^2$.

Applying Lemma 15 to the problem (3.7) with a fixed coordinate i , we specify $\mathbf{p}$ by $[\lambda_k]_{k \in N_i}$ and $\mathbf{v}$ by $[u_i^k]_{k \in N_i}$. It follows immediately that $u_i^k = \lambda_k$ for all $k \in N_i$, which implies that

$$\frac{u_i^k}{\lambda_k} = 1, \quad \forall k \in N_i.$$

Moreover, due to $s_i^* = 0$, we have

$$u_i^k = 0, \quad \forall k \in K_0.$$

Together with the fact (3.5) and applying Lemma 15 to the problems (3.7) across all coordinates $i \in [1:n]$, we complete the proof. $\square$

Theorem 17. *Assume that (3.1) holds. Then, for all $\gamma \in \mathbb{R}$, the conic relaxation (CP-soft(γ)) is equivalent to the problem (MIQP-soft(γ)) in the sense of objective value.*

Proof. By Lemma 14, for any (CP-soft(γ))-feasible $\mathbf{Y}$, we have a CP decomposition

$$\begin{aligned} \mathbf{Y} &= \begin{pmatrix} 1 & \mathbf{x}^\top & \mathbf{u}^\top \\ \mathbf{x} & \mathbf{Y}^{\mathbf{x}\mathbf{x}} & \mathbf{Y}^{\mathbf{x}\mathbf{u}} \\ \mathbf{u} & (\mathbf{Y}^{\mathbf{x}\mathbf{u}})^\top & \mathbf{Y}^{\mathbf{u}\mathbf{u}} \end{pmatrix} = \sum_{k \in K} \begin{pmatrix} \lambda_k \\ \mathbf{x}^k \\ \mathbf{u}^k \end{pmatrix} \begin{pmatrix} \lambda_k \\ \mathbf{x}^k \\ \mathbf{u}^k \end{pmatrix}^\top = \mathbf{M}_1 + \mathbf{M}_2, \quad \text{with} \\ \mathbf{M}_1 &:= \sum_{k \in K_+} \beta_k \begin{pmatrix} 1 \\ \frac{\mathbf{x}^k}{\lambda_k} \\ \frac{\mathbf{u}^k}{\lambda_k} \end{pmatrix} \begin{pmatrix} 1 \\ \frac{\mathbf{x}^k}{\lambda_k} \\ \frac{\mathbf{u}^k}{\lambda_k} \end{pmatrix}^\top \quad \text{and} \quad \mathbf{M}_2 := \sum_{k \in K_0} \begin{pmatrix} 0 \\ \mathbf{x}^k \\ \mathbf{u}^k \end{pmatrix} \begin{pmatrix} 0 \\ \mathbf{x}^k \\ \mathbf{u}^k \end{pmatrix}^\top \end{aligned} \quad (3.8)$$

where $\lambda_k^2 = \beta_k$ for all $k \in K_+$ and $\sum_{k \in K_+} \beta_k = 1$. Moreover, it holds that $\mathbf{A}\mathbf{x}^k = \lambda_k \mathbf{b}$ for all $k \in K_+$ while $\mathbf{A}\mathbf{x}^k = \mathbf{o}$ for all $k \in K_0$.

Now consider a (CP-soft(γ))-feasible $\mathbf{Y}^*$ which is (nearly)³ optimal, $z_\gamma(\mathbf{Y}^*) \approx z_\gamma^*$.

If $K_0 = \emptyset$, then by default $\mathbf{M}_2 = \mathbf{O}$, so that it is easy to see the equivalence. Indeed, Lemma 16 gives

$$z_\gamma(\mathbf{Y}^*) = z_\gamma(\mathbf{M}_1) = \sum_{k \in K_+} \beta_k \left(\frac{(\mathbf{x}^k)^\top}{\lambda_k} \mathbf{Q} \frac{\mathbf{x}^k}{\lambda_k} + 2\mathbf{c}^\top \frac{\mathbf{x}^k}{\lambda_k} - \gamma \mathbf{e}^\top \frac{\mathbf{u}^k}{\lambda_k} \right)$$

with $\left(\frac{\mathbf{x}^k}{\lambda_k} \right)^\top \frac{\mathbf{u}^k}{\lambda_k} = 0$ and $\frac{\mathbf{u}^k}{\lambda_k} \in \{0, 1\}^n$ for all $k \in K_+$, hence all $\left(\frac{\mathbf{x}^k}{\lambda_k}, \frac{\mathbf{u}^k}{\lambda_k} \right)$ are (MIQP-soft(γ))-feasible, so

$$z_\gamma(\mathbf{Y}^*) \geq \min_{k \in K_+} \left[\frac{(\mathbf{x}^k)^\top}{\lambda_k} \mathbf{Q} \frac{\mathbf{x}^k}{\lambda_k} + 2\mathbf{c}^\top \frac{\mathbf{x}^k}{\lambda_k} - \lambda \mathbf{e}^\top \frac{\mathbf{u}^k}{\lambda_k} \right] \geq \zeta_\gamma^*,$$

and we get $\zeta_\gamma^* \leq z_\gamma^*$; the reverse inequality $\zeta_\gamma^* \geq z_\gamma^*$ follows by relaxation.

If $K_0 \neq \emptyset$ for $\mathbf{Y} = \mathbf{Y}^*$, then Lemma 16 ensures $\mathbf{u}^k = \mathbf{o}$ and $\mathbf{A}\mathbf{x}^k = \mathbf{o}$ for all $k \in K_0$. It follows that $\mathbf{x}(t) := \mathbf{x} + t\mathbf{x}^k$ is (MIQP-soft(γ))-feasible for all $t > 0$ and all $k \in K_0$, hence Assumption (3.1) guarantees $(\mathbf{x}^k)^\top \mathbf{Q}\mathbf{x}^k \geq 0$ for all $k \in K_0$. Recalling that

$$\mathbf{Y}^* = \sum_{k \in K_+} \beta_k \begin{pmatrix} 1 \\ \frac{\mathbf{x}^k}{\lambda_k} \\ \frac{\mathbf{u}^k}{\lambda_k} \end{pmatrix} \begin{pmatrix} 1 \\ \frac{\mathbf{x}^k}{\lambda_k} \\ \frac{\mathbf{u}^k}{\lambda_k} \end{pmatrix}^\top + \sum_{k \in K_0} \begin{pmatrix} 0 \\ \mathbf{x}^k \\ \mathbf{o} \end{pmatrix} \begin{pmatrix} 0 \\ \mathbf{x}^k \\ \mathbf{o} \end{pmatrix}^\top,$$

this implies

$$z_\gamma^* \approx z_\gamma(\mathbf{Y}^*) \geq z_\gamma(\bar{\mathbf{Y}}) \quad \text{with} \quad \bar{\mathbf{Y}} = \sum_{k \in K_+} \beta_k \begin{pmatrix} 1 \\ \frac{\mathbf{x}^k}{\lambda_k} \\ \frac{\mathbf{u}^k}{\lambda_k} \end{pmatrix} \begin{pmatrix} 1 \\ \frac{\mathbf{x}^k}{\lambda_k} \\ \frac{\mathbf{u}^k}{\lambda_k} \end{pmatrix}^\top,$$

and $z_\gamma(\bar{\mathbf{Y}}) \geq \zeta_\gamma^*$ as discussed for the case $K_0 = \emptyset$ above. This completes the proof. $\square$

³ We need to argue by approximate optimality here since attainability of the optimal value of conic problems is not guaranteed *a priori*. But see Corollary 18 below.

Corollary 18. *The optimal value z_γ^* of $(\text{CP-soft}(\gamma))$ is attained at an optimal solution $\mathbf{Y}^*$ of the form*

$$\mathbf{Y}^* = \sum_{k \in K_+} \beta_k \begin{pmatrix} 1 \\ \frac{\mathbf{x}^k}{\lambda_k} \\ \frac{\mathbf{u}^k}{\lambda_k} \end{pmatrix} \begin{pmatrix} 1 \\ \frac{\mathbf{x}^k}{\lambda_k} \\ \frac{\mathbf{u}^k}{\lambda_k} \end{pmatrix}^\top,$$

where $\beta_k = \lambda_k^2 > 0$ and $\sum_{k \in K_+} \beta_k = 1$. Moreover, we have

$$\frac{\mathbf{u}^k}{\lambda_k} \in \{0, 1\}^n,$$

and $u_i^k x_i^k = 0$ for all $i \in [1:n]$ and $k \in K_+$.

Proof. As argued in Theorem 17, there exists a nearly $(\text{CP-soft}(\gamma))$ -optimal solution $\bar{\mathbf{Y}}$ of the form:

$$\bar{\mathbf{Y}} = \sum_{k \in K_+} \beta_k \begin{pmatrix} 1 \\ \frac{\mathbf{x}^k}{\lambda_k} \\ \frac{\mathbf{u}^k}{\lambda_k} \end{pmatrix} \begin{pmatrix} 1 \\ \frac{\mathbf{x}^k}{\lambda_k} \\ \frac{\mathbf{u}^k}{\lambda_k} \end{pmatrix}^\top,$$

where $\beta_k = \lambda_k^2 > 0$ and $\sum_{k \in K_+} \beta_k = 1$. Moreover, we have

$$\frac{\mathbf{u}^k}{\lambda_k} \in \{0, 1\}^n,$$

and $u_i^k x_i^k = 0$ for all $i \in [1:n]$ and $k \in K_+$. In each component of the decomposition of $\bar{\mathbf{Y}}$, if the support of $\mathbf{u}^k$ is given, to find the corresponding $\mathbf{x}^k$ reduces to find an optimal solution of a QP, which is attainable. On the other hand, the number of supports of $\mathbf{u}^k$ is finite ($\leq 2^n$), which implies that the system of possible sets K_+ is finite.⁴ Therefore, the claim follows. $\square$

3.2.3. The decomposed CP reformulation

In this subsection, we propose another exact CP relaxation for (MIQP-hard) based on the sparse pattern of the problem.

We notice that the problem $(\text{CP-soft}(\gamma))$ has a sparse structure, namely, only the submatrices of the underlying variable $\mathbf{Y}$ are involved in both the objective function and the constraints. To be more precise, the specification graph of the variable $\mathbf{Y}$ involved in the problem $(\text{CP-soft}(\gamma))$ is chordal. It is known that a partial SDP matrix is SDP-completable if its specification graph is chordal, and a partial CP matrix is CP-completable if this graph is block-clique [SMB21]. Inspired by this fact, we consider the following decomposed CP/DNN relaxation:

$$\begin{aligned} \hat{z}_\gamma^* &= \inf_{\mathbf{X}, \mathbf{x}, \mathbf{u}} \quad \hat{z}_\gamma(\mathbf{X}, \mathbf{x}, \mathbf{u}) := \langle \mathbf{Q}, \mathbf{X} \rangle + 2\mathbf{c}^\top \mathbf{x} - \gamma \mathbf{e}^\top \mathbf{u} \\ \text{s.t.} \quad & \mathbf{A}\mathbf{x} = \mathbf{b}, \\ & \text{diag}(\mathbf{A}\mathbf{X}\mathbf{A}^\top) = \mathbf{b} \circ \mathbf{b}, \\ & \begin{pmatrix} 1 & x_i & u_i \\ x_i & X_{ii} & 0 \\ u_i & 0 & u_i \end{pmatrix} \in \mathcal{DNN}^3, \quad \forall i \in [1:n], \\ & \begin{pmatrix} 1 & \mathbf{x} \\ \mathbf{x} & \mathbf{X} \end{pmatrix} \in \mathcal{CP}^{n+1}. \end{aligned} \quad (\text{CP-soft-decomposed}(\gamma))$$

⁴ Note that $|K_+|$ may exceed the CP-rank of $\bar{\mathbf{Y}}$ so that we cannot claim a tighter bound on $|K_+|$ like those in [BSU15]. In other words, decomposition (3.8) may not be a minimal CP decomposition of $\mathbf{Y}^*$, so we cannot even apply Caratheodory's theorem. But fortunately only finiteness matters here.

Some comments related to $(\text{CP-soft-decomposed}(\gamma))$:

- $(\text{CP-soft-decomposed}(\gamma))$ is a direct consequence of a procedure that relaxes the cone constraint in $(\text{CP-soft}(\gamma))$ to its specified principal submatrices, which is known as chordal decomposition. Generally speaking, after the chordal decomposition for the CP constraint, the two relaxations are not necessarily equivalent in the sense of the objective value. However, we will show the equivalence in our cases.
- It is worth noticing that when $n \leq 4$, we have $\mathcal{DN}\mathcal{N}^n = \mathcal{CP}^n$ [SMB21], hence we replace $\mathcal{CP}^3$ by $\mathcal{DN}\mathcal{N}^3$ in the formulation. Note that this condition is equivalent to the two conditions $u_i \geq 0$ and $X_{ii}(1 - u_i) \geq x_i^2$, given that $X_{ii} \geq 0$.
- In the decomposed CP relaxation, we notice that the dimension of the cone constraint reduces dramatically from $2n + 1$ to $n + 1$, meanwhile the complexity of the problem decreases from $O(2^{2n+1})$ to $O(2^{n+1})$.

In the following, we will show that the relaxation $(\text{CP-soft-decomposed}(\gamma))$ is also exact.

Theorem 19. *The CP relaxation $(\text{CP-soft-decomposed}(\gamma))$ is exact.*

Proof. In analogy to Lemma 16, we study the necessary condition of a point $(\mathbf{X}, \mathbf{x}, \mathbf{u})$ being $(\text{CP-soft-decomposed}(\gamma))$ -optimal. First, similar to the proof of Lemma 16, for a given point $(\mathbf{X}, \mathbf{x}, \mathbf{u})$ which is $(\text{CP-soft-decomposed}(\gamma))$ -feasible with

$$\begin{pmatrix} 1 & \mathbf{x} \\ \mathbf{x} & \mathbf{X} \end{pmatrix} = \sum_{k \in K_+} \beta_k \begin{pmatrix} 1 \\ \frac{\mathbf{x}^k}{\lambda_k} \end{pmatrix} \begin{pmatrix} 1 \\ \frac{\mathbf{x}^k}{\lambda_k} \end{pmatrix}^\top + \sum_{k \in K_0} \begin{pmatrix} 0 \\ \mathbf{x}^k \end{pmatrix} \begin{pmatrix} 0 \\ \mathbf{x}^k \end{pmatrix}^\top,$$

we construct another $(\text{CP-soft-decomposed}(\gamma))$ -feasible point $(\mathbf{X}, \mathbf{x}, \hat{\mathbf{u}})$ with

$$\hat{z}_\gamma(\mathbf{X}, \mathbf{x}, \hat{\mathbf{u}}) \leq \hat{z}_\gamma(\mathbf{X}, \mathbf{x}, \mathbf{u}) \quad (3.9)$$

by defining $\hat{u}_i^k := 1 - \text{sign}(x_i^k) \in \{0, 1\}$. Since $\hat{u}_i = \sum_{k \in K_+} \beta_k \hat{u}_i^k \geq u_i$, the claimed relation (3.9) follows. Furthermore, the following CP decomposition across all $i \in [1:n]$ holds:

$$\begin{pmatrix} 1 & x_i & \hat{u}_i \\ x_i & X_{ii} & 0 \\ \hat{u}_i & 0 & \hat{u}_i \end{pmatrix} = \sum_{k \in K_+} \beta_k \begin{pmatrix} 1 \\ \frac{x_i^k}{\lambda_k} \\ \hat{u}_i^k \end{pmatrix} \begin{pmatrix} 1 \\ \frac{x_i^k}{\lambda_k} \\ \hat{u}_i^k \end{pmatrix}^\top.$$

Here $\mathbf{x}^k \in \mathbb{R}_+^n$ for all $k \in K_+ \cup K_0$, which is a point of the projection of the feasible region of $(\text{CP-soft}(\gamma))$ onto $(\mathbf{X}, \mathbf{x}, \mathbf{u})$, which implies that $\hat{z}_\gamma^* \geq z_\gamma^*$. On the other hand, it is clear that the feasible region of $(\text{CP-soft-decomposed}(\gamma))$ contains the corresponding projected set of the feasible region of $(\text{CP-soft}(\gamma))$. Therefore, together with Theorem 17 we have

$$\hat{z}_\gamma^* = z_\gamma^* = \zeta_\gamma,$$

with which we complete the proof. $\square$

Corollary 20. *The optimal value $\hat{z}_\gamma^*$ of $(\text{CP-soft}(\gamma))$ is attained at an optimal solution $(\mathbf{X}^*, \mathbf{x}^*, \mathbf{u}^*)$ of the form*

$$\begin{pmatrix} 1 & x_i & u_i \\ x_i & X_{ii} & 0 \\ u_i & 0 & u_i \end{pmatrix} = \sum_{k \in K_+} \beta_k \begin{pmatrix} 1 \\ \frac{x_i^k}{\lambda_k} \\ u_i^k \end{pmatrix} \begin{pmatrix} 1 \\ \frac{x_i^k}{\lambda_k} \\ u_i^k \end{pmatrix}^\top, \quad \forall i \in [1:n],$$

where $\beta_k = \lambda_k^2 > 0$ and $\sum_{k \in K_+} \beta_k = 1$. Moreover, we have

$$\mathbf{u}^k \in \{0, 1\}^n,$$

and $u_i^k x_i^k = 0$ for all $i \in [1:n]$ and $k \in K_+$.

Proof. Directly from Theorem 19, there exists a nearly (CP-soft-decomposed(γ))-optimal solution having the properties stated above. As the rest is quite identical to the proof of Corollary 18, we omit it for the sake of brevity. $\square$

3.3. CP reformulation for hard sparse QPs

In this section, we propose two compact CP relaxations of (MIQP-hard) and prove the exactness of the relaxations.

3.3.1. New and more compact CP reformulation for hard sparse QPs

Recall the hard sparse QP is of the form:

$$\begin{aligned} \ell_\rho^* &:= \min_{\mathbf{x} \in \mathbb{R}_+^n} \quad \mathbf{x}^\top \mathbf{Q} \mathbf{x} + 2\mathbf{c}^\top \mathbf{x} \\ \text{s.t.} \quad & \mathbf{A} \mathbf{x} = \mathbf{b}, \\ & \|\mathbf{x}\|_0 \leq \rho, \end{aligned} \tag{Hard-QP}$$

where $\rho \in [0, n]$ is a parameter restricting the number of nonzero elements of $\mathbf{x}$. Notice that we only assume the feasible region of (QP) is nonempty as shown in (3.1), therefore it may happen for small ρ , the problem (Hard-QP) is infeasible. Indeed, if the feasible region of (Hard-QP) is nonempty for $\rho = \mu \in [0, n]$, then the set remains nonempty for $\rho \in [\mu, n]$. To accommodate the above case, by $\underline{\rho} \in [0, n]$ we define the l_0 -pseudonorm of the sparsest feasible point of the problem (QP). It follows immediately that the problem (Hard-QP) is feasible for $\rho \in [\underline{\rho}, n]$. As explained above, for all these ρ , the optimal solution of (Hard-QP) is attainable due to the Frank/Wolfe Theorem.

Proposition 21. *For a given optimization problem (Hard-QP), we have $\ell_{\rho_1}^* \geq \ell_{\rho_2}^*$ if $\rho_1 \leq \rho_2$.*

Proof. The claim follows since the feasible region of (Hard-QP) enlarges when the parameter ρ increases. By default, if an optimization problem is infeasible, the optimal value is $+\infty$. So the relation remains true even if $\rho_1 < \underline{\rho}$. $\square$

Before loading to the CP reformulation, we present a useful property on the optimal value of (Hard-QP) when the sparsity parameter ρ varies: the objective value is not only decreasing, but even convex decreasing in ρ .

Lemma 22. *The problem (Hard-QP) has the property:*

$$\alpha \ell_\rho^* + (1 - \alpha) \ell_{\rho'}^* \geq \ell_{\alpha\rho + (1-\alpha)\rho'}^* \quad \text{for all } \{\rho, \rho'\} \subset [0, n] \text{ and } \alpha \in [0, 1].$$

Proof. In the case of $\min\{\rho, \rho'\} < \underline{\rho}$, it is clear that the $\max\{\ell_\rho^*, \ell_{\rho'}^*\} = +\infty$, therefore

$$\alpha \ell_\rho^* + (1 - \alpha) \ell_{\rho'}^* \geq \ell_{\alpha\rho + (1-\alpha)\rho'}^*$$

holds for all $\alpha \in [0, 1]$.

In the following, we show the property when $\min\{\rho, \rho'\} \geq \underline{\rho}$. First, by $\mathbf{x}_\rho^*$ we denote an optimal solution to the problem (Hard-QP) and consider any strictly positive vector $\mathbf{d} \in \text{int } \mathbb{R}_+^n$. Letting

$$h := \max_{\rho \in [\underline{\rho}, n]} \mathbf{d}^\top \mathbf{x}_\rho^* = \max_{\rho \in [\underline{\rho}, n]} \mathbf{d}^\top \mathbf{x}_\rho^*,$$

we add the redundant constraint $\mathbf{d}^\top \mathbf{x} + t = h$ to (Hard-QP):

$$\begin{aligned} \ell_\rho^* = \min_{\mathbf{x} \in \mathbb{R}_+^n, t \in \mathbb{R}_+} \quad & \mathbf{x}^\top \mathbf{Q} \mathbf{x} + 2\mathbf{c}^\top \mathbf{x} \\ \text{s.t.} \quad & \mathbf{A} \mathbf{x} = \mathbf{b}, \\ & \mathbf{d}^\top \mathbf{x} + t = h, \\ & \|\mathbf{x}\|_0 \leq \rho. \end{aligned} \tag{3.10}$$

Due to the construction of h , it is clear that the problem (3.10) is equivalent to the problem (Hard-QP) for all $\rho \in [\underline{\rho}, n]$. By introducing auxiliary variables and slack variables $\{\mathbf{u}, \mathbf{s}\} \subset \mathbb{R}_+^n$, the problem (3.10) can be rewritten:

$$\begin{aligned} \ell_\rho^* = \min_{\{\mathbf{x}, \mathbf{u}, \mathbf{s}\} \subset \mathbb{R}_+^n, t \in \mathbb{R}_+} \quad & \mathbf{x}^\top \mathbf{Q} \mathbf{x} + 2\mathbf{c}^\top \mathbf{x} \\ \text{s.t.} \quad & \mathbf{A} \mathbf{x} = \mathbf{b}, \\ & \mathbf{d}^\top \mathbf{x} + t = h, \\ & \mathbf{e}^\top \mathbf{u} = n - \rho, \\ & \mathbf{u} + \mathbf{s} = \mathbf{e}, \\ & \mathbf{x}^\top \mathbf{u} = \mathbf{o}. \end{aligned} \tag{3.11}$$

Since the feasible region of problem (3.11) is bounded, it satisfies the condition ensuring exactness of the CP relaxation proposed in [Bur09], it admits an exact CP reformulation of the form

$$\begin{aligned} \ell_\rho^* = \inf_{\mathbf{Y}} \quad & \langle \mathbf{Q}, \mathbf{Y}^{\mathbf{xx}} \rangle + \mathbf{c}^\top \mathbf{x} \\ \text{s.t.} \quad & \mathbf{Y}^{00} = 1, & [y_1] \\ & \mathbf{A} \mathbf{x} = \mathbf{b}, & [\mathbf{z}_1] \\ & \mathbf{d}^\top \mathbf{x} + t = h, & [y_2] \\ & \mathbf{e}^\top \mathbf{u} = n - \rho, & [y_3] \\ & \mathbf{u} + \mathbf{s} = \mathbf{e}, & [\mathbf{z}_2] \\ & \text{diag}(\mathbf{A} \mathbf{Y}^{\mathbf{xx}} \mathbf{A}^\top) = \mathbf{b} \circ \mathbf{b}, & [\mathbf{z}_3] \\ & \langle \mathbf{d} \mathbf{d}^\top, \mathbf{Y}^{\mathbf{xx}} \rangle + 2\mathbf{d}^\top \mathbf{Y}^{\mathbf{xt}} + \mathbf{Y}^{tt} = h^2, & [y_4] \\ & \text{diag}(\mathbf{Y}^{\mathbf{uu}} + 2\mathbf{Y}^{\mathbf{us}} + \mathbf{Y}^{\mathbf{ss}}) = \mathbf{e}, & [\mathbf{z}_4] \\ & \mathbf{e}^\top \text{diag}(\mathbf{Y}^{\mathbf{xu}}) = 0, & [y_5] \\ & \langle \mathbf{E}, \mathbf{Y}^{\mathbf{uu}} \rangle = (n - \rho)^2, & [y_6] \\ & \mathbf{Y} := \begin{pmatrix} \mathbf{Y}^{00} & \mathbf{x}^\top & \mathbf{u}^\top & \mathbf{s}^\top & t \\ \mathbf{x} & \mathbf{Y}^{\mathbf{xx}} & \mathbf{Y}^{\mathbf{xu}} & \mathbf{Y}^{\mathbf{xs}} & \mathbf{Y}^{\mathbf{xt}} \\ \mathbf{u} & (\mathbf{Y}^{\mathbf{xu}})^\top & \mathbf{Y}^{\mathbf{uu}} & \mathbf{Y}^{\mathbf{us}} & \mathbf{Y}^{\mathbf{ut}} \\ \mathbf{s} & (\mathbf{Y}^{\mathbf{xs}})^\top & (\mathbf{Y}^{\mathbf{us}})^\top & \mathbf{Y}^{\mathbf{ss}} & \mathbf{Y}^{\mathbf{st}} \\ t & (\mathbf{Y}^{\mathbf{xt}})^\top & (\mathbf{Y}^{\mathbf{ut}})^\top & (\mathbf{Y}^{\mathbf{st}})^\top & \mathbf{Y}^{tt} \end{pmatrix} \in \mathcal{CP}^{3n+2}, \end{aligned} \tag{3.12}$$

in which the associated dual variable for each constraint is introduced on the right. Notice that $\mathbf{z}_i \in \mathbb{R}^n$ for all $i \in [1:4]$ and $y_j \in \mathbb{R}$ for all $j \in [1:6]$. Further observe that the equality constraint $\langle \mathbf{E}, \mathbf{Y}^{uu} \rangle = (n - \rho)^2$ can be equivalently relaxed to the inequality $\langle \mathbf{E}, \mathbf{Y}^{uu} \rangle \leq (n - \rho)^2$. To show this, for any (3.12)-feasible $\mathbf{Y}^{uu}$, we notice that $\mathbf{Y}^{uu} \succeq \mathbf{u}\mathbf{u}^\top$ and $\mathbf{e}^\top \mathbf{u} = n - \rho$, which implies that $\langle \mathbf{E}, \mathbf{Y}^{uu} \rangle \geq (n - \rho)^2$. Therefore, such replacement does not change the feasible set of (3.12). An immediate implication is that the dual variable y_6 has to be non-negative.

The conic dual of (3.12) reads as

$$\begin{aligned} d_\rho^* &:= \sup_{\mathbf{S}} \quad (n - \rho)y_3 + (n - \rho)^2 y_6 + f(y_i, \mathbf{z}_j) \\ \text{s.t.} \quad &y_6 \geq 0 \\ &\mathbf{S}(y_i, \mathbf{z}_j) \in \mathcal{COP}^{3n+2}, \end{aligned} \tag{3.13}$$

where $f(y_i, \mathbf{z}_j) := y_1 + \mathbf{b}^\top \mathbf{z}_1 + h y_2 + \mathbf{e}^\top \mathbf{z}_2 + (\mathbf{b} \circ \mathbf{b})^\top \mathbf{z}_3 + h^2 y_4 + \mathbf{e}^\top \mathbf{z}_4$ and the upper triangular part of the symmetric slack matrix $\mathbf{S}(y_i, \mathbf{z}_j)$ is defined as an affine function of all dual variables $(y_i, \mathbf{z}_j)$ as follows:

$$\begin{pmatrix} -2y_1 & (\mathbf{c} - \mathbf{A}^\top \mathbf{z}_1 - y_2 \mathbf{d})^\top & -(y_3 \mathbf{e} + \mathbf{z}_2)^\top & -\mathbf{z}_2^\top & -y_2 \\ * & 2\mathbf{B}(y_4, \mathbf{z}_3) & -y_5 \mathbf{I} & \mathbf{O} & -2y_4 \mathbf{d} \\ * & * & -2(\text{Diag}(\mathbf{z}_4) + y_6 \mathbf{E}) & -\text{Diag}(\mathbf{z}_4) & \mathbf{o} \\ * & * & * & -2\text{Diag}(\mathbf{z}_4) & \mathbf{o} \\ * & * & * & * & -2y_4 \end{pmatrix},$$

with $\mathbf{B}(y_4, \mathbf{z}_3) := \mathbf{Q} - \text{Diag}(\mathbf{z}_3)\mathbf{A}\mathbf{A}^\top - y_4 \mathbf{d}\mathbf{d}^\top$. Let

$$\bar{y}_2 < \min\{0\} \cup \left\{ \frac{c_i}{d_i} : i \in [1:n] \right\} \quad \text{and} \quad \bar{y}_4 < \min\{0\} \cup \left\{ \frac{Q_{ij}}{d_i d_j} : (i, j) \in [1:n]^2 \right\}.$$

Furthermore, let $\bar{y}_1 = -0.5$, $\bar{y}_3 = \bar{y}_5 = \bar{y}_6 = 0$, $\bar{\mathbf{z}}_1 = \bar{\mathbf{z}}_2 = \bar{\mathbf{z}}_3 = \mathbf{o}$ while $\bar{\mathbf{z}}_4 = -\mathbf{e}$. Then the slack matrix $\bar{\mathbf{S}} := \mathbf{S}(\bar{y}_i, \bar{\mathbf{z}}_j)$ generates a quadratic form in $[r, \mathbf{x}^\top, \mathbf{u}^\top, \mathbf{s}^\top, t]^\top \in \mathbb{R}_+^{3n+2} \setminus \{\mathbf{o}\}$ satisfying

$$[r, \mathbf{x}^\top, \mathbf{u}^\top, \mathbf{s}^\top, t] \bar{\mathbf{S}} [r, \mathbf{x}^\top, \mathbf{u}^\top, \mathbf{s}^\top, t]^\top \geq r^2 + 2 \left(\mathbf{x}^\top \mathbf{B}(\bar{y}_4, \bar{\mathbf{z}}_3) \mathbf{x} + \|\mathbf{u}\|^2 + \|\mathbf{s}\|^2 \right) + t^2,$$

which implies that $\bar{\mathbf{S}}$ is a strictly copositive matrix, since $\mathbf{B}(\bar{y}_4, \bar{\mathbf{z}}_3) = \mathbf{Q} - \bar{y}_4 \mathbf{d}\mathbf{d}^\top$ is so, having only positive entries by choice of $\bar{y}_4$. Therefore strong duality holds for the pair (3.12) and (3.13): $\ell_\rho^* = d_\rho^*$. But by definition in (3.13), the function d_ρ^* is convex in ρ as a pointwise supremum of convex functions. We conclude that ℓ_ρ^* is also convex in ρ , which completes the proof. $\square$

Now we turn to the mixed-integer reformulation of (Hard-QP). Again, by introducing auxiliary (now binary) variables $\mathbf{u}$, we arrive at the reformulation:

$$\begin{aligned} \ell_\rho^* &= \min_{\mathbf{x} \in \mathbb{R}_+^n, \mathbf{u} \in \{0,1\}^n} \quad \mathbf{x}^\top \mathbf{Q} \mathbf{x} + \mathbf{c}^\top \mathbf{x} \\ \text{s.t.} \quad &\mathbf{A} \mathbf{x} = \mathbf{b}, \\ &\mathbf{e}^\top \mathbf{u} = n - \rho, \\ &\mathbf{x}^\top \mathbf{u} = 0. \end{aligned} \tag{MIQP-hard}$$

Consider the following CP relaxation of (MIQP-hard):

$$\begin{aligned}
\ell_\rho^{\text{CP}} &:= \inf \quad \langle \mathbf{Q}, \mathbf{Z}^{\text{xx}} \rangle + \mathbf{c}^\top \mathbf{x} \\
&\text{s.t.} \quad \mathbf{A}\mathbf{x} = \mathbf{b} \\
&\quad \text{diag}(\mathbf{A}\mathbf{Z}^{\text{xx}}\mathbf{A}^\top) = \mathbf{b} \circ \mathbf{b}, \\
&\quad \mathbf{e}^\top \mathbf{u} = n - \rho \\
&\quad \mathbf{e}^\top \text{diag}(\mathbf{Z}^{\text{xu}}) = 0 \\
&\quad \text{diag}(\mathbf{Z}^{\text{uu}}) = \mathbf{u} \\
&\quad \mathbf{Z} := \begin{pmatrix} 1 & \mathbf{x}^\top & \mathbf{u}^\top \\ \mathbf{x} & \mathbf{Z}^{\text{xx}} & \mathbf{Z}^{\text{xu}} \\ \mathbf{u} & (\mathbf{Z}^{\text{xu}})^\top & \mathbf{Z}^{\text{uu}} \end{pmatrix} \in \mathcal{CP}^{2n+1}.
\end{aligned} \tag{CP-hard(\rho)}$$

We notice that the relaxation has fewer variables, with only $3n + 2$ linear constraints and a much smaller conic one, compared to CP reformulations in [BP23, Bur09].

In the following part, we will show, under the assumption (3.1), the equivalence between (MIQP-hard) and (CP-hard(ρ)). First of all, we establish an important result similar to Lemma 15.

Lemma 23. *For any (CP-hard(ρ))-feasible point $\mathbf{Z}$, the following decomposition holds:*

$$\mathbf{Z} = \begin{pmatrix} 1 & \mathbf{x}^\top & \mathbf{u}^\top \\ \mathbf{x} & \mathbf{Z}^{\text{xx}} & \mathbf{Z}^{\text{xu}} \\ \mathbf{u} & (\mathbf{Z}^{\text{xu}})^\top & \mathbf{Z}^{\text{uu}} \end{pmatrix} = \sum_{k \in K_+} \beta_k \begin{pmatrix} 1 \\ \frac{\mathbf{x}^k}{\lambda_k} \\ \frac{\mathbf{u}^k}{\lambda_k} \end{pmatrix} \begin{pmatrix} 1 \\ \frac{\mathbf{x}^k}{\lambda_k} \\ \frac{\mathbf{u}^k}{\lambda_k} \end{pmatrix}^\top + \sum_{k \in K_0} \begin{pmatrix} 0 \\ \mathbf{x}^k \\ \mathbf{u}^k \end{pmatrix} \begin{pmatrix} 0 \\ \mathbf{x}^k \\ \mathbf{u}^k \end{pmatrix}^\top,$$

where $\beta_k = \lambda_k^2 > 0$ for all $k \in K_+$ and $\sum_{k \in K_+} \beta_k = 1$. Furthermore, we have $\sum_{k \in K_+} \beta_k \|\mathbf{x}^k\|_0 \leq \rho$ and $\mathbf{u}^k = \mathbf{0}$ for all $k \in K_0$.

Proof. For a given (CP-hard(ρ))-feasible point $\mathbf{Z}$, there is always a CP decomposition of the form (3.8). Therefore, we can present a feasible point either by $\mathbf{Z}$ or $(\lambda_k, \mathbf{x}^k, \mathbf{u}^k)_{k \in K}$. Let us fix $\mathbf{Z}^{\text{xx}}$ and $\mathbf{x}$ (or equivalently λ_k and $\mathbf{x}^k$ for all $k \in K$), and try to find the necessary condition on $\mathbf{u}^k$ for the point $(\lambda_k, \mathbf{x}^k, \mathbf{u}^k)_{k \in K}$ being (CP-hard(ρ))-feasible. Again, we need some index sets which play a significant role in the proof: for all $i \in [1:n]$, consider

$$P_i = \{k \in K_+ : x_i^k > 0\} \quad \text{and} \quad N_i = K_+ \setminus P_i = \{k \in K_+ : x_i^k = 0\}.$$

Now consider the following optimization problem:

$$\begin{aligned}
\ell &:= \max_{\mathbf{u}^k \in \mathbb{R}_+^n} \quad \sum_{k \in K_+} \lambda_k \mathbf{e}^\top \mathbf{u}^k \\
&\text{s.t.} \quad \sum_{k \in K} (u_i^k)^2 = \sum_{k \in K_+} \lambda_k u_i^k, \quad \forall i \in [1:n], \\
&\quad u_i^k x_i^k = 0, \quad \forall i \in [1:n], \quad \forall k \in K.
\end{aligned} \tag{3.14}$$

It is clear that one necessary condition of the point $(\lambda_k, \mathbf{x}^k, \mathbf{u}^k)_{k \in K}$ being (CP-hard(ρ))-feasible is that the optimal value of (3.14) satisfying

$$\ell \geq n - \rho. \tag{3.15}$$

Due to the fully decomposed structure in both the objective and the constraints regarding $\mathbf{u}^k$, and by the defined sets P_i and N_i , (3.14) can be equivalently splitted into a set of more simple problems with index $i \in [1:n]$:

$$\begin{aligned} \max_{u_i^k \in \mathbb{R}_+} \quad & \sum_{k \in K_+} \lambda_k u_i^k \\ \text{s.t.} \quad & \sum_{k \in N_i} (u_i^k)^2 + \sum_{k \in K_0} (u_i^k)^2 = \sum_{k \in K_+} \lambda_k u_i^k. \end{aligned} \quad (3.16)$$

It follows from Lemma 15 that the optimal solution to the problem (3.16) is

$$(u_i^k)^* = \begin{cases} \lambda_k, & \text{if } k \in N_i; \\ 0, & \text{else.} \end{cases}$$

The optimal value of (3.16) is

$$\ell = \sum_{k \in K_+} \beta_k \mathbf{e}^\top \left(\frac{(\mathbf{u}^k)^*}{\lambda_k} \right) = \sum_{k \in K_+} \beta_k (n - \|\mathbf{x}^k\|_0) = n - \sum_{k \in K_+} \beta_k \|\mathbf{x}^k\|_0.$$

Together with (3.15), we have $\sum_{k \in K_+} \beta_k \|\mathbf{x}^k\|_0 \leq \rho$, which is exactly the first half of the statement. The second half follows immediately from the structure of the optimal solution to (3.16). i.e., $(u_i^k)^* = 0$ for all $k \in K_0$ and $i \in [1:n]$. $\square$

Theorem 24. *Under the assumption (3.1), the CP relaxation (CP-hard(ρ)) is equivalent to the problem (MIQP-hard).*

Proof. It is clear that the problem (CP-hard(ρ)) is a relaxation of the problem (MIQP-hard), hence

$$\ell_\rho^* \geq \ell_\rho^{\text{CP}}.$$

To show the converse, we take a (CP-hard(ρ)) feasible point $\mathbf{Z}$ and compare the objective values between (CP-hard(ρ)) and (MIQP-hard). Let us consider

$$\begin{aligned} & \langle \mathbf{Q}, \mathbf{Z}^{\text{xx}} \rangle \\ \stackrel{(3.8)}{=} & \sum_{k \in K_+} \beta_k \left(\frac{\mathbf{x}^k}{\lambda_k} \right)^\top \mathbf{Q} \frac{\mathbf{x}^k}{\lambda_k} \\ \geq & \sum_{k \in K_+} \beta_k \ell_{\|\mathbf{x}^k\|_0}^* \\ \geq & \ell_{\sum_{k \in K_+} \beta_k \|\mathbf{x}^k\|_0}^* \\ \geq & \ell_\rho^*, \end{aligned} \quad (3.17)$$

where the last inequality follows from Lemma 23 and Proposition 21, while the previous one follows from Lemma 22. So (3.17) holds which implies $\ell_\rho^{\text{CP}} \geq \ell_\rho^*$. $\square$

Corollary 25. *The optimal solution of (CP-hard(ρ)) is attained.*

Proof. Directly from Theorem 24 via the Frank/Wolfe theorem. $\square$

3.4. Conclusion

In this chapter, we introduced new, significantly more compact CP reformulations for QPs with a sparsity term under mild conditions. By leveraging the sparse structure of these CP reformulations, we also proved the equivalence of their decomposed forms, demonstrating that the complexity of the problems increases only mildly with dimension. These compact relaxations facilitate the development of tight yet tractable relaxations in the area of sparse QPs.

4. Conic relaxations for sparse StQPs

In this chapter, we consider standard quadratic optimization problems (StQPs) under a hard sparsity constraint and propose novel relaxations based on the corresponding CP reformulations with significant dimensional reduction, which is essential for the tractability of these relaxations when the problem size grows. Moreover, an instance-generation procedure of StQPs is proposed, which systematically avoids too easy instances. Extensive computational results illustrate the high quality of the relaxation bounds in a significant number of instances. The content of this chapter is based on the submitted paper [BPQY24]. A precursor of the material presented in the previous chapter, which also complements the study in this chapter, is published as [BPQY25].

4.1. The problem

An StQP consists of minimizing a (not necessarily convex) quadratic form over the standard simplex:

$$\ell_n(\mathbf{Q}) := \min_{\mathbf{x} \in \mathbb{R}^n} \left\{ \mathbf{x}^\top \mathbf{Q} \mathbf{x} : \mathbf{x} \in F_n \right\}, \quad (\text{StQP})$$

where

$$F_n := \left\{ \mathbf{x} \in \mathbb{R}_+^n : \mathbf{e}^\top \mathbf{x} = 1 \right\} \quad (4.1)$$

is the standard simplex, the simplest polytope in the n -dimensional Euclidean space $\mathbb{R}^n$. Here, $\mathbf{e} \in \mathbb{R}^n$ denotes the vector of all ones.

Despite its simplicity, this problem class serves to model manifold real-world applications, ranging from the analysis of social networks (community detection via dominant-set-clustering) [BKL21, MS65] through biology, game theory to economy and finance [Bom02, Mar52], to cite just a few. For a survey on mixed-integer convex quadratic optimization approaches to portfolio selection, see, e.g., [MD19].

One motivation to introduce sparsity constraints comes from high-dimensional portfolio selection. Recently, the significant benefits of controlling the cardinality of a portfolio have been widely recognized in the literature, e.g., [DGW23], [HV19].

Therefore, we propose to study the StQP with an exact, hard sparsity constraint, with the aim of developing tight yet tractable relaxations. Such a problem can be expressed as follows:

$$\ell_\rho(\mathbf{Q}) := \min_{\mathbf{x} \in \mathbb{R}^n} \left\{ \mathbf{x}^\top \mathbf{Q} \mathbf{x} : \mathbf{x} \in F_\rho \right\}, \quad (\text{StQP}(\rho))$$

where

$$F_\rho := \{ \mathbf{x} \in F_n : \|\mathbf{x}\|_0 \leq \rho \} = \left\{ \mathbf{x} \in \mathbb{R}_+^n : \mathbf{e}^\top \mathbf{x} = 1, \quad \|\mathbf{x}\|_0 \leq \rho \right\}. \quad (4.2)$$

The sparse StQP problem appears in numerous applications, notably in portfolio selection with asset limits [Bie96, Per84, XTYP24]. `MATLAB`'s latest release [The24] includes a MILP framework to solve MIQP formulations of convex portfolio problems, handling lower and upper cardinality bounds and semi-continuous variables. However, tests show that even for $n = 30$ and extreme sparsity $\rho = 1$, solution time was around 40 seconds, and no solution was found for $\rho = 2$ within

two hours. Additional applications include selecting regression variables [Mil84] and compressive sampling [CW08].

The sparse StQP problem has been studied extensively and can be categorized into two main approaches. The first addresses sparsity constraints indirectly through regularization methods such as the ℓ_1 -norm or the ℓ_p -pseudonorms with $p \in (0, 1)$ [BDDM⁺09, CPZ13]. The second approach directly tackles cardinality constraints using various methods. Semidefinite programming-based reformulations for generating perspective cuts have been explored in [FG07, ZSL14], while nonconvex reformulations using regularization methods are presented in [BKS16]. Heuristic methods based on semidefinite programming that provide upper bounds for the sparse StQP problem are proposed in [BM05], and tight semidefinite relaxations under convexity assumptions are discussed in [WZ21]. Additionally, mixed-integer formulations are solved through discrete optimization techniques, either to optimality or as approximations [BS09, Bie96, DLLR⁺12, MS12, RTGMS10, SZL13, ZSLS14]. In this work, our focus is on finding tight but tractable relaxations that scale well with the problem dimension.

4.2. Two exact completely positive formulations

In this section, we present two exact formulations of the sparse StQP problem ($\text{StQP}(\rho)$) as mixed-binary quadratic optimization problems, a direct one covered by Section 4.2.1, whereas Section 4.2.2 is devoted to a formulation involving a complementarity constraint.

4.2.1. Conic formulation of direct MIQP

By introducing binary variables, the sparse StQP can be reformulated as the following mixed-binary quadratic optimization problem:

$$\begin{aligned} \ell_\rho(Q) = \min_{(\mathbf{x}, \mathbf{u}) \in \mathbb{R}^n \times \mathbb{R}^n} \quad & \mathbf{x}^\top Q \mathbf{x} \\ \text{s.t.} \quad & \mathbf{e}^\top \mathbf{x} = 1 \\ & \mathbf{e}^\top \mathbf{u} = \rho \\ & \mathbf{x} \leq \mathbf{u} \\ & \mathbf{u} \in \{0, 1\}^n \\ & \mathbf{x} \geq \mathbf{o}. \end{aligned} \tag{P1(\rho)}$$

By introducing the redundant constraints $\mathbf{u} \leq \mathbf{e}$, slack variables $\mathbf{y} \geq \mathbf{o}$ and $\mathbf{v} \geq \mathbf{o}$ so that $\mathbf{x} + \mathbf{y} = \mathbf{u}$ and $\mathbf{u} + \mathbf{v} = \mathbf{e}$, we arrive at the following reformulation of (P1(\rho)):

$$\begin{aligned} \ell_\rho(Q) = \min_{(\mathbf{x}, \mathbf{u}, \mathbf{v}, \mathbf{y}) \in \mathbb{R}^n \times \mathbb{R}^n \times \mathbb{R}^n \times \mathbb{R}^n} \quad & \mathbf{x}^\top Q \mathbf{x} \\ \text{s.t.} \quad & \mathbf{e}^\top \mathbf{x} = 1 \\ & \mathbf{e}^\top \mathbf{u} = \rho \\ & \mathbf{x} + \mathbf{y} = \mathbf{u} \\ & \mathbf{u} + \mathbf{v} = \mathbf{e} \\ & \mathbf{u} \in \{0, 1\}^n \\ & \mathbf{x}, \mathbf{y}, \mathbf{v} \geq \mathbf{o} \end{aligned} \tag{P1A(\rho)}$$

By [Bur09], (P1A(ρ)) admits an equivalent reformulation as the following conic optimization problem:

$$\begin{aligned}
\ell_\rho(\mathbf{Q}) = \min_{\mathbf{Z} \in \mathcal{S}^{4n+1}} \quad & \langle \mathbf{Q}, \mathbf{Z}^{\mathbf{x}\mathbf{x}} \rangle \\
\text{s.t.} \quad & \mathbf{e}^\top \mathbf{x} = 1 \\
& \mathbf{e}^\top \mathbf{u} = \rho \\
& \mathbf{x} + \mathbf{y} = \mathbf{u} \\
& \mathbf{u} + \mathbf{v} = \mathbf{e} \\
& \langle \mathbf{E}, \mathbf{Z}^{\mathbf{x}\mathbf{x}} \rangle = 1 \\
& \langle \mathbf{E}, \mathbf{Z}^{\mathbf{u}\mathbf{u}} \rangle = \rho^2 \\
& \text{diag}(\mathbf{Z}^{\mathbf{u}\mathbf{u}}) = \mathbf{u} \\
& \text{diag}(\mathbf{Z}^{\mathbf{x}\mathbf{x}}) + \text{diag}(\mathbf{Z}^{\mathbf{y}\mathbf{y}}) + \text{diag}(\mathbf{Z}^{\mathbf{u}\mathbf{u}}) + 2[\text{diag}(\mathbf{Z}^{\mathbf{x}\mathbf{y}}) - \text{diag}(\mathbf{Z}^{\mathbf{x}\mathbf{u}}) - \text{diag}(\mathbf{Z}^{\mathbf{u}\mathbf{y}})] = \mathbf{o} \\
& \text{diag}(\mathbf{Z}^{\mathbf{u}\mathbf{u}}) + 2\text{diag}(\mathbf{Z}^{\mathbf{u}\mathbf{v}}) + \text{diag}(\mathbf{Z}^{\mathbf{v}\mathbf{v}}) = \mathbf{e} \\
& \mathbf{Z} := \begin{pmatrix} 1 & \mathbf{x}^\top & \mathbf{u}^\top & \mathbf{v}^\top & \mathbf{y}^\top \\ \mathbf{x} & \mathbf{Z}^{\mathbf{x}\mathbf{x}} & \mathbf{Z}^{\mathbf{x}\mathbf{u}} & \mathbf{Z}^{\mathbf{x}\mathbf{v}} & \mathbf{Z}^{\mathbf{x}\mathbf{y}} \\ \mathbf{u} & (\mathbf{Z}^{\mathbf{x}\mathbf{u}})^\top & \mathbf{Z}^{\mathbf{u}\mathbf{u}} & \mathbf{Z}^{\mathbf{u}\mathbf{v}} & \mathbf{Z}^{\mathbf{u}\mathbf{y}} \\ \mathbf{v} & (\mathbf{Z}^{\mathbf{x}\mathbf{v}})^\top & (\mathbf{Z}^{\mathbf{u}\mathbf{v}})^\top & \mathbf{Z}^{\mathbf{v}\mathbf{v}} & \mathbf{Z}^{\mathbf{v}\mathbf{y}} \\ \mathbf{y} & (\mathbf{Z}^{\mathbf{x}\mathbf{y}})^\top & (\mathbf{Z}^{\mathbf{u}\mathbf{y}})^\top & (\mathbf{Z}^{\mathbf{v}\mathbf{y}})^\top & \mathbf{Z}^{\mathbf{y}\mathbf{y}} \end{pmatrix} \in \mathcal{CP}^{4n+1}.
\end{aligned}
\tag{CP1A(\rho)}$$

By [Bur09], (P1A(ρ)) and (CP1A(ρ)) are equivalent in the following sense: Both of their optimal values are equal to $\ell_\rho(\mathbf{Q})$ and for any optimal solution $\mathbf{Z} \in \mathcal{S}_+^{4n+1}$ of (CP1A(ρ)), $(\mathbf{x}, \mathbf{u}, \mathbf{y}, \mathbf{v}) \in \mathbb{R}^n \times \mathbb{R}^n \times \mathbb{R}^n \times \mathbb{R}^n$ is in the convex hull of the set of optimal solutions of (P1A(ρ)). Therefore, (CP1A(ρ)) is an exact convex reformulation of the nonconvex optimization problem (P1A(ρ)).

4.2.2. Conic formulation of MIQP with complementarity constraint

In (P1(ρ)), sparsity constraints are handled by big-M constraints. Alternatively, sparsity can be enforced by dropping the big-M constraints $\mathbf{x} \leq \mathbf{u}$ and replacing them with the complementarity constraints $\mathbf{x}_j(1 - \mathbf{u}_j) = 0$, $j = 1, \dots, n$ (see, e.g., [BKS16]). Defining $\mathbf{v} = \mathbf{e} - \mathbf{u} \in \{0, 1\}^n$, we obtain the following quadratic optimization problem with complementarity constraints:

$$\begin{aligned}
\ell_\rho(\mathbf{Q}) = \min_{(\mathbf{x}, \mathbf{v}) \in \mathbb{R}^n \times \mathbb{R}^n} \quad & \mathbf{x}^\top \mathbf{Q} \mathbf{x} \\
\text{s.t.} \quad & \mathbf{e}^\top \mathbf{x} = 1 \\
& \mathbf{e}^\top \mathbf{v} = n - \rho \\
& \mathbf{x}^\top \mathbf{v} = 0 \\
& \mathbf{v} \in \{0, 1\}^n \\
& \mathbf{x} \geq \mathbf{o}.
\end{aligned}
\tag{P2(\rho)}$$

Similarly, adding the redundant constraints $\mathbf{v} \leq \mathbf{e}$ and reintroducing the slack variables $\mathbf{u} \geq \mathbf{o}$

so that $\mathbf{u} + \mathbf{v} = \mathbf{e}$, $(\text{P2}(\rho))$ can be reformulated as follows:

$$\begin{aligned}
\ell_\rho(\mathbf{Q}) = \min_{(\mathbf{x}, \mathbf{u}, \mathbf{v}) \in \mathbb{R}^n \times \mathbb{R}^n \times \mathbb{R}^n} \quad & \mathbf{x}^\top \mathbf{Q} \mathbf{x} \\
\text{s.t.} \quad & \mathbf{e}^\top \mathbf{x} = 1 \\
& \mathbf{e}^\top \mathbf{u} = \rho \\
& \mathbf{u} + \mathbf{v} = \mathbf{e} \\
& \mathbf{x}^\top \mathbf{v} = 0 \\
& \mathbf{u} \in \{0, 1\}^n \\
& \mathbf{x}, \mathbf{v}, \mathbf{u} \geq \mathbf{0}.
\end{aligned} \tag{P2A(\rho)}$$

By [BP23, Bur09], $(\text{P2A}(\rho))$ admits an equivalent reformulation as the following copositive optimization problem:

$$\begin{aligned}
\ell_\rho(\mathbf{Q}) = \min_{\mathbf{Y} \in \mathcal{S}^{3n+1}} \quad & \langle \mathbf{Q}, \mathbf{Y}^{\mathbf{x}\mathbf{x}} \rangle \\
\text{s.t.} \quad & \mathbf{e}^\top \mathbf{x} = 1 \\
& \mathbf{e}^\top \mathbf{u} = \rho \\
& \mathbf{u} + \mathbf{v} = \mathbf{e} \\
& \langle \mathbf{E}, \mathbf{Y}^{\mathbf{x}\mathbf{x}} \rangle = 1 \\
& \langle \mathbf{E}, \mathbf{Y}^{\mathbf{u}\mathbf{u}} \rangle = \rho^2 \\
& \text{diag}(\mathbf{Y}^{\mathbf{u}\mathbf{u}}) = \mathbf{u} \\
& \text{diag}(\mathbf{Y}^{\mathbf{u}\mathbf{u}}) + 2 \text{diag}(\mathbf{Y}^{\mathbf{u}\mathbf{v}}) + \text{diag}(\mathbf{Y}^{\mathbf{v}\mathbf{v}}) = \mathbf{e} \\
& \mathbf{e}^\top \text{diag}(\mathbf{Y}^{\mathbf{x}\mathbf{v}}) = 0 \\
& \mathbf{Y} := \begin{pmatrix} 1 & \mathbf{x}^\top & \mathbf{u}^\top & \mathbf{v}^\top \\ \mathbf{x} & \mathbf{Y}^{\mathbf{x}\mathbf{x}} & \mathbf{Y}^{\mathbf{x}\mathbf{u}} & \mathbf{Y}^{\mathbf{x}\mathbf{v}} \\ \mathbf{u} & (\mathbf{Y}^{\mathbf{x}\mathbf{u}})^\top & \mathbf{Y}^{\mathbf{u}\mathbf{u}} & \mathbf{Y}^{\mathbf{u}\mathbf{v}} \\ \mathbf{v} & (\mathbf{Y}^{\mathbf{x}\mathbf{v}})^\top & (\mathbf{Y}^{\mathbf{u}\mathbf{v}})^\top & \mathbf{Y}^{\mathbf{v}\mathbf{v}} \end{pmatrix} \in \mathcal{CP}^{3n+1}.
\end{aligned} \tag{CP2A(\rho)}$$

Once again, the equivalence between $(\text{P2A}(\rho))$ and $(\text{CP2A}(\rho))$ is understood similarly to that between $(\text{P1A}(\rho))$ and $(\text{CP1A}(\rho))$.

4.3. Two convex relaxations and their smaller equivalents

In this section, we consider two convex relaxations arising from the conic formulations $(\text{CP1A}(\rho))$ and $(\text{CP2A}(\rho))$.

Despite the fact that $(\text{CP1A}(\rho))$ and $(\text{CP2A}(\rho))$ are convex optimization problems, they are, in general, NP-hard [MK87]. A well-known tractable convex relaxation of completely positive conic problems is obtained by replacing the intractable cone $\mathcal{CP}^d$ with the larger cone $\mathcal{DNN}^d$ of doubly nonnegative (DNN) matrices, i.e., the convex cone of symmetric $d \times d$ matrices that are both PSD and componentwise nonnegative. It follows that the optimal value of the DNN relaxation is a lower bound on the optimal value of the completely positive conic optimization problem. Here, we present the DNN relaxations arising from conic optimization problems $(\text{CP1A}(\rho))$ and $(\text{CP2A}(\rho))$. Furthermore, we establish that each of the two DNN relaxations can, in fact, be reformulated in smaller dimensions without sacrificing the strength of the corresponding lower bound. Finally, we present a theoretical comparison of the strength of the two DNN relaxations.

4.3.1. DNN relaxations of the direct MIQP formulation

In this section, we consider the DNN relaxation of the problem (CP1A(ρ)). By replacing the intractable conic constraint $Z \in \mathcal{CP}^{4n+1}$ in (CP1A(ρ)) by $Z \in \mathcal{DNN}^{4n+1}$, where $\mathcal{DNN}^{4n+1}$ denotes the cone of doubly nonnegative (DNN) matrices in $\mathcal{S}_+^{4n+1}$, we obtain the following tractable convex relaxation of (P1A(ρ)):

$$\begin{aligned}
\nu(\text{D1A}(\rho)) &:= \min_{Z \in \mathcal{S}_+^{4n+1}} \langle Q, Z^{\text{xx}} \rangle \\
\text{s.t. } & \mathbf{e}^\top \mathbf{x} = 1 \\
& \mathbf{e}^\top \mathbf{u} = \rho \\
& \mathbf{x} + \mathbf{y} = \mathbf{u} \\
& \mathbf{u} + \mathbf{v} = \mathbf{e} \\
& \langle \mathbf{E}, Z^{\text{xx}} \rangle = 1 \\
& \langle \mathbf{E}, Z^{\text{uu}} \rangle = \rho^2 \\
& \text{diag}(Z^{\text{uu}}) = \mathbf{u} \\
& \text{diag}(Z^{\text{xx}}) + \text{diag}(Z^{\text{yy}}) + \text{diag}(Z^{\text{uu}}) + \\
& 2[\text{diag}(Z^{\text{xy}}) - \text{diag}(Z^{\text{xu}}) - \text{diag}(Z^{\text{uy}})] = \mathbf{o} \\
& \text{diag}(Z^{\text{uu}}) + 2\text{diag}(Z^{\text{uv}}) + \text{diag}(Z^{\text{vv}}) = \mathbf{e} \\
& Z := \begin{pmatrix} 1 & \mathbf{x}^\top & \mathbf{u}^\top & \mathbf{v}^\top & \mathbf{y}^\top \\ \mathbf{x} & Z^{\text{xx}} & Z^{\text{xu}} & Z^{\text{xv}} & Z^{\text{xy}} \\ \mathbf{u} & (Z^{\text{xu}})^\top & Z^{\text{uu}} & Z^{\text{uv}} & Z^{\text{uy}} \\ \mathbf{v} & (Z^{\text{xv}})^\top & (Z^{\text{uv}})^\top & Z^{\text{vv}} & Z^{\text{vy}} \\ \mathbf{y} & (Z^{\text{xy}})^\top & (Z^{\text{uy}})^\top & (Z^{\text{vy}})^\top & Z^{\text{yy}} \end{pmatrix} \in \mathcal{DNN}^{4n+1}.
\end{aligned} \tag{D1A(\rho)}$$

Note that (D1A(ρ)) consists of $5n + 4$ linear equality constraints and a doubly nonnegative constraint in $\mathcal{S}_+^{4n+1}$. First, we make several observations about feasible solutions of (D1A(ρ)).

Lemma 26. *Let $Z \in \mathcal{S}_+^{4n+1}$ be (D1A(ρ))-feasible. Then, the following relations hold:*

$$\mathbf{x}\mathbf{e}^\top - Z^{\text{xu}} = Z^{\text{xv}} \geq \mathbf{O}, \tag{4.3}$$

$$-Z^{\text{xx}} + Z^{\text{xu}} = Z^{\text{xy}} \geq \mathbf{O}, \tag{4.4}$$

$$Z^{\text{xx}} - Z^{\text{xu}} - (Z^{\text{xu}})^\top + Z^{\text{uu}} = Z^{\text{yy}} \geq \mathbf{O}, \tag{4.5}$$

$$\mathbf{e}\mathbf{e}^\top - \mathbf{e}\mathbf{u}^\top - \mathbf{u}\mathbf{e}^\top + Z^{\text{uu}} = Z^{\text{vv}} \geq \mathbf{O}, \tag{4.6}$$

$$-\mathbf{e}\mathbf{x}^\top + (Z^{\text{xu}})^\top + \mathbf{e}\mathbf{u}^\top - Z^{\text{uu}} = Z^{\text{vy}} \geq \mathbf{O}, \tag{4.7}$$

$$\text{diag}(Z^{\text{xu}}) = \mathbf{x}, \tag{4.8}$$

$$\text{diag}(Z^{\text{xv}}) = \mathbf{o}, \tag{4.9}$$

$$\text{diag}(Z^{\text{vy}}) = \mathbf{o}, \tag{4.10}$$

$$\text{diag}(Z^{\text{vv}}) = \mathbf{v}, \tag{4.11}$$

$$\text{diag}(Z^{\text{uv}}) = \mathbf{o}. \tag{4.12}$$

Proof. Let $Z \in \mathcal{S}_+^{4n+1}$ be (D1A(ρ))-feasible. By using the Schur complementation, we obtain

$$\begin{pmatrix} Z^{\text{xx}} & Z^{\text{xu}} & Z^{\text{xv}} & Z^{\text{xy}} \\ (Z^{\text{xu}})^\top & Z^{\text{uu}} & Z^{\text{uv}} & Z^{\text{uy}} \\ (Z^{\text{xv}})^\top & (Z^{\text{uv}})^\top & Z^{\text{vv}} & Z^{\text{vy}} \\ (Z^{\text{xy}})^\top & (Z^{\text{uy}})^\top & (Z^{\text{vy}})^\top & Z^{\text{yy}} \end{pmatrix} = \begin{pmatrix} \mathbf{x} \\ \mathbf{u} \\ \mathbf{v} \\ \mathbf{y} \end{pmatrix} \begin{pmatrix} \mathbf{x} \\ \mathbf{u} \\ \mathbf{v} \\ \mathbf{y} \end{pmatrix}^\top + \sum_{k \in K} \begin{pmatrix} \mathbf{a}^k \\ \mathbf{b}^k \\ \mathbf{c}^k \\ \mathbf{d}^k \end{pmatrix} \begin{pmatrix} \mathbf{a}^k \\ \mathbf{b}^k \\ \mathbf{c}^k \\ \mathbf{d}^k \end{pmatrix}^\top,$$

where K is a finite set and $\{\mathbf{a}^k, \mathbf{b}^k, \mathbf{c}^k, \mathbf{d}^k\} \subset \mathbb{R}^n$ for each $k \in K$. Using $\text{diag}(\mathbf{Z}^{\mathbf{u}\mathbf{u}}) + 2 \text{diag}(\mathbf{Z}^{\mathbf{u}\mathbf{v}}) + \text{diag}(\mathbf{Z}^{\mathbf{v}\mathbf{v}}) = \mathbf{e}$, we obtain

$$\begin{aligned} u_i^2 + \sum_{k \in K} (b_i^k)^2 + 2 \left[u_i v_i + \sum_{k \in K} (b_i^k) (c_i^k) \right] + v_i^2 + \sum_{k \in K} (c_i^k)^2 &= (u_i + v_i)^2 + \sum_{k \in K} (b_i^k + c_i^k)^2 \\ &= 1 + \sum_{k \in K} (b_i^k + c_i^k)^2 = 1 \end{aligned}$$

for each $i \in \{1, \dots, n\}$, where we used $\mathbf{u} + \mathbf{v} = \mathbf{e}$ in the second equality. We conclude that

$$\mathbf{b}^k = -\mathbf{c}^k, \quad \text{for all } k \in K. \quad (4.13)$$

In a similar manner, by using $\text{diag}(\mathbf{Z}^{\mathbf{x}\mathbf{x}}) + \text{diag}(\mathbf{Z}^{\mathbf{y}\mathbf{y}}) + \text{diag}(\mathbf{Z}^{\mathbf{u}\mathbf{u}}) + 2 [\text{diag}(\mathbf{Z}^{\mathbf{x}\mathbf{y}}) - \text{diag}(\mathbf{Z}^{\mathbf{x}\mathbf{u}}) - \text{diag}(\mathbf{Z}^{\mathbf{u}\mathbf{y}})] = \mathbf{o}$, we arrive at

$$(x_i + y_i - u_i)^2 + \sum_{k \in K} (d_i^k + d_i^k - b_i^k)^2 = 0$$

for each $i \in \{1, \dots, n\}$, which, together with $\mathbf{x} + \mathbf{y} = \mathbf{u}$, implies that

$$\mathbf{a}^k + \mathbf{d}^k = \mathbf{b}^k, \quad \text{for all } k \in K. \quad (4.14)$$

Therefore,

$$\mathbf{x}\mathbf{e}^\top - \mathbf{Z}^{\mathbf{x}\mathbf{u}} = \mathbf{x}\mathbf{e}^\top - \mathbf{x}\mathbf{u}^\top - \sum_{k \in K} \mathbf{a}^k (\mathbf{b}^k)^\top = \mathbf{x}\mathbf{v}^\top + \sum_{k \in K} \mathbf{a}^k (\mathbf{c}^k)^\top = \mathbf{Z}^{\mathbf{x}\mathbf{v}} \geq \mathbf{O},$$

where we used $\mathbf{u} + \mathbf{v} = \mathbf{e}$, (4.13), and $\mathbf{Z} \in \mathcal{DN}\mathcal{N}^{4n+1}$. This establishes (4.3). Arguing similarly,

$$-\mathbf{Z}^{\mathbf{x}\mathbf{x}} + \mathbf{Z}^{\mathbf{x}\mathbf{u}} = -\mathbf{x}\mathbf{x}^\top - \sum_{k \in K} \mathbf{a}^k (\mathbf{a}^k)^\top + \mathbf{x}\mathbf{u}^\top + \sum_{k \in K} \mathbf{a}^k (\mathbf{b}^k)^\top = \mathbf{x}\mathbf{y}^\top + \sum_{k \in K} \mathbf{a}^k (\mathbf{d}^k)^\top = \mathbf{Z}^{\mathbf{x}\mathbf{y}} \geq \mathbf{O},$$

where we used $\mathbf{x} + \mathbf{y} = \mathbf{u}$, (4.14), and $\mathbf{Z} \in \mathcal{DN}\mathcal{N}^{4n+1}$. This establishes (4.4). Next,

$$\begin{aligned} \mathbf{Z}^{\mathbf{x}\mathbf{x}} - \mathbf{Z}^{\mathbf{x}\mathbf{u}} - (\mathbf{Z}^{\mathbf{x}\mathbf{u}})^\top + \mathbf{Z}^{\mathbf{u}\mathbf{u}} &= \mathbf{x}\mathbf{x}^\top + \sum_{k \in K} \mathbf{a}^k (\mathbf{a}^k)^\top - \mathbf{x}\mathbf{u}^\top - \sum_{k \in K} \mathbf{a}^k (\mathbf{b}^k)^\top \\ &\quad - \mathbf{u}\mathbf{x}^\top - \sum_{k \in K} \mathbf{b}^k (\mathbf{a}^k)^\top + \mathbf{u}\mathbf{u}^\top + \sum_{k \in K} \mathbf{b}^k (\mathbf{b}^k)^\top \\ &= \mathbf{y}\mathbf{y}^\top + \sum_{k \in K} \mathbf{d}^k (\mathbf{d}^k)^\top \\ &= \mathbf{Z}^{\mathbf{y}\mathbf{y}} \geq \mathbf{O}, \end{aligned}$$

where we used $\mathbf{x} + \mathbf{y} = \mathbf{u}$, (4.14), and $\mathbf{Z} \in \mathcal{DN}\mathcal{N}^{4n+1}$. Similarly,

$$\mathbf{e}\mathbf{e}^\top - \mathbf{e}\mathbf{u}^\top - \mathbf{u}\mathbf{e}^\top + \mathbf{Z}^{\mathbf{u}\mathbf{u}} = \mathbf{e}\mathbf{e}^\top - \mathbf{e}\mathbf{u}^\top - \mathbf{u}\mathbf{e}^\top + \mathbf{u}\mathbf{u}^\top + \sum_{k \in K} \mathbf{b}^k (\mathbf{b}^k)^\top = \mathbf{v}\mathbf{v}^\top + \sum_{k \in K} \mathbf{c}^k (\mathbf{c}^k)^\top = \mathbf{Z}^{\mathbf{v}\mathbf{v}} \geq \mathbf{O},$$

where we used $\mathbf{u} + \mathbf{v} = \mathbf{e}$, (4.13), and $\mathbf{Z} \in \mathcal{DN}\mathcal{N}^{4n+1}$, establishing (4.6). Next,

$$\begin{aligned} -\mathbf{e}\mathbf{x}^\top + (\mathbf{Z}^{\mathbf{x}\mathbf{u}})^\top + \mathbf{e}\mathbf{u}^\top - \mathbf{Z}^{\mathbf{u}\mathbf{u}} &= -\mathbf{e}\mathbf{x}^\top + \mathbf{u}\mathbf{x}^\top + \sum_{k \in K} \mathbf{b}^k (\mathbf{a}^k)^\top + \mathbf{e}\mathbf{u}^\top - \mathbf{u}\mathbf{u}^\top - \sum_{k \in K} \mathbf{b}^k (\mathbf{b}^k)^\top \\ &= -\mathbf{v}\mathbf{x}^\top + \mathbf{v}\mathbf{u}^\top + \sum_{k \in K} \mathbf{c}^k (\mathbf{d}^k)^\top \\ &= \mathbf{Z}^{\mathbf{v}\mathbf{y}} \geq \mathbf{O}, \end{aligned}$$

where we used $\mathbf{u} + \mathbf{v} = \mathbf{e}$, $\mathbf{x} + \mathbf{y} = \mathbf{u}$, (4.13), (4.14), and $\mathbf{Z} \in \mathcal{DN}^{4n+1}$, establishing (4.7).

Using $\text{diag}(\mathbf{Z}^{\mathbf{uu}}) = \mathbf{u}$, (4.3), and (4.7),

$$\begin{aligned} \mathbf{o} \leq \text{diag}(\mathbf{x}\mathbf{e}^\top - \mathbf{Z}^{\mathbf{xu}}) &= \mathbf{x} - \text{diag}(\mathbf{Z}^{\mathbf{xu}}) \quad \text{and} \\ \mathbf{o} \leq \text{diag}(-\mathbf{e}\mathbf{x}^\top + (\mathbf{Z}^{\mathbf{xu}})^\top + \mathbf{e}\mathbf{u}^\top - \mathbf{Z}^{\mathbf{uu}}) &= -\mathbf{x} + \text{diag}(\mathbf{Z}^{\mathbf{xu}}) + \mathbf{u} - \text{diag}(\mathbf{Z}^{\mathbf{uu}}) = -\mathbf{x} + \text{diag}(\mathbf{Z}^{\mathbf{xu}}), \end{aligned}$$

which together yield (4.8). By (4.8) and $\text{diag}(\mathbf{Z}^{\mathbf{uu}}) = \mathbf{u}$, it is easy to see that (4.9) and (4.10) follow from (4.3) and (4.7), respectively. Using (4.6), $\text{diag}(\mathbf{Z}^{\mathbf{uu}}) = \mathbf{u}$ and $\mathbf{u} + \mathbf{v} = \mathbf{e}$, we obtain

$$\text{diag}(\mathbf{e}\mathbf{e}^\top - \mathbf{e}\mathbf{u}^\top - \mathbf{u}\mathbf{e}^\top + \mathbf{Z}^{\mathbf{uu}}) = \mathbf{e} - 2\mathbf{u} + \mathbf{u} = \mathbf{v} = \text{diag}(\mathbf{Z}^{\mathbf{vv}}),$$

which establishes (4.11). Finally, (4.12) follows from (4.11), $\text{diag}(\mathbf{Z}^{\mathbf{uu}}) = \mathbf{u}$, $\text{diag}(\mathbf{Z}^{\mathbf{uu}}) + 2\text{diag}(\mathbf{Z}^{\mathbf{uv}}) + \text{diag}(\mathbf{Z}^{\mathbf{vv}}) = \mathbf{e}$, and $\mathbf{u} + \mathbf{v} = \mathbf{e}$. This completes the proof. $\square$

In view of Lemma 26, let us now consider the following optimization problem:

$$\begin{aligned} \nu(\text{D1B}(\rho)) &:= \min_{\mathbf{W} \in \mathcal{S}_+^{2n+1}} \langle \mathbf{Q}, \mathbf{W}^{\mathbf{xx}} \rangle \\ \text{s.t.} \quad &\mathbf{e}^\top \mathbf{x} = 1 \\ &\mathbf{e}^\top \mathbf{u} = \rho \\ &\langle \mathbf{E}, \mathbf{W}^{\mathbf{xx}} \rangle = 1 \\ &\langle \mathbf{E}, \mathbf{W}^{\mathbf{uu}} \rangle = \rho^2 \\ &\text{diag}(\mathbf{W}^{\mathbf{uu}}) = \mathbf{u} \\ &\mathbf{x}\mathbf{e}^\top - \mathbf{W}^{\mathbf{xu}} \geq \mathbf{O} \\ &-\mathbf{W}^{\mathbf{xx}} + \mathbf{W}^{\mathbf{xu}} \geq \mathbf{O} \\ &\mathbf{W}^{\mathbf{xx}} - \mathbf{W}^{\mathbf{xu}} - (\mathbf{W}^{\mathbf{xu}})^\top + \mathbf{W}^{\mathbf{uu}} \geq \mathbf{O} \\ &\mathbf{e}\mathbf{e}^\top - \mathbf{e}\mathbf{u}^\top - \mathbf{u}\mathbf{e}^\top + \mathbf{W}^{\mathbf{uu}} \geq \mathbf{O} \\ &-\mathbf{e}\mathbf{x}^\top + (\mathbf{W}^{\mathbf{xu}})^\top + \mathbf{e}\mathbf{u}^\top - \mathbf{W}^{\mathbf{uu}} \geq \mathbf{O} \\ &\mathbf{W}^{\mathbf{xx}} \geq \mathbf{O} \\ &\mathbf{W} := \begin{pmatrix} 1 & \mathbf{x}^\top & \mathbf{u}^\top \\ \mathbf{x} & \mathbf{W}^{\mathbf{xx}} & \mathbf{W}^{\mathbf{xu}} \\ \mathbf{u} & (\mathbf{W}^{\mathbf{xu}})^\top & \mathbf{W}^{\mathbf{uu}} \end{pmatrix} \succeq \mathbf{O}. \end{aligned} \tag{D1B}(\rho)$$

It is easy to see that $(\text{D1B}(\rho))$ is a convex relaxation of $(\text{P1}(\rho))$ since, for any feasible solution $(\mathbf{x}, \mathbf{u}) \in \mathbb{R}^n \times \mathbb{R}^n$ of $(\text{P1}(\rho))$,

$$\mathbf{W} = \begin{pmatrix} 1 \\ \mathbf{x} \\ \mathbf{u} \end{pmatrix} \begin{pmatrix} 1 \\ \mathbf{x} \\ \mathbf{u} \end{pmatrix}^\top \in \mathcal{S}_+^{2n+1}$$

is a feasible solution of $(\text{D1B}(\rho))$ with the same objective function value. Note that $(\text{D1B}(\rho))$ has $n + 4$ linear equality constraints, $(9/2)n^2 + (3/2)n$ linear inequality constraints, and a PSD constraint in $\mathcal{S}^{2n+1}$. It is worth noticing that the inequality constraints of $(\text{D1B}(\rho))$ are precisely given by all the RLT (reformulation-linearization technique) constraints obtained from the inequalities $\mathbf{x} \geq \mathbf{o}$, $\mathbf{x} \leq \mathbf{u}$, and $\mathbf{u} \leq \mathbf{e}$ (see, e.g., [SA99]).

First, we present some properties of feasible solutions of $(\text{D1B}(\rho))$.

Lemma 27. Let $W \in \mathcal{S}_+^{2n+1}$ be $(D1B(\rho))$ -feasible. Then,

$$W \in \mathcal{DNN}^{2n+1}, \quad (4.15)$$

$$-(W^{xu})^\top + W^{uu} \geq O, \quad (4.16)$$

$$ue^\top - W^{uu} \geq O, \quad (4.17)$$

$$u \leq e, \quad (4.18)$$

$$x \leq u. \quad (4.19)$$

Proof. Let $W \in \mathcal{S}_+^{2n+1}$ be $(D1B(\rho))$ -feasible. Since $W^{xx} \geq O$ and $-W^{xx} + W^{xu} \geq O$, we obtain $W^{xu} \geq O$. By $xe^\top - W^{xu} \geq O$, we arrive at $x \geq o$. Since $\text{diag}(W^{uu}) = u$ and $W \succeq O$, we have $u \geq o$. Then, adding the inequalities $-W^{xx} + W^{xu} \geq O$ and $W^{xx} - W^{xu} - (W^{xu})^\top + W^{uu} \geq O$ yields (4.16), which, in turn, establishes (4.15) since $W^{uu} \geq (W^{xu})^\top \geq O$. Similarly, adding the inequalities $(xe^\top - W^{xu})^\top = ex^\top - (W^{xu})^\top \geq O$ and $-ex^\top + (W^{xu})^\top + eu^\top - W^{uu} \geq O$, we arrive at (4.17). Adding (4.17) to $ee^\top - eu^\top - ue^\top + W^{uu} \geq O$, we get $e(e - u)^\top \geq O$, which yields (4.18). Similarly, using $\text{diag}(W^{uu}) = u$,

$$o \leq \text{diag}(xe^\top - W^{xu}) = x - \text{diag}(W^{xu}) \quad \text{and}$$

$$o \leq \text{diag}(-ex^\top + (W^{xu})^\top + eu^\top - W^{uu}) = -x + \text{diag}(W^{xu}) + u - \text{diag}(W^{uu}) = -x + \text{diag}(W^{xu}),$$

which together yield $\text{diag}(W^{xu}) = x$. Combining this with (4.16), we get

$$o \leq \text{diag}(-(W^{xu})^\top + W^{uu}) = -\text{diag}(W^{xu}) + \text{diag}(W^{uu}) = -x + u,$$

where we used $\text{diag}(W^{uu}) = u$. This establishes (4.19) and completes the proof. $\square$ $\square$

Our next result shows that $(D1B(\rho))$ is equivalent to the DNN relaxation $(D1A(\rho))$ of $(CP1A(\rho))$.

Proposition 28. $(D1A(\rho))$ and $(D1B(\rho))$ are equivalent to each other. Therefore, $\nu(D1A(\rho)) = \nu(D1B(\rho))$.

Proof. Let $W \in \mathcal{S}_+^{2n+1}$ be $(D1B(\rho))$ -feasible. Let us define $Z \in \mathcal{S}_+^{4n+1}$ as follows:

$$Z = \begin{pmatrix} 1 & o^\top & o^\top \\ o & I & o \\ o & O & I \\ e & O & -I \\ o & -I & I \end{pmatrix} \begin{pmatrix} 1 & x^\top & u^\top \\ x & W^{xx} & W^{xu} \\ u & (W^{xu})^\top & W^{uu} \end{pmatrix} \begin{pmatrix} 1 & o^\top & o^\top \\ o & I & o \\ o & O & I \\ e & O & -I \\ O & -I & I \end{pmatrix}^\top. \quad (4.20)$$

Therefore,

$$Z = \begin{pmatrix} 1 & x^\top & u^\top & e^\top - u^\top & -x^\top + u^\top \\ x & W^{xx} & W^{xu} & xe^\top - W^{xu} & -W^{xx} + W^{xu} \\ u & (W^{xu})^\top & W^{uu} & ue^\top - W^{uu} & -(W^{xu})^\top + W^{uu} \\ e - u & ex^\top - (W^{xu})^\top & eu^\top - W^{uu} & ee^\top - eu^\top - ue^\top + W^{uu} & -ex^\top + (W^{xu})^\top + eu^\top - W^{uu} \\ -x + u & -W^{xx} + (W^{xu})^\top & -W^{xu} + W^{uu} & -xe^\top + W^{xu} + ue^\top - W^{uu} & W^{xx} - W^{xu} - (W^{xu})^\top + W^{uu} \end{pmatrix}. \quad (4.21)$$

We claim that $Z \in \mathcal{S}_+^{4n+1}$ is $(D1A(\rho))$ -feasible. By Lemma 27, (4.20), and (4.21), we conclude that $Z \in \mathcal{DNN}^{4n+1}$. Since $v = e - u$, $y = u - x$, $Z^{xx} = W^{xx}$, and $Z^{uu} = W^{uu}$, the first seven sets of

equality constraints of (D1A(ρ)) are satisfied. Considering the eighth set of equality constraints of (D1A(ρ)) and substituting the corresponding submatrices in (4.21), we obtain

$$\begin{aligned}\text{diag} (Z^{\text{xx}}) &= \text{diag} (W^{\text{xx}}) \\ \text{diag} (Z^{\text{yy}}) &= \text{diag} (W^{\text{xx}}) - 2 \text{diag} (W^{\text{xu}}) + \text{diag} (W^{\text{uu}}) \\ \text{diag} (Z^{\text{uu}}) &= \text{diag} (W^{\text{uu}}) \\ \text{diag} (Z^{\text{xy}}) &= -\text{diag} (W^{\text{xx}}) + \text{diag} (W^{\text{xu}}) \\ \text{diag} (Z^{\text{xu}}) &= \text{diag} (W^{\text{xu}}) \\ \text{diag} (Z^{\text{uy}}) &= -\text{diag} (W^{\text{xu}}) + \text{diag} (W^{\text{uu}}),\end{aligned}$$

which implies that

$$\text{diag} (Z^{\text{xx}}) + \text{diag} (Z^{\text{yy}}) + \text{diag} (Z^{\text{uu}}) + 2 [\text{diag} (Z^{\text{xy}}) - \text{diag} (Z^{\text{xu}}) - \text{diag} (Z^{\text{uy}})] = \mathbf{0}.$$

Finally, consider the last set of equality constraints of (D1A(ρ)):

$$\text{diag} (Z^{\text{uu}}) + 2 \text{diag} (Z^{\text{uv}}) + \text{diag} (Z^{\text{vv}}) = \text{diag} (W^{\text{uu}}) + 2 (\mathbf{u} - \text{diag} (W^{\text{uu}})) + \mathbf{e} - 2 \mathbf{u} + \text{diag} (W^{\text{uu}}) = \mathbf{e},$$

which implies that $Z \in \mathcal{S}^{4n+1}$ is (D1A(ρ))-feasible and achieves the same objective function value as W in (D1B(ρ)).

Conversely, let $Z \in \mathcal{S}^{4n+1}$ be (D1A(ρ))-feasible. Let $W \in \mathcal{S}^{2n+1}$ be given by the top left 3×3 -block of Z , i.e.,

$$W = \begin{pmatrix} 1 & \mathbf{x}^\top & \mathbf{u}^\top \\ \mathbf{x} & Z^{\text{xx}} & Z^{\text{xu}} \\ \mathbf{u} & (Z^{\text{xu}})^\top & Z^{\text{uu}} \end{pmatrix}. \quad (4.22)$$

Clearly, $W \in \mathcal{DN}^{2n+1}$ and satisfies all linear equality constraints and $W^{\text{xx}} = Z^{\text{xx}} \geq \mathbf{O}$ in (D1B(ρ)). By Lemma 26, we conclude that W satisfies each of the remaining linear inequality constraints and achieves the same objective function value as Z in (D1A(ρ)). It follows that $\nu(\text{D1A}(\rho)) = \nu(\text{D1B}(\rho))$. $\square$ $\square$

Recall that (D1A(ρ)) consists of $5n + 4$ linear equality constraints and a doubly nonnegative constraint in $\mathcal{S}^{4n+1}$. In contrast, (D1B(ρ)) has $n + 4$ linear equality constraints, $(9/2)n^2 + (3/2)n$ linear inequality constraints, and a PSD constraint in $\mathcal{S}^{2n+1}$. By Proposition 28, we conclude that the lower bound arising from the DNN relaxation of (CP1A(ρ)) can be computed by solving a conic optimization problem in a much smaller dimension.

4.3.2. DNN relaxations of the complementarity constraint formulation

In this section, we focus on the DNN relaxation of the copositive optimization problem (CP2A(ρ)). By replacing the intractable conic constraint $Z \in \mathcal{CP}^{3n+1}$ in (CP2A(ρ)) by $Z \in \mathcal{DN}^{3n+1}$, we obtain the following convex optimization problem:

$$\begin{aligned}
\nu(\text{D2A}(\rho)) &:= \min_{\mathbf{Y} \in \mathcal{S}_+^{3n+1}} \langle \mathbf{Q}, \mathbf{Y}^{\text{xx}} \rangle \\
\text{s.t. } & \mathbf{e}^\top \mathbf{x} = 1 \\
& \mathbf{e}^\top \mathbf{u} = \rho \\
& \mathbf{u} + \mathbf{v} = \mathbf{e} \\
& \langle \mathbf{E}, \mathbf{Y}^{\text{xx}} \rangle = 1 \\
& \langle \mathbf{E}, \mathbf{Y}^{\text{uu}} \rangle = \rho^2 \\
& \text{diag}(\mathbf{Y}^{\text{uu}}) = \mathbf{u} \\
& \mathbf{e}^\top \text{diag}(\mathbf{Y}^{\text{xv}}) = 0 \\
& \text{diag}(\mathbf{Y}^{\text{uu}}) + 2 \text{diag}(\mathbf{Y}^{\text{uv}}) + \text{diag}(\mathbf{Y}^{\text{vv}}) = \mathbf{e} \\
& \mathbf{Y} := \begin{pmatrix} 1 & \mathbf{x}^\top & \mathbf{u}^\top & \mathbf{v}^\top \\ \mathbf{x} & \mathbf{Y}^{\text{xx}} & \mathbf{Y}^{\text{xu}} & \mathbf{Y}^{\text{xv}} \\ \mathbf{u} & (\mathbf{Y}^{\text{xu}})^\top & \mathbf{Y}^{\text{uu}} & \mathbf{Y}^{\text{uv}} \\ \mathbf{v} & (\mathbf{Y}^{\text{xv}})^\top & (\mathbf{Y}^{\text{uv}})^\top & \mathbf{Y}^{\text{vv}} \end{pmatrix} \in \mathcal{DN}^{3n+1}.
\end{aligned} \tag{D2A(\rho)}$$

Once again, it is easy to verify that $(\text{D2A}(\rho))$ is a tractable convex relaxation of $(\text{P2A}(\rho))$. Note that $(\text{D2A}(\rho))$ consists of $3n + 5$ linear equality constraints and a doubly nonnegative constraint in $\mathcal{S}^{3n+1}$.

First, we establish several properties of feasible solutions of $(\text{D2A}(\rho))$.

Lemma 29. *Let $\mathbf{Y} \in \mathcal{S}^{3n+1}$ be $(\text{D2A}(\rho))$ -feasible. Then, the following relations hold:*

$$\mathbf{x}\mathbf{e}^\top - \mathbf{Y}^{\text{xu}} = \mathbf{Y}^{\text{xv}} \geq \mathbf{O}, \tag{4.23}$$

$$\mathbf{e}\mathbf{e}^\top - \mathbf{e}\mathbf{u}^\top - \mathbf{u}\mathbf{e}^\top + \mathbf{Y}^{\text{uu}} = \mathbf{Y}^{\text{vv}} \geq \mathbf{O}, \tag{4.24}$$

$$\mathbf{u}\mathbf{e}^\top - \mathbf{Y}^{\text{uu}} = \mathbf{Y}^{\text{uv}} \geq \mathbf{O}, \tag{4.25}$$

$$\text{diag}(\mathbf{Y}^{\text{xv}}) = \mathbf{0}, \tag{4.26}$$

$$\text{diag}(\mathbf{Y}^{\text{xu}}) = \mathbf{x}, \tag{4.27}$$

$$\text{diag}(\mathbf{Y}^{\text{vv}}) = \mathbf{v}, \tag{4.28}$$

$$\text{diag}(\mathbf{Y}^{\text{uv}}) = \mathbf{0}. \tag{4.29}$$

Proof. Let $\mathbf{Y} \in \mathcal{S}^{3n+1}$ be $(\text{D2A}(\rho))$ -feasible. By using the Schur complementation, we obtain

$$\begin{pmatrix} \mathbf{Y}^{\text{xx}} & \mathbf{Y}^{\text{xu}} & \mathbf{Y}^{\text{xv}} \\ (\mathbf{Y}^{\text{xu}})^\top & \mathbf{Y}^{\text{uu}} & \mathbf{Y}^{\text{uv}} \\ (\mathbf{Y}^{\text{xv}})^\top & (\mathbf{Y}^{\text{uv}})^\top & \mathbf{Y}^{\text{vv}} \end{pmatrix} = \begin{pmatrix} \mathbf{x} \\ \mathbf{u} \\ \mathbf{v} \end{pmatrix} \begin{pmatrix} \mathbf{x} \\ \mathbf{u} \\ \mathbf{v} \end{pmatrix}^\top + \sum_{k \in K} \begin{pmatrix} \mathbf{a}^k \\ \mathbf{b}^k \\ \mathbf{c}^k \end{pmatrix} \begin{pmatrix} \mathbf{a}^k \\ \mathbf{b}^k \\ \mathbf{c}^k \end{pmatrix}^\top,$$

where K is a finite set and $\{\mathbf{a}^k, \mathbf{b}^k, \mathbf{c}^k\} \subset \mathbb{R}^n$ for each $k \in K$. Arguing similarly to the proof of Lemma 26, we obtain $\mathbf{b}^k = -\mathbf{c}^k$ for each $k \in K$. The relations (4.23), (4.24), and (4.25) can be established in a very similar manner. Since $\mathbf{e}^\top \text{diag}(\mathbf{Y}^{\text{xv}}) = 0$ and $\mathbf{Y} \in \mathcal{DN}^{3n+1}$, we obtain (4.26). By using $\mathbf{u} + \mathbf{v} = \mathbf{e}$, $\mathbf{b}^k = -\mathbf{c}^k$ for each $k \in K$, and the decomposition above, it is easy to verify that $\text{diag}(\mathbf{Y}^{\text{xu}}) + \text{diag}(\mathbf{Y}^{\text{xv}}) = \mathbf{x}$, which establishes (4.27) due to (4.26). Similarly, we obtain $\text{diag}(\mathbf{Y}^{\text{vv}}) = \mathbf{e} - 2\mathbf{u} + \text{diag}(\mathbf{Y}^{\text{uu}}) = \mathbf{e} - \mathbf{u} = \mathbf{v}$ since $\text{diag}(\mathbf{Y}^{\text{uu}}) = \mathbf{u}$, which yields (4.28). Finally, we obtain (4.29) from $\text{diag}(\mathbf{Y}^{\text{uu}}) = \mathbf{u}$, $\text{diag}(\mathbf{Y}^{\text{uu}}) + 2 \text{diag}(\mathbf{Y}^{\text{uv}}) + \text{diag}(\mathbf{Y}^{\text{vv}}) = \mathbf{e}$, $\mathbf{u} + \mathbf{v} = \mathbf{e}$, and (4.28). This completes the proof. $\square$

Once again, we can utilize Lemma 29 to obtain the following optimization problem:

$$\begin{aligned}
\nu(\text{D2B}(\rho)) := \min_{\mathbf{S} \in \mathcal{S}^{2n+1}} \quad & \langle \mathbf{Q}, \mathbf{S}^{\mathbf{x}\mathbf{x}} \rangle \\
\text{s.t.} \quad & \mathbf{e}^\top \mathbf{x} = 1 \\
& \mathbf{e}^\top \mathbf{u} = \rho \\
& \langle \mathbf{E}, \mathbf{S}^{\mathbf{x}\mathbf{x}} \rangle = 1 \\
& \langle \mathbf{E}, \mathbf{S}^{\mathbf{u}\mathbf{u}} \rangle = \rho^2 \\
& \text{diag}(\mathbf{S}^{\mathbf{u}\mathbf{u}}) = \mathbf{u} \\
& \text{diag}(\mathbf{S}^{\mathbf{x}\mathbf{u}}) = \mathbf{x} \\
& \mathbf{x}\mathbf{e}^\top - \mathbf{S}^{\mathbf{x}\mathbf{u}} \geq \mathbf{O} \\
& \mathbf{e}\mathbf{e}^\top - \mathbf{e}\mathbf{u}^\top - \mathbf{u}\mathbf{e}^\top + \mathbf{S}^{\mathbf{u}\mathbf{u}} \geq \mathbf{O} \\
& \mathbf{u}\mathbf{e}^\top - \mathbf{S}^{\mathbf{u}\mathbf{u}} \geq \mathbf{O} \\
& \mathbf{S} := \begin{pmatrix} 1 & \mathbf{x}^\top & \mathbf{u}^\top \\ \mathbf{x} & \mathbf{S}^{\mathbf{x}\mathbf{x}} & \mathbf{S}^{\mathbf{x}\mathbf{u}} \\ \mathbf{u} & (\mathbf{S}^{\mathbf{x}\mathbf{u}})^\top & \mathbf{S}^{\mathbf{u}\mathbf{u}} \end{pmatrix} \in \mathcal{DN}\mathcal{N}_{2n+1}^+.
\end{aligned} \tag{D2B(\rho)}$$

Note that $(\text{D2B}(\rho))$ has $2n + 4$ linear equality constraints, $(5/2)n^2 + (1/2)n$ linear inequality constraints, and a doubly nonnegative constraint in $\mathcal{S}^{2n+1}$. Once again, it is worth noticing that $(\text{D2B}(\rho))$ contains all the RLT constraints obtained from the inequalities $\mathbf{x} \geq \mathbf{o}$, $\mathbf{u} \geq \mathbf{o}$, and $\mathbf{u} \leq \mathbf{e}$.

Our next result establishes the equivalence between $(\text{D2A}(\rho))$ and $(\text{D2B}(\rho))$.

Proposition 30. $(\text{D2A}(\rho))$ and $(\text{D2B}(\rho))$ are equivalent to each other. Therefore, $\nu(\text{D2A}(\rho)) = \nu(\text{D2B}(\rho))$.

Proof. The proof is very similar to that of Proposition 28. Let $\mathbf{S} \in \mathcal{S}^{2n+1}$ be $(\text{D2B}(\rho))$ -feasible. Let us define $\mathbf{Y} \in \mathcal{S}^{3n+1}$ as follows:

$$\mathbf{Y} = \begin{pmatrix} 1 & \mathbf{o}^\top & \mathbf{o}^\top \\ \mathbf{o} & \mathbf{I} & \mathbf{o} \\ \mathbf{o} & \mathbf{O} & \mathbf{I} \\ \mathbf{e} & \mathbf{O} & -\mathbf{I} \end{pmatrix} \begin{pmatrix} 1 & \mathbf{x}^\top & \mathbf{u}^\top \\ \mathbf{x} & \mathbf{S}^{\mathbf{x}\mathbf{x}} & \mathbf{S}^{\mathbf{x}\mathbf{u}} \\ \mathbf{u} & (\mathbf{S}^{\mathbf{x}\mathbf{u}})^\top & \mathbf{S}^{\mathbf{u}\mathbf{u}} \end{pmatrix} \begin{pmatrix} 1 & \mathbf{o}^\top & \mathbf{o}^\top \\ \mathbf{o} & \mathbf{I} & \mathbf{o} \\ \mathbf{o} & \mathbf{O} & \mathbf{I} \\ \mathbf{e} & \mathbf{O} & -\mathbf{I} \end{pmatrix}^\top. \tag{4.30}$$

Therefore,

$$\mathbf{Y} = \begin{pmatrix} 1 & \mathbf{x}^\top & \mathbf{u}^\top & \mathbf{e}^\top - \mathbf{u}^\top \\ \mathbf{x} & \mathbf{S}^{\mathbf{x}\mathbf{x}} & \mathbf{S}^{\mathbf{x}\mathbf{u}} & \mathbf{x}\mathbf{e}^\top - \mathbf{S}^{\mathbf{x}\mathbf{u}} \\ \mathbf{u} & (\mathbf{S}^{\mathbf{x}\mathbf{u}})^\top & \mathbf{S}^{\mathbf{u}\mathbf{u}} & \mathbf{u}\mathbf{e}^\top - \mathbf{S}^{\mathbf{u}\mathbf{u}} \\ \mathbf{e} - \mathbf{u} & \mathbf{e}\mathbf{x}^\top - (\mathbf{S}^{\mathbf{x}\mathbf{u}})^\top & \mathbf{e}\mathbf{u}^\top - \mathbf{S}^{\mathbf{u}\mathbf{u}} & \mathbf{e}\mathbf{e}^\top - \mathbf{e}\mathbf{u}^\top - \mathbf{u}\mathbf{e}^\top + \mathbf{S}^{\mathbf{u}\mathbf{u}} \end{pmatrix}. \tag{4.31}$$

We claim that $\mathbf{Y} \in \mathcal{S}^{3n+1}$ is $(\text{D2A}(\rho))$ -feasible. Since $\text{diag}(\mathbf{S}^{\mathbf{u}\mathbf{u}}) = \mathbf{u}$, and $\mathbf{S}^{\mathbf{u}\mathbf{u}} - \mathbf{u}\mathbf{u}^\top \succeq \mathbf{O}$, we obtain $\mathbf{u} \leq \mathbf{e}$. Therefore, by (4.30) and (4.31), we conclude that $\mathbf{Y} \in \mathcal{DN}\mathcal{N}^{4n+1}$. Since $\mathbf{v} = \mathbf{e} - \mathbf{u}$, $\mathbf{Y}^{\mathbf{x}\mathbf{x}} = \mathbf{S}^{\mathbf{x}\mathbf{x}}$, $\mathbf{Y}^{\mathbf{x}\mathbf{u}} = \mathbf{S}^{\mathbf{x}\mathbf{u}}$, and $\mathbf{Y}^{\mathbf{u}\mathbf{u}} = \mathbf{S}^{\mathbf{u}\mathbf{u}}$, the first six sets of equality constraints of $(\text{D2A}(\rho))$ are satisfied. Consider the seventh set of equality constraints of $(\text{D2A}(\rho))$:

$$\mathbf{e}^\top \text{diag}(\mathbf{Y}^{\mathbf{x}\mathbf{v}}) = \mathbf{e}^\top \text{diag}(\mathbf{x}\mathbf{e}^\top - \mathbf{S}^{\mathbf{x}\mathbf{u}}) = \mathbf{e}^\top [\mathbf{x} - \text{diag}(\mathbf{S}^{\mathbf{x}\mathbf{u}})] = 0,$$

where we used $\text{diag}(\mathbf{S}^{\mathbf{x}\mathbf{u}}) = \mathbf{x}$. Finally, the last set of equality constraints of $(\text{D2A}(\rho))$ can be directly verified using (4.31). Therefore, $\mathbf{Y} \in \mathcal{S}^{3n+1}$ is $(\text{D2A}(\rho))$ -feasible and achieves the same objective function value as $\mathbf{S}$ in $(\text{D2B}(\rho))$.

Conversely, let $\mathbf{Y} \in \mathcal{S}^{3n+1}$ be $(D2A(\rho))$ -feasible. Let $\mathbf{S} \in \mathcal{S}^{2n+1}$ be given by the top left 3×3 -block of $\mathbf{Z}$, i.e.,

$$\mathbf{S} = \begin{pmatrix} 1 & \mathbf{x}^\top & \mathbf{u}^\top \\ \mathbf{x} & \mathbf{Y}^{\mathbf{x}\mathbf{x}} & \mathbf{Y}^{\mathbf{x}\mathbf{u}} \\ \mathbf{u} & (\mathbf{Y}^{\mathbf{x}\mathbf{u}})^\top & \mathbf{Y}^{\mathbf{u}\mathbf{u}} \end{pmatrix}. \quad (4.32)$$

Clearly, $\mathbf{S} \in \mathcal{DN}^{2n+1}$ and satisfies all linear equality and inequality constraints in $(D2B(\rho))$ by Lemma 29. Furthermore, $\mathbf{S}$ has the same objective function value in $(D2B(\rho))$ as that of $\mathbf{Y}$ in $(D2A(\rho))$. It follows that $\nu(D2A(\rho)) = \nu(D2B(\rho))$. $\square$ $\square$

A comparison of $(D2A(\rho))$ and $(D2B(\rho))$ reveals that the former consists of $3n+5$ linear equality constraints and a doubly nonnegative constraint in $\mathcal{S}^{3n+1}$ whereas the latter has $2n+4$ linear equality constraints, $(5/2)n^2 + (1/2)n$ linear inequality constraints, and a doubly nonnegative constraint in $\mathcal{S}^{2n+1}$. Proposition 30 implies that the lower bound arising from the DNN relaxation of $(CP2A(\rho))$ can be computed by solving a conic optimization problem in a smaller dimension.

4.3.3. Comparing the two DNN relaxations

Now we compare the lower bounds arising from the two DNN relaxations of $(CP1A(\rho))$ and $(CP2A(\rho))$.

Proposition 31. *$(D1A(\rho))$ (or equivalently, $(D1B(\rho))$) is at least as tight as $(D2A(\rho))$ (or equivalently, $(D2B(\rho))$), i.e.,*

$$\nu(D2A(\rho)) = \nu(D2B(\rho)) \leq \nu(D1A(\rho)) = \nu(D1B(\rho)) \leq \ell_\rho(\mathbf{Q}).$$

Proof. Let $\mathbf{Z} \in \mathcal{S}^{4n+1}$ be $(D1A(\rho))$ -feasible. Let $\mathbf{Y} \in \mathcal{S}^{3n+1}$ be given by the top left 4×4 -block of $\mathbf{Z}$, i.e.,

$$\mathbf{Y} = \begin{pmatrix} 1 & \mathbf{x}^\top & \mathbf{v}^\top & \mathbf{u}^\top \\ \mathbf{x} & \mathbf{Z}^{\mathbf{x}\mathbf{x}} & \mathbf{Z}^{\mathbf{x}\mathbf{v}} & \mathbf{Z}^{\mathbf{x}\mathbf{u}} \\ \mathbf{u} & (\mathbf{Z}^{\mathbf{x}\mathbf{u}})^\top & \mathbf{Z}^{\mathbf{u}\mathbf{u}} & \mathbf{Z}^{\mathbf{u}\mathbf{v}} \\ \mathbf{v} & (\mathbf{Z}^{\mathbf{x}\mathbf{v}})^\top & (\mathbf{Z}^{\mathbf{u}\mathbf{v}})^\top & \mathbf{Z}^{\mathbf{v}\mathbf{v}} \end{pmatrix}.$$

We claim that $\mathbf{Y} \in \mathcal{S}^{3n+1}$ is $(D2A(\rho))$ -feasible. Clearly, $\mathbf{Y} \in \mathcal{DN}^{3n+1}$ and all constraints of $(D2A(\rho))$ are satisfied by Lemma 26. Therefore, for every $(D1A(\rho))$ -feasible solution, there exists a corresponding $(D2A(\rho))$ -feasible solution with the same objective function value. By Propositions 28 and 30, we conclude that $\nu(D2A(\rho)) = \nu(D2B(\rho)) \leq \nu(D1A(\rho)) = \nu(D1B(\rho)) \leq \ell_\rho(\mathbf{Q})$. $\square$ $\square$

Our next example illustrates that we can have $\nu(D2A(\rho)) < \nu(D1A(\rho))$. i.e., $(D2A(\rho))$ can be strictly weaker than $(D1A(\rho))$.

Example 6. *Consider the following instance, where $n = 6$ and $\rho = 3$:*

$$\mathbf{Q} = \begin{pmatrix} 2.6947 & -0.2028 & -1.1144 & -2.4230 & -2.1633 & 0.7710 \\ -0.2028 & 5.1998 & 0.5005 & -2.0941 & 0.7828 & -3.3611 \\ -1.1144 & 0.5005 & 5.2918 & 1.4119 & -1.1526 & -1.9723 \\ -2.4230 & -2.0941 & 1.4119 & 4.1140 & -0.2025 & 0.3132 \\ -2.1633 & 0.7828 & -1.1526 & -0.2025 & 7.6645 & -0.0170 \\ 0.7710 & -3.3611 & -1.9723 & 0.3132 & -0.0170 & 3.3040 \end{pmatrix}.$$

For this instance, we have $\nu(D2A(\rho)) = \nu(D2B(\rho)) \approx 0.1320 < \nu(D1A(\rho)) = \nu(D1B(\rho)) \approx 0.1333$. In addition, we observe that the computed optimal solution $\mathbf{S}^$ of $(D2B(\rho))$ is not feasible for $(D1B(\rho))$. For instance, $(S_{31}^{\mathbf{x}\mathbf{u}})^* - (S_{31}^{\mathbf{x}\mathbf{x}})^* \approx -0.00137 < 0$.*

4.4. Generating nontrivial instances of sparse StQPs

In this section, we discuss procedures for generating nontrivial instances of the sparse StQP problem $(\text{StQP}(\rho))$. To this end, we say that an instance of $(\text{StQP}(\rho))$ is *nontrivial* if the sparsity constraint is violated by each optimal solution of the corresponding instance of StQP given by (StQP) (or, equivalently, $(\text{StQP}(\rho))$ with $\rho = n$). For a nontrivial instance of the sparse StQP, it follows that the sparsity constraint cuts off all optimal solutions of the corresponding StQP instance without the sparsity constraints.

We aim to construct a nontrivial instance of $(\text{StQP}(\rho))$ in two steps. First, we discuss how to construct an instance of (StQP) with a prespecified *unique* optimal solution. We then choose the sparsity parameter ρ to be strictly smaller than the sparsity of the designated unique optimal solution of (StQP) . It follows that the resulting instance of $(\text{StQP}(\rho))$ is nontrivial since the unique optimal solution of (StQP) is cut off by the sparsity constraint.

In Section 4.4.1, we consider generating an instance of (StQP) with a unique optimal solution. We then present simple algorithms for generating nontrivial instances of $(\text{StQP}(\rho))$ in Section 4.4.2.

4.4.1. Generating StQP instances with a unique optimal solution

Let us next focus on generating an instance of (StQP) with a unique optimal solution. To that end, we first recall that an StQP can be reformulated as the following copositive optimization problem [BDDK⁺00]:

$$\ell_n(\mathbf{Q}) = \min_{\mathbf{x} \in \mathbb{R}^n} \{ \mathbf{x}^\top \mathbf{Q} \mathbf{x} : \mathbf{x} \in F_n \} = \min_{\mathbf{X} \in \mathcal{S}^n} \{ \langle \mathbf{Q}, \mathbf{X} \rangle : \langle \mathbf{E}, \mathbf{X} \rangle = 1, \mathbf{X} \in \mathcal{CP}_n \},$$

where F_n is given by (4.1) and $\mathbf{E} = \mathbf{e}\mathbf{e}^\top \in \mathcal{S}^n$.

The doubly nonnegative (DNN) relaxation of an StQP is therefore given by

$$\mu(\mathbf{Q}) = \min_{\mathbf{X} \in \mathcal{S}^n} \{ \langle \mathbf{Q}, \mathbf{X} \rangle : \langle \mathbf{E}, \mathbf{X} \rangle = 1, \mathbf{X} \in \mathcal{DN}^n \}. \quad (4.33)$$

Note that $\mu(\mathbf{Q}) \leq \ell_n(\mathbf{Q})$. For a given instance of StQP, we say that its DNN relaxation is *exact* if $\mu(\mathbf{Q}) = \ell_n(\mathbf{Q})$ and *inexact* otherwise.

We will review the following useful results, which will play a fundamental role in our instance construction procedures.

Theorem 32. *Let $\mathbf{x} \in F_n$, where F_n is given by (4.1).*

- (i) *$\mathbf{x}$ is a globally optimal solution of (StQP) if and only if there exist $\mathbf{M} \in \mathbf{bd} \mathcal{COP}^n$ and $\lambda \in \mathbb{R}$ such that $\mathbf{x} \in \mathbf{V}^{\mathbf{M}}$ and $\mathbf{Q} = \mathbf{M} + \lambda \mathbf{E}$, where $\mathbf{bd} \mathcal{COP}^n$ and $\mathbf{V}^{\mathbf{M}}$ are given by*

$$\mathbf{bd} \mathcal{COP}^n = \{ \mathbf{M} \in \mathcal{COP}^n : \exists \mathbf{u} \in F_n \text{ s.t. } \mathbf{u}^\top \mathbf{M} \mathbf{u} = 0 \}, \quad (4.34)$$

and

$$\mathbf{V}^{\mathbf{M}} = \{ \mathbf{u} \in F_n : \mathbf{u}^\top \mathbf{M} \mathbf{u} = 0 \}. \quad (4.35)$$

respectively. Furthermore, in this case, λ is the optimal value of (StQP) , and the set of optimal solutions of (StQP) is given by $\mathbf{V}^{\mathbf{M}}$.

- (ii) *$\mathbf{x}$ is a globally optimal solution of (StQP) and its doubly nonnegative relaxation is exact if and only if there exists $\mathbf{R} \succeq \mathbf{O}$, $\mathbf{N} \succeq \mathbf{O}$, and $\lambda \in \mathbb{R}$ such that $\mathbf{x}^\top \mathbf{N} \mathbf{x} = 0$ and $\mathbf{Q} = (\mathbf{I} - \mathbf{e}\mathbf{x}^\top)\mathbf{R}(\mathbf{I} - \mathbf{x}\mathbf{e}^\top) + \mathbf{N} + \lambda \mathbf{E}$.*

(iii) $\mathbf{x}$ is a globally optimal solution of (StQP) and its doubly nonnegative relaxation is inexact if and only if there exist $\mathbf{M} \in \mathbf{bd} \mathcal{COP}^n \setminus \mathcal{SPN}^n$ and $\lambda \in \mathbb{R}$ such that $\mathbf{x} \in \mathbf{V}^{\mathbf{M}}$ and $\mathbf{Q} = \mathbf{M} + \lambda \mathbf{E}$, where $\mathbf{bd} \mathcal{COP}^n$ and $\mathbf{V}^{\mathbf{M}}$ are given by (4.34) and (4.35), respectively.

Proof. The reader is referred to [Bom97] for (i), and to [GY22] for (ii) and (iii). $\square$ $\square$

In view of Theorem 32, we first give a simple recipe to construct an instance of (StQP) such that it has a unique optimal solution and its doubly nonnegative relaxation is exact.

Proposition 33. *Let $\mathbf{x} \in F_n$, where F_n is given by (4.1). Let $\mathbb{A} = \{j \in \{1, \dots, n\} : x_j > 0\}$ and $\mathbb{B} = \{j \in \{1, \dots, n\} : x_j = 0\}$. Then, for any $\mathbf{R} \succ \mathbf{O}$, any $\mathbf{N} \in \mathcal{S}^n$ is such that $\mathbf{N}_{\mathbb{A}\mathbb{A}} = \mathbf{O}$, $\mathbf{N}_{\mathbb{A}\mathbb{B}} \geq \mathbf{O}$, and $\mathbf{N}_{\mathbb{B}\mathbb{B}} \geq \mathbf{O}$, and any $\lambda \in \mathbb{R}$, if $\mathbf{Q} = (\mathbf{I} - \mathbf{e}\mathbf{x}^\top)\mathbf{R}(\mathbf{I} - \mathbf{x}\mathbf{e}^\top) + \mathbf{N} + \lambda \mathbf{E}$, then $\mathbf{x}$ is the unique globally optimal solution of (StQP) and its DNN relaxation is exact.*

Proof. It follows from the hypothesis and Theorem 32(ii) that $\mathbf{x}$ is a globally optimal solution of (StQP) and its DNN relaxation is exact. We next show the uniqueness. For any $\mathbf{y} \in F_n$, where F_n is given by (4.1), we have

$$\mathbf{y}^\top \mathbf{Q} \mathbf{y} = (\mathbf{y} - \mathbf{x})^\top \mathbf{R}(\mathbf{y} - \mathbf{x}) + \mathbf{y}^\top \mathbf{N} \mathbf{y} + \lambda \geq \lambda = \mathbf{x}^\top \mathbf{Q} \mathbf{x},$$

where we used $\mathbf{x} \in F_n$, $\mathbf{y} \in F_n$, $\mathbf{R} \succ \mathbf{O}$, $\mathbf{y} \geq \mathbf{o}$, and $\mathbf{N} \geq \mathbf{O}$. Since $\mathbf{R} \succ \mathbf{O}$, the inequality above holds with equality if and only if $\mathbf{y} = \mathbf{x}$. We conclude that $\mathbf{x}$ is the unique globally optimal solution of (StQP). $\square$ $\square$

In Proposition 33, if one chooses $\mathbf{N} = \mathbf{O}$ and $\lambda \geq \mathbf{o}$, we obtain $\mathbf{Q} \succeq \mathbf{O}$, which implies that the resulting instance of (StQP) is a convex optimization problem. On the other hand, by choosing $\mathbf{N}$ and λ in such a way that $\mathbf{Q} \not\succeq \mathbf{O}$, Proposition 33 allows us to construct a nonconvex instance of StQP that admits an exact DNN relaxation. We will utilize both of these observations in our computational experiments.

Next, we consider constructing an instance of StQP with a unique optimal solution but an inexact DNN relaxation. In contrast with the previous case, such a construction is more involved and requires additional care. Our next result presents such a procedure.

Proposition 34. *Let $\mathbf{x} \in F_n$, where F_n is given by (4.1). Let $\mathbb{A} = \{j \in \{1, \dots, n\} : x_j > 0\}$ and $\mathbb{B} = \{j \in \{1, \dots, n\} : x_j = 0\}$. Let $\mathbf{Q} = (\mathbf{I} - \mathbf{e}\mathbf{x}^\top)\mathbf{R}(\mathbf{I} - \mathbf{x}\mathbf{e}^\top) + \mathbf{N} + \lambda \mathbf{E}$, where $\lambda \in \mathbb{R}$, $\mathbf{R} \in \mathcal{S}^n$ is such that $\mathbf{R}_{\mathbb{A}\mathbb{A}} \succeq \mathbf{O}$, $\mathbf{R}_{\mathbb{B}\mathbb{B}} \in \mathcal{COP}^{|\mathbb{B}|}$, $\mathbf{R}_{\mathbb{A}\mathbb{B}} = \mathbf{O}$, and $\mathbf{N} \in \mathcal{S}^n$ is such that $\mathbf{N}_{\mathbb{A}\mathbb{A}} = \mathbf{O}$, $\mathbf{N}_{\mathbb{A}\mathbb{B}} \geq \mathbf{O}$, and $\mathbf{N}_{\mathbb{B}\mathbb{B}} \geq \mathbf{O}$.*

(i) $\mathbf{x}$ is a globally optimal solution of (StQP). Furthermore, if $\mathbf{R}_{\mathbb{A}\mathbb{A}} \succ \mathbf{O}$, then $\mathbf{x}$ is the unique globally optimal solution.

(ii) If $\mathbf{R}_{\mathbb{B}\mathbb{B}} \in \mathcal{COP}^{|\mathbb{B}|} \setminus \mathcal{SPN}^{|\mathbb{B}|}$ and $\mathbf{N} = \mathbf{O}$, there exists an $\epsilon > 0$ such that, for all $\mathbf{R}_{\mathbb{A}\mathbb{A}} \succ \mathbf{O}$ with $\|\mathbf{R}_{\mathbb{A}\mathbb{A}}\| < \epsilon$, where $\|\cdot\|$ denotes the operator norm, $\mathbf{x}$ is the unique globally optimal solution of (StQP) and its doubly nonnegative relaxation is inexact.

Proof. (i) Let $\mathbf{M} = (\mathbf{I} - \mathbf{e}\mathbf{x}^\top)\mathbf{R}(\mathbf{I} - \mathbf{x}\mathbf{e}^\top) + \mathbf{N}$. By Theorem 32(i), it suffices to show that $\mathbf{M} \in \mathbf{bd} \mathcal{COP}^n$ and $\mathbf{x} \in \mathbf{V}^{\mathbf{M}}$, where $\mathbf{V}^{\mathbf{M}}$ is given by (4.35). For any $\mathbf{y} \in F_n$,

$$\mathbf{y}^\top \mathbf{M} \mathbf{y} = (\mathbf{y} - \mathbf{x})^\top \mathbf{R}(\mathbf{y} - \mathbf{x}) + \mathbf{y}^\top \mathbf{N} \mathbf{y} = (\mathbf{y}_{\mathbb{A}} - \mathbf{x}_{\mathbb{A}})^\top \mathbf{R}_{\mathbb{A}\mathbb{A}}(\mathbf{y}_{\mathbb{A}} - \mathbf{x}_{\mathbb{A}}) + \mathbf{y}_{\mathbb{B}}^\top \mathbf{R}_{\mathbb{B}\mathbb{B}} \mathbf{y}_{\mathbb{B}} + \mathbf{y}^\top \mathbf{N} \mathbf{y} \geq 0 = \mathbf{x}^\top \mathbf{M} \mathbf{x}, \quad (4.36)$$

where we used $\mathbf{y} \in F_n$ in the first equality, $\mathbf{R}_{\mathbb{A}\mathbb{B}} = \mathbf{O}$ and $\mathbf{x}_{\mathbb{B}} = \mathbf{o}$ in the second equality, $\mathbf{R}_{\mathbb{A}\mathbb{A}} \succeq \mathbf{O}$, $\mathbf{R}_{\mathbb{B}\mathbb{B}} \in \mathcal{COP}^{|\mathbb{B}|}$, $\mathbf{N} \geq \mathbf{O}$, and $\mathbf{y} \geq \mathbf{o}$ to derive the inequality, and $\mathbf{x} \in F_n$ in the

last equality. We conclude that $\mathbf{M} \in \mathbf{bd} \mathcal{COP}^n$ and $\mathbf{x} \in \mathbf{V}^{\mathbf{M}}$. By Theorem 32(i), $\mathbf{x}$ is a globally optimal solution of (StQP). If, in addition, $\mathbf{R}_{\mathbb{A}\mathbb{A}} \succ \mathbf{O}$, then it follows from (4.36) that $\mathbf{y}^\top \mathbf{M} \mathbf{y} = 0$ if and only if $(\mathbf{y}_{\mathbb{A}} - \mathbf{x}_{\mathbb{A}})^\top \mathbf{R}_{\mathbb{A}\mathbb{A}} (\mathbf{y}_{\mathbb{A}} - \mathbf{x}_{\mathbb{A}}) = \mathbf{y}_{\mathbb{B}}^\top \mathbf{R}_{\mathbb{B}\mathbb{B}} \mathbf{y}_{\mathbb{B}} = \mathbf{y}^\top \mathbf{N} \mathbf{y} = 0$. Since $\mathbf{R}_{\mathbb{A}\mathbb{A}} \succ \mathbf{O}$, we conclude that $\mathbf{y}_{\mathbb{A}} = \mathbf{x}_{\mathbb{A}}$. Since $\mathbf{x} \in F_n$, $\mathbf{x}_{\mathbb{B}} = \mathbf{o}$, and $\mathbf{y}_{\mathbb{A}} = \mathbf{x}_{\mathbb{A}}$, we obtain $\mathbf{1} = \mathbf{e}^\top \mathbf{x} = \mathbf{e}_{\mathbb{A}}^\top \mathbf{x}_{\mathbb{A}} + \mathbf{e}_{\mathbb{B}}^\top \mathbf{x}_{\mathbb{B}} = \mathbf{e}_{\mathbb{A}}^\top \mathbf{x}_{\mathbb{A}} = \mathbf{e}_{\mathbb{A}}^\top \mathbf{y}_{\mathbb{A}}$, which implies that $\mathbf{y}_{\mathbb{B}} = \mathbf{o}$ since $\mathbf{y} \in F_n$. It follows that $\mathbf{y} = \mathbf{x}$. Therefore, $\mathbf{V}^{\mathbf{M}} = \{\mathbf{x}\}$. By Theorem 32(i), $\mathbf{x}$ is the unique globally optimal solution of (StQP).

- (ii) Let $\mathbf{M} = (\mathbf{I} - \mathbf{e}\mathbf{x}^\top)\mathbf{R}(\mathbf{I} - \mathbf{x}\mathbf{e}^\top)$. By Theorem 32(ii), it suffices to show that $\mathbf{M} \in \mathbf{bd} \mathcal{COP}^n \setminus \mathcal{SPN}^n$ and $\mathbf{x} \in \mathbf{V}^{\mathbf{M}}$. By (i), $\mathbf{M} \in \mathbf{bd} \mathcal{COP}_n$ and $\mathbf{x} \in \mathbf{V}^{\mathbf{M}}$. Since $\mathbf{R}_{\mathbb{B}\mathbb{B}} \in \mathcal{COP}^{|\mathbb{B}|} \setminus \mathcal{SPN}^{|\mathbb{B}|}$, $\mathcal{CP}_{|\mathbb{B}|}$ and $\mathcal{DN}^{|\mathbb{B}|}$ are the dual cones of $\mathcal{COP}^{|\mathbb{B}|}$ and $\mathcal{SPN}^{|\mathbb{B}|}$, respectively, it follows from (4.36) that there exists $\mathbf{T} \in \mathcal{DN}^{|\mathbb{B}|} \setminus \mathcal{CP}^{|\mathbb{B}|}$ such that $\langle \mathbf{R}_{\mathbb{B}\mathbb{B}}, \mathbf{T} \rangle = -\delta < 0$. Let $\mathbf{U} \in \mathcal{S}^n$ be such that $\mathbf{U}_{\mathbb{A}\mathbb{A}} = \mathbf{O}$, $\mathbf{U}_{\mathbb{A}\mathbb{B}} = \mathbf{O}$, and $\mathbf{U}_{\mathbb{B}\mathbb{B}} = \mathbf{T}$. Clearly, $\mathbf{U} \in \mathcal{DN}^n$. Furthermore,

$$\begin{aligned} \langle \mathbf{M}, \mathbf{U} \rangle &= \langle \mathbf{M}_{\mathbb{B}\mathbb{B}}, \mathbf{T} \rangle \\ &= \left\langle \left(\mathbf{x}_{\mathbb{A}}^\top \mathbf{R}_{\mathbb{A}\mathbb{A}} \mathbf{x}_{\mathbb{A}} \right) \mathbf{e}_{\mathbb{B}} \mathbf{e}_{\mathbb{B}}^\top + \mathbf{R}_{\mathbb{B}\mathbb{B}}, \mathbf{T} \right\rangle \\ &= \left(\mathbf{x}_{\mathbb{A}}^\top \mathbf{R}_{\mathbb{A}\mathbb{A}} \mathbf{x}_{\mathbb{A}} \right) \left(\mathbf{e}_{\mathbb{B}}^\top \mathbf{T} \mathbf{e}_{\mathbb{B}} \right) - \delta \\ &\leq \|\mathbf{R}_{\mathbb{A}\mathbb{A}}\| \mu - \delta, \end{aligned}$$

where $\mu = \mathbf{e}_{\mathbb{B}}^\top \mathbf{T} \mathbf{e}_{\mathbb{B}} > 0$ and we used $\|\mathbf{x}_{\mathbb{A}}\| \leq 1$ to derive the last inequality. We conclude that $\langle \mathbf{M}, \mathbf{U} \rangle < 0$ provided that $\|\mathbf{R}_{\mathbb{A}\mathbb{A}}\| < \epsilon$, where $\epsilon = \delta/\mu > 0$. Therefore, for any $\mathbf{R}_{\mathbb{A}\mathbb{A}} \succeq \mathbf{O}$ such that $\|\mathbf{R}_{\mathbb{A}\mathbb{A}}\| < \epsilon$, we obtain $\mathbf{M} \in \mathbf{bd} \mathcal{COP}^n \setminus \mathcal{SPN}^n$ since $\mathbf{U} \in \mathcal{DN}^n$. By Theorem 32(ii), it follows that the DNN relaxation of (StQP) is inexact. If, in addition, $\mathbf{R}_{\mathbb{A}\mathbb{A}} \succ \mathbf{O}$, then $\mathbf{x}$ is the unique globally optimal solution of (StQP) by part (i). This concludes the proof. $\square$

4.4.2. Two algorithms for generating nontrivial sparse StQP instances

In this section, we present two algorithms for generating nontrivial instances of (StQP(ρ)). Given $\mathbf{x} \in F_n$ with $\|\mathbf{x}\|_0 \geq 2$, both algorithms initially generate an instance of (StQP) such that $\mathbf{x}$ is the unique optimal solution. Then, a nontrivial instance of (StQP(ρ)) is constructed by choosing $\rho \in \{1, \dots, \|\mathbf{x}\|_0 - 1\}$.

The details of our algorithms are presented in Algorithm 2 and Algorithm 3. The two algorithms differ only in terms of whether the StQP instance generated in the first stage admits an exact or inexact DNN relaxation.

Data: $n \geq 2$, $\mathbf{x} \in F_n$ such that $\|\mathbf{x}\|_0 \geq 2$, $\lambda \in \mathbb{R}$

Result: A nontrivial instance of (StQP(ρ))

$\mathbb{A} \leftarrow \{j \in \{1, \dots, n\} : x_j > 0\}$, $\mathbb{B} \leftarrow \{j \in \{1, \dots, n\} : x_j = 0\}$

Choose an arbitrary $\mathbf{R} \in \mathcal{S}^n$ such that $\mathbf{R} \succ \mathbf{O}$

Choose an arbitrary $\mathbf{N} \in \mathcal{S}^n$ such that $\mathbf{N}_{\mathbb{A}\mathbb{A}} = \mathbf{O}$, $\mathbf{N}_{\mathbb{A}\mathbb{B}} \geq \mathbf{O}$, and $\mathbf{N}_{\mathbb{B}\mathbb{B}} \geq \mathbf{O}$

$\mathbf{Q} \leftarrow (\mathbf{I} - \mathbf{e}\mathbf{x}^\top)\mathbf{R}(\mathbf{I} - \mathbf{x}\mathbf{e}^\top) + \mathbf{N} + \lambda \mathbf{E}$

Choose an arbitrary $\rho \in \{1, \dots, \|\mathbf{x}\|_0 - 1\}$.

Algorithm 2: Nontrivial instance of (StQP(ρ))
with (StQP) admitting an exact DNN relaxation

Our first algorithm, given by Algorithm 2, requires as input $n \geq 2$, $\mathbf{x} \in F_n$ such that $\|\mathbf{x}\|_0 \geq 2$, and $\lambda \in \mathbb{R}$, and generates an instance of (StQP(ρ)). By Proposition 33, the instance generated

by Algorithm 2 is a nontrivial instance of $(\text{StQP}(\rho))$ such that $\mathbf{x}$ is the unique optimal solution of (StQP) , which admits an exact DNN relaxation. We remark that each step of Algorithm 2 can be implemented in a fairly straightforward way.

Data: $n \geq 7$, $\mathbf{x} \in F_n$ such that $2 \leq \|\mathbf{x}\|_0 \leq n - 5$, $\lambda \in \mathbb{R}$
Result: A nontrivial instance of $(\text{StQP}(\rho))$
 $\mathbb{A} \leftarrow \{j \in \{1, \dots, n\} : x_j > 0\}$, $\mathbb{B} \leftarrow \{j \in \{1, \dots, n\} : x_j = 0\}$
Construct $\mathbf{R} \in \mathcal{S}^n$ using the following steps.
 $\mathbf{R}_{\mathbb{A}\mathbb{B}} \leftarrow \mathbf{O}$
Choose an arbitrary $\mathbf{R}_{\mathbb{B}\mathbb{B}} \in \mathcal{COP}^{|\mathbb{B}|} \setminus \mathcal{SPN}^{|\mathbb{B}|}$.
Choose an arbitrary $\mathbf{T} \in \mathcal{DNN}^{|\mathbb{B}|} \setminus \mathcal{CP}^{|\mathbb{B}|}$ such that $\langle \mathbf{R}_{\mathbb{B}\mathbb{B}}, \mathbf{T} \rangle < 0$.
 $\delta \leftarrow -\langle \mathbf{R}_{\mathbb{B}\mathbb{B}}, \mathbf{T} \rangle$, $\mu \leftarrow \mathbf{e}_{\mathbb{B}}^\top \mathbf{T} \mathbf{e}_{\mathbb{B}}$, $\epsilon \leftarrow \frac{\delta}{\mu}$
Choose any $\mathbf{R}_{\mathbb{A}\mathbb{A}} \succ \mathbf{O}$ such that $\|\mathbf{R}_{\mathbb{A}\mathbb{A}}\| < \epsilon$.
 $\mathbf{Q} \leftarrow (\mathbf{I} - \mathbf{e}\mathbf{x}^\top)\mathbf{R}(\mathbf{I} - \mathbf{x}\mathbf{e}^\top) + \lambda \mathbf{E}$
Choose an arbitrary $\rho \in \{1, \dots, \|\mathbf{x}\|_0 - 1\}$.
Algorithm 3: Nontrivial instance of $(\text{StQP}(\rho))$
with (StQP) admitting an inexact DNN relaxation

Our second algorithm, given by Algorithm 3, takes as input $n \geq 7$, $\mathbf{x} \in F_n$ such that $2 \leq \|\mathbf{x}\|_0 \leq n - 5$, and $\lambda \in \mathbb{R}$, and outputs an instance of $(\text{StQP}(\rho))$. Similarly, it follows from Proposition 34 that the instance generated by Algorithm 3 is a nontrivial instance of $(\text{StQP}(\rho))$ such that $\mathbf{x}$ is the unique optimal solution of (StQP) , which admits an inexact DNN relaxation.

In contrast with Algorithm 2, the practical implementation of Algorithm 3 merits further discussion. Step 4 of Algorithm 3 requires $\mathbf{R}_{\mathbb{B}\mathbb{B}} \in \mathcal{COP}^{|\mathbb{B}|} \setminus \mathcal{SPN}^{|\mathbb{B}|}$. Such matrices are called *exceptional*. Recall that $\mathcal{SPN}^n \subseteq \mathcal{COP}^n$, and $\mathcal{SPN}^n = \mathcal{COP}^n$ if and only if $n \leq 4$ (see [Dia62]). Therefore, such an exceptional matrix exists if and only if $|\mathbb{B}| \geq 5$, which is ensured by the choices of the input parameters of Algorithm 3. A well-known exceptional 5×5 matrix is the Horn matrix given by

$$\mathbf{H} = \begin{pmatrix} 1 & -1 & 1 & 1 & -1 \\ -1 & 1 & -1 & 1 & 1 \\ 1 & -1 & 1 & -1 & 1 \\ 1 & 1 & -1 & 1 & -1 \\ -1 & 1 & 1 & -1 & 1 \end{pmatrix} \in \mathcal{COP}_5 \setminus \mathcal{SPN}_5. \quad (4.37)$$

Furthermore, $\mathbf{H}$ is an extreme ray of $\mathcal{COP}^5$ [HN63].

We next present a practical procedure for constructing an exceptional matrix in higher dimensions. Our procedure is inspired by a similar procedure outlined in [GY22, Section 6.2].

Let $n \geq 7$ and let $\mathbf{x} \in F_n$ satisfy $\|\mathbf{x}\|_0 \leq n - 5$, which ensures that $|\mathbb{B}| \geq 5$. We now construct $\mathbf{R}_{\mathbb{B}\mathbb{B}} \in \mathcal{COP}^{|\mathbb{B}|} \setminus \mathcal{SPN}^{|\mathbb{B}|}$ as follows. Choose an arbitrary $\mathbf{B} \in \mathcal{COP}^{|\mathbb{B}|-5}$, which can, for instance, be ensured by choosing $\mathbf{B} \in \mathcal{SPN}^{|\mathbb{B}|-5}$ (or even $\mathbf{B} \geq \mathbf{O}$ or $\mathbf{B} \succeq \mathbf{O}$), and an arbitrary $\mathbf{C} \in \mathbb{R}^{(|\mathbb{B}|-5) \times |\mathbb{B}|}$ such that $\mathbf{C} \geq \mathbf{O}$. Let us define

$$\mathbf{R}_{\mathbb{B}\mathbb{B}} = \begin{bmatrix} \mathbf{B} & \mathbf{C} \\ \mathbf{C}^\top & \mathbf{H} \end{bmatrix} \in \mathcal{S}^{|\mathbb{B}|}. \quad (4.38)$$

It is easy to see that $\mathbf{R}_{\mathbb{B}\mathbb{B}} \in \mathcal{COP}^{|\mathbb{B}|}$ (see also [SM16, Lemma 3.4 (a)]). Since $\mathbf{H} \in \mathcal{COP}^5 \setminus \mathcal{SPN}^5$, there exists $\mathbf{F} \in \mathcal{DNN}^5 \setminus \mathcal{CP}^5$ such that $\langle \mathbf{H}, \mathbf{F} \rangle = -\delta < 0$. Let

$$\mathbf{T} = \begin{pmatrix} \mathbf{O} & \mathbf{O} \\ \mathbf{O} & \mathbf{F} \end{pmatrix}. \quad (4.39)$$

Note that $\mathbf{T} \in \mathcal{DN}\mathcal{N}^{|\mathbb{B}|}$ and $\langle \mathbf{R}_{\mathbb{B}\mathbb{B}}, \mathbf{T} \rangle = \langle \mathbf{H}, \mathbf{F} \rangle = -\delta < 0$, which implies that $\mathbf{R}_{\mathbb{B}\mathbb{B}} \notin \mathcal{SPN}^{|\mathbb{B}|}$. We conclude that $\mathbf{R}_{\mathbb{B}\mathbb{B}} \in \mathcal{COP}^{|\mathbb{B}|} \setminus \mathcal{SPN}^{|\mathbb{B}|}$. We then define $\mu = \mathbf{e}_{\mathbb{B}}^{\top} \mathbf{T} \mathbf{e}_{\mathbb{B}} > 0$ and $\epsilon = \delta/\mu$.

While any choice of $\mathbf{F} \in \mathcal{DN}\mathcal{N}^5 \setminus \mathcal{CP}^5$ such that $\langle \mathbf{H}, \mathbf{F} \rangle = -\delta < 0$ works in the above construction, it may be preferable in practice to choose $\mathbf{F}$ such that the corresponding value of ϵ is as large as possible. This can be easily achieved in practice by solving the following small semidefinite programming (SDP) problem:

$$\epsilon := \max_{\mathbf{F} \in \mathcal{DN}\mathcal{N}^5 \setminus \{\mathbf{O}\}} \frac{-\langle \mathbf{H}, \mathbf{F} \rangle}{\mathbf{e}_{\mathbb{B}}^{\top} \mathbf{F} \mathbf{e}_{\mathbb{B}}} = \max\{-\langle \mathbf{H}, \mathbf{F} \rangle : \mathbf{e}_{\mathbb{B}}^{\top} \mathbf{F} \mathbf{e}_{\mathbb{B}} = 1, \mathbf{F} \in \mathcal{DN}\mathcal{N}_5\}, \quad (4.40)$$

where the second equality follows from the homogeneity of the objective function. An approximate optimal solution of (4.40) is given by

$$\mathbf{F} = \frac{1}{78.2} \begin{pmatrix} 7 & 4.32 & 0 & 0 & 4.32 \\ 4.32 & 7 & 4.32 & 0 & 0 \\ 0 & 4.32 & 7 & 4.32 & 0 \\ 0 & 0 & 4.32 & 7 & 4.32 \\ 4.32 & 0 & 0 & 4.32 & 7 \end{pmatrix}, \quad (4.41)$$

which yields $\epsilon \approx 0.1049$.

Remark 1. *The above construction works for any exceptional $\mathbf{R}_{\mathbb{B}\mathbb{B}}$ once we can control $\mathbf{T}$, the normal separating $\mathbf{R}_{\mathbb{B}\mathbb{B}}$ (corresponding to $\mathbf{H}$ above) from $\mathcal{SPN}^{|\mathbb{B}|}$, in the sense to solve the (effectively small) SDP problem*

$$\epsilon := \max_{\mathbf{T} \in \mathcal{DN}\mathcal{N}^{|\mathbb{B}|} \setminus \{\mathbf{O}\}} \frac{-\langle \mathbf{R}_{\mathbb{B}\mathbb{B}}, \mathbf{T} \rangle}{\mathbf{e}_{\mathbb{B}}^{\top} \mathbf{T} \mathbf{e}_{\mathbb{B}}} = \max\{-\langle \mathbf{R}_{\mathbb{B}\mathbb{B}}, \mathbf{T} \rangle : \mathbf{e}_{\mathbb{B}}^{\top} \mathbf{T} \mathbf{e}_{\mathbb{B}} = 1, \mathbf{T} \in \mathcal{DN}\mathcal{N}^{|\mathbb{B}|}\}. \quad (4.42)$$

Observe that any extremal ray of $\mathcal{COP}^{|\mathbb{B}|}$ different from a rank-one matrix or a symmetric componentwise nonnegative matrix with exactly two positive entries – which constitute the extremal rays of the cone of PSD matrices and the cone of componentwise nonnegative matrices – is necessarily exceptional, by extremality. While for larger $|\mathbb{B}|$, these extremal rays are (yet) unknown, we can imitate the above extension strategy and build upon the known ones for $|\mathbb{B}| \in \{5, 6\}$, e.g., using Hildebrand matrices [AHD21, Hil12]. All these used in $\mathbf{R}_{\mathbb{B}\mathbb{B}}$ would render (4.42) feasible (recommended only for extremal rays of moderate order to keep this effort low), and proceeding as above with $\mathbf{H}$ and $\mathbf{F}$, this approach provides more nontrivial examples. Finally, note that extremality is used in this context only to ensure exceptionality; any non-extremal but exceptional matrix would work as well.

4.4.3. Set of instances

To assess the impact of the choice of the instance set on the solution time of each exact model and each convex relaxation as well as on the quality of the lower bound, we conduct extensive experiments on a variety of carefully constructed instances of $(\text{StQP}(\rho))$. Let us now describe the set of instances in detail.

Using Algorithm 2 and Algorithm 3, we generated three sets of nontrivial instances of $(\text{StQP}(\rho))$ with different characteristics.

- (i) *PSD Instances:* This set of instances of $(\text{StQP}(\rho))$ is constructed using Algorithm 2 in such a way that $\mathbf{Q} \succeq \mathbf{O}$, i.e., the corresponding instance of (StQP) is a convex optimization

problem. For instance, such instances may arise in portfolio optimization problems. We remark that the resulting instances of $(\text{StQP}(\rho))$ are, in general, still NP-hard. In Step 2 of Algorithm 2, we construct $\mathbf{R} \succ \mathbf{O}$ so that all of its eigenvalues are uniformly distributed in $(0, 3)$. By choosing $\mathbf{N} = \mathbf{O}$ and $\lambda = 0$, we ensure that $\mathbf{Q} = (\mathbf{I} - \mathbf{e}\mathbf{x}^\top)\mathbf{R}(\mathbf{I} - \mathbf{x}\mathbf{e}^\top) \succeq \mathbf{O}$. Since $\mathbf{x}^\top\mathbf{Q}\mathbf{x} = 0$, $\mathbf{Q}$ is PSD but not positive definite. By Proposition 33, each PSD instance satisfies

$$\ell_\rho(\mathbf{Q}) > \ell_n(\mathbf{Q}) = \mu(\mathbf{Q}) = 0, \quad (4.43)$$

where $\mu(\mathbf{Q})$ is defined as in (4.33).

- (ii) *SPN Instances:* This set of instances of $(\text{StQP}(\rho))$ is generated using Algorithm 2 in a similar manner to PSD instances, i.e., we set $\lambda = 0$ and use the same procedure to construct $\mathbf{R}$. However, instead of choosing $\mathbf{N} = \mathbf{O}$ as in PSD instances, each entry of $\mathbf{N}_{\mathbb{A}\mathbb{B}} \in \mathbb{R}^{|\mathbb{A}| \times |\mathbb{B}|}$ and $\mathbf{N}_{\mathbb{B}\mathbb{B}} \in \mathcal{S}^{|\mathbb{B}|}$ is uniformly generated in $(0, 3)$. In contrast with PSD instances, $\mathbf{Q} = (\mathbf{I} - \mathbf{e}\mathbf{x}^\top)\mathbf{R}(\mathbf{I} - \mathbf{x}\mathbf{e}^\top) + \mathbf{N}$ is, in general, not PSD. In fact, we numerically verified that $\mathbf{Q}$ was an indefinite matrix for each SPN instance. By Proposition 33, each SPN instance satisfies

$$\ell_\rho(\mathbf{Q}) > \ell_n(\mathbf{Q}) = \mu(\mathbf{Q}) = 0. \quad (4.44)$$

- (iii) *COP Instances:* This set of instances of $(\text{StQP}(\rho))$ is generated by Algorithm 3. In Step 2, we construct $\mathbf{R} \in \mathcal{S}^n$ using the procedure outlined in Section 4.4.2. In particular, $\mathbf{R}_{\mathbb{B}\mathbb{B}}$ is constructed as in (4.38), where $\mathbf{B} \succ \mathbf{O}$ is generated with eigenvalues uniformly distributed in $(0, 3)$, and each entry of $\mathbf{C} \geq \mathbf{O}$ is uniformly distributed in $(0, 1)$. Our choices are, in part, motivated by the observation that the eigenvalues of the Horn matrix $\mathbf{H}$ given by (4.37) lie in $[-1.236, 3.236]$ so that the eigenvalues of $\mathbf{B} \succ \mathbf{O}$ have similar magnitude. The separating matrix $\mathbf{T}$ is constructed as in (4.39), where $\mathbf{F}$ is given by (4.41), which yields $\epsilon \approx 0.1049$. Finally, $\mathbf{R}_{\mathbb{A}\mathbb{A}} \succ \mathbf{O}$ is constructed with eigenvalues uniformly distributed in $(0, 0.99\epsilon)$, ensuring that $\|\mathbf{R}_{\mathbb{A}\mathbb{A}}\| < \epsilon$. We then choose $\lambda = 0$ and set $\mathbf{Q} = (\mathbf{I} - \mathbf{e}\mathbf{x}^\top)\mathbf{R}(\mathbf{I} - \mathbf{x}\mathbf{e}^\top)$. By Proposition 34, this procedure ensures that $\mathbf{Q} \not\succeq \mathbf{O}$ and each COP instance satisfies

$$\ell_\rho(\mathbf{Q}) > \ell_n(\mathbf{Q}) = 0 > \mu(\mathbf{Q}). \quad (4.45)$$

In summary, each of the three sets of instances of $(\text{StQP}(\rho))$ is comprised of nontrivial instances that exhibit different characteristics. While each PSD instance and each SPN instance has a corresponding StQP instance that admits an exact doubly nonnegative (DNN) relaxation, they differ in terms of the convexity of the objective function. On the other hand, each COP instance has a nonconvex objective function with the additional property that the DNN relaxation of the corresponding StQP instance is inexact.

4.5. Computational results

In this section, we report and discuss the results of our computational experiments. After describing the experimental setup in Section 4.5.1, we report our results in detail in Sections 4.5.2 and 4.5.3.

4.5.1. Experimental setup

The set of instances was already described in Section 4.4.3. Here, we outline the details of our experimental setup.

In our experiments, we chose $n \in \{25, 50\}$. For each choice of n , we identified three sparsity levels for the designated optimal solution $\mathbf{x} \in F_n$ of (StQP), denoted by $\rho_0 = \|\mathbf{x}\|_0$. In particular, we considered $\rho_0 \in \{\lfloor 0.25n \rfloor, \lfloor 0.5n \rfloor, \lfloor 0.75n \rfloor\}$, where $\lfloor \cdot \rfloor$ denotes the nearest integer function. Finally, for each choice of ρ_0 , we chose $\rho \in \{\lfloor 0.25\rho_0 \rfloor, \lfloor 0.5\rho_0 \rfloor, \lfloor 0.75\rho_0 \rfloor\}$. The parameters are summarized in Table 4.1.

n	ρ_0	ρ
25	6	$\{2,3,4\}$
	12	$\{3,6,9\}$
	19	$\{5,10,14\}$
50	12	$\{3,6,9\}$
	25	$\{6,12,19\}$
	38	$\{10,19,28\}$

Table 4.1.: Choices of the parameters n , ρ_0 , and ρ

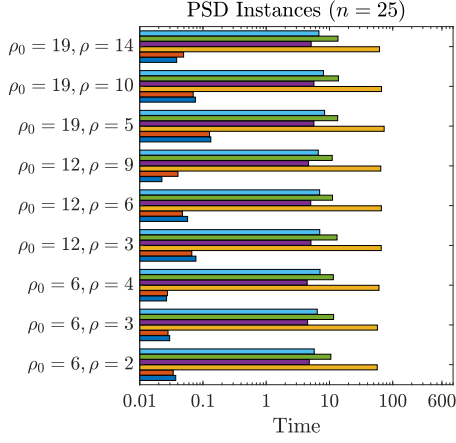
For each choice of n , ρ_0 , and ρ presented in Table 4.1, we generated 25 instances from each of the three sets of instances described in Section 4.4.3, which gave rise to a total of 1,350 nontrivial instances of (StQP(ρ)) comprising of 450 PSD instances, 450 SPN instances, and 450 COP instances. For each instance of (StQP(ρ)), we solved (P1(ρ)), (P2(ρ)), (D1A(ρ)), (D1B(ρ)), (D2A(ρ)), and (D2B(ρ)).

Algorithm 2 and Algorithm 3 were implemented in Julia version 1.8.5 [BEKS17]. Each instance of (P1(ρ)) and (P2(ρ)) was solved by **Gurobi** version 11.0.0 [Gur23] via the modeling language **JuMP** v1.19.0 [LDD⁺23]. We solved (D1A(ρ)), (D1B(ρ)), (D2A(ρ)), and (D2B(ρ)) by **MOSEK** version 10.1.24 [ApS24] via **JuMP**. Our computational experiments were carried out on a 64-bit HP workstation with 24 threads (2 sockets, 6 cores per socket, 2 threads per core) running **Ubuntu Linux** with 96 GB of RAM and Intel Xeon CPU E5-2667 processors with a clock speed of 2.90 GHz. We imposed a time limit of 600 seconds on each solve. The default settings were used for all of the other parameters of **Gurobi** and **MOSEK**.

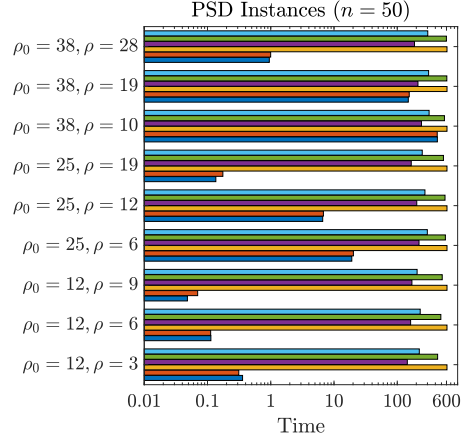
4.5.2. Solution time

Now we focus on the solution time of each of the six models (P1(ρ)), (P2(ρ)), (D1A(ρ)), (D1B(ρ)), (D2A(ρ)), and (D2B(ρ)).

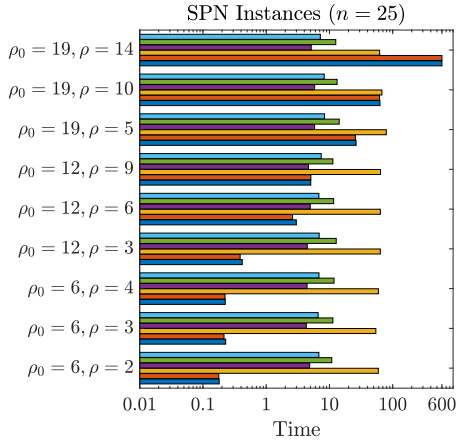
Figure 4.1 depicts the average solution times of each of the six models on each of the three instance sets. The results for PSD instances, SPN instances, and COP instances are presented in the first, second, and third rows of Figure 4.1, respectively. In each row, the results for $n = 25$ and $n = 50$ are given in the first and second column, respectively. In each graph, the horizontal axis denotes the average solution time (in seconds) presented in logarithmic scale. We used identical axis limits in each graph to facilitate a better comparison across different graphs. The vertical axis is comprised of nine sets of bar charts, each of which represents a particular choice of the tuple (ρ_0, ρ) corresponding to the choice of n as outlined in Table 4.1. Finally, for each choice of (ρ_0, ρ) , the corresponding bar chart represents the average solution times of the six models (P1(ρ)), (P2(ρ)), (D1A(ρ)), (D1B(ρ)), (D2A(ρ)), and (D2B(ρ)) over 25 instances generated using the procedure in Section 4.4.3. Note that the solution times of instances that were terminated due to the time limit of 600 seconds were also included in the average solution times. Therefore, each bar chart represents the average computational requirement of the corresponding model. We



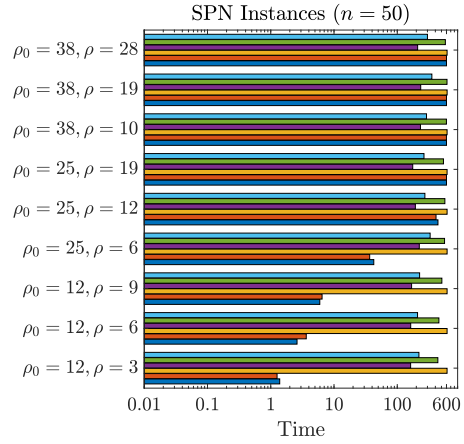
(a) PSD Instances ($n = 25$)



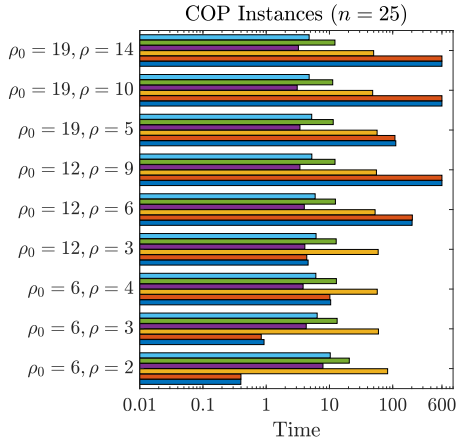
(b) PSD Instances ($n = 50$)



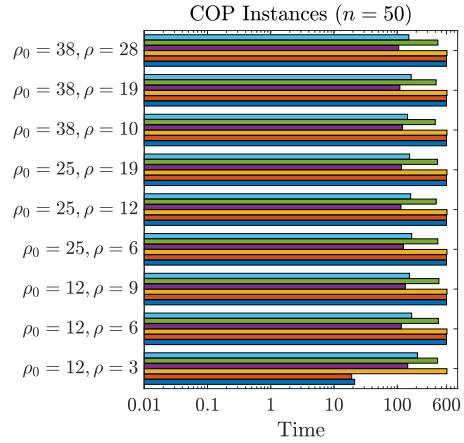
(c) SPN Instances ($n = 25$)



(d) SPN Instances ($n = 50$)



(e) COP Instances ($n = 25$)



(f) COP Instances ($n = 50$)

Figure 4.1.: Average solution times (in seconds) of (P1(ρ)), (P2(ρ)), (D1A(ρ)), (D1B(ρ)), (D2A(ρ)), and (D2B(ρ)) for PSD, SPN, and COP instances

recall that (P1(ρ)) and (P2(ρ)) were solved by **Gurobi** whereas (D1A(ρ)), (D1B(ρ)), (D2A(ρ)), and (D2B(ρ)) by **MOSEK**.

Figure 4.1 reveals several interesting relations about average solution times. We first outline our observations concerning the exact MIQP models $(P1(\rho))$ and $(P2(\rho))$:

- (i) Average solution times of the two exact models exhibit very similar behavior for each fixed choice of the instance set and the triple (n, ρ_0, ρ) . However, there seems to be a very strong correlation between the average solution time and the choices of (n, ρ_0, ρ) as well as the instance set.
- (ii) Considering each row of Figure 4.1 separately, the average solution time of each of the two exact models across all choices of the tuple (ρ_0, ρ) increases as n increases. However, the rate of increase seems to be highly dependent on the choice of the instance set.

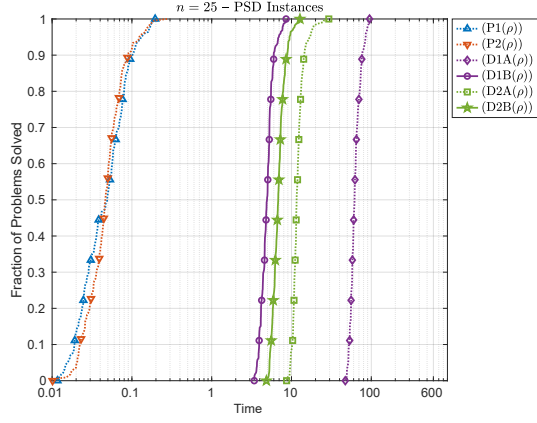
We next present the corresponding observations for the average solution times of the convex relaxations $(D1A(\rho))$, $(D1B(\rho))$, $(D2A(\rho))$, and $(D2B(\rho))$:

- (i) In contrast with the exact models, for each choice of n , Figure 4.1 reveals that the average solution times of the convex relaxations exhibit very similar behavior across all three sets of instances and all choices of (ρ_0, ρ) .
- (ii) The average solution time of our reduced formulation $(D1B(\rho))$ consistently achieves the lowest average solution times, followed by the reduced formulation $(D2B(\rho))$, which, in turn, is followed by $(D2A(\rho))$ and $(D1A(\rho))$, respectively.
- (iii) The average solution time of our reduced formulation $(D1B(\rho))$ exhibits at least an order of magnitude improvement over that of $(D1A(\rho))$. On the other hand, the corresponding improvement of $(D2B(\rho))$ over $(D2A(\rho))$ seems to be less pronounced, especially in smaller dimensions.
- (iv) As expected, the average solution time of each relaxation increases with n . While the rate of increase does not seem to be influenced by the specific instance set, we observe that it does depend on the specific relaxation.

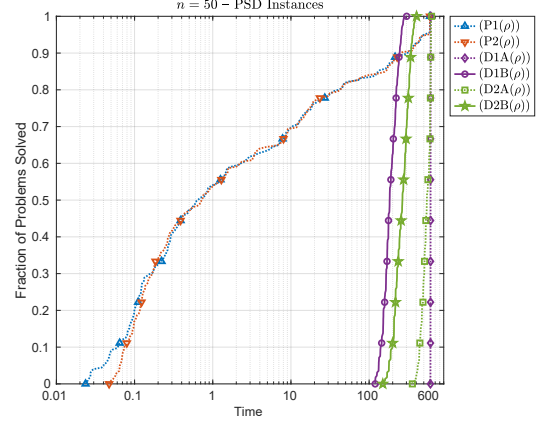
Since Figure 4.1 is based only on average solution times, it cannot capture the variability in the data. In an attempt to shed more light on the distribution of the solution times across all instance sets, we present the empirical cumulative distribution functions of all six models in Figure 4.2, which is organized similarly to Figure 4.1. The three rows display the results for PSD, SPN, and COP instances, respectively, whereas the two columns are devoted to the results for $n = 25$ and $n = 50$, respectively. Each of the six graphs presents six plots corresponding to the empirical cumulative distribution functions of solution times of each of the six models on all 225 instances for the corresponding choice of the instance set and n (see Table 4.1). The horizontal axis denotes the solution time (in seconds) on a logarithmic scale, and the vertical axis represents the fraction of the instances solved. The markers represent the data points for every 25th instance. Once again, we employed identical axis limits in each of the six graphs.

Figure 4.2 clearly reveals the difference between the distributions of solution times of the exact models $(P1(\rho))$ and $(P2(\rho))$ and those of the convex relaxations $(D1A(\rho))$, $(D1B(\rho))$, $(D2A(\rho))$, and $(D2B(\rho))$. We outline our observations below:

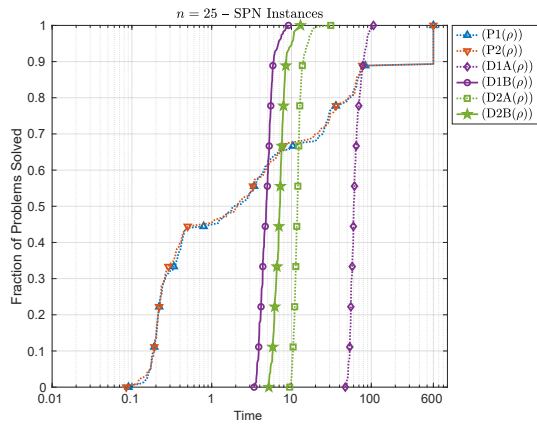
- (i) The solution times of the two exact models, in general, exhibit a very similar distribution but a significant variability across different instance sets.



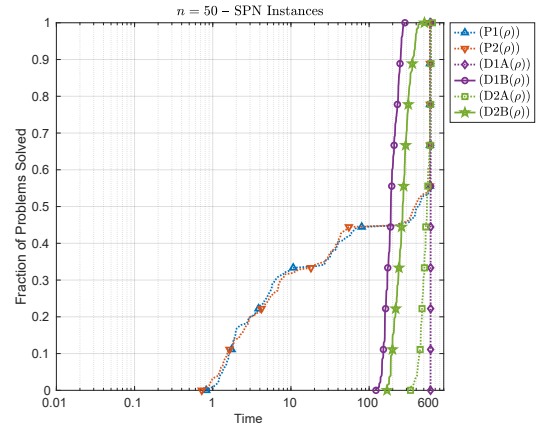
(a) PSD Instances ($n = 25$)



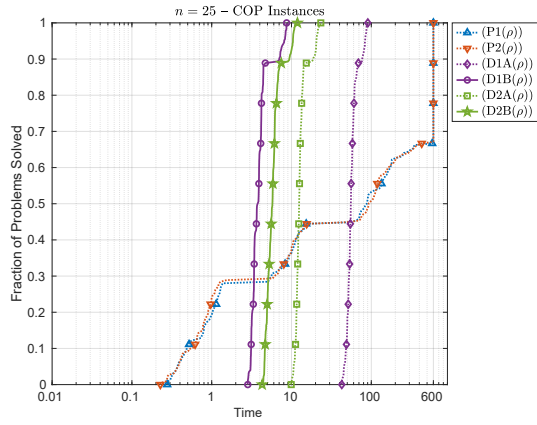
(b) PSD Instances ($n = 50$)



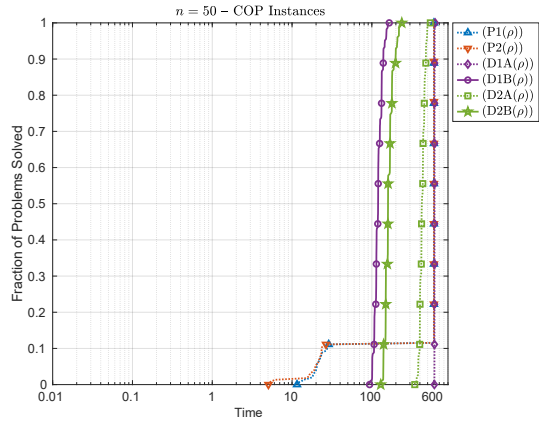
(c) SPN Instances ($n = 25$)



(d) SPN Instances ($n = 50$)



(e) COP Instances ($n = 25$)



(f) COP Instances ($n = 50$)

Figure 4.2.: Empirical cumulative distribution functions of solution times of $(P1(\rho))$, $(P2(\rho))$, $(D1A(\rho))$, $(D1B(\rho))$, $(D2A(\rho))$, and $(D2B(\rho))$

- (ii) In contrast, we observe that the solution times of the convex relaxations tend to have a much smaller variance and exhibit very similar distributions across different instance sets.

Therefore, in contrast with the exact models $(P1(\rho))$ and $(P2(\rho))$, we conclude that the solution time of each convex relaxation seems to be very robust with respect to the choice of the instance set and the choice of (ρ_0, ρ) in our setting.

In the sequel, we present a more detailed discussion of the behavior of the solution times of the exact models and convex relaxations.

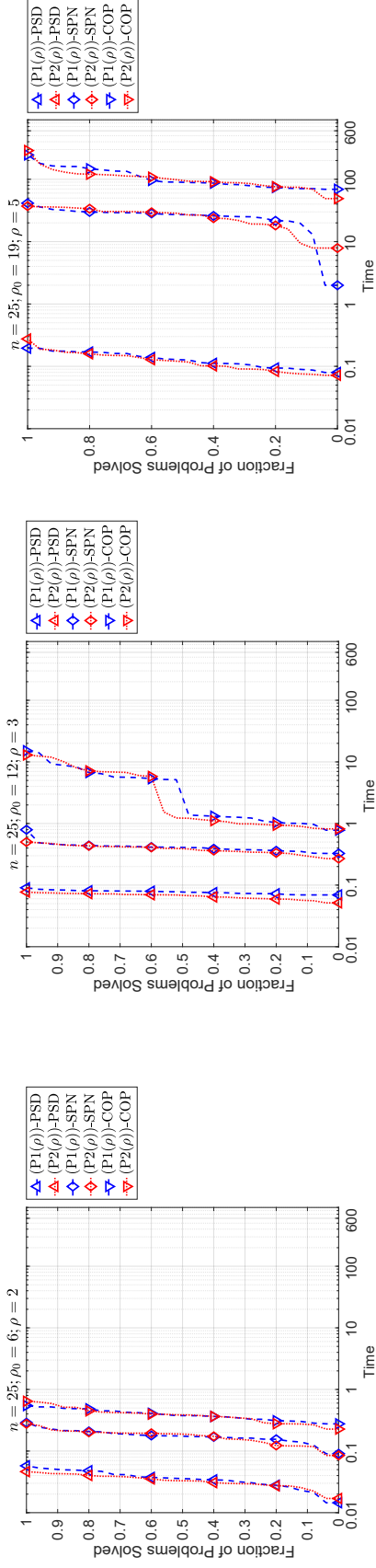
Exact MIQP models

As illustrated by Figures 4.1 and 4.2, the solution time of the exact models is highly dependent on the choices of the instance set and on the triple (n, ρ_0, ρ) . Here, we aim to shed more light on this dependence.

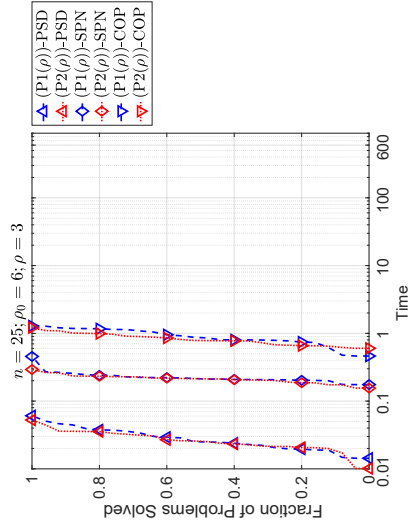
In Figures 4.3 and 4.4, we present the empirical cumulative distribution functions of solution times of $(P1(\rho))$ and $(P2(\rho))$ for all instances with $n = 25$ and $n = 50$, respectively. Each graph in each figure consists of six plots, each of which corresponds to the solution times of $(P1(\rho))$ and $(P2(\rho))$ on each of PSD, SPN, and COP instances for a fixed choice of (n, ρ_0, ρ) . In each column of each figure, we fix the tuple (n, ρ_0) and present the empirical cumulative distributions of the solution times of 25 instances corresponding to the three different choices of ρ as presented in Table 4.1. On the other hand, each row corresponds to a different ratio of $\rho/\rho_0 \in \{0.25, 0.5, 0.75\}$. Again, we use a logarithmic scale for the solution time and ensure that all axis limits are identical for a meaningful comparison across different graphs. Finally, we note that the markers represent the data points at every 5th instance.

A close examination of Figures 4.3 and 4.4 reveal the following observations about the two exact models $(P1(\rho))$ and $(P2(\rho))$:

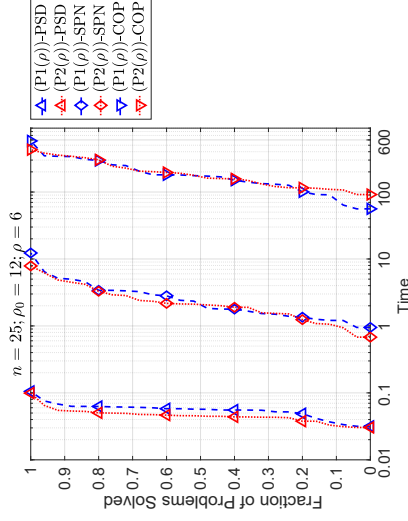
- (i) The two models, in general, exhibit very similar distributions across all parameter choices, except for some small differences for a few choices of (n, ρ_0, ρ) on PSD instances, where $(P1(\rho))$ seems to be solved slightly faster than $(P2(\rho))$.
- (ii) We clearly observe an empirical first-order stochastic dominance among PSD, SPN, and COP instances. For each choice of (n, ρ_0, ρ) , PSD instances tend to achieve the lowest solution times, followed by SPN instances, whereas COP instances exhibit the highest solution times.
- (iii) For each fixed (n, ρ_0) , we observe that the solution time tends to decrease as ρ increases for PSD instances, especially for $n = 50$, whereas it increases for each of SPN and COP instances.
- (iv) If we fix n and the ratio $\rho/\rho_0 \in \{0.25, 0.5, 0.75\}$, the solution time increases as ρ_0 increases across all three sets of instances. For fixed n , we therefore conclude that the solution time of the exact models seems to be also highly influenced by the choice of the ratio ρ/ρ_0 .
- (v) For $n = 25$, both exact models can be solved very quickly for all PSD instances. In contrast, for SPN and COP instances, each of $(P1(\rho))$ and $(P2(\rho))$ is terminated due to the time limit of 600 seconds on certain subsets of the instances. Therefore, we observe that SPN and COP instances can be particularly challenging for $(P1(\rho))$ and $(P2(\rho))$, even for $n = 25$ with particular choices of (ρ_0, ρ) . For $n = 50$, we observe that even some PSD instances were terminated due to the time limit. Recall that each instance in this set has a convex objective function. For each choice of (n, ρ_0, ρ) , a comparison of PSD, SPN, and COP instances reveals that the number of instances terminated due to the time limit exhibits a monotonically increasing behavior.



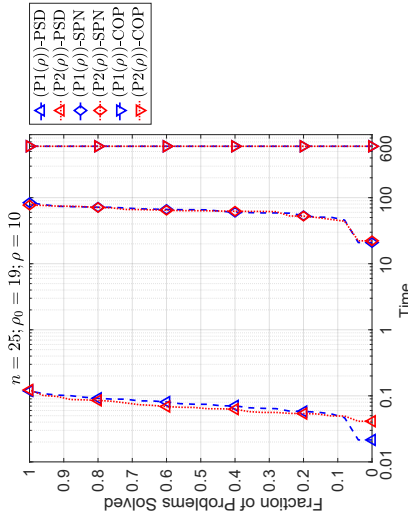
(a) $n = 25, \rho_0 = 6, \rho = 2$



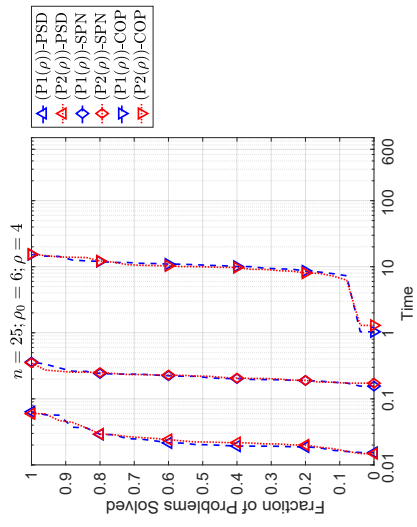
(b) $n = 25, \rho_0 = 12, \rho = 3$



(c) $n = 25, \rho_0 = 19, \rho = 5$

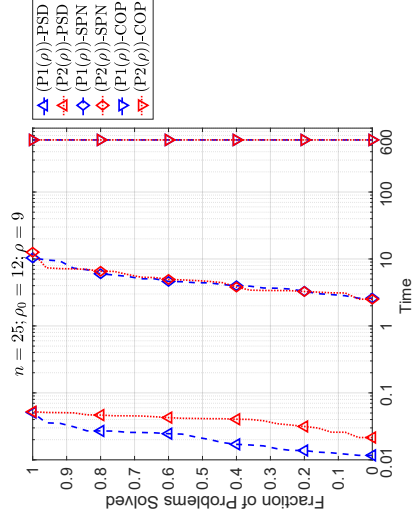


(d) $n = 25, \rho_0 = 6, \rho = 3$



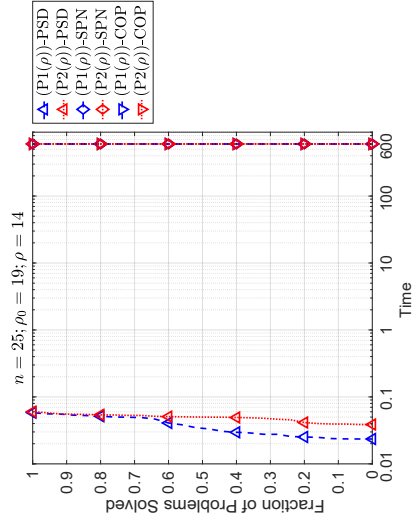
(e) $n = 25, \rho_0 = 6, \rho = 4$

(f) $n = 25, \rho_0 = 19, \rho = 10$



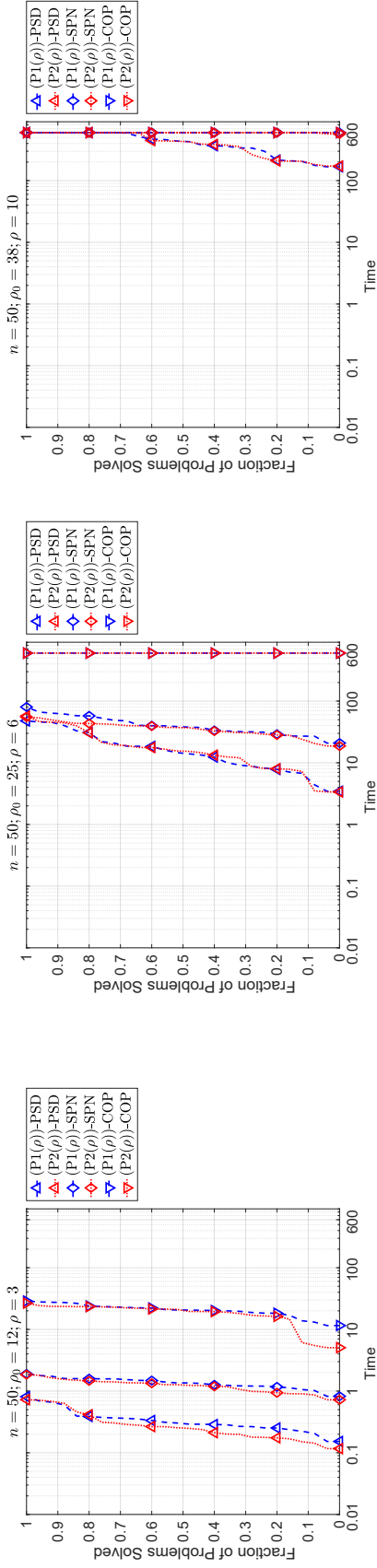
(g) $n = 25, \rho_0 = 12, \rho = 9$

(h) $n = 25, \rho_0 = 19, \rho = 14$

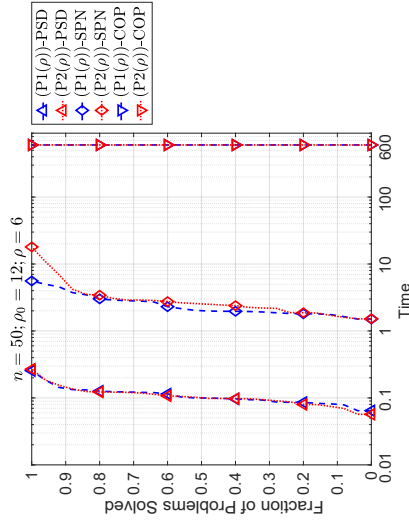


(i) $n = 25, \rho_0 = 19, \rho = 14$

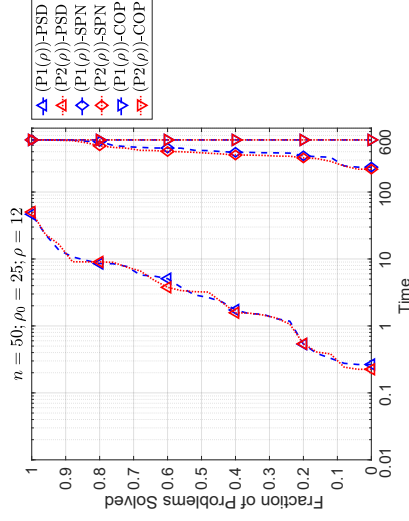
Figure 4.3.: Empirical cumulative distribution functions of solution times of $(P1(\rho))$ and $(P2(\rho))$ for all instances with $n = 25$



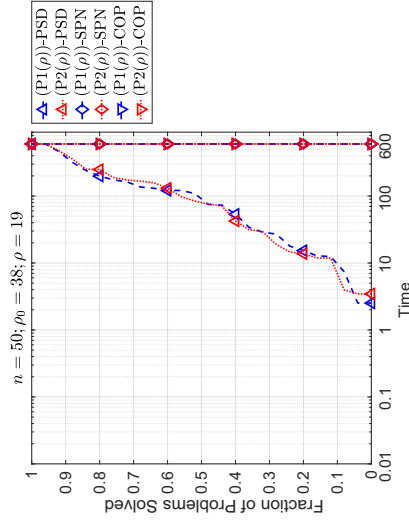
(a) $n = 50, \rho_0 = 12, \rho = 3$



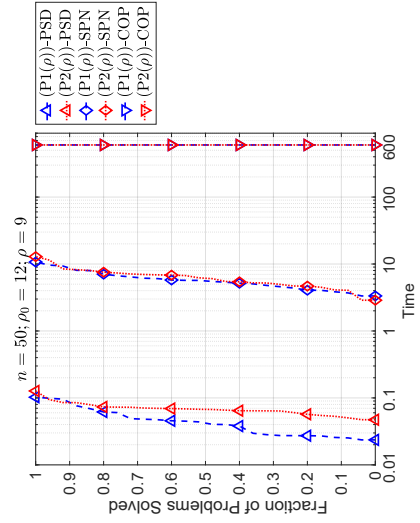
(b) $n = 50, \rho_0 = 25, \rho = 6$



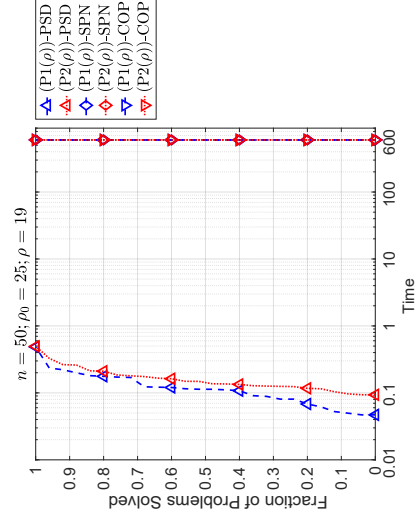
(c) $n = 50, \rho_0 = 38, \rho = 10$



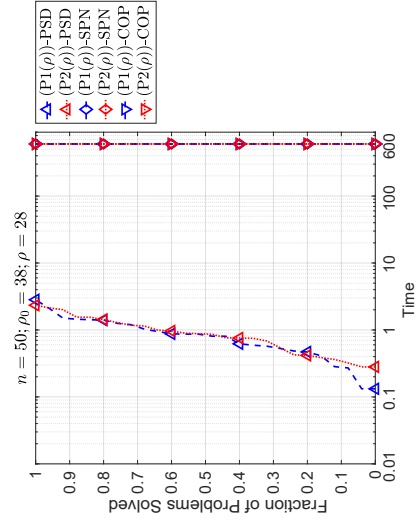
(d) $n = 50, \rho_0 = 12, \rho = 6$



(e) $n = 50, \rho_0 = 25, \rho = 12$



(f) $n = 50, \rho_0 = 38, \rho = 19$



(g) $n = 50, \rho_0 = 12, \rho = 9$

(h) $n = 50, \rho_0 = 25, \rho = 19$

(i) $n = 50, \rho_0 = 38, \rho = 28$

Figure 4.4.: Empirical cumulative distribution functions of solution times of $(P1(\rho))$ and $(P2(\rho))$ for all instances with $n = 50$

In summary, the solution time of exact models ($P1(\rho)$) and ($P2(\rho)$) seem to be highly dependent on each of the choices of the instance set, (n, ρ_0, ρ) , as well as on the ratio ρ/ρ_0 . COP instances tend to be the most challenging for the exact models, followed by SPN instances, which, in turn, are followed by PSD instances. These results illustrate that Algorithms 2 and 3 are capable of generating instances of ($StQP(\rho)$) that are particularly challenging for **Gurobi**, even in relatively small dimensions.

Convex DNN relaxations

Now let us focus on solution times of the four convex relaxations ($D1A(\rho)$), ($D1B(\rho)$), ($D2A(\rho)$), and ($D2B(\rho)$).

In contrast with the exact models, Figure 4.2 illustrates that the solution times of convex relaxations do not exhibit a strong correlation with the choices of the instance set and the tuple (ρ_0, ρ) . As such, we do not present the counterparts of Figures 4.3 and 4.4 for the convex relaxations due to space considerations.

Based on Figure 4.2, we make the following observations about the convex relaxations ($D1A(\rho)$), ($D1B(\rho)$), ($D2A(\rho)$), and ($D2B(\rho)$):

- (i) The distributions of the solution times clearly illustrate the computational advantages of the reduced formulations ($D2A(\rho)$) and ($D2B(\rho)$) over their counterparts.
- (ii) We observe an empirical first-order stochastic dominance among the distributions of the solution times of ($D1A(\rho)$), ($D1B(\rho)$), ($D2A(\rho)$), and ($D2B(\rho)$) on each of the six graphs. In particular, we observe that the reduced formulation ($D1B(\rho)$) consistently achieves the lowest solution times, followed by the reduced formulation ($D2B(\rho)$), which, in turn, is followed by ($D2A(\rho)$) and ($D1A(\rho)$), respectively. In addition to being the theoretically tightest relaxation, it is worth noticing that ($D1B(\rho)$) also outperforms ($D2B(\rho)$) in terms of the solution time.
- (iii) For $n = 25$, each of the four relaxations can be solved to optimality within the time limit for each instance set and each choice of ρ_0 and ρ . For $n = 50$, we observe that ($D1A(\rho)$) was terminated due to the time limit on every instance whereas ($D2A(\rho)$) was terminated due to the time limit on some subsets of the PSD and SPN instances. In contrast, each of the two reduced formulations ($D1B(\rho)$) and ($D2B(\rho)$) was solved to optimality within the time limit on all instances. These observations clearly demonstrate the computational advantages of the reduced formulations ($D2A(\rho)$) and ($D2B(\rho)$) over their counterparts, especially in higher dimensions.
- (iv) For PSD instances with $n = 25$, the distributions of the solution times of each of the two exact models exhibit a first-order stochastic dominance on the corresponding distribution of each of the four convex relaxations. In contrast, we observe that the solution times of the convex relaxations achieve better performance on a larger subset of instances for each instance set with a larger n . For fixed n , a comparison of PSD, SPN, and COP instances reveals an increasingly better performance of the convex relaxations in comparison with the exact models.

Table 4.2 reports the number of instances on which ($D2A(\rho)$) is terminated due to the time limit. Note that we omit ($D1A(\rho)$) since all instances with $n = 50$ were terminated due to the time limit. It is worth noticing that ($D2A(\rho)$) was solved to optimality within the time limit on all COP instances with $n = 50$. In contrast with $n = 25$, we therefore conclude that the solution

time of $(D2A(\rho))$ for $n = 50$ seems to be somewhat sensitive to the set of instances and the choices of the parameters (ρ_0, ρ) .

Instance Set	n	ρ_0	ρ	$(D2A(\rho))$
PSD	50	12	9	1
		25	6	14
		25	12	10
		25	19	3
		38	10	8
		38	19	22
		38	28	18
SPN	50	12	3	1
		25	6	12
		25	12	8
		25	19	2
		38	10	20
		38	19	24
		38	28	10
COP	50	–	–	–

Table 4.2.: Number of instances (out of 25) terminated due to the time limit (excluding $(D1A(\rho))$ that was terminated due to the time limit on *all* instances with $n = 50$)

Finally, we recall that optimality gaps are not reported on instances terminated due to the time limit since they are not particularly meaningful in our setting.

In conclusion, the solution time of each relaxation seems to be very robust with respect to the choice of the instance set and the choice of (ρ_0, ρ) in our setting. The reduced formulations $(D2A(\rho))$ and $(D2B(\rho))$ yield significant computational advantages over their counterparts, especially in higher dimensions.

4.5.3. Quality of lower bounds

In this section, we discuss the quality of the lower bounds arising from the convex relaxations. To that end, we compare the absolute gaps of the exact models with those of the convex relaxations. For an instance, we define the MIQP gap as the difference between the best upper bound and the best lower bound obtained from $(P1(\rho))$ and $(P2(\rho))$. Similarly, we define the DNN gap to be the difference between the best upper bound obtained from the two exact models, $(P1(\rho))$ and $(P2(\rho))$, and the best lower bound obtained from the reduced formulations $(D2A(\rho))$ and $(D2B(\rho))$. Note that we use the lower bounds from the reduced formulations since all of them can be solved to optimality on all instances.

If at least one of $(P1(\rho))$ and $(P2(\rho))$ can solve an instance to optimality, then the MIQP gap will be very close to zero. On such an instance, we deem that our relaxation is *exact* if the DNN gap is of similar magnitude to that of the MIQP gap. If each of $(P1(\rho))$ and $(P2(\rho))$ is terminated due to the limit on an instance, then the MIQP gap will be sufficiently away from zero. In this case, if the DNN gap is smaller than or equal to the MIQP gap, we say that our relaxations are better than the exact modes. Otherwise, our relaxations are *worse* than the exact models.

In Figure 4.5, which is organized similarly to Figure 4.2, we plot the DNN and MIQP gaps. To ease reading, we have accumulated all instances of the same dimensions from each instance set across all different parameter constellations (ρ_0, ρ) , yielding 225 such instances in total for every graph (a)-(f), depicting the situation in the different hardness classes of the generated instances. In each graph, the horizontal axis represents the instances ordered in nondecreasing MIQP gaps, which are represented by the blue curve, and the vertical axis denotes the absolute gap. Finally, the markers indicate the data points at every 25th instance.

We outline our observations:

- (i) For PSD instances, our relaxations are exact on the vast majority of all instances solved to optimality by an exact model and are better than the exact models on all instances with a positive MIQP gap.
- (ii) For SPN instances, our relaxations are exact on a relatively smaller proportion of all instances solved to optimality by an exact model and are better than the exact models on all instances with a positive MIQP gap. Furthermore, it is worth noticing that our relaxations continue to be exact on several instances with a positive MIQP gap.
- (iii) On COP instances, our relaxations are worse than the exact models on the vast majority of instances. However, we remark that the DNN gap is smaller on a small subset of instances with $n = 50$ that admit a positive MIQP gap.

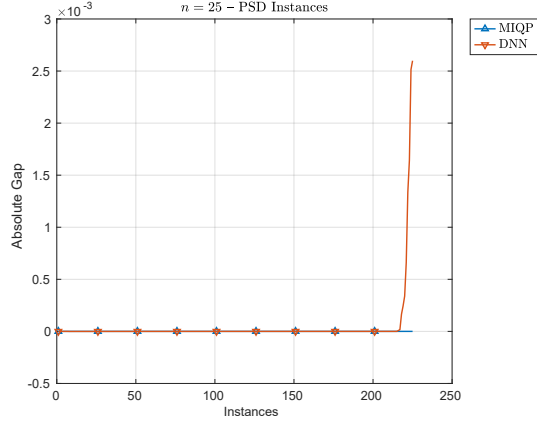
The results show clearly that the DNN relaxations provide tight lower bounds on a large number of PSD and SPN instances and that there is a significant number of instances where the MIQP gap is larger than the gaps achieved by the DNN relaxations, sometimes drastically so (one or even two orders of magnitude). The dominance of MIQP models on the extremely difficult COP instances in graphs (e) and (f) may be explained by the fact that additional cuts may be needed to tighten the DNN gaps without resorting to higher-order relaxations of the (intractable) exact conic reformulations.

We close this part by reporting that we have not observed a significant difference between the quality of the lower bounds arising from the provably tighter relaxation (D1B(ρ)) and the weaker (D2B(ρ)). However, we still recommend using the tighter (D1B(ρ)) since it also outperforms (D2B(ρ)) in terms of the solution time on our instance sets.

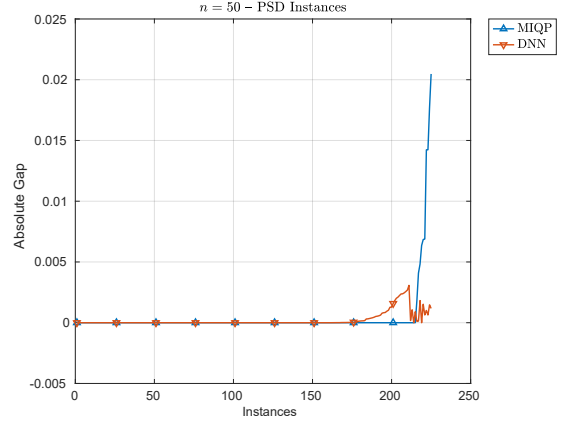
In this section, we studied the sparse StQP. By utilizing exact copositive reformulations of two mixed integer quadratic optimization models, we proposed two tractable convex relaxations. For both relaxations, we established equivalent tractable reformulations in significantly smaller dimensions. We also presented a theoretical comparison of the two relaxations. Finally, we proposed an instance generation scheme that is guaranteed to construct nontrivial instances. Our computational results clearly illustrate the computational advantages of our reduced relaxations as well as the quality of the relaxation bounds.

4.6. Conclusion

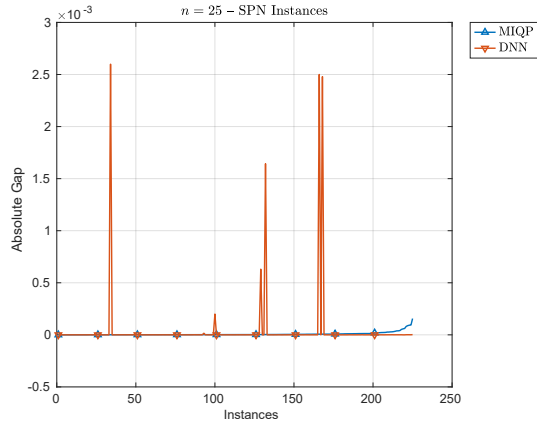
In this chapter, we studied the sparse StQP. By utilizing exact copositive reformulations of two mixed integer quadratic optimization models, we proposed two tractable convex relaxations. For both relaxations, we established equivalent tractable reformulations in significantly smaller dimensions. We also presented a theoretical comparison of the two relaxations. Finally, we proposed an instance generation scheme that is guaranteed to construct nontrivial instances. Our computational results clearly illustrate the computational advantages of our reduced relaxations as well as the quality of the relaxation bounds.



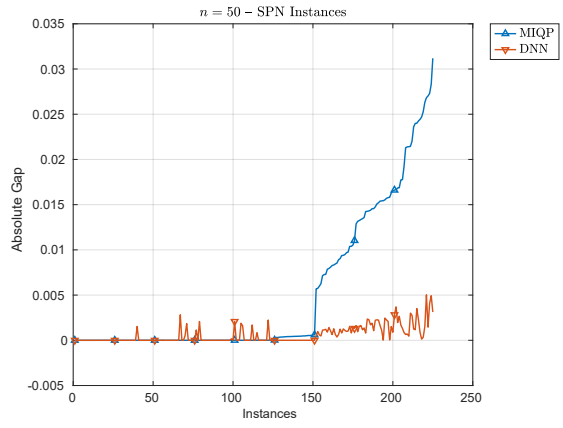
(a) PSD Instances ($n = 25$)



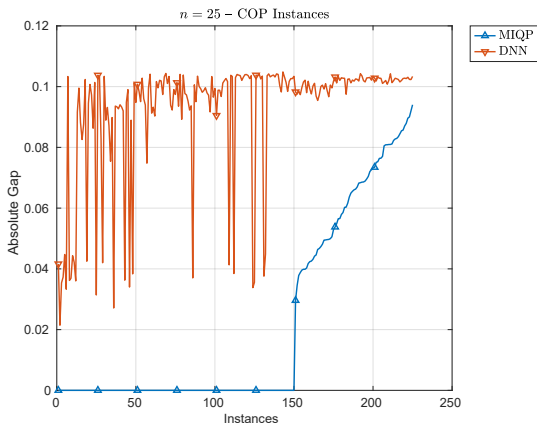
(b) PSD Instances ($n = 50$)



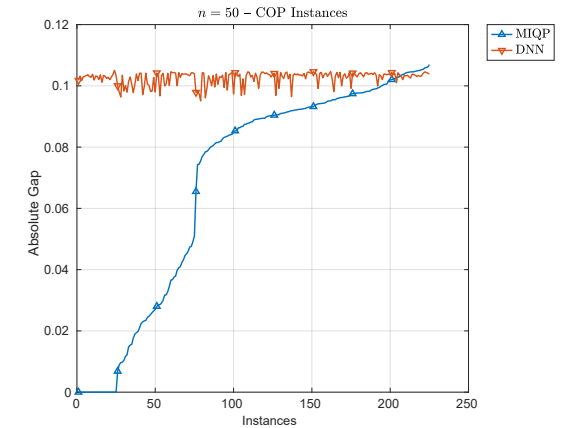
(c) SPN Instances ($n = 25$)



(d) SPN Instances ($n = 50$)



(e) COP Instances ($n = 25$)



(f) COP Instances ($n = 50$)

Figure 4.5.: Absolute gaps of the best of $(P1(\rho))$ and $(P2(\rho))$ versus the best of $(D1B(\rho))$ and $(D2B(\rho))$

5. A scalable conic approach for feature selection in linear SVMs

In this chapter, we study the embedded feature selection problem in linear Support Vector Machines (SVMs), in which a cardinality constraint is employed, leading to an explainable selection model. To handle the hard problem, we first introduce two mixed-integer formulations for which novel semidefinite relaxations are proposed. In addition, we propose both heuristics and exact algorithms using the information of its optimal solution. The content of this chapter is based on the submitted paper [BDPP24] for which revisions were requested by the editor after the first round of evaluations.

5.1. The problem

Given a dataset with m samples, and n features, $\{(\mathbf{x}^i, y^i) \in \mathbb{R}^n \times \{-1, 1\} : i \in [1:m]\}$, the linear SVM classification problem defines a classifier as the function $f : \mathbb{R}^n \rightarrow \{-1, 1\}$,

$$f(\mathbf{x}) = \text{sign}(\mathbf{w}^{*\top} \mathbf{x} + b^*).$$

In traditional SVMs [CV95], coefficients $(\mathbf{w}^*, b^*) \in \mathbb{R}^n \times \mathbb{R}$ identify a separating hyperplane (if it exists) that maximizes its margin, i.e. the distance between the hyperplane and the closest data points of each class. Accounting for classes not linearly separable, the tuple $(\mathbf{w}^*, b^*)$ is found by solving the following convex quadratic problem:

$$\begin{aligned} (\text{SVM}) \quad & \min_{\mathbf{w}, b, \xi} \quad \frac{1}{2} \|\mathbf{w}\|_2^2 + C \sum_{i=1}^m \xi_i \\ & \text{s.t.} \quad y^i (\mathbf{w}^\top \mathbf{x}^i + b) \geq 1 - \xi_i, \quad i \in [1:m], \\ & \quad \xi \geq \mathbf{0}. \end{aligned} \tag{5.1}$$

The non-negative variables ξ allow points to be on the wrong side of their “soft margin”, as well as being “too close” to the decision boundary. The sum $\sum_{i=1}^m \xi_i$ in the objective function represents an upper bound on the total misclassification error, which is weighted by the hyperparameter $C > 0$ so to balance the maximization of the margin, which is inversely proportional to $\|\mathbf{w}\|_2$.

The optimal solution $\mathbf{w}^*$ of (5.1) is usually dense, in the sense that usually all components of $\mathbf{w}^*$ are nonzero and thus every feature contributes to the definition of the classifier. Indeed, through duality theory, $\mathbf{w}^*$ can be expressed as the linear combination of a subset of samples, called support vectors. More precisely, $\mathbf{w}^* = \sum_{i=1}^m \alpha_i^* y^i \mathbf{x}^i$, with $\alpha_i^* \geq 0$ being the dual optimal solution of problem (5.1) with nonzero value only if sample i represents a support vector. Thus, it is unlikely to obtain $w_j^* = 0$ for some feature j , and sparsity in the vector $\mathbf{w}^*$ must be forced. In this work, we aim to strictly control sparsity by imposing hard cardinality constraints on the number of nonzero components in $\mathbf{w}^*$ and this leads to the following problem:

$$\begin{aligned} (\text{FS-SVM}) \quad & \min_{\mathbf{w}, b, \xi} \quad \frac{1}{2} \|\mathbf{w}\|_2^2 + C \sum_{i=1}^m \xi_i \\ & \text{s.t.} \quad y^i (\mathbf{w}^\top \mathbf{x}^i + b) \geq 1 - \xi_i, \quad i \in [1:m], \\ & \quad \|\mathbf{w}\|_0 \leq B \\ & \quad \xi \geq \mathbf{0}, \end{aligned} \tag{5.2}$$

where the user-defined parameter $B \in \mathbb{Z}_+$ constitutes a budget on the number of features used, indicating that the number of features involved in the classifier is at most B . Even though objective function and all other constraints are convex quadratic or linear, the cardinality constraint $\|\mathbf{w}\|_0 \leq B$ is of combinatorial type, which renders (5.2) NP-hard.

The integration of feature selection in linear SVM has been intensively studied over the past few decades. There are essentially two types of models for feature selection. The first approach involves using surrogates to pursue sparse solutions, such as well-known ℓ_1 regularization and its variants, such as [FM04, HRTZ04, NDT10, YCHL10, WZZ06, ZH05]. The second approach directly addresses the cardinality constraint. Several studies [MPWL14, LMMRC19] focus on the representing the constraint through mixed-integer modelling, incorporating a large constant (referred to as big-M), and solving the problem using LP-based branch and bound methods. However, the numerical performance of this approach is highly dependent on the choice of the big-M constant. Estimating this constant is often as challenging as solving the original problem itself, leading to inefficiencies, particularly when limited information is available regarding the optimal value of big-M.

In this chapter, we propose several novel semidefinite relaxations based on the complementarity reformulation of the problem (5.2). Exploiting the sparsity pattern of the relaxations, we decompose the problems and obtain equivalent relaxations in a much smaller cone, making the conic approaches scalable. To make the best usage of the decomposed relaxations, we propose heuristics using the information of its optimal solution. Moreover, an exact procedure is proposed by solving a sequence of mixed-integer decomposed semidefinite optimization problems. Numerical results on classical benchmarking datasets are reported, showing the efficiency and effectiveness of our approach. Before proceeding to the details, we introduce several sets that will be frequently used throughout the remainder of the text.

The margin set is denoted by

$$\Xi = \{(\mathbf{w}, b, \boldsymbol{\xi}) \in \mathbb{R}^{n+1+m} : y^i(\mathbf{w}^\top \mathbf{x}^i + b) \geq 1 - \xi_i, i \in [1:m], \boldsymbol{\xi} \geq \mathbf{0}\}, \quad (5.3)$$

so that problem (5.2) reads as:

$$\begin{aligned} \min_{\mathbf{w}, b, \boldsymbol{\xi}} \quad & \frac{1}{2} \|\mathbf{w}\|_2^2 + C \sum_{i=1}^m \xi_i \\ \text{s.t.} \quad & (\mathbf{w}, b, \boldsymbol{\xi}) \in \Xi \\ & \|\mathbf{w}\|_0 \leq B. \end{aligned}$$

In the next sections, we will also make use of big-M reformulations, so it is convenient to introduce the two sets

$$\Omega_M = \{(\mathbf{w}, \mathbf{v}) \in \mathbb{R}^{2n} : -M\mathbf{v} \leq \mathbf{w} \leq M\mathbf{v}, \mathbf{e}^\top \mathbf{v} = B\} \quad (5.4)$$

and

$$\mathcal{B}_M = \{(\mathbf{w}, \mathbf{u}) \in \mathbb{R}^{2n} : -M(1 - \mathbf{u}) \leq \mathbf{w} \leq M(1 - \mathbf{u}), \mathbf{e}^\top \mathbf{u} = n - B\}, \quad (5.5)$$

where the big-M parameter $M > 0$ is large enough (see Remark 4 below) and the variables $\mathbf{u} = 1 - \mathbf{v}$ are the negated $\mathbf{v}$, sometimes more convenient in notation.

5.2. Exact Mixed-Integer approaches

In this section, we propose two exact mixed-integer formulations of the original problem (5.2), using two different mixed-integer systems to present the cardinality constraint with auxillary

binary variables introduced as an indicator, namely a big-M strategy and a complementarity constraint strategy.

Differently from [LMMRC19], we study the following MIQP model incorporating the standard ℓ_2 -norm on the hyperplane weights as the regularization term:

$$\begin{aligned}
(\text{BigMP}) \quad & \min_{\mathbf{w}, b, \boldsymbol{\xi}, \mathbf{v}} \quad \frac{1}{2} \|\mathbf{w}\|_2^2 + C \sum_{i=1}^m \xi_i \\
& \text{s.t.} \quad (\mathbf{w}, b, \boldsymbol{\xi}) \in \Xi \\
& \quad (\mathbf{w}, \mathbf{v}) \in \Omega_M \\
& \quad \mathbf{v} \in \{0, 1\}^n,
\end{aligned} \tag{5.6}$$

where Ξ and Ω_M are defined in (5.3) and (5.4).

To show the equivalence between (5.2) and (5.6), we consider a B -sparse $\mathbf{w}$, e.g., $\|\mathbf{w}\|_0 \leq B$. If $\|\mathbf{w}\|_0 = B$, there exists a $\mathbf{v} \in \{0, 1\}^n$ such that $(\mathbf{w}, \mathbf{v}) \in \Omega_M$, where $\mathbf{v}$ is given by the rule that $v_j = 0$ if $w_j = 0$ and $v_j = 1$ if $w_j \neq 0$; If $\|\mathbf{w}\|_0 < B$, there also exists a $\mathbf{v} \in \{0, 1\}^n$ such that $(\mathbf{w}, \mathbf{v}) \in \Omega_M$, where $\mathbf{v}$ is given by the rule that $v_j = 1$ for $j \in I$ and $v_j = 0$ otherwise, where I is a index set covers all of features with $w_j \neq 0$, satisfying $|I| = B$.

Remark 2. *It is worth noticing that the optimal value of model (5.6) remains the same if we substitute the equality constraint $\mathbf{e}^\top \mathbf{v} = B$ by the budget constraint $\mathbf{e}^\top \mathbf{v} \leq B$, which is derived directly from the targeted problem (5.2).*

Alternatively, instead of expressing the cardinality constraint by the big-M formulation, we employ a set of binary variables $\mathbf{u} \in \{0, 1\}^n$ and a complementarity constraint to count the nonzero elements of $\mathbf{w}$ and end up with another mixed-integer reformulation of (5.2):

$$\begin{aligned}
(\text{CoP}) \quad & \min_{\mathbf{w}, b, \boldsymbol{\xi}, \mathbf{u}} \quad \frac{1}{2} \|\mathbf{w}\|_2^2 + C \sum_{i=1}^m \xi_i \\
& \text{s.t.} \quad (\mathbf{w}, b, \boldsymbol{\xi}) \in \Xi \\
& \quad u_j w_j = 0, \quad j \in [1:n] \\
& \quad \mathbf{e}^\top \mathbf{u} = n - B \\
& \quad \mathbf{u} \in \{0, 1\}^n.
\end{aligned} \tag{5.7}$$

Unlike (5.6), the binary variables u_j count the number of zero elements of $\mathbf{w}$. For a (5.7)-feasible point $(\mathbf{w}, b, \boldsymbol{\xi}, \mathbf{u})$, we let $\mathbf{v} = 1 - \mathbf{u}$ and the point $(\mathbf{w}, b, \boldsymbol{\xi}, \mathbf{v})$ is (5.6)-feasible. To show this, we only need to show $(\mathbf{w}, \mathbf{v}) \in \Omega_M$, namely, that $(\mathbf{w}, \mathbf{v})$ satisfies the big-M constraint in (5.4). It is clear that the constraint holds for the coordinates of $\mathbf{w}$ with $w_j = 0$. On the other hand, the complementarity constraint $(1 - v_j)w_j = 0$ implies $v_j = 1$ if $w_j \neq 0$, and therefore the point $(\mathbf{w}, \mathbf{v}) \in \Omega_M$. Hence, the problem (5.7) is equivalent to the problem (5.6).

Remark 3. *In the above problem (5.7), the constraint $\mathbf{e}^\top \mathbf{u} = n - B$ can be equivalently replaced by the constraint $\mathbf{e}^\top \mathbf{u} \geq n - B$ for the same reason as Remark 2, and the binary constraint $u_j \in \{0, 1\}$ can be equivalently reduced to the box constraint $0 \leq u_j \leq 1$ due to the existence of the complementarity constraint. To show this, we consider the model with the box constraint and assume that $\mathbf{w}$ is ρ -sparse with $\rho > B$. Due to the constraint $u_j w_j = 0$, $u_j = 0$ if w_j is nonzero and $0 \leq u_j \leq 1$ if w_j is zero. Hence, $\mathbf{e}^\top \mathbf{u} \leq n - \rho < n - B$, which contradicts the budget constraint. Therefore all $\mathbf{w}$ vectors with cardinality strictly larger than B are rendered infeasible to the model (5.7) and its relaxation (5.8) below.*

In the first formulation, we need a big-M parameter M for the reformulation. It is known that the choice of M has a profound influence on the computational performance if we solve

the problem with off-the-shell global solvers. In a branch-and-bound framework, we solve LP relaxation at each node, and in many cases, the LP relaxation can be as weak as the SVM problem without feature selection, if too large M is chosen, as shown later in Theorem 35. In contrast, in the second formulation, the difficulty of the problem mainly goes to the complementarity constraint, and no big-M needed.

5.3. Tractable Relaxations

In this section, we propose several tractable relaxations, e.g., LP relaxations, and SDP relaxations, of the FS-SVM model, and a comparison is made among some of them. Moreover, a decomposition strategy is proposed by exploring the sparse pattern involved in the both objective and feasible region of the original problem, leading to a much smaller and scalable SDP-based relaxation.

5.3.1. Box relaxation

In this subsection, we study the box relaxation (also known as a type of LP relaxation) of two exact mixed-integer formulations presented in the last chapter and point out the crucial role of M that it plays in the first formulation.

After relaxing the binarity constraints to the box constraints in the problem (5.6), we arrive at:

$$\begin{aligned}
(\text{BoxMP}) \quad & \min_{\mathbf{w}, b, \boldsymbol{\xi}, \mathbf{v}} \quad \frac{1}{2} \|\mathbf{w}\|_2^2 + C \sum_{i=1}^m \xi_i \\
& (\mathbf{w}, b, \boldsymbol{\xi}) \in \Xi \\
& (\mathbf{w}, \mathbf{v}) \in \Omega_M \\
& \mathbf{0} \leq \mathbf{v} \leq \mathbf{e}.
\end{aligned} \tag{5.8}$$

The relaxation is a convex quadratic optimization problem solvable in polynomial time, therefore it has been widely used in many previous works, see references in [LMMRC19]. However, the tightness of the relaxation is highly dependent on the choice of M and there is a possibility that it can be as weak as the original SVM model without feature selection (5.1). To be more precise, if M is larger than a data-based threshold, the relaxation (5.8) is equivalent to the problem (5.1), as presented in the following theorem.

Theorem 35. *If the parameter M in the problem (5.8) satisfies*

$$M \geq \frac{\|\mathbf{w}^*\|_1}{B},$$

where $(\mathbf{w}^, b^*, \boldsymbol{\xi}^*)$ is an (5.1)-optimal solution, then (5.8) is equivalent to (5.1).*

Proof. It is clear to see that (5.1) is a relaxation of (5.8), as the former includes the constraint $(\mathbf{w}, b, \boldsymbol{\xi}) \in \Xi$. Hence the optimal value of (5.1) is not larger than that of (5.8). To show the converse, we assume that $(\mathbf{w}^*, b^*, \boldsymbol{\xi}^*)$ is optimal to the problem (5.1). Let $v_j = \frac{|w_j^*|}{M}$ for all $j \in [1:n]$. Then the point $(\mathbf{w}^*, b^*, \boldsymbol{\xi}^*, \mathbf{v})$ is (5.8)-feasible. Indeed, due to the construction of v_j , the constraint $-Mv_j \leq w_j \leq Mv_j$ is satisfied automatically for all $j \in [1:n]$. Since $\mathbf{e}^\top \mathbf{v} = \frac{\|\mathbf{w}^*\|_1}{M}$, the constraint $\mathbf{e}^\top \mathbf{v} \leq B$ is implied by the assumption we made. Since the objective functions in both problems coincide, the claim follows. $\square$

On the other hand, M can not be chosen smaller than:

$$M \geq \|\bar{\mathbf{w}}^*\|_\infty,$$

where $\bar{\mathbf{w}}^*$ is part of an optimal solution to (5.2). Otherwise, the optimal solution will be cut off in the sense that the problem (5.6) is not a reformulation for an insufficient choice of M . Hence, the proper range of the parameter M has to be:

$$M \in \left[\|\bar{\mathbf{w}}^*\|_\infty, \frac{\|\mathbf{w}^*\|_1}{B} \right). \quad (5.9)$$

However, this set can be empty depending on the problem data itself, which means no matter what M we chose, the box relaxation (5.8) ends up with one of the following situations: (i) no improvement compared to the original SVM problem (5.1); (ii) the optimal solution of the targeted problem (5.2) is cut off, because M is too small to render (5.8) a reformulation of (5.2). This is quite problematic in practice and it encourages us to explore the other formulation (5.7).

Remark 4. *Since estimating the range of $\bar{\mathbf{w}}^*$ is as difficult as solving the problem itself, we usually choose M under a safe threshold, which is much larger than the tightest upper bound. In most practical cases, the condition $M \geq \frac{\|\mathbf{w}^*\|_1}{B}$ holds then as well, which means the LP relaxation (5.8) does not provide any improvement compared to the plain problem (5.1).*

Due to the presence of complementarity constraints, the box relaxation of (5.7) is equivalent to (5.7) itself, as mentioned in Remark 3. However, the combinatorial nature of a complementarity constraint is the main difficulty in both the unrelaxed and the relaxed version, so this continuous reformulation does not solve the problem immediately without any further processing.

5.3.2. The Shor relaxation

In this subsection, we study the Shor relaxation, a type of SDP relaxation, of the two proposed mixed-integer reformulations.

Letting $\mathbf{W} = \mathbf{w}\mathbf{w}^\top$, $\mathbf{V} = \mathbf{v}\mathbf{v}^\top$ and dropping the rank one constraint, the Shor relaxation of (5.6) reads as:

$$\begin{aligned} (\text{SMP}) \quad & \min_{\mathbf{w}, \mathbf{W}, b, \boldsymbol{\xi}, \mathbf{v}, \mathbf{V}} \quad \frac{1}{2} \text{trace}(\mathbf{W}) + C \sum_{i=1}^m \xi_i \\ & \text{s.t.} \quad (\mathbf{w}, b, \boldsymbol{\xi}) \in \Xi \\ & \quad (\mathbf{w}, \mathbf{v}) \in \Omega_M \\ & \quad \text{diag } \mathbf{V} = \mathbf{v} \\ & \quad \begin{bmatrix} 1 & \mathbf{w}^\top \\ \mathbf{w} & \mathbf{W} \end{bmatrix} \in \mathcal{S}_+^{n+1}, \quad \begin{bmatrix} 1 & \mathbf{v}^\top \\ \mathbf{v} & \mathbf{V} \end{bmatrix} \in \mathcal{S}_+^{n+1}. \end{aligned} \quad (5.10)$$

Theorem 36. *The Shor relaxation (5.10) is equivalent to the box relaxation (5.8).*

Proof. For a (5.10)-feasible point $(\mathbf{w}, \mathbf{W}, b, \boldsymbol{\xi}, \mathbf{v}, \mathbf{V})$, the projected point $(\mathbf{w}, b, \boldsymbol{\xi}, \mathbf{v})$ satisfies all of the constraints of (5.8). Moreover, we have $\text{trace}(\mathbf{W}) \geq \|\mathbf{w}\|_2^2$, hence the relaxation (5.8) is weaker than the relaxation (5.10). To show the other direction, we suppose that the point $(\mathbf{w}, b, \boldsymbol{\xi}, \mathbf{v})$ is feasible to the problem (5.8), and let $\mathbf{W} = \mathbf{w}\mathbf{w}^\top$ and $\mathbf{V} = \mathbf{v}\mathbf{v}^\top + \mathbf{Z}$, where $\mathbf{Z} \in \mathcal{S}_+^n$ is a diagonal matrix with $Z_{jj} = v_j - v_j^2 \geq 0$ for all $j \in [1:n]$. The constructed point $(\mathbf{w}, \mathbf{W}, b, \boldsymbol{\xi}, \mathbf{v}, \mathbf{V})$ then satisfies all constraints of (5.10). Indeed, the last positive-semidefiniteness constraint follows by the construction of $\mathbf{V}$ due to the Schur complement, which therefore completes the proof. $\square$

Letting $\mathbf{W} = \mathbf{w}\mathbf{w}^\top$, $\mathbf{R} = \mathbf{w}\mathbf{u}^\top$, $\mathbf{U} = \mathbf{u}\mathbf{u}^\top$ and dropping the rank-one constraint, we arrive at the Shor relaxation of (5.7):

$$\begin{aligned}
(\text{SCoP}) \quad & \min_{\substack{\mathbf{w}, \mathbf{W}, b, \boldsymbol{\xi}, \\ \mathbf{R}, \mathbf{u}, \mathbf{U}}} & \frac{1}{2} \text{trace}(\mathbf{W}) + C \sum_{i=1}^m \xi_i \\
& \text{s.t.} & (\mathbf{w}, b, \boldsymbol{\xi}) \in \Xi \\
& & \mathbf{e}^\top \mathbf{u} = n - B \\
& & \text{diag } \mathbf{R} = \mathbf{o}, \quad \text{diag } \mathbf{U} = \mathbf{u} \\
& & \begin{bmatrix} 1 & \mathbf{w}^\top & \mathbf{u}^\top \\ \mathbf{w} & \mathbf{W} & \mathbf{R} \\ \mathbf{u} & \mathbf{R}^\top & \mathbf{U} \end{bmatrix} \in \mathcal{S}_+^{2n+1}
\end{aligned} \tag{5.11}$$

Theorem 37. *If there exists part $\mathbf{w}$ of an optimal solution to (5.1) with $\|\mathbf{w}\|_0 > B$, then the optimal value of (5.11) is strictly larger than the optimal value of (5.1).*

Proof. For a given (5.1)-feasible point $(\mathbf{w}, b, \boldsymbol{\xi})$ with $\|\mathbf{w}\|_0 > B$, we restrict the corresponding variables of (5.11) to the same value and estimate the optimal value of the restricted version of (5.11). The matrix

$$\begin{bmatrix} 1 & \mathbf{w}^\top & \mathbf{u}^\top \\ \mathbf{w} & \mathbf{W} & \mathbf{R} \\ \mathbf{u} & \mathbf{R}^\top & \mathbf{U} \end{bmatrix} \in \mathcal{S}_+^{2n+1}, \quad \text{with} \quad \text{diag } \mathbf{R} = \mathbf{o} \text{ and } \text{diag } \mathbf{U} = \mathbf{u},$$

has a principal submatrix

$$\begin{bmatrix} 1 & w_j & u_j \\ w_j & W_{jj} & 0 \\ u_j & 0 & u_j \end{bmatrix} \in \mathcal{S}_+^3, \quad \text{for all } j \in [1:n],$$

which implies that

$$\begin{cases} 0 \leq u_j \leq 1, & W_{jj} \geq w_j^2, & j \in [1:n] \\ (W_{jj} - w_j^2)(u_j - u_j^2) - w_j^2 u_j^2 \geq 0, & j \in [1:n]. \end{cases}$$

It follows that

$$\begin{cases} 0 \leq u_j \leq 1, & \text{if } j \in S_0 := \{j : w_j = 0\}, \\ 0 \leq u_j \leq \frac{W_{jj} - w_j^2}{W_{jj}} < 1 & \text{if } j \in S_+ := [1:n] \setminus S_0 \end{cases}$$

(observe that $w_j > 0$ implies $W_{jj} > 0$). Recalling the assumption $\|\mathbf{w}\|_0 > B$, it follows that

$$\sum_{j \in S_0} u_j < n - B.$$

Together with the constraint

$$\mathbf{e}^\top \mathbf{u} = n - B,$$

we have $u_j > 0$ for some $j \in S_+$, which implies that $W_j > w_j^2$ for some $j \in S_+$. Hence, the optimal value of (5.11) is strictly larger than the optimal value of (5.1). $\square$

There is no direct comparison between the relaxation (5.10) and the relaxation (5.11). But given the case we have a valid estimate of M at hand satisfying (5.9), the relaxation (5.11) can be tightened with a small computational cost by adding the big-M constraint via $\mathcal{B}_M$, which reads as:

$$\begin{aligned}
(\text{SCoMP}) \quad & \min_{\substack{\mathbf{w}, \mathbf{W}, b, \boldsymbol{\xi}, \\ \mathbf{R}, \mathbf{u}, \mathbf{U}}} & \frac{1}{2} \text{trace}(\mathbf{W}) + C \sum_{i=1}^m \xi_i \\
& \text{s.t.} & (\mathbf{w}, b, \boldsymbol{\xi}) \in \Xi \\
& & (\mathbf{w}, \mathbf{u}) \in \mathcal{B}_M \\
& & \text{diag } \mathbf{R} = \mathbf{o}, \quad \text{diag } \mathbf{U} = \mathbf{u}, \\
& & \begin{bmatrix} 1 & \mathbf{w}^\top & \mathbf{u}^\top \\ \mathbf{w} & \mathbf{W} & \mathbf{R} \\ \mathbf{u} & \mathbf{R}^\top & \mathbf{U} \end{bmatrix} \in \mathcal{S}_+^{2n+1}.
\end{aligned} \tag{5.12}$$

Theorem 38. *The relaxation (5.12) is tighter than the relaxations (5.10) and (5.11).*

Proof. It is clear that the relaxation (5.12) is tighter than the relaxation (5.11) due to the additional big-M constraint. To show the other claim, we construct a (5.10)-feasible point, given a (5.12)-feasible point. Suppose that the point $(\mathbf{w}, \mathbf{W}, b, \boldsymbol{\xi}, \mathbf{R}, \mathbf{u}, \mathbf{U})$ is (5.12)-feasible, let $\mathbf{v} = \mathbf{e} - \mathbf{u}$ and $\mathbf{V} = \mathbf{v}\mathbf{v}^\top + \mathbf{Z}$, where $\mathbf{Z} \in \mathcal{S}_+^n$ is a diagonal matrix with $Z_{jj} = v_j - v_j^2 \geq 0$ for all $j \in [1:n]$. Then $(\mathbf{w}, \mathbf{W}, b, \boldsymbol{\xi}, \mathbf{v}, \mathbf{V})$ is (5.10)-feasible with the same objective value. $\square$

5.3.3. Decomposed SDP-based relaxations

In this subsection, we propose a procedure for obtaining equivalent and much smaller SDP-based relaxations by exploiting the sparsity patterns of the Shor relaxations.

We notice that, for example: in problem (5.11), some sparsity patterns appear both in the objective and the constraints. e.g., only diagonal elements of matrices $\mathbf{W}$, $\mathbf{R}$, and $\mathbf{U}$ are involved. Based on this observation, we write down the following decomposed version of (5.11):

$$\begin{aligned}
(\text{DSCoP}) \quad & \min_{\substack{\mathbf{w}, \text{diag } \mathbf{W}, b, \boldsymbol{\xi}, \mathbf{u}}} & \frac{1}{2} \text{trace}(\mathbf{W}) + C \sum_{i=1}^m \xi_i \\
& \text{s.t.} & (\mathbf{w}, b, \boldsymbol{\xi}) \in \Xi \\
& & \mathbf{e}^\top \mathbf{u} = n - B \\
& & \begin{bmatrix} 1 & w_j & u_j \\ w_j & W_{jj} & 0 \\ u_j & 0 & u_j \end{bmatrix} \in \mathcal{S}_+^3, \quad j \in [1:n]
\end{aligned} \tag{5.13}$$

Theorem 39. *The relaxation (5.11) is equivalent to the relaxation (5.13).*

Proof. Suppose that $(\mathbf{w}, \mathbf{W}, b, \boldsymbol{\xi}, \mathbf{u})$ is (5.11)-feasible. It is clear that $(\mathbf{w}, \text{diag } \mathbf{W}, b, \boldsymbol{\xi}, \mathbf{u})$ is (5.13)-feasible. In the opposite direction, we assume that the point $(\mathbf{w}, \text{diag } \mathbf{W}, b, \boldsymbol{\xi}, \mathbf{u})$ is (5.13)-feasible. It is clear that all linear constraints in (5.11) are satisfied if we let $\text{diag } \mathbf{R} = \mathbf{o}$ and $\text{diag } \mathbf{U} = \mathbf{u}$. The only thing that needs to be checked is if the following partial positive-semidefinite matrix is

positive-semidefinite completable:

$$\begin{bmatrix} 1 & w_1 & \dots & w_n & u_1 & \dots & u_n \\ w_1 & W_{11} & * & * & 0 & * & * \\ \vdots & * & \ddots & * & * & \ddots & * \\ w_n & * & * & W_{nn} & * & * & 0 \\ u_1 & 0 & * & * & u_1 & * & * \\ \vdots & * & \ddots & * & * & \ddots & * \\ u_n & * & * & 0 & * & * & u_n \end{bmatrix}.$$

By $\mathcal{G}$ we denote the specification graph of the above partial positive-semidefinite matrix. Since the vertex w_i is not connected with w_j and u_j for all $i \neq j$, and the vertex u_i is not connected with u_j and w_j for all $i \neq j$, the only cycle we can find in the graph is the triangle $1 \rightarrow w_j \rightarrow u_j \rightarrow 1$ for all $[1:n]$. Therefore $\mathcal{G}$ is chordal. Due to [BSM03, Theorem 1.39], if the specification graph of a partial positive-semidefinite matrix is chordal, it is always completable to a positive-semidefinite matrix. $\square$

Following a similar fashion, we propose the decomposed version of the relaxation (5.12):

$$\begin{aligned} (\text{DSCoMP}) \quad & \min_{\mathbf{w}, \text{diag } \mathbf{W}, \mathbf{b}, \boldsymbol{\xi}, \mathbf{u}} \quad \frac{1}{2} \text{trace}(\mathbf{W}) + C \sum_{i=1}^m \xi_i \\ & \text{s.t.} \quad (\mathbf{w}, \mathbf{b}, \boldsymbol{\xi}) \in \Xi \\ & \quad (\mathbf{w}, \mathbf{u}) \in \mathcal{B}_M \\ & \quad \begin{bmatrix} 1 & w_j & u_j \\ w_j & W_{jj} & 0 \\ u_j & 0 & u_j \end{bmatrix} \in \mathcal{S}_+^3, \quad j \in [1:n]. \end{aligned} \tag{5.14}$$

Corollary 40. *The relaxation (5.14) is equivalent to the relaxation (5.12).*

Proof. See the proof of Theorem 39. $\square$

For the formulation (5.6), the decomposed version of its Shor relaxation (5.15) reads as:

$$\begin{aligned} (\text{DSMP}) \quad & \min_{\mathbf{w}, \text{diag } \mathbf{W}, \mathbf{b}, \boldsymbol{\xi}, \mathbf{v}} \quad \frac{1}{2} \text{trace}(\mathbf{W}) + C \sum_{i=1}^m \xi_i \\ & \text{s.t.} \quad (\mathbf{w}, \mathbf{b}, \boldsymbol{\xi}) \in \Xi \\ & \quad (\mathbf{w}, \mathbf{v}) \in \Omega_M \\ & \quad \begin{bmatrix} 1 & w_j \\ w_j & W_{jj} \end{bmatrix} \in \mathcal{S}_+^2, \quad \begin{bmatrix} 1 & v_j \\ v_j & v_j \end{bmatrix} \in \mathcal{S}_+^2, \quad j \in [1:n]. \end{aligned} \tag{5.15}$$

Corollary 41. *The relaxation (5.15) is equivalent to the relaxation (5.10).*

Proof. See the proof of Theorem 39. $\square$

5.4. The algorithms

In this section, we introduce both a heuristic and an exact algorithm for solving the feature selection problem we are interested in. In particular, we will first present two heuristic strategies

to find good upper bound values for our problem and a heuristic strategy to estimate a tight value for the big-M parameter M for the problem. All strategies are based on the resolution of the decomposed SDP-based relaxations presented in the previous section. These strategies will then be embedded in a heuristic algorithm which can ultimately improve the heuristic upper bound. Finally, we present an exact algorithm that implements one of the two upper bound procedures as a subroutine and consists of the resolution of a sequence of Mixed-Integer Second-Order Cone Optimization problems (MISOCPs).

5.4.1. Two upper bounding strategies

In this subsection, we propose two strategies for obtaining an upper bound, namely a feasible objective value of a feasible solution $(\widehat{\mathbf{w}}, \widehat{\mathbf{b}}, \widehat{\boldsymbol{\xi}})$ to (5.2), based on the optimal solution of relaxations.

Local Search

The first strategy, denoted as Local Search, consists of three main steps and employs a user-specified excess parameter k which will temporarily allow to work on just $k + B$ instead of all features. We first solve the relaxation by off-the-shelf software such as **Mosek**, an interior point solver. With the optimal solution $\mathbf{u}^*$ in hand, we want to search for a feasible point nearby. We next sort features according to increasing u^* . Namely, if u_j^* is close to 1, then the feature j is unlikely to be selected because of the complementarity constraint $w_j u_j = 0$, while if u_j^* is closer to 0, then the feature j is more likely to be selected. Based on it, we select the first $k + B$ features with small corresponding u , denoted by a set $K \subseteq [1:n]$ with $|K| = k + B$. In the end, we solve the restricted mixed-integer problem:

$$\begin{aligned}
 (\text{CoP}(K)) \quad & \min_{\mathbf{w}, \mathbf{b}, \boldsymbol{\xi}, \mathbf{u}} \quad \frac{1}{2} \|\mathbf{w}\|_2^2 + C \sum_{i=1}^m \xi_i \\
 & \text{s.t.} \quad (\mathbf{w}, \mathbf{b}, \boldsymbol{\xi}) \in \Xi \\
 & \quad \sum_{j \in K} u_j = |K| - B \\
 & \quad u_j w_j = 0, \quad j \in K, \\
 & \quad u_j \in \{0, 1\}, \quad j \in K.
 \end{aligned} \tag{5.16}$$

which is much smaller than the original problem (5.7), as it only involves $|K|$ binary variables u_j . The problem is solved by the global solver **Gurobi**. This first upper bounding strategy is described in Algorithm 4.

Data: $(\mathbf{x}^i, y^i) \in \mathbb{R}^n \times \{1, -1\}$, B, C , parameter $k \in [0:n - B]$

Result: The set of selected features and an upper bound UB.

Solve the relaxation (5.13) and return the optimal solution $\mathbf{u}^*$;

Generate a subset $K \subseteq [1:n]$ with $|K| = k + B$ such that $u_i^* \leq u_j^*$ for all $i \in K$ and for all $j \in [1:n] \setminus K$

Solve the mixed-integer problem (5.16)

Algorithm 4: Local Search

Kernel Search

Slightly enlarging the search set at each iteration, we propose the following search strategy for a feasible solution of (5.2). Our Kernel Search strategy consists in solving a sequence of

small MIQPs considering an initial ranking on the features. This strategy takes inspiration from a heuristic proposed by [AMS10] for the general multi-dimensional knapsack problem which has been applied to different kinds of problems such as location problems [GS12] and portfolio optimization [AMS12]. Recently, [LMMRC19] also applied it to solve their embedded feature selection for ℓ_1 -SVMs. In this strategy, based on some ranking, a restricted set of promising features is kept and updated throughout each iteration. Usually, for ranking variables in general MILPs, LP theory is exploited using both relaxed solution values and reduced costs, a strategy followed by [LMMRC19] for ranking features.

By contrast, in our case we decided to use as a ranking criterion the coordinates of an optimal solution $\mathbf{u}^*$ to the decomposed SDP relaxation (5.13). Features are sorted with respect to this vector $\mathbf{u}^*$: the smaller the value of u_j^* , the higher the relevance of feature j . Since in practice (5.13) leads to much tighter lower bound solutions in contrast to the plain relaxation (5.8) similar to the model used by [LMMRC19], there is a justified hope that the ranking provided by $\mathbf{u}^*$ is closer to the relevance of features.

Let K be the ordered set of features, ρ a user-defined parameter and $N := \lceil \frac{n-\rho}{\rho} \rceil$. Then K is divided into N subsets denoted as K_i for $i \in [1:N]$. In particular, each subset K_i , $i \in [1:N-1]$, will be composed of ρ features, and K_N will contain the remaining features.

An initial value of the UB is set to ∞ . Also define a feature subset $\overline{K} \subseteq [1:n]$ as the *kernel set*, containing features that are kept at the next iteration. To initialize, set $\overline{K} = \emptyset$. At each iteration k , the heuristic considers set $\mathcal{K}_k = \overline{K} \cup K_k$ i.e. the union of the features in the kernel set and the features in the set K_k . To update the UB, at each iteration we solve problem $\text{CoP}(\mathcal{K}_k)$ plus the following two constraints:

$$\frac{1}{2} \|\mathbf{w}\|_2^2 + C \sum_{i=1}^m \xi_i \leq \text{UB}, \quad (5.17)$$

$$\sum_{j \in K_k} u_j \leq |K_k| - 1. \quad (5.18)$$

Constraint (5.17) restricts the objective function to take a value smaller than or equal to the current upper bound, while constraint (5.18) makes sure that at least one feature from the new set K_k is selected. Note that, due to the addition of these constraints, the optimization problems may potentially be infeasible. If the optimization problems are feasible, the new features selected are then added to the kernel set $\overline{K}$ used in the next iteration since adding these features obtains an identical or better upper bound. Conversely, the set of features of $\overline{K}$ that have not been chosen in the optimal solution in the previous iterations is removed from the kernel set $\overline{K}$. The removal of some of the features from the kernel set is decisive in that it does not excessively increase the number of binary variables considered in each iteration. We decided to remove the features that were not selected in the previous two iterations. The set of added features is denoted as K^+ and the set of removed features as K^- . The resulting kernel set for the next iteration is $\overline{K} \cup K^+ \setminus K^-$. Conversely, if the problem is infeasible, the kernel set is not modified and the procedure skips to the next iteration. The Kernel Search strategy is described in Algorithm 5.

Data: $(x^i, y^i) \in \mathbb{R}^n \times \{1, -1\}$, B , C , parameter $\rho \in [1:n]$.

Result: The set of selected features and an upper bound UB.

Solve the relaxation (5.13) and return the optimal solution u^* ;

Sort features according to the numerical size of corresponding u^* increasingly, denoted by a vector $K \in \mathbb{Z}_+^n$;

Separate K into $N := \lfloor \frac{n}{\rho} \rfloor$ groups evenly such that $K = [K_1, \dots, K_N]$ and set $\bar{K} = \emptyset$, $k = 1$, $UB = \infty$;

while $k \leq N$ **do**

$\mathcal{K}_k \leftarrow K_k \cup \bar{K}$;

Solve the problem $\text{CoP}(\mathcal{K}_k) + (5.17) + (5.18)$;

if *problem is feasible* **then**

Return solution u^* and optimal value UB^* ;

$UB \leftarrow UB^*$;

Build $K^+ := \{j \in K_k : j \text{ selected at current iteration}\}$;

Build $K^- := \{j \in \bar{K} : j \text{ not selected in the solution of the last two iterations}\}$;

Update $\bar{K} \leftarrow \bar{K} \cup K^+ \setminus K^-$;

end

$k \leftarrow k + 1$;

end

Algorithm 5: Kernel Search

5.4.2. A strategy to tighten the big-M parameter

Given an upper bound UB of problem (5.2), one can estimate the upper bound of w_j by solving the following problem:

$$\begin{aligned}
 \bar{M}_j := & \max_{\substack{\mathbf{w}, \text{diag } W, b, \boldsymbol{\xi}, \mathbf{u} \\ \text{s.t.}}} & w_j \\
 & (\mathbf{w}, b, \boldsymbol{\xi}) \in \Xi \\
 & \frac{1}{2} \|\mathbf{w}\|_2^2 + C \sum_{i=1}^m \xi_i \leq UB \\
 & \mathbf{e}^\top \mathbf{u} = n - B \\
 & \begin{bmatrix} 1 & w_j & u_j \\ w_j & W_{jj} & 0 \\ u_j & 0 & u_j \end{bmatrix} \in \mathcal{S}_+^3, \quad j \in [1:n].
 \end{aligned} \tag{5.19}$$

Similarly, the lower bound of w_j can be obtained by solving the same optimization problem with the max replaced by min and we denote the optimal value by $\underline{M}_j$. The proper choice of M is $\max_{j \in [1:n]} \{\max\{|\bar{M}_j|, |\underline{M}_j|\}\}$.

Proposition 42. Let $[\hat{\mathbf{w}}, \hat{b}, \hat{\boldsymbol{\xi}}]$ be a feasible solution of (5.2) and $UB := \frac{1}{2} \|\hat{\mathbf{w}}\|_2^2 + C \sum_{i=1}^m \hat{\xi}_i$ be the corresponding objective value. Then

$$M = \max_{j \in [1:n]} \{\max\{|\bar{M}_j|, |\underline{M}_j|\}\}$$

is a valid bound, where $\bar{M}_j$ and $\underline{M}_j$ are obtained by (5.19).

Proof. obvious. □

Remark 5. *The proposed bound M is tighter than most of the currently used big- M s. Nevertheless, empirically, the tightening of any such bound is of little relevance unfortunately, because they still fall outside of the range specified in (5.9).*

5.4.3. The heuristic algorithm

In this subsection, we propose a procedure for generating a good feasible solution of (5.2) taking one of the two upper bound strategies and the big- M estimation strategy into account. At first, a feasible solution is computed with the chosen upper bound strategy. Then, according to Proposition 42, we solve a series of $2n$ many SDP problems to find, for each feature j , values $\overline{M}_j, \underline{M}_j$ and then obtain the big- M parameter M . If this value satisfies (5.9), we know that the SDP formulation (5.14) is better than (5.13), thus we compute a new relaxed solution $\mathbf{u}^*$ of (5.14). We thus use vector $\mathbf{u}^*$ as a ranking criterion for the chosen upper bound strategy, we run the strategy again and we update the heuristic solution if a better one is found.

Data: $(\mathbf{x}^i, y^i) \in \mathbb{R}^n \times \{1, -1\}$, B , C , the parameter $k \in [0:n-B]$ or the parameter $\rho \in [1:n]$.
Result: $(\mathbf{w}^*, b^*, \boldsymbol{\xi}^*)$ solution of (5.2).
Run Algorithm 4 or Algorithm 5 to obtain a feasible solution $(\widehat{\mathbf{w}}, \widehat{b}, \widehat{\boldsymbol{\xi}})$ to (5.2) with upper bound value $\widehat{UB}$;
Set $(\mathbf{w}^*, b^*, \boldsymbol{\xi}^*) \leftarrow (\widehat{\mathbf{w}}, \widehat{b}, \widehat{\boldsymbol{\xi}})$;
Obtain M according to Proposition 42;
Solve problem (5.1) and obtain $\mathbf{w}^*$;
if $M < \frac{\|\mathbf{w}^*\|_1}{B}$ **then**
 Solve problem (5.14) and obtain $\mathbf{u}^*$;
 Run Algorithm 4 or Algorithm 5 starting with $\mathbf{u}^*$, get $(\overline{\mathbf{w}}, \overline{b}, \overline{\boldsymbol{\xi}})$ feasible to (5.2) and upper bound $\overline{UB}$;
 if $\overline{UB} \leq \widehat{UB}$ **then**
 $(\mathbf{w}^*, b^*, \boldsymbol{\xi}^*) \leftarrow (\overline{\mathbf{w}}, \overline{b}, \overline{\boldsymbol{\xi}})$
 end
end

Algorithm 6: Heuristic procedure

5.4.4. The exact algorithm

In this subsection, we propose an exact procedure for (5.2) by solving a sequence of MISOCPs. With a slight abuse of notation, let us define, given a subset of features $K \subseteq [1:n]$, the set

$$\mathcal{B}_M(K) = \{(\mathbf{w}, \mathbf{u}) \in \mathbb{R}^{2n} : -M(1 - u_j) \leq w_j \leq M(1 - u_j), j \in K, \mathbf{e}^\top \mathbf{u} = n - B\}.$$

The subproblem considered in the exact procedure is

$$\begin{aligned}
(\text{SR-DSCoP}(K)) \quad & \min_{\mathbf{w}, \text{diag } \mathbf{W}, b, \boldsymbol{\xi}, \mathbf{u}} \quad \frac{1}{2} \text{trace } (\mathbf{W}) + C \sum_{i=1}^m \xi_i \\
\text{s.t.} \quad & (\mathbf{w}, b, \boldsymbol{\xi}) \in \Xi \\
& (\mathbf{w}, \mathbf{u}) \in \mathcal{B}_M(K) \\
& \begin{bmatrix} 1 & w_j & u_j \\ w_j & W_{jj} & 0 \\ u_j & 0 & u_j \end{bmatrix} \in \mathcal{S}_+^3, \quad j \in [1:n] \setminus K, \\
& u_j \in \{0, 1\}, \quad j \in K, \\
& \begin{bmatrix} 1 & w_j \\ w_j & W_{jj} \end{bmatrix} \in \mathcal{S}_+^2, \quad j \in K,
\end{aligned} \tag{5.20}$$

where SR is the abbreviation for semi-relaxation. In the problem (5.20), we essentially relax the complementarity constraints associated with features in the set $[1:n] \setminus K$ via decomposed SDP relaxation, and express the other by the big-M mixed-integer formulation. We note that the problem (5.20) is a mixed-integer positive semi-definite optimization problem (MISDP) and there is no off-the-shelf solver dealing with it. In the following part, we show a way to express the problem in the form of MISOCP, which can be handled by the off-the-shelf solvers like `copt`.

Before specifying the MISOCP reformulation, we present a useful result: some special slices $\mathcal{F}$ of the positive-semidefinite cone are second-order cone representable, due to Lemma 43 below.

Lemma 43.

$$\mathcal{F} := \left\{ (w, u, W) \in \mathbb{R}^3 : \begin{bmatrix} 1 & w & u \\ w & W & 0 \\ u & 0 & u \end{bmatrix} \in \mathcal{S}_+^3 \right\} = \left\{ (w, u, W) \in \mathbb{R}^3 : \begin{array}{l} \sqrt{(W-s)^2 + 4w^2} \leq W + s \\ u - t \leq 0, \quad u \geq 0 \\ s + t = 1, \quad \text{some } (s, t) \in \mathbb{R}^2. \end{array} \right\}$$

Proof. Noticing that the matrix in the left set has a chordal sparsity pattern, due to [AHMR88], we have the decomposition:

$$\begin{bmatrix} 1 & w & u \\ w & W & 0 \\ u & 0 & u \end{bmatrix} \in \mathcal{S}_+^3 \iff \exists (s, t) \in \mathbb{R}^2 \quad \text{s.t.} \quad \begin{bmatrix} s & w \\ w & W \end{bmatrix} \in \mathcal{S}_+^2, \begin{bmatrix} t & u \\ u & u \end{bmatrix} \in \mathcal{S}_+^2, \\
s + t = 1.$$

It follows that

$$\begin{bmatrix} s & w \\ w & W \end{bmatrix} \in \mathcal{S}_+^2 \iff s + W \geq 0 \quad \text{and} \quad sW \geq w^2 \iff \sqrt{(W-s)^2 + 4w^2} \leq W + s.$$

Finally, noticing the fact that

$$\begin{bmatrix} t & u \\ u & u \end{bmatrix} \in \mathcal{S}_+^2 \iff u - t \leq 0, \quad u \geq 0,$$

we complete the proof. □

Applying Lemma 43, the problem (5.20) can be rewritten of the form

$$\begin{aligned}
(\text{SR-DLMP}(K)) \quad & \min_{\mathbf{w}, \text{diag } \mathbf{W}, \mathbf{b}, \boldsymbol{\xi}, \mathbf{u}, \mathbf{s}, \mathbf{t}} \quad \frac{1}{2} \text{trace}(\mathbf{W}) + C \sum_{i=1}^m \xi_i \\
\text{s.t.} \quad & (\mathbf{w}, \mathbf{b}, \boldsymbol{\xi}) \in \Xi \\
& (\mathbf{w}, \mathbf{u}) \in \mathcal{B}_M(K) \\
& \mathbf{d}_j := \begin{bmatrix} W_{jj} - s_j \\ 2w_j \end{bmatrix} \in \mathbb{R}^2, \quad j \in [1:n] \setminus K, \\
& \|\mathbf{d}_j\|_2 \leq s_j + W_{jj}, \quad j \in [1:n] \setminus K, \\
& u_j - t_j \leq 0, \quad j \in [1:n] \setminus K, \\
& s_j + t_j = 1, \quad j \in [1:n] \setminus K, \\
& u_j \in \{0, 1\}, \quad j \in K, \\
& w_j^2 \leq W_{jj}, \quad j \in K,
\end{aligned} \tag{5.21}$$

where $2n$ (actually, $2(n - |K|)$) more variables $\mathbf{s} \in \mathbb{R}^n$ and $\mathbf{t} \in \mathbb{R}^n$ are introduced.

Now we can introduce the exact algorithm reported as Algorithm 7. In this exact procedure, we solve a sequence of semi-relaxed problems $\text{SR-DLMP}(K)$ associated with a subset $K \subseteq [1:n]$ of features. In such problems, only variables u_j with $j \in K$ will be considered as binary and the remaining ones are relaxed. To obtain initial bounds on the objective value, the exact procedure exploits the upper bound strategies detailed above. The main step of the exact procedure consists of solving the sequence of semi-relaxed problems (5.21) to improve the lower bound of the objective value.

To start with, we must select a subset of features $K \subseteq [1:n]$ whose associated $\mathbf{u}$ variables will be considered as binary variables in the first semi-relaxed problem. The upper-bound strategy run at the beginning provides a subset of features that allows us to obtain a good bound on the optimal objective value. Therefore, the exact procedure will consider the set provided by the heuristic as the initial K and it will obtain an initial LB solving $\text{SR-DLMP}(K)$. Then the set K is updated by adding and removing some of the features, in order to improve the lower bound of the objective value. To this end, two sets (denoted by K^+ and K^-) are built in each iteration. Set K^+ consists of some of the features in $[1:n] \setminus K$ whose associated $\mathbf{u}$ variables will be considered as binary in the next iteration, i.e. features of K^+ will be added to K . Similarly, K^- consists of features in K that will not be considered as binary in the next iteration.

In addition and if possible, we will update the UB in the main step running the upper bound strategies and using vector $\tilde{\mathbf{u}}$ as the initial ranking. There could be many different ways of defining set K^+ and updating K . In the strategy we adopted, we set K^+ as the set of s features j such that the relaxed solution $\tilde{\mathbf{u}}$ of $\text{SR-DLMP}(K)$ was "less binary". This is why we first create a ranking on the features based on terms $\tilde{u}_j - \tilde{u}_j^2$ in descending order, and then select the first s , which is a user-defined parameter.

Data: $(x^i, y^i) \in \mathbb{R}^n \times \{1, -1\}$, $B, C, s \in [1:n]$, $LB = -\infty$, $UB = \infty$, $K = \emptyset$.

Result: (w^*, b^*, ξ^*) solution of (5.2).

Run Algorithm 4 or Algorithm 5, obtain set K as the set of the selected features and update the UB value;

```

while  $\frac{UB-LB}{UB} \cdot 100 \geq 0.01$  do
  Solve problem SR-DLMP( $K$ );
  if a time limit of 1800 seconds is reached then
    | STOP
  end
  Obtain optimal value  $\widehat{LB}$  and optimal solution  $\tilde{u}$ ;
  Run Algorithm 4 or Algorithm 5 with the first step skipped vector  $\tilde{u}$  used for initial ranking;
  Obtain the solution  $(\widehat{w}, \widehat{b}, \widehat{\xi}, \widehat{u})$  and upper bound value  $\widehat{UB}$ ;
  Sort features using terms  $\tilde{u}_j - \tilde{u}_j^2$  in descending order;
  Build set  $K^+$  with the first  $s$  features;
  Build  $K^- := \{j \in K : j \text{ not selected in the solution of the last two iterations}\}$ ;
   $K \leftarrow K \cup K^+ \setminus K^-$ ;
  if  $\widehat{UB} < UB$  then
    |  $UB \leftarrow \widehat{UB}$ ;
    |  $(w^*, b^*, \xi^*, u^*) \leftarrow (\widehat{w}, \widehat{b}, \widehat{\xi}, \widehat{u})$ ;
  end
  if  $\widehat{LB} > LB$  then
    |  $LB \leftarrow \widehat{LB}$ ;
  end
end

```

Algorithm 7: Exact procedure

5.5. Numerical experiments

The various computational experiments were performed on an Apple M1 CPU 16 GB RAM computer. All models and algorithms were implemented in **Python**. Results were performed using **Gurobi 11.0** for the resolution of the MIQPs, **Mosek 10.1** for the resolution of the SDP relaxations and **Copt** for the resolution of the MISOCP models solved in the Exact algorithm.

The experiments were carried out on fifteen different datasets. Nine of them can be found in the UCI repository [DG17], (see Table 5.1), where m is the number of data points, n is the number of features and the last column shows the percentage of samples in each class. As can be observed, they contain a small number of features. The other six datasets used in the experiments have a larger number of features, as shown in Table 5.2. The Arrhythmia, Madelon, and MFeat datasets are available in the UCI repository as well. The remaining three datasets, Colorectal, DLBCL, and Lymphoma, are microarray datasets containing thousands of features but smaller sample sizes. Further descriptions of these last datasets can be found in [ABN⁺99, GWBV02, SRW⁺02].

Dataset	m	n	Class (%)
Breast Cancer D.	569	30	63/37
Breast Cancer W.	683	9	65/35
Breast Cancer P.	198	33	65/35
Cleveland	297	13	54/46
Diabetes	768	8	65/35
German	1000	24	70/30
Ionosphere	351	33	36/64
Parkinsons	195	22	25/75
Wholesale	440	7	68/32

Table 5.1.: First datasets

Dataset	m	n	Class (%)
Arrhythmia	452	274	54/46
Colorectal	62	2000	65/35
DLBCL	77	7129	25/75
Lymphoma	96	4026	35/65
Madelon	2000	500	50/50
Mfeat	2000	649	10/90

Table 5.2.: Second datasets

5.5.1. Computational analysis on the MIQP models

As starting results, we solved the two proposed MIQP models BigMP and CoP with **Gurobi**. Initially, we created different instances of the problems for each dataset in Table 5.1 ranging B values among all the possible values from 1 to n and ranging C values in the set $\{10^r : r \in [-3:3]\}$. For this first set of results, **Gurobi** was able to reach global optimum in every case but formulation CoP was the one which was easier to solve. Figure 5.1 shows average computation times for solving the two MIQPs, CoP and BigMP. On average, among datasets in Table 5.1, solving model CoP with **Gurobi** was 1.4 times faster.

Dataset	B	BigMP		CoP	
		ObjFun	MipGap	ObjFun	MipGap
Arrhythmia	10	2398.06	0.64	2151.56	0.82
	20	2047.67	0.67	1712.64	0.81
	30	1733.28	0.67	1441.63	0.78
Colorectal	10	2.99	0.64	2.12	0.98
	20	0.64	0.13	0.68	0.93
	30	0.41	0.04	0.43	0.89
DLBCL	10	0.63	1.00	0.56	1.00
	20	0.18	0.99	0.15	0.99
	30	0.16	0.99	0.09	0.98
Lymphoma	10	18.07	1.00	1.16	0.99
	20	0.71	0.99	0.29	0.98
	30	0.34	0.98	0.16	0.96
Madelon	10	16680.48	0.32	16800.07	0.33
	20	16419.87	0.32	18558.98	0.40
	30	16075.70	0.30	18425.44	0.39
Mfeat	10	0.92	0.59	0.68	0.93
	20	0.26	0.21	0.26	0.81
	30	0.18	0.15	0.18	0.71

Table 5.3.: Comparison of BigMP and CoP on the large datasets

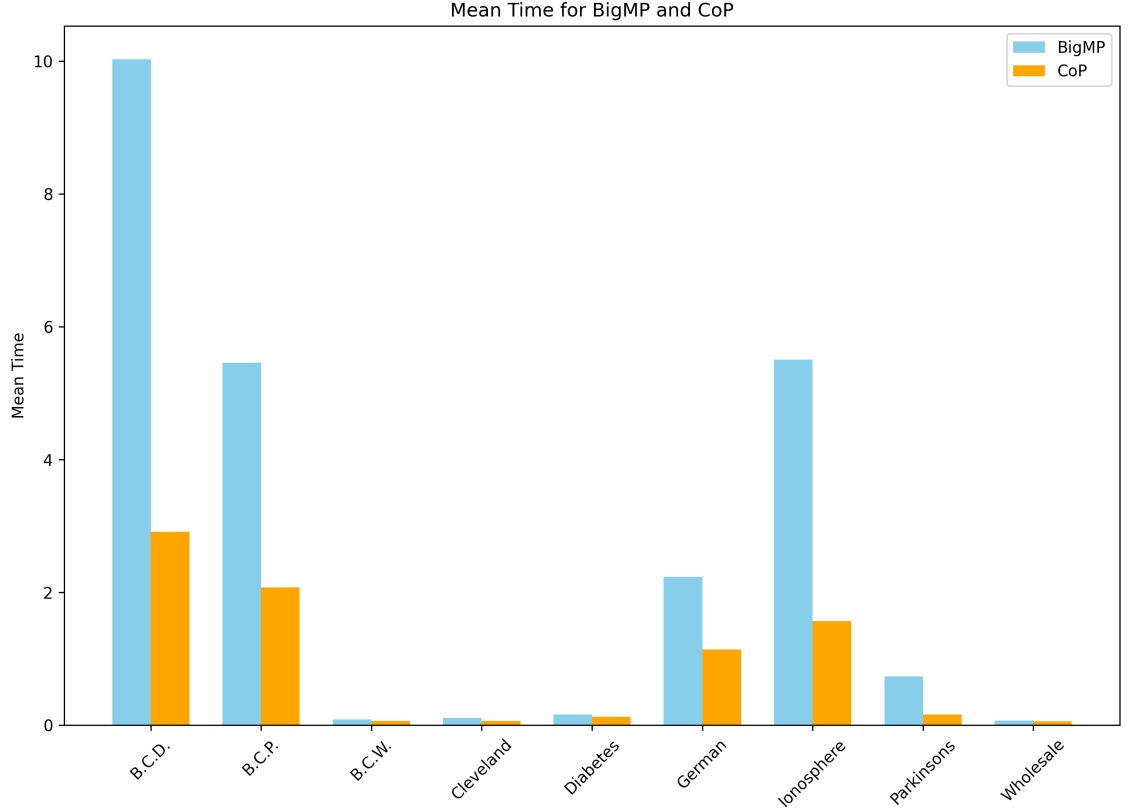


Figure 5.1.: Comparison of solving BigMP and CoP with **Gurobi** for the first set of datasets: mean computational times (in seconds)

The reason behind this first comparisons is that **Gurobi** is used in the two upper bound strategies we proposed. The two strategies, the Local Search and the Kernel Search, require the resolution of the FS-SVM model on instances created on a smaller subset of features, and this can be either done by using the BigMP or the CoP formulation. From these initial tests, we understand that CoP is the model that is faster to solve with **Gurobi** when dealing with datasets with a small number of features, thus, we will use the resolution of model CoP as the subroutine used by the heuristic algorithms to find and update the upper bounds values.

In Table 5.3 we show the computational performances of **Gurobi** when tested on the larger instances related to the datasets with the larger number of features. For these tests, we set the hyperparameter $C = 10$, and, similarly to [LMMRC19], we ranged B values in $\{10, 20, 30\}$. We can notice how the resulting optimization problems are much harder to solve in contrast to the problem related to the smaller datasets. A time limit of 1 hour was set, and the final MipGap values for the majority of these dataset are very high. This is certainly due to the fact that the lower bound values computed in the Branch & Bound tree of **Gurobi** are not very effective for these type of formulations. This is particularly the case for model CoP whose final MipGap values tend to be worse than the ones of model BigMP. In contrast, **Gurobi** manages to find much better incumbent solutions when solving model CoP for all datasets, apart from Colorectal and Madelon.

5.5.2. Results on the heuristics

In this section we analyse the performances of the heuristic algorithm we proposed and we compare with **Gurobi**. We have two versions of the heuristic algorithm depending on whether the Local Search or the Kernel Search upper bound strategies are used. We denote by H-LS the heuristic algorithm which implements the Local Search strategy, and by H-KS the heuristic implementing Kernel Search.

Table 5.4 shows the optimization performance of H-LS and H-KS. Both heuristics were given a time limit of 600 seconds. In particular for H-LS we set the parameter $k = 10$, which means that after getting the relaxed solution we solved model CoP on the $B + k$ most promising features. For H-KS we also set parameter $s = 10$, which means that the overall set of features was divided in subsets of ten features each. A smaller time limit of 60 seconds was given for the resolution of each CoP subproblem solved in the Kernel Search strategy.

As expected, H-KS manages to find better upper bounds compared to H-LS, at the cost of generally longer computational times. In some instances like in Lymphoma, Madelon and Mfeat, both heuristics found the same solutions. The time limit of 600 seconds was reached by H-KS for the instances related to datasets Madelon and Mfeat. For the rest of the datasets, the average computational times do not exceed 300 seconds for both heuristics. In general, compared to the objective function values of the solutions found by **Gurobi** in 1 hour (Table 5.3), the solutions found by the heuristics are much better.

Dataset	B	H-LS		H-KS	
		ObjFun	Time	ObjFun	Time
Arrhythmia	10	2491.04	5.37	2142.67	41.24
	20	1982.08	9.55	1652.83	168.83
	30	1729.55	7.59	1380.19	311.62
Colorectal	10	2.13	15.38	2.04	15.88
	20	0.67	20.12	0.65	57.51
	30	0.411	37.26	0.406	69.46
DLBCL	10	0.25	51.01	0.23	26.56
	20	0.104	46.97	0.101	67.74
	30	0.065	134.95	0.063	117.29
Lymphoma	10	0.69	28.73	0.69	98.64
	20	0.24	40.87	0.23	54.42
	30	0.14	98.25	0.13	109.33
Madelon	10	16660.85	25.07	16615.43	607.32
	20	16236.87	48.57	16236.87	612.62
	30	15917.52	176.92	15917.99	605.93
Mfeat	10	0.70	21.27	0.66	341.87
	20	0.25	76.68	0.25	621.16
	30	0.17	276.90	0.17	606.96

Table 5.4.: Results on large datasets of the Heuristic algorithm implementing either the Local search (H-LS) or the Kernel Search (H-KS) as the upper bound strategies

Another set of results we carried out was the one of using H-KS as a warm start heuristic in order to evaluate if **Gurobi** was able to find better solution. In Table 5.5 we compare the optimal value of the solution of **Gurobi** after 1 hour for model CoP, the value of the H-KS heuristic,

and the value of the solution found by **Gurobi** using as a warm start solution the one of H-KS (denoted by CoP*).

Regarding Table 5.5, we can see that **Gurobi** was unable to find a better solution than the heuristic after 1 hour for all datasets, apart from Arrhythmia, and this is probably due to the fact that the heuristic solutions are already optimal or nearly-optimal. In this table, column d_f indicate the percentage of improvement of the objective function found by **Gurobi** using H-KS starting solution, against the objective function of the solution found by **Gurobi** alone.

Dataset	B	ObjFun			
		CoP	H-KS	CoP*	d_f
Arrhythmia	10	2151.56	2142.67	2098.55	2%
	20	1712.64	1652.83	1652.83	3%
	30	1441.63	1380.19	1366.94	5%
Colorectal	10	2.12	2.04	2.04	4%
	20	0.68	0.65	0.65	5%
	30	0.43	0.41	0.41	6%
DLBCL	10	0.56	0.23	0.23	60%
	20	0.15	0.10	0.10	35%
	30	0.09	0.06	0.06	29%
Lymphoma	10	1.16	0.69	0.69	41%
	20	0.29	0.22	0.22	21%
	30	0.16	0.13	0.13	16%
Madelon	10	16800.07	16615.43	16615.43	1%
	20	18558.98	16236.86	16236.86	13%
	30	18425.44	15917.51	15917.51	14%
Mfeat	10	0.68	0.66	0.66	3%
	20	0.26	0.25	0.25	4%
	30	0.17	0.17	0.17	3%

Table 5.5.: Comparison of the objective functions of CoP returned by **Gurobi** with and without using the Heuristic Kernel Search algorithm as a warm start for large datasets

5.5.3. Results on the exact approach

As shown in Table 5.7, the proposed exact algorithm performs much better than **Gurobi** with a warm start heuristic, both on optimal value and computation time. In particular, the exact algorithm is able to globally solve four problems under a time limit of 3600 seconds. For the unsolved problems, the exact algorithm provides a better objective value and a smaller MipGap.

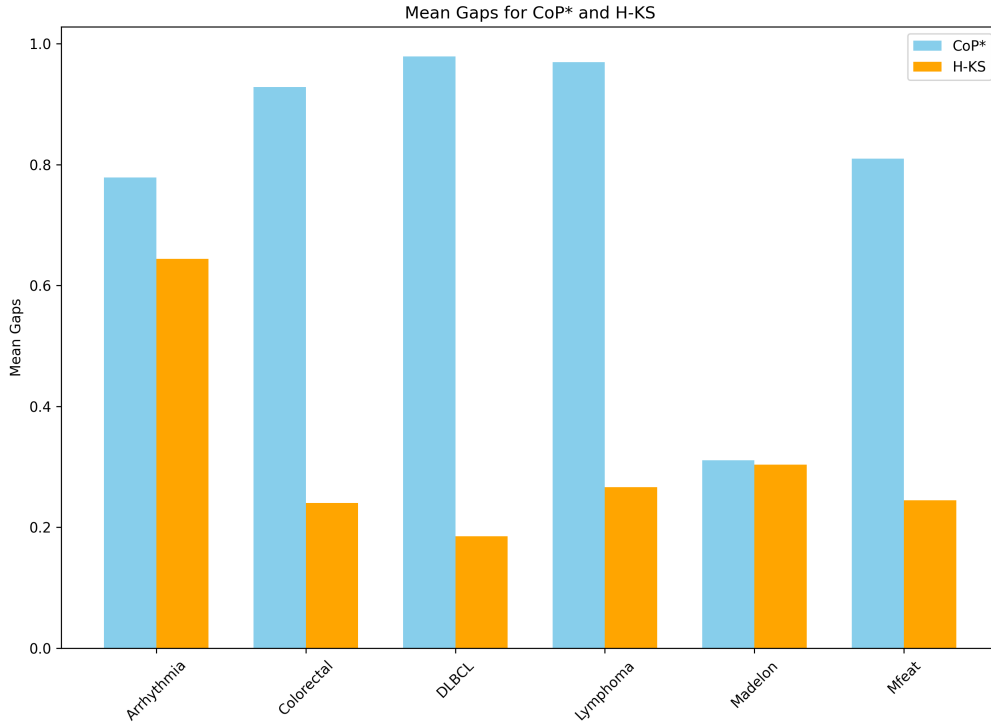


Figure 5.2.: Comparison of the mean MipGap values of the solution found by **Gurobi** solving CoP after 1 hour, and the Relaxed Gap of the Heuristic Kernel Search solution using the solution of DSCoP as a lower bound

5.6. Conclusion

In this chapter, we addressed the NP-hard feature selection problem in linear SVMs under a cardinality constraint, which ensures that the model remains interpretable by selecting a limited number of features. Our approach involves formulating the problem as an MIQP problem and introducing novel SDP relaxations to handle its complexity and enhance scalability.

We developed two MIQP formulations: one employing a big-M method and another one using a complementarity constraint. These formulations accommodate the cardinality constraint by integrating binary variables to select features. To solve these challenging formulations, we proposed several tractable SDP relaxations and a decomposed SDP approach, exploiting the sparsity pattern inherent to the problems. This decomposition significantly reduces computational complexity, making the conic relaxations scalable even for datasets with a large number of features.

Numerical experiments on various benchmark datasets demonstrate the effectiveness of our approach. The proposed methods not only outperform **Gurobi** but also provide competitive classification accuracy. By allowing a flexible adjustment of the number of features, our approach enhances interpretability without significantly compromising predictive performance.

Initial Bounds				
Dataset	B	DSCoP	BigMP	CoP
Arrhythmia	10	754.46	451.10	274.79
	20	622.44	395.92	274.79
	30	538.86	303.66	274.79
Colorectal	10	1.07	0.05	0.05
	20	0.54	0.05	0.05
	30	0.38	0.05	0.05
DLBCL	10	0.16	0.00	0.00
	20	0.08	0.00	0.00
	30	0.06	0.00	0.00
Lymphoma	10	0.35	0.00	0.01
	20	0.18	0.00	0.01
	30	0.12	0.00	0.01
Madelon	10	11243.14	0.00	0.00
	20	11209.44	0.00	0.00
	30	11197.73	0.00	0.00
Mfeat	10	0.36	0.05	0.05
	20	0.20	0.05	0.05
	30	0.15	0.05	0.05

Table 5.6.: Lower bounds comparison between the SDP one and the starting bound found by **Gurobi** for the BigMP and CoP models for different datasets

Dataset	B	CoP*			Exact KS		
		ObjFun	Time	MipGap	ObjFun	Time	MipGap
Arrhythmia	10	2098.55	3600	0.82	2003.98	3600	0.22
	20	1652.83	3600	0.80	1632.32	3600	0.27
	30	1366.94	3600	0.77	1362.75	3600	0.32
Colorectal	10	2.04	3600	0.98	2.04	407.2	0
	20	0.65	3600	0.93	0.65	431.5	0
	30	0.41	3600	0.88	0.41	407.6	0
DLBCL	10	0.23	3600	0.99	0.23	451.3	0
	20	0.10	3600	0.98	0.10	542.7	0
	30	0.06	3600	0.97	0.06	746.2	0
Lymphoma	10	0.69	3600	0.99	0.69	314.6	0
	20	0.22	3600	0.97	0.22	532.3	0
	30	0.13	3600	0.95	0.13	823.4	0
Madelon	10	16615.43	3600	0.33	16499.38	3600	0.17
	20	16236.86	3600	0.31	16142.54	3600	0.18
	30	15917.51	3600	0.30	15765.83	3600	0.21
Mfeat	10	0.66	3600	0.93	0.64	397.5	0
	20	0.25	3600	0.80	0.24	573.5	0
	30	0.17	3600	0.71	0.17	1045.5	0

Table 5.7.: Comparison of the **Gurobi** solution to the CoP model using the Heuristic Kernel Search solution as a warm start heuristic, and the Exact algorithm using the Kernel Search procedure as the upper bound strategy.

Bibliography

- [ABN⁺99] Uri Alon, Naama Barkai, Daniel A Notterman, Kurt Gish, Suzanne Ybarra, Daniel Mack, and Arnold J. Levine. Broad patterns of gene expression revealed by clustering analysis of tumor and normal colon tissues probed by oligonucleotide arrays. *Proceedings of the National Academy of Sciences*, 96(12):6745–6750, 1999.
- [AHD21] Andrey Afonin, Roland Hildebrand, and Peter J.C. Dickinson. The extreme rays of the 6×6 copositive cone. *Journal of Global Optimization*, 79(1):153–190, 2021.
- [AHMR88] Jim Agler, William Helton, Scott McCullough, and Leiba Rodman. Positive semi-definite matrices with a given sparsity pattern. *Linear Algebra and its Applications*, 107:101–149, 1988.
- [AMS10] Enrico Angelelli, Renata Mansini, and M. Grazia Speranza. Kernel search: A general heuristic for the multi-dimensional knapsack problem. *Computers & Operations Research*, 37(11):2017–2026, 2010.
- [AMS12] Enrico Angelelli, Renata Mansini, and M. Grazia Speranza. Kernel search: A new heuristic framework for portfolio selection. *Computational Optimization and Applications*, 51:345–361, 2012.
- [Ans21] Kurt M. Anstreicher. Testing copositivity via mixed-integer linear programming. *Linear Algebra and Its Applications*, 609:218–230, 2021.
- [AOHE19] Tiago Andrade, Fabricio Oliveira, Silvio Hamacher, and Andrew Eberhard. Enhancing the normalized multiparametric disaggregation technique for mixed-integer quadratic programming. *Journal of Global Optimization*, 73:701–722, 2019.
- [ApS24] MOSEK ApS. *The MOSEK optimization toolbox for MATLAB manual. Version 10.1.*, 2024.
- [BDDK⁺00] Immanuel M. Bomze, Mirjam Dür, Etienne De Klerk, Cornelis Roos, Arie J. Quist, and Tamás Terlaky. On copositive programming and standard quadratic optimization problems. *Journal of Global Optimization*, 18(4):301–320, 2000.
- [BDDM⁺09] Joshua Brodie, Ingrid Daubechies, Christine De Mol, Domenico Giannone, and Ignace Loris. Sparse and stable Markowitz portfolios. *Proceedings of the National Academy of Sciences*, 106(30):12267–12272, 2009.
- [BDPP24] Immanuel M. Bomze, Federico D’Onofrio, Laura Palagi, and Bo Peng. Feature selection in linear SVMs via hard cardinality constraint: a scalable SDP decomposition approach. *arXiv preprint arXiv:2404.10099*, 2024.
- [BEKS17] Jeff Bezanson, Alan Edelman, Stefan Karpinski, and Viral B. Shah. Julia: A fresh approach to numerical computing. *SIAM Review*, 59(1):65–98, 2017.

- [Bie96] Daniel Bienstock. Computational study of a family of mixed-integer quadratic programming problems. *Mathematical Programming*, 74:121–140, 1996.
- [BKL21] Immanuel M. Bomze, Michael Kahr, and Markus Leitner. Trust your data or not – StQP remains StQP: Community detection via robust Standard Quadratic Optimization. *Mathematics of Operations Research*, 46:301–316, 2021.
- [BKS16] Oleg P. Burdakov, Christian Kanzow, and Alexandra Schwartz. Mathematical programs with cardinality constraints: reformulation by complementarity-type conditions and a regularization method. *SIAM Journal on Optimization*, 26(1):397–425, 2016.
- [BLT08] Immanuel M. Bomze, Marco Locatelli, and Fabio Tardella. New and old bounds for standard quadratic optimization: dominance, equivalence and incomparability. *Mathematical Programming*, 115(1):31–64, 2008.
- [BM05] Stephen Braun and John E. Mitchell. A semidefinite programming heuristic for quadratic programming problems with complementarity constraints. *Computational Optimization and Applications*, 31(1):5–29, 2005.
- [BMP16] Lijie Bai, John E. Mitchell, and Jong-Shi Pang. On conic QPCCs, conic QCQPs and completely positive programs. *Mathematical Programming*, 159:109–136, 2016.
- [Bom97] Immanuel M. Bomze. Evolution towards the maximum clique. *Journal of Global Optimization*, 10(2):143–164, 1997.
- [Bom02] Immanuel M. Bomze. Regularity versus degeneracy in dynamics, games, and optimization: A unified approach to different aspects. *SIAM Review*, 44(3):394–414, 2002.
- [Bom12] Immanuel M. Bomze. Copositive optimization—recent developments and applications. *European Journal of Operational Research*, 216(3):509–520, 2012.
- [BP23] Immanuel M. Bomze and Bo Peng. Conic formulation of QPCCs applied to truly sparse QPs. *Computational Optimization and Applications*, 84:703–735, 2023.
- [BPQY24] Immanuel M. Bomze, Bo Peng, Yuzhou Qiu, and E. Alper Yıldırım. Tighter yet more tractable relaxations and nontrivial instance generation for sparse standard quadratic optimization. *arXiv preprint arXiv:2406.01239*, 2024.
- [BPQY25] Immanuel M. Bomze, Bo Peng, Yuzhou Qiu, and E. Alper Yıldırım. On tractable convex relaxations of standard quadratic optimization problems under sparsity constraints. *Journal of Optimization Theory and Applications*, to appear 2025.
- [BS09] Dimitris Bertsimas and Romy Shioda. Algorithm for cardinality-constrained quadratic optimization. *Computational Optimization and Applications*, 43(1):1–22, 2009.
- [BSM03] Abraham Berman and Naomi Shaked-Monderer. *Completely Positive Matrices*. World Scientific, 2003.
- [BSU12] Immanuel M. Bomze, Werner Schachinger, and Gabriele Uchida. Think co(mpletely)positive! Matrix properties, examples and a clustered bibliography on copositive optimization. *Journal of Global Optimization*, 52(3):423–445, 2012.

- [BSU15] Immanuel M. Bomze, Werner Schachinger, and Reinhard Ullrich. New lower bounds and asymptotics for the cp-rank. *SIAM Journal on Matrix Analysis and Applications*, 36(1):20–37, 2015.
- [Bur09] Samuel Burer. On the copositive representation of binary and continuous nonconvex quadratic programs. *Mathematical Programming*, 120(2):479–495, 2009.
- [Bur15] Samuel Burer. A gentle, geometric introduction to copositive optimization. *Mathematical Programming*, 151:89–116, 2015.
- [CPS92] Richard W. Cottle, Jong-Shi Pang, and Richard E. Stone. *The Linear Complementarity Problem*, volume 60. SIAM, 1992.
- [CPZ13] Xin Chen, Jiming Peng, and Shuzhong Zhang. Sparse solutions to random standard quadratic optimization problems. *Mathematical Programming*, 141:273–293, 2013.
- [CRT06] Emmanuel J. Candès, Justin Romberg, and Terence Tao. Robust uncertainty principles: Exact signal reconstruction from highly incomplete frequency information. *IEEE Transactions on Information Theory*, 52(2):489–509, 2006.
- [CT06] Emmanuel J. Candes and Terence Tao. Near-optimal signal recovery from random projections: Universal encoding strategies? *IEEE transactions on Information Theory*, 52(12):5406–5425, 2006.
- [CV95] Corinna Cortes and Vladimir Vapnik. Support-vector networks. *Machine Learning*, 20(3):273–297, 1995.
- [CW08] Emmanuel J. Candès and Michael B. Wakin. An introduction to compressive sampling. *IEEE Signal Processing Magazine*, 25(2):21–30, 2008.
- [DDM17] Alberto Del Pia, Santanu S. Dey, and Marco Molinaro. Mixed-integer quadratic programming is in NP. *Mathematical Programming*, 162:225–240, 2017.
- [DG17] Dheeru Dua and Casey Graff. UCI Machine Learning Repository. <http://archive.ics.uci.edu/ml>, 2017.
- [DGW23] Jin-Hong Du, Yifeng Guo, and Xueqin Wang. High-dimensional portfolio selection with cardinality constraints. *Journal of the American Statistical Association*, 118(542):779–791, 2023.
- [Dia62] Palahenedi H. Diananda. On non-negative forms in real variables some or all of which are non-negative. *Mathematical Proceedings of the Cambridge Philosophical Society*, 58(1):17–25, 1962.
- [Dic19] Peter J.C. Dickinson. A new certificate for copositivity. *Linear Algebra and its Applications*, 569:15–37, 2019.
- [DLLR⁺12] David Di Lorenzo, Giampaolo Liuzzi, Francesco Rinaldi, Fabio Schoen, and Marco Scia drone. A concave optimization-based approach for sparse portfolio selection. *Optimization Methods and Software*, 27(6):983–1000, 2012.
- [DR21] Mirjam Dür and Franz Rendl. Conic optimization: a survey with special focus on copositive optimization and binary quadratic problems. *EURO Journal on Computational Optimization*, 9:100021, 2021.

- [FG07] Antonio Frangioni and Claudio Gentile. SDP diagonalizations and perspective cuts for a class of nonseparable MIQP. *Operations Research Letters*, 35(2):181–185, 2007.
- [FM04] Glenn Fung and Olvi L. Mangasarian. A feature selection Newton method for support vector machine classification. *Computational Optimization and Applications*, 28:185–202, 2004.
- [FMP⁺18] Mingbin Feng, John J. Mitchell, Jong Shi Pang, Xin Shen, and Andreas Waechter. Complementarity formulations of l0-norm optimization. *Pacific Journal of Optimization*, 14(2):273–305, 2018.
- [FW56] Marguerite Frank and Philip Wolfe. An algorithm for quadratic programming. *Naval Research Logistics Quarterly*, 3(1-2):95–110, 1956.
- [GS12] Gianfranco Guastaroba and M. Grazia Speranza. Kernel search for the capacitated facility location problem. *Journal of Heuristics*, 18(6):877–917, 2012.
- [Gur23] Gurobi Optimization, LLC. Gurobi Optimizer Reference Manual, 2023.
- [GWBV02] Isabelle Guyon, Jason Weston, Stephen Barnhill, and Vladimir Vapnik. Gene selection for cancer classification using support vector machines. *Machine Learning*, 46:389–422, 2002.
- [GY21] Jacek Gondzio and E. Alper Yildirim. Global solutions of nonconvex standard quadratic programs via mixed integer linear programming reformulations. *Journal of Global Optimization*, 81(2):293–321, 2021.
- [GY22] Y. Gökem Gökmen and E. Alper Yildirim. On standard quadratic programs with exact and inexact doubly nonnegative relaxations. *Mathematical Programming*, 193(Ser. A):365–403, 2022.
- [Hil12] Roland Hildebrand. The extreme rays of the 5×5 copositive cone. *Linear Algebra and its Applications*, 437(7):1538–1547, 2012.
- [HN63] Marshall Hall and Morris Newman. Copositive and completely positive quadratic forms. *Mathematical Proceedings of the Cambridge Philosophical Society*, 59(2):329–339, 1963.
- [HP73] Alan J. Hoffman and Francisco Pereira. On copositive matrices with $-1, 0, 1$ entries. *Journal of Combinatorial Theory, Series A*, 14(3):302–309, 1973.
- [HRTZ04] Trevor Hastie, Saharon Rosset, Robert Tibshirani, and Ji Zhu. The entire regularization path for the support vector machine. *Journal of Machine Learning Research*, 5(Oct):1391–1415, 2004.
- [HV19] Nikolaus Hautsch and Stefan Voigt. Large-scale portfolio allocation under transaction costs and model uncertainty. *Journal of Econometrics*, 212(1):221–240, 2019.
- [Jol02] Ian T. Jolliffe. *Principal Component Analysis*. Springer, 2002.

- [LDD⁺23] Miles Lubin, Oscar Dowson, Joaquim Dias Garcia, Joey Huchette, Benoît Legat, and Juan Pablo Vielma. JuMP 1.0: Recent improvements to a modeling language for mathematical optimization. *Mathematical Programming Computation*, 2023.
- [LMMRC19] Martine Labbé, Luisa I. Martínez-Merino, and Antonio M. Rodríguez-Chía. Mixed integer linear programming for feature selection in support vector machine. *Discrete Applied Mathematics*, 261:276–304, 2019.
- [LPR96] Zhi-Quan Luo, Jong-Shi Pang, and Daniel Ralph. *Mathematical programs with equilibrium constraints*. Cambridge University Press, 1996.
- [Mar52] Harry Markowitz. Portfolio selection. *The Journal of Finance*, 7(1):77–91, 1952.
- [MD19] Luca Mencarelli and Claudia D’Ambrosio. Complex portfolio selection via convex mixed-integer quadratic programming: A survey. *International Transactions in Operational Research*, 26(2):389–414, 2019.
- [Mil84] Alan J. Miller. Selection of subsets of regression variables. *Journal of the Royal Statistical Society Series A*, 147(3):389–410, 1984.
- [MK87] Katta G. Murty and Santosh N. Kabadi. Some NP-complete problems in quadratic and nonlinear programming. *Mathematical Programming*, 39(2):117–129, 1987.
- [Mot52] Theodore S. Motzkin. Copositive quadratic forms. *National Bureau of Standards Report 1818*, pages 11–12, 1952.
- [MPWL14] Sebastián Maldonado, Juan Pérez, Richard Weber, and Martine Labbé. Feature selection for support vector machines via mixed integer linear programming. *Information Sciences*, 279:163–175, 2014.
- [MS65] Theodore S. Motzkin and Ernst G. Straus. Maxima for graphs and a new proof of a theorem of Turán. *Canadian Journal of Mathematics*, 17:533–540, 1965.
- [MS12] Walter Murray and Howard Shek. A local relaxation method for the cardinality constrained portfolio optimization problem. *Computational Optimization and Applications*, 53:681–709, 2012.
- [NDIT10] Minh Hoai Nguyen and Fernando De la Torre. Optimal feature selection for support vector machines. *Pattern Recognition*, 43(3):584–591, 2010.
- [Pen22] Bo Peng. Performance comparison of two recently proposed copositivity tests. *EURO Journal on Computational Optimization*, 10:100037, 2022.
- [Per84] Andre F. Perold. Large-scale portfolio optimization. *Management Science*, 30(10):1143–1160, 1984.
- [PVZ15] Javier Pena, Juan C. Vera, and Luis F. Zuluaga. Completely positive reformulations for polynomial optimization. *Mathematical Programming*, 151:405–431, 2015.
- [QdKRT98] Arie J. Quist, Etienne de Klerk, Cornelis Roos, and Tamas Terlaky. Copositive relaxation for general quadratic programming. *Optimization Methods and Software*, 9(1-3):185–208, 1998.

- [RTGMS10] Rubén Ruiz-Torrubiano, Sergio García-Moratilla, and Alberto Suárez. Optimization problems with cardinality constraints. In *Computational Intelligence in Optimization: Applications and Implementations*, pages 105–130. Springer, 2010.
- [SA99] Hanif D. Sherali and Warren P. Adams. *A Reformulation-Linearization Technique for Solving Discrete and Continuous Nonconvex Problems*. Springer US, Boston, MA, 1999.
- [Sha48] Claude Elwood Shannon. A mathematical theory of communication. *The Bell System Technical Journal*, 27(3):379–423, 1948.
- [SM16] Naomi Shaked-Monderer. SPN graphs: when copositive = SPN. *Linear Algebra and its Applications*, 509:82–113, 2016.
- [SMB21] Naomi Shaked-Monderer and Abraham Berman. *Copositive and completely positive matrices*. World Scientific Publishing Co. Pte. Ltd., Hackensack, NJ, 2021.
- [SRW⁺02] Margaret A. Shipp, Pablo Ross, Ken N. and Tamayo, Andrew P. Weng, Jeffery L. Kutok, Ricardo C.T. Aguiar, Michelle Gaasenbeek, Michael Angelo, Michael Reich, Geraldine S. Pinkus, et al. Diffuse large B-cell lymphoma outcome prediction by gene-expression profiling and supervised machine learning. *Nature Medicine*, 8(1):68–74, 2002.
- [SZL13] Xiaoling Sun, Xiaojin Zheng, and Duan Li. Recent advances in mathematical programming with semi-continuous variables and cardinality constraint. *Journal of the Operations Research Society of China*, 1(1):55–77, 2013.
- [The24] The MathWorks Inc. MATLAB Optimization Toolbox version 24.1 (R2024a), 2024.
- [WZ21] Angelika Wiegele and Shudian Zhao. Tight SDP relaxations for cardinality-constrained problems. In Norbert Trautmann and Mario Gnägi, editors, *Operations Research Proceedings 2021*, Lecture Notes in Operations Research, pages 167–172. Springer, 2021.
- [WZZ06] Li Wang, Ji Zhu, and Hui Zou. The doubly regularized support vector machine. *Statistica Sinica*, 16(2):589–615, 2006.
- [XTYP24] Wei Xu, Jie Tang, Ka Fai Cedric Yiu, and Jian Wen Peng. An efficient global optimal method for cardinality constrained portfolio optimization. *INFORMS Journal on Computing*, 36(2):690–704, 2024.
- [YCHL10] Guo-Xun Yuan, Kai-Wei Chang, Cho-Jui Hsieh, and Chih-Jen Lin. A comparison of optimization methods and software for large-scale l1-regularized linear classification. *Journal of Machine Learning Research*, 11:3183–3234, 2010.
- [ZH05] Hui Zou and Trevor Hastie. Regularization and variable selection via the elastic net. *Journal of the Royal Statistical Society Series B*, 67(2):301–320, 2005.
- [ZSL14] Xiaojin Zheng, Xiaoling Sun, and Duan Li. Improving the performance of miqp solvers for quadratic programs with cardinality and minimum threshold constraints: A semidefinite program approach. *INFORMS Journal on Computing*, 26(4):690–703, 2014.

- [ZSLS14] Xiaojin Zheng, Xiaoling Sun, Duan Li, and Jie Sun. Successive convex approximations to cardinality-constrained convex programs: a piecewise-linear DC approach. *Computational Optimization and Applications*, 59(1):379–397, 2014.

A. Appendix

A.1. Notation

Throughout the thesis, by serif font lower case, e.g. $\mathbf{x}$, and serif font upper case letter, e.g. $\mathbf{X}$, we denote a vector and a matrix respectively. For instance, $\mathbf{o}$ and $\mathbf{O}$ stand for zero vectors and matrices of suitable size, respectively. For $\mathbf{x} \in \mathbb{R}^n$, the n -dimensional Euclidean space, and any nonempty $I \subset [1:n]$, we denote by $\mathbf{x}_I = [x_i]_{i \in I}$ the restricted subvector. We also define sets by:

$$\begin{aligned}\Delta^n &= \{\mathbf{x} \in \mathbb{R}_+^n : \|\mathbf{x}\|_1 = 1\}, \\ \mathcal{S}^n &= \{\mathbf{X} \in \mathbb{R}^{n \times n} : \mathbf{X} = \mathbf{X}^\top\}, \\ \mathcal{N}^n &= \mathcal{S}^n \cap \mathbb{R}_+^{n \times n}, \\ \mathcal{S}_+^n &= \{\mathbf{X} \in \mathcal{S}^n : \mathbf{v}^\top \mathbf{X} \mathbf{v} \geq 0 \text{ for all } \mathbf{v} \in \mathbb{R}^n\}, \\ \mathcal{COP}^n &= \{\mathbf{X} \in \mathcal{S}^n : \mathbf{v}^\top \mathbf{X} \mathbf{v} \geq 0 \text{ for all } \mathbf{v} \in \mathbb{R}_+^n\}, \\ \mathcal{CP}^n &= \left\{ \sum_{i \in I} \mathbf{x}_i \mathbf{x}_i^\top \in \mathcal{S}^n : \mathbf{x}_i \in \mathbb{R}_+^n, \text{ for all } i \in I \right\}, \\ \mathcal{DNN}^n &= \mathcal{S}_+^n \cap \mathcal{N}^n, \\ \mathcal{SPN}^n &= \mathcal{S}_+^n + \mathcal{N}^n.\end{aligned}$$