



universität
wien

MASTER'S THESIS

Title of the Master's Thesis

„Semidefinite Programming:
Theory and Applications in Finance “

submitted by

Markus Gerald Alexander Gabl Bakk.phil. BSc.

in partial fulfilment of the requirements for the degree of
Master of Science (MSc)

Vienna 2016

degree programme code as it appears on
the student record sheet:

A 066 915

degree programme as it appears on
the student record sheet:

Betriebswirtschaft

Supervisor:

Univ.-Prof. Mag. Dr. Immanuel Bomze

Contents

List of Figures	3
Abstract	4
Zusammenfassung	5
The Author (CV)	6
Acknowledgments	6
1 Basic Concepts	7
1.1 The Convex Cone of Positive Semidefinite Matrices	7
1.2 Linear Matrix Inequalities and Constraints using a Frobenius Product	8
2 Semidefinite Programs (SDPs)	10
2.1 Statment of the Problem and its Dual	10
2.2 Duality, Optimality, Feasibility and Orthogonality	13
2.3 The Central Path	15
2.3.1 Logarithmic Barrier Function for a Linear Matrix Inequality .	15
2.3.2 Objective and Duality Gap Parametrization	16
2.4 Algorithms for solving SDPs	17
2.4.1 Primal-Dual Potential-Reduction Methods	18
2.4.2 Primal Logarithmic Barrier Methods	21
2.4.3 Primal-Dual Affine-Scaling Methods	21
2.4.4 Primal-Dual Path-Following Methods	22
2.5 Some useful remarks on formulating SDPs	23
2.5.1 Linear Programs	23
2.5.2 Schur Complements and other Nonlinear Problems	24
2.5.3 Quadratic Programms and Quadratic constraints	25
2.5.4 Second Order Cone Programms (SOCP)	26
2.5.5 Relaxations of some NP-Hard Problems	28
2.5.6 S-Lemma and S-Procedure	29
2.5.7 Sufficient Conditions for Positive Polynomials	30
3 Asset Pricing Approaches using Conic and Semidefinite Optimization	32
3.1 Arbitrage Pricing Theory: Recovering Risk Neutral Probability Mea- sures ("RNPM") from Option Prices	32

3.1.1	Arbitrage Pricing : Basic Theory	32
3.1.2	Model for Recovering the RNPM	34
3.2	Option Pricing	37
3.2.1	Merton-Black-Scholes Model (MBSM)	37
3.2.2	Tight Bounds for Option Prices via an SDP-Approach	40
4	Sparse Principal Component Analysis: SDP-Relaxations and Application	44
4.1	Standard Principal Component Analysis	44
4.2	Sparse Principal Component Analysis: Two SDP-Relaxation Approaches	46
4.2.1	SDP-Relaxation based on the cardinality constrained Eigen- value Problem	47
4.2.2	SDP-Relaxation based on a convex Restriction	49
4.3	Applications of SPCA in Finance	51
5	Portfolio selection under data certainty	52
5.1	Classical Portfolio Theory	52
5.2	Tracking error	53
5.3	SDP-Relaxation for Portfolio Selection with Higher Order Moments .	54
5.4	The Knapsack Problem: SDP-Relaxations and Applications in Finance	57
6	Portfolio Selection with uncertain data	61
6.1	Modeling Uncertainty using Ellipsoidal Sets	61
6.1.1	Ellipsoidal uncertainty for linear constraints	62
6.1.2	Ellipsoidal uncertainty for quadratic constraints	63
6.1.3	A Robust Multi-Period Portfolio Selection Model	65
6.2	Distributional Robust Modelling	67
6.2.1	A General Framework	67
6.2.2	Applying the theory to the Knapsack Problem	74
	Conclusion	87
	References	87

List of Figures

1	Figure 1	7
2	Figure 2	9
3	Figure 3	10
4	Figure 4	15

Abstract

This thesis deals with semidefinite programming, the basic theory and the respective algorithms, as well as with its many applications in finance, ranging from asset pricing to portfolio optimization. Semidefinite programming is a convex optimization technique, that seeks to optimize the a linear objective function over a convex set that is described via linear matrix inequalities or via intersections of affine sets with the cone of positive definite matrices. The power of this problem formulation arises not only from the fact, that it generalizes linear programming as well as second order cone programming and quadratic programming, but also from the fact, that many non-convex and combinatorial problems can be relaxed as semidefinite programs and solved approximately, where the so obtained bounds can often be shown to be comparatively tight. Though the concept is known at least since the 1980's (with the earliest reference even dating back to 1961), interest only started to excel in the 1990's, when it became clear, that interior point methods (previously used to solve linear and quadratic programs) could also be used to solve semidefinite programs efficiently. Since then a plethora of algorithms have been developed with applications ranging from engineering, over statistics and logistics to, of course, finance.

The Thesis will start out by briefly reviewing some important concepts from linear algebra and convex analysis. After that, there will be a more extensive discussion of the model itself, as well as of its properties, as far as they will be useful for the understanding of the applications and algorithms. The later will be discussed in the subsequent chapters, where the focus will be on Primal-Dual Potential Reduction Methods. The theoretical part will close with a survey of basic ideas, that may guide the formulation of generic problems as semidefinite programs.

Following that, there will be an extensive discussion of models, that are addressing problems in finance. The survey contains applications in asset pricing, sparse principal component analysis and portfolio selection under data certainty as well as with uncertain data. If necessary for the understanding of the respective model, the chapters will provide some background in financial theory or other fields that might be of interest, in order to contextualize the models discussion as well as to highlight their significance. In the final two chapters we will first introduce the reader to a relatively new framework for distributionally robust convex programming, recently established by Wiesmann, Kuhn and Sim. Preceding that, their theory will illustrated by applying it to the famous Knapsack problem. Although other authors have proposed formulations for the distributionally robust Knapsack problem, the approach presented here is solely based on the aforementioned framework.

Zusammenfassung

Die vorliegende Arbeit beschäftigt sich mit semidefiniter Programmierung, der ihr zugrundeliegenden Theorie, den entsprechenden Algorithmen, sowie die wichtigsten Anwendungen in der Finanzwirtschaft, von Preisfindung bis zur Portfolio Optimierung. Semidefinite Programmierung ist eine Technik der konvexen Optimierung. Dabei wird eine lineare Funktion über eine Konvexe Menge maximiert oder minimiert. Diese Menge wird durch lineare Matrixungleichungen, sowie über Schnittmengen von Affinen Mengen mit dem konvexen Kegel der positive Semidefinite Matrizen beschrieben. Die Stärke der semidefinite Programmierung ruht einerseits auf der Tatsache, dass es sich bei ihr um eine Verallgemeinerung der linearen als auch der quadratischen und Second-Order-Cone Programmierung handelt, andererseits erlaubt sie die Relaxierung vieler nicht konvexen und kombinatorischen Problemstellungen, wobei die so erhaltenen Schranken vergleichsweise eng sind. Obwohl das grundlegende Konzept bereits seit den 1980-ern bekannt ist, gelangte diese Technik erst in den 1990-ern zu größerer Popularität, nachdem man herausfand, dass Innere-Punkt Methoden, die ursprünglich zur Lösung linearer und quadratischer Programme entwickelt wurden, auch zur effizienten Lösung semidefiniter Programme geeignet sind. Seitdem wurde eine Vielzahl an Algorithmen entwickelt, wobei die Anwendungsgebiete von Maschinenbau über Statistik und Logistik bis hin zur Finanzwirtschaft reichen. Am Beginn dieser Arbeit, werden die wichtigsten Konzepte aus der linearen Algebra sowie der Konvexen Analysis kurz aufbereiten. Dem folgen eine ausführlichere Diskussion der Grundlagen der semidefiniten Programmierung, sowie wichtige Eigenschaften des Models sofern sie für das Verständnis der Anwendungen und Algorithmen maßgeblich sind. Letztere werden in dem darauf folgenden Kapitel dargestellt, wobei der Fokus auf Primal-Dual-Potential-Reduction-Methods liegen wird. Der theoretische Teil endet mit einer Darstellung von wichtigen Modellierungstechniken. Anschließend werden wichtige Anwendungen in der Finanzwirtschaft präsentiert, wobei Probleme wie die Preisfindung von Finanztiteln, Hauptkomponentenanalyse und Portfoliooptimierung sowie robuste Portfoliooptimierungen behandelt werden. Soweit es dem besseren Verständnis der Modelle dient, werden Grundlagen aus dem jeweiligen Bereich kurz dargestellt. In den letzten beiden Kapiteln wird zuerst ein relativ neuer theoretischer Rahmen für die Verteilungsrobuste ("distributionally robust") Optimierung wiedergegeben der von Wiesmann, Kuhn und Sim ausgearbeitet wurde. Danach wird zum Zwecke der Illustration diese Theorie, auf das Rucksack Problem angewandt. Zwar haben andere Autoren bereits verteilungsrobuste Formulierungen dieses Problems ausgearbeitet, der hier präsentierte Ansatz baut jedoch alleine auf der besagten Rahmentheorie auf.

The Author

Markus Gabl, Curriculum Vitae

1986. Geboren in Linz, Austria

1993-1997. Volksschule: Jahnschule Linz

1997-2005. Bundesrealgymnasium Aufhof, Linz

2005-2006 Zivildients : Rotes Kreuz, Linz

An der Universität Wien:

2006-2012. Publizistik und Kommunikationswissenschaften (Bakk.phil.),

2009-2013. Betriebswirtschaftslehre(BSc.),

2014-2016. Betriebswirtschaftslehre(MSc. noch nicht abgeschlossen),

Acknowledgments

I would like to thank ao. Univ.-Prof. Mag. Dr. Gabriele Uchida for introducing me to this interesting topic of semidefinite programming and Univ.-Prof. Mag. Dr. Immanuel Bomze for his great support and effort in supervising this thesis.

1 Basic Concepts

For the sake of self-containment, various basic concepts and some of their implications, as far as they are useful in later chapters, shall be revisited in the following. These include basic concepts of linear Algebra as well as of convex analysis of which an extensive outlay can be found in [48]. A layout more focused on SDPs is given in [15].

1.1 The Convex Cone of Positive Semidefinite Matrices

A real, square, symmetric matrix (which is the only type of matrix of interest here), say X , is an $n \times n$ matrix, with real entries, where $X^T = X$. It is said to be positive semidefinite (p.s.d.) if its quadratic form $v^T X v$ is greater or equal to zero for any vector $v \in \mathbb{R}^n$. Thus we can define the set of real p.s.d. matrices as

$$S_+^n = \{X \in \mathbb{R}^{n \times n} \mid X^T = X, v^T X v \geq 0 \ \forall v \in \mathbb{R}^n\}. \quad (1)$$

Positive definite (*p.d.*) matrices fulfill the condition on the quadratic form with strict inequality, and the respective set shall be denoted by S_{++}^n . The (partial) order induced by the notion of definiteness is called the *Loewner partial order*. It allows us to write $A \geq 0$ whenever $A \in S_+^n$, and $A > 0$ whenever $A \in S_{++}^n$. Further $A \geq B$ means that $A - B \geq 0$. Not all properties of scalar inequalities translate into the Loewner partial order, but we will skip a discussion here.

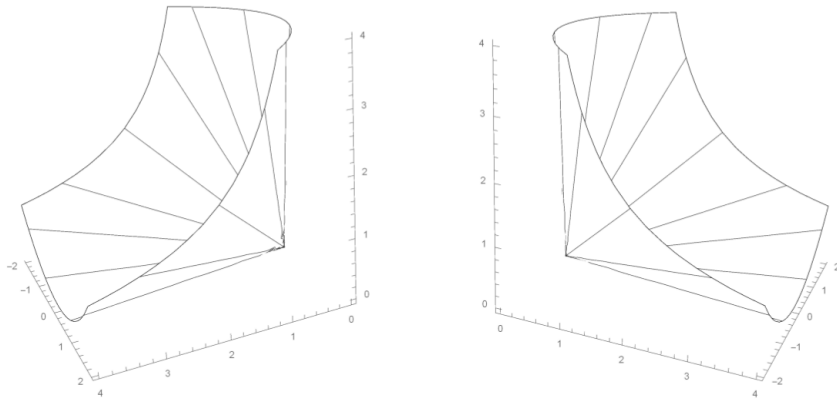


Figure 1: The graphs show the locus of all Matrices of the form $\begin{pmatrix} x & y \\ y & z \end{pmatrix}$ that are p.s.d.

One can easily show that the convex combination $Z = X\alpha + Y(1-\alpha)$ with $X, Y \in S_+^n$ and $\alpha \in [0, 1]$ is itself p.s.d. (i.e. $Z \in S_+^n$, thus S_+^n is closed under convex combinations).

It is obvious in that calculating the quadratic form of Y amounts to adding up two non-negative real numbers. This establishes S_+^n and (by a similar argument) S_{++}^n as convex sets. S_+^n as well as S_{++}^n are also closed under multiplication with a positive scalar (i.e. $Y = X\alpha$ with $X \in S_+^n$ and $\alpha \in \mathfrak{R}^+$ implies $Y \in S_+^n$, analogously with S_{++}^n ; to see this, again observe that the quadratic form of Y amounts to the product of two non negative/positive scalars, respectively). Thus both sets are also cones by the standard definition, so that one can indeed speak of the convex cone of *p.s.d/p.d.* matrices. The case where $X \in S_+^2$ allows for a nice graphical representation of that fact (see Table 1).

The eigenvalues of *p.d.* matrices are strictly positive, whereas in the case of *p.s.d.* matrices that are not also *p.d.*, at least one is equal to zero. Since the determinant of a matrix can be represented by the product of its eigenvalues (i.e. $\det(X) = \prod_{i=1}^n \lambda_i$), *p.s.d.* matrices are singular unless they are also *p.d.*, in which case they have a positive determinant. By definition, S_+^n represents the closure of S_{++}^n . Thus all singular p.s.d. matrices lie on the boundary of the p.s.d cone, which one has to cross when trying to "leave" the cone. More formally $\exists \alpha \in [0, 1] : \det(A\alpha + B(1 - \alpha)) = 0$ if $A \in S_+^n$ and $B \notin S_+^n$. So, if the p.s.d. cone is "left" along a line between $A \geq 0$ and $B \not\geq 0$, at least one of the eigenvalues of that matrix converges to 0 as the boundary of S_+^n is approached.

1.2 Linear Matrix Inequalities and Constraints using a Frobenius Product

Let's define a matrix $F(x)$ as follows:

$$F(x) = F_0 + \sum_{i=1}^m x_i F_i \quad (2)$$

where $x \in \mathfrak{R}^m$, $F_i \in \mathfrak{R}^{n \times n}$ with $F_i^T = F_i$. Then the constraint

$$F(x) \geq 0 \quad (3)$$

is called a *Linear Matrix Inequality*, which by restricting $F(x)$ to S_+^n also restricts x to a convex subset of $\mathfrak{R}^n$. To see this, observe that for any $x, y \in \mathfrak{R}^n$ for which (3) is fulfilled one has

$$F(x\alpha + y(1 - \alpha)) = \alpha F(x) + (1 - \alpha)F(y) \geq 0. \quad (4)$$

The equality in (4) is true by linearity of $F(x)$; the inequality is due to the convexity of S_+^n . Figure X gives an illustration of an example where $m = 2$ and $n = 3$.

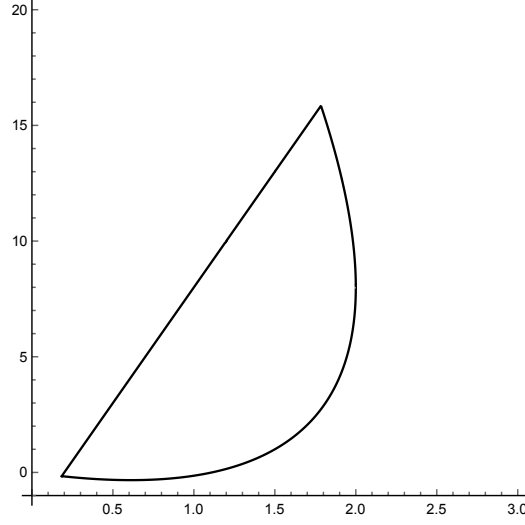


Figure 2: The graphic shows the set of alle x and y for which $\begin{pmatrix} 5x-0.5y-1 & 0 & 0 \\ 0 & 10 & 3x+0.5y \\ 0 & 3x+0.5y & x+y \end{pmatrix} \geq 0$

Another key way of using the p.s.d. cone to confine variables to convex sets is invoked by the so called *FrobeniusProduct*. It is defined as follows:

$$A \bullet B = \text{Tr}(A^T B) = \sum_{i=1}^n \sum_{j=1}^n a_{ij} b_{ij} \quad (5)$$

The matrix norm that is induced by the Frobenius Product is called the Frobenius Norm and is given for any $A = A^T \in \mathfrak{R}^{n \times n}$ by

$$\|A\| = A \bullet A = \sum_{i=1}^n \sum_{j=1}^n a_{ij}^2 \quad (6)$$

.

Clearly, the equation

$$A \bullet X = c \quad (7)$$

with symmetric $A, X \in \mathfrak{R}^{n \times n}$ and $c \in \mathfrak{R}$ restricts the entries of X to a hyperplane in $\mathfrak{R}^{n \frac{n+1}{2}}$ for given A and c . If one adds to that the restriction $X \geq 0$, the set of entries of X so described is the intersection between that hyperplane and the convex p.s.d. cone. Since both represent convex sets, their intersection itself is again convex. For a graphic example with $n = 2$, see Figure 3.

It is worth mentioning that any constraint expressed via a linear matrix inequality can also be stated in terms of intersections of hyperplanes with S_+^n . Consider a system of m equations of the (7) form, where $A_1 \dots A_m$ are linearly independent. Then the set

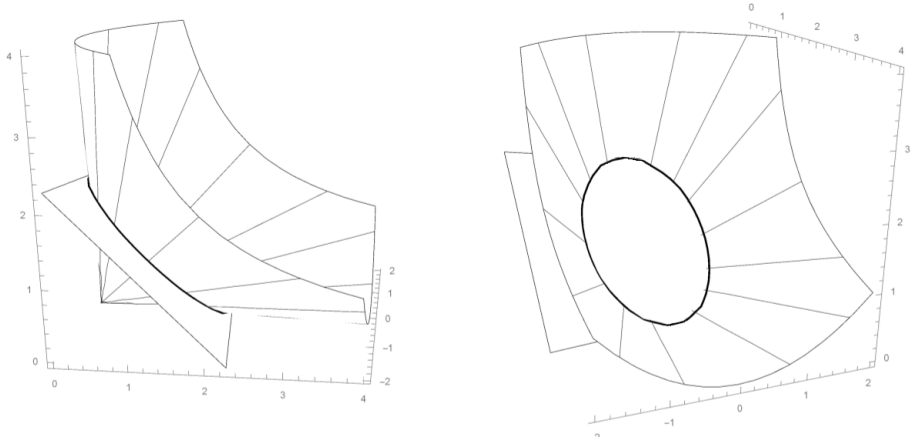


Figure 3: The intersection of the p.s.d. cone with a hyperplane

of solutions is an affine set expressed by

$$\{X \mid X^T = X \in \mathfrak{R}^{n \times n}, A_i \bullet X = c_i, i = 1, \dots, m\}.$$

By choosing appropriate F_i one can describe the same set by

$$\{F(y) = F_0 + \sum_{i=1}^p y_i F_i \mid y \in \mathfrak{R}^p\}$$

where $p = n \frac{n+1}{2}$.

2 Semidefinite Programs (SDPs)

In this chapter, a summary of the basic literature on SDPs is given. The main sources are [45], [13] and [24]. First, the statement of the problem is presented along with the derivation of its dual. Then some important concepts are explained in more detail, before discussing the algorithms. Finally, the chapter provides some remarks on problem formulation, which will also give an idea of the plethora of problems that can be cast as SDPs.

2.1 Statment of the Problem and its Dual

The problem considered in semidefinite programming is to minimize a linear function subject to the constraint given by a linear matrix inequality:

$$\begin{aligned} \min_x \quad & c^T x \\ \text{subject to} \quad & F(x) \geq 0 \end{aligned} \tag{8}$$

where $x, c \in \mathfrak{R}^m$, $F(x) = F_0 + \sum_{i=1}^m x_i F_i$ and $F_i^T = F_i \in \mathfrak{R}^{n \times n}$ for $i = 0, \dots, m$. The problem in its primal form (8) will henceforth be denoted as $\mathcal{P}$, the set of its feasible solutions as

$$P = \{ x \mid F(x) \geq 0, x \in \mathfrak{R}^m \}$$

and its optimal solutions as $p^* = \inf_{x \in P} (c^T x)$, with the respective set denoted as

$$P^* = \{ x \mid F(x) \geq 0, x \in \mathfrak{R}^m, c^T x = p^* \}.$$

In order for p^* to always have a value, it is defined to be equal ∞ whenever $\mathcal{P}$ is infeasible, i.e. if there is no solution that satisfies the constraint. Elements of P that fulfill the constraint with strict inequality, i.e. solutions that satisfy $F(x) > 0$, are called strictly feasible solutions.

Since the objective function is linear and the feasible set convex, semidefinite programming presents itself as an instance of convex programming. Further, the restriction can be reformulated as

$$\begin{aligned} F(x) - Y &= 0 \\ Y &\geq 0 \quad (\text{or } Y \in S_+^n) \end{aligned}$$

in order to emphasize the fact that SPDs are special cases of cone programming, where the respective cone is S_+^n .

The dual may be established by the standard procedure. One turns the constrained problem into an unconstrained problem with more variables that act as (positive) penalties to the objective function. These penalties are incurred whenever a constraint is violated (and forgiven, whenever the constraint is exceeded), in the case of the primal problem that is:

$$\min_{x, Z} L(x, Z) := c^T x - Z \bullet F(x) \tag{9}$$

where $L(x, Z)$ is called the *Lagrangian Function*. Note that $Z \bullet F(x)$ is non-positive if $F(x) \leq 0$, since $Z \geq 0$ (the proof is due to Fejer, see [13] pp. 233). Thus Z can be interpreted as a price that has to be paid in order to perturb $F(x) \geq 0$. Further, notice that for a given $Z \geq 0$ the minimum of the Lagrangian gives a lower bound on the

optimal value of (8). This naturally raises the question of what the "best" (meaning the highest) lower bound might be. A simple reformulation shows that

$$\begin{aligned} L(x, Z) &= \sum_{i=1}^m c_i x_i - Z \bullet (F_0 + \sum_{i=1}^m x_i F_i) \\ &= -Z \bullet F_0 + \sum_{i=1}^m x_i (c_i - Z \bullet F_i). \end{aligned}$$

Thus

$$\inf_x L(x, Z) = \begin{cases} -Z \bullet F_0 & Z \bullet F_i = c_i \quad \forall i \\ -\infty & \text{else} \end{cases} \quad (10)$$

with the largest lower bound being given by

$$g(x, Z) = \sup_{Z \geq 0} \inf_x L(x, Z)$$

where $g(x, Z)$ is called the *Lagrangian Dual* of $\mathcal{P}$. It can be stated as an optimization problem of the form

$$\begin{aligned} \max_Z \quad & -Z \bullet F_0 \\ \text{subject to} \quad & Z \bullet F_i = c_i \quad \forall i = 1, \dots, m \\ & Z \geq 0 \end{aligned} \quad (11)$$

. This problem is further referred to as the *Dual* of $\mathcal{P}$, and will be denoted by $\mathcal{D}$. The set of its feasible solutions is given by

$$D = \{ Z \mid Z \bullet F_i = c_i \quad \forall i = 1, \dots, m, \quad Z \geq 0 \}$$

and its set of optimal solutions $d^* = \sup_Z D$ by

$$D^* = \{ Z \mid Z \bullet F_i = c_i \quad \forall i = 1, \dots, m, \quad Z \geq 0, \quad -Z \bullet F_0 = d^* \}.$$

Again, if D is empty, i.e. $\mathcal{D}$ is infeasible, then $d^* = -\infty$, which preserves the lower bound property in any case. Again, solutions for which $Z \succ 0$ holds are called strictly feasible solutions. The constraints can be reformulated as a linear matrix Inequality. Also, by reformulating the objective function, one can state $\mathcal{D}$ in the same form as

$\mathcal{P}$. So $\mathcal{D}$ is also an SDP.

2.2 Duality, Optimality, Feasibility and Orthogonality

As already mentioned, the solutions for $\mathcal{D}$ provide a lower bound for $\mathcal{P}$. Let's consider the difference in their objective function values for any $x \in P$ and $Z \in D$:

$$c^T x + Z \bullet F_0 = \sum_{i=1}^m (Z \bullet F_i) x_i + Z \bullet F_0 = Z \bullet F(x) \geq 0 \quad (12)$$

This quantity is called the *Duality Gap*, and we can easily see that it is positive, since by feasibility, $F(x)$ and Z are both p.s.d ($A \bullet B \geq 0$ if $A, B \geq 0$). Of course, this is also clear from the construction of the Lagrangian dual. The fact that the duality gap is positive is called *Weak Duality*.

If, however $p^* = d^*$, one speaks of *Perfect Duality*, which does not necessarily imply that P^* or D^* are nonempty. Indeed one can construct examples, where perfect duality holds, while either P^* or D^* is empty, but these are of little interest here. Rather, we consider the case in which both sets are nonempty, in which case one speaks of *Strong Duality*. The following theorem can be proved for Strong Duality:

Theorem 2.2.1.: Given $p^* > -\infty$ and $d^* < \infty$, further P and D contain strictly feasible solutions. Then it holds that $P^* \neq \emptyset$, $D^* \neq \emptyset$ and $p^* = d^*$ (i.e. $Z \bullet F(x) = 0$)

Proofs can be found in [48], [13], [53]. Given that Strong Duality holds, we find for $x \in P^*$ and $Z \in D^*$ that $F(x)Z = 0$ which is implied by $Z \bullet F(x) = 0$. Adding to that the problems constraints we can characterize any optimal solution by

$$\begin{aligned} F(x) &\geq 0 \\ Z &\geq 0 \\ Z \bullet F_i &= c_i \\ F(x)Z &= 0. \end{aligned} \quad (13)$$

It is intuitively clear that if $\mathcal{P}$ is unbounded, i.e. the constraints are such that the target function value can be minimized arbitrarily, $\mathcal{D}$ is by definition infeasible since it has the lower bound property. The converse is true if $\mathcal{D}$ is unbounded. In the following, analytic conditions that characterize such cases are disclosed in more detail.

We say $\mathcal{P}$ has a *Primal Improving Ray* if there exists a vector Δx such that $\sum_{i=1}^m F_i \Delta x \geq$

0 and $c^T \Delta x < 0$. In such a case, one could always add Δx to any feasible solution and maintain feasibility while improving the value of the target function. Analogously, we say that $\mathcal{D}$ has a *Dual Improving Ray* if there is a symmetric $\Delta Z \geq 0$, such that $\Delta Z \bullet F_i = 0 \ \forall i = 1, \dots, m$ and $-\Delta Z \bullet F_0 > 0$.

We can now show that the existence of a primal improving ray is incompatible with dual feasibility. Assume $\mathcal{D}$ is feasible and that there exists a primal improving ray. Then there is a Δx such that

$$0 > c^T \Delta x = \sum_{i=1}^m Z \bullet F_i \Delta x_i = Z \bullet \sum_{i=1}^m F_i \Delta x_i > 0 \quad (14)$$

which is a contradiction. Assuming primal feasibility and a Dual Improving Ray, an analogous argument can be made. The concept of infeasibility discussed above is called *Strong Infeasibility*. We can state the following Theorem:

Theorem 2.2.2.: If $\mathcal{D}$ has an improving ray, then $\mathcal{P}$ is strongly infeasible, and vice versa.

Finally an important concept in this regard is the fact that P and D are orthogonal in the following sense. Note that for any primal step $\delta x = x^0 - x$ we have

$$\delta F = \sum_{i=1}^m (x_i^0 - x_i) F_i \quad (15)$$

such that

$$\delta F \in \text{span}\{F_1, \dots, F_m\}.$$

So any dual step given by $\delta Z = Z_0 - Z$ has it that

$$F_i \bullet \delta Z = F_i \bullet Z_0 - F_i \bullet Z = c_i - c_i = 0 \quad \forall i \quad (16)$$

which further implies

$$\delta Z \in \text{span}\{F_1, \dots, F_m\}^\perp.$$

In other words, δF and δS lie in subspaces of S^n that are perpendicular to each other, i.e. $\delta F \bullet \delta Z = 0$

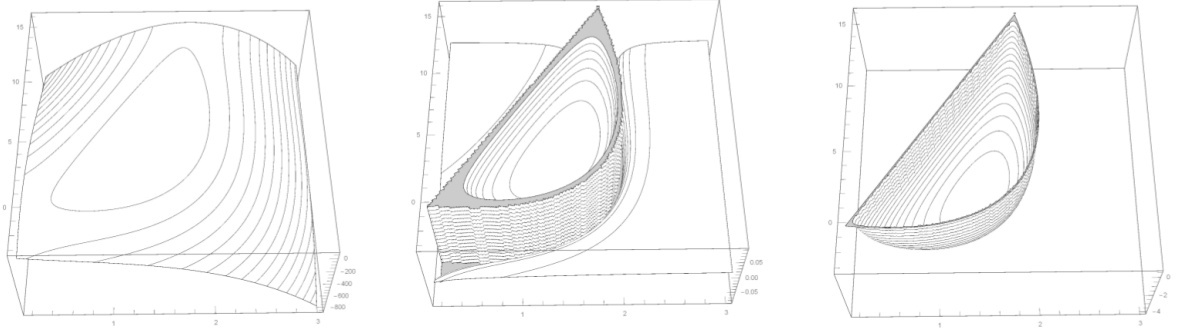


Figure 4: The graphics show the graphs of $\det F(x)$, $\det F(x)^{-1}$ and finally $\phi(x)$ respectively, for the example from Figure 3

2.3 The Central Path

2.3.1 Logarithmic Barrier Function for a Linear Matrix Inequality

Many interior point methods rely heavily on functions that give a strong "signal" in an analytical sense whenever the algorithm in question is evaluating how close the current iterate is to the boundaries of the feasible set. A natural way to provide such a function is given by

$$\phi(x) = \begin{cases} \log \det F(x)^{-1} & \text{if } F(x) > 0 \\ +\infty & \text{else} \end{cases}. \quad (17)$$

Note that the quantity inside the log-function is one over the product of the eigenvalues of $F(x)$, which will go to infinity as at least one of these eigenvalues approaches zero, i.e. as x approaches the boundaries of the feasible region. The log-function itself provides for the barrier function to stay fairly flat inside the feasible region, and to become very steep at its borders. Figure 4 provides the graph of ϕ for the example from Figure 3.

Besides being intuitively reasonable, it is also analytic, strictly convex and self concordant as long as $F(x) > 0$ (for a proof, see [53]). The entries of the gradient and the Hessian of ϕ are given by

$$\begin{aligned} (\nabla \phi(x))_i &= \text{Tr } F(x)^{-1} F_i \\ (\nabla^2 \phi(x))_{ij} &= \text{Tr } F(x)^{-1} F_i F(x)^{-1} F_j. \end{aligned} \quad (18)$$

The minimizer of the barrier function is called the *analytic centre*, which can be

computed efficiently using Newton's method ([53] provides a complete analysis of the speed of convergence of such an algorithm).

2.3.2 Objective and Duality Gap Parametrization

In the following we want to convey a picture of what the idea behind the central path really is. To this end, consider the linear matrix inequality given by

$$\begin{aligned} F(x) &\geq 0 \\ c^T x &= \gamma \end{aligned} \tag{19}$$

where $p^* < \gamma < \sup\{c^T x \mid F(x) \geq 0\}$. Minimizing $\phi(x)$ over this set along all values of γ within the constraints yields a curve called the *Central Path*. One can analyze the respective minimization problem given by

$$\begin{aligned} \min_x \quad & \log \det F(x)^{-1} \\ \text{subject to} \quad & F(x) \geq 0 \\ & c^T x = \gamma \end{aligned} \tag{20}$$

and find that for any given γ the minimizer $x^*(\gamma)$ satisfies

$$\text{Tr } F(x^*(\gamma))^{-1} F_i = \lambda c_i, \quad \forall i = 1, \dots, m \tag{21}$$

where λ is a Lagrange multiplier. This implies that $F(x^*(\gamma))^{-1} \lambda^{-1}$ is a dual feasible solution, with the associated duality gap given by

$$\eta = F(x) \bullet Z = F(x^*(\gamma)) \bullet F(x^*(\gamma))^{-1} \lambda^{-1} = \frac{n}{\lambda} \tag{22}$$

Further one can consider the optimization problem that yields $\mathcal{D}$'s central path, parameterized using the above duality gap:

$$\begin{aligned} \min_Z \quad & \log \det Z^{-1} \\ \text{subject to} \quad & \text{Tr } F_i Z = c_i, \quad \forall i = 1, \dots, m \\ & F_0 \bullet Z = \gamma - \frac{n}{\lambda} \\ & Z \geq 0 \end{aligned} \tag{23}$$

and find that $F(x^*(\gamma))^{-1} \lambda^{-1}$ is its minimizer, i.e. $Z^*(\gamma - \frac{n}{\lambda}) = F(x^*(\gamma))^{-1} \lambda^{-1}$. This gives a nice pairing between points on the dual and the primal Central Path.

Finally, one can simultaneously solve both problems by considering

$$\begin{aligned}
& \min_{x, Z} \quad \log \det F(x)^{-1} + \log \det Z^{-1} \\
& \text{subject to} \quad \text{Tr } F_i Z = c_i, \quad \forall i = 1, \dots, m \\
& \quad \quad \quad c^T x + F_0 \bullet Z = \eta \\
& \quad \quad \quad F(x) \geq 0 \\
& \quad \quad \quad Z \geq 0
\end{aligned} \tag{24}$$

where the parameter now is the duality gap η . From the optimality conditions we have that the optimal pair $x^*(\eta), Z^*(\eta)$ satisfies

$$Z^*(\eta)F(x^*(\eta)) = \left(\frac{\eta}{n}\right)I. \tag{25}$$

Note that (25) converges to zero as $\eta \rightarrow 0$, so that eventually $x^*(\eta), Z^*(\eta)$ satisfy the optimality conditions in (13). The principal idea behind many algorithms that solve SPDs is to follow the central path on its way to the optimal primal dual solution pair. Now we can construct a quantity that will serve as a measure of deviation from the central path. Let x and Z be a feasible pair with duality gap $\eta = c^T x + F_0 \bullet Z$. Then $x^*(\eta), Z^*(\eta)$ is the pair on the Central Path with the same duality gap and we have

$$-\log \det Z^*(\eta)F(x^*(\eta)) = -n \log \left(\frac{\eta}{n}\right) = n \log n - n \log F(x) \bullet Z. \tag{26}$$

Thus we can define

$$\begin{aligned}
\psi(x, Z) &:= -\log \det ZF(x) + \log \det Z^*(\eta)F(x^*(\eta)) \\
&= -\log \det ZF(x) + n \log n - n \log F(x) \bullet Z.
\end{aligned} \tag{27}$$

This quantity is zero if x and Z are central, and is positive and increasing as they diverge from centrality. Note that there is no necessity to actually compute the respective central pair in order to evaluate $\psi(x, Z)$.

2.4 Algorithms for solving SDPs

Extensive literature exists that covers all different kinds of algorithms. However, the discussion here will be restricted to one approach that is provided in [45] since this approach appears to be the most comprehensive given the material covered thus far.

Next, a short survey of other algorithms will be given that lays out the basic ideas and the respective literature without going into great detail.

2.4.1 Primal-Dual Potential-Reduction Methods

So far we have studied two decisive ideas that will guide the construction of this type of algorithm. The first being Strong Duality (i.e. the fact that the duality gap is zero at the optimum), and the second being the primal-dual central path. The aim of potential-reduction methods is to reduce the duality gap at every iteration, while maintaining proximity to the central path in order to ensure feasible convergence. To this end, consider the following *Potential Function*:

$$\begin{aligned}\varphi(x, Z) &:= v \sqrt{n} \log(F(x) \bullet Z) + \psi(x, Z) \\ &= (n + v \sqrt{n}) \log(F(x) \bullet Z) - \log \det F(x) - \log \det Z - n \log n\end{aligned}\tag{28}$$

which is a combination of the Duality Gap and the aforementioned measure of deviance from the Central Path (it was first introduced in [23] and [54], hence it is also called *Tanabe – Todd – Ye potential function*). The constant $v \geq 1$ is a parameter that gives relative weight to the term involving the duality gap.

Note that since $\psi(x, Z) \geq 0$ for all $x \in P$ and $Z \in D$, it holds that

$$\eta \leq e^{\frac{\varphi}{v \sqrt{n}}}.\tag{29}$$

Thus, minimizing the φ over the feasible set eventually results in optimal Solutions (x^*, Z^*) of $\mathcal{P}$ and $\mathcal{D}$.

The basic outline of potential-reduction algorithms is as follows:

Initialization

give strictly feasible $(x^{(0)}, Z^{(0)})$

give v

repeat

1. Find a suitable direction δx and a suitable dual feasible direction δZ
2. Find $(p, q) \in \mathbb{R}^2$ that minimize $\varphi(x + p\delta x, Z + q\delta Z)$ (plane search)
3. Update: $x^{(k)} = x^{(k-1)} + p\delta x$ and $Z^{(k)} = Z^{(k-1)} + q\delta Z$

until Duality Gap $\leq \epsilon$

There are many different ways to achieve step 1., but they all have it in common

that the directions δx and δZ are obtained by solving system equations of the form

$$\begin{aligned} S\delta ZS + \sum_{i=1}^m \delta x_i F_i &= -D \\ F_j \bullet \delta Z &= 0, \quad j = 1, \dots, m. \end{aligned} \quad (30)$$

Matrices $D = D^T$ and $S = S^T \geq 0$ depend on the method of choice, one of which will be introduced below. Others can be found in [45], which are originally due to [53], [14], [32] and [26]. One can solve this system by calculating δZ from the first equation, and substituting into the m remaining equations. The system so obtained is given by

$$\sum_{i=1}^m \delta x_i (F_j \bullet S^{-1})(F_i \bullet S^{-1}) = -(F_j \bullet S^{-1})(D \bullet S^{-1}) \quad (31)$$

While there are many ways to specify (31), we restrict ourselves to the basic example given in [45]. For convenience, x and Z here denote the current iterates and $F = F_0 + \sum_{i=1}^m x_i F_i$. The specification of (31) is obtained by setting $S = F$ and $D = F - \rho FZF$ so that we have

$$\begin{aligned} F\delta Z^p F + \sum_{i=1}^m \delta x_i^p F_i &= -\rho FZF + F \\ F_j \bullet Z^p &= 0, \quad j = 1, \dots, m \end{aligned} \quad (32)$$

$$\begin{aligned} Z^{-1}\delta Z^d Z^{-1} + \sum_{i=1}^m \delta x_i^d F_i &= -\rho F + Z^{-1} \\ F_j \bullet Z^d &= 0, \quad j = 1, \dots, m. \end{aligned} \quad (33)$$

The primal step is computed from (32) and the dual step from (33). Both arise from a modified Newton's Method. Recall that the aim of Newton's Method is to iteratively minimize a function by considering its second order approximation from which a suitable direction for updating a current iterate is obtained. This might only work with functions that are convex, since otherwise there is no guarantee that the direction obtained actually points to the minimum. Unfortunately, the first term of $\varphi(x, Z)$, involving the Duality Gap, is concave. This is simply accounted for by skipping the second derivative of that term when calculating the approximation. There are obvious downsides to this approach. The first is that two systems of equations have to be solved in every iteration of the algorithm. Second, solving these two systems requires the calculation of inverses of Z and F , which poses a

problem especially as they approach the optimal solution, thus becoming singular. The aforementioned alternatives for specifying (31) tackle this issues to some extent.

Once suitable directions have been calculated, the step-length is determined by plane search, i.e. by minimizing $\varphi(x + p\delta x, Z + q\delta Z)$ over p and d . To this end, the potential function can be reformulated (recall, that the term $\delta F \delta Z = 0$ and thus vanishes):

$$\begin{aligned} \varphi(x + p\delta x, Z + q\delta Z) &= \varphi(x, Z) + (n + v \sqrt{n}) \log(1 + c_1 p + c_2 q) \\ &\quad - \log \det(I + pF^{-1/2} \delta F F^{-1/2}) - \log \det(I + qZ^{-1/2} \delta Z Z^{-1/2}) \end{aligned} \quad (34)$$

where again $F = F(x)$, $\delta F = \sum_{i=1}^m \delta x_i F_i$ and

$$c_1 = \frac{c^T \delta x}{F(x) \bullet Z}, \quad c_2 = \frac{F_0 \delta Z}{F(x) \bullet Z}. \quad (35)$$

This can be minimized quite efficiently by diagonalizing $F^{-1/2} \delta F F^{-1/2}$ and $Z^{-1/2} \delta Z Z^{-1/2}$. Once the respective eigenvalues are obtained, (35) can be further reformulated into

$$\begin{aligned} \varphi(x + p\delta x, Z + q\delta Z) &= \varphi(x, Z) + (n + v \sqrt{n}) \log(1 + c_1 p + c_2 q) \\ &\quad - \sum_{i=1}^n \log(1 + p\mu_i) - \sum_{i=1}^n \log(1 + qv_i). \end{aligned} \quad (36)$$

From this it is clear that step lengths p and q are restricted to the rectangle given by the intervals $p_{\min} < p < p_{\max}$ and $q_{\min} < q < q_{\max}$, where

$$\begin{aligned} p_{\min} &= \max_i \left\{ \frac{-1}{\mu_i} \middle| \mu_i > 0 \right\} \\ p_{\max} &= \min_i \left\{ \frac{-1}{\mu_i} \middle| \mu_i < 0 \right\} \\ q_{\min} &= \max_i \left\{ \frac{-1}{v_i} \middle| v_i > 0 \right\} \\ q_{\max} &= \min_i \left\{ \frac{-1}{v_i} \middle| v_i < 0 \right\}. \end{aligned} \quad (37)$$

Two facts are worth mentioning in regard to plane search. The first thing is that $\varphi(x + p\delta x, Z + q\delta Z)$ is quasi-concave and has a unique local minimum in the feasible region, which makes standard algorithms, such as the guarded Newton's Method or the Nelder-Mead Method, viable options for solving this problem. Second, note that trivially the minimization of duality gap is achieved by one of the corner solutions of

the aforementioned rectangle. Thus, if the duality gap can at that stage be reduced to a size meeting the stopping criterion of the algorithm, it can terminate.

2.4.2 Primal Logarithmic Barrier Methods

These methods are quite similar to the potential reduction method discussed above with the difference being that they only consider the primal (or the dual) and follow the respective primal (or dual) central path. The method goes back to [36], who presented the principle idea for linear programs, which was later extended to SPDs. The basic approach is to incorporate the barrier into the objective function to obtain

$$f(x) = \frac{1}{\eta} c^T x + \log \det F(x)^{-1}. \quad (38)$$

This function is then minimized in an iterative process, where the steps for updating current iterates comprise an affine scaling component that aims to minimize the primal target function, and a centering component that minimizes the barrier function. The various methods that follow this general design differ in what is considered the primal problem (either $\mathcal{P}$, or $\mathcal{D}$ in which case the target function looks slightly different) and the step length. A detailed discussion of these algorithms can be found in [41] and [13].

2.4.3 Primal-Dual Affine-Scaling Methods

Since optimality conditions for the minimization of the duality gap are given by (12), it comes as an immediate idea to find some iterative scheme to solve this non linear system of equations. This is achieved by Affine-Scaling Methods, which aim to minimize the duality gap over some ellipsoid inscribed in $P \times D$. They can be viewed as 'greedy' algorithms, which seek to reach an optimum in a single step. This general idea has been introduced by [46] and [31] whose results differ in the computation of the step direction. First, note that $F(x)$ and Z can be scaled to the same symmetric matrix such that

$$V^2 = D^{-\frac{1}{2}} F(x) Z D^{\frac{1}{2}} \quad (39)$$

where D is an appropriate matrix. This type of scaling can be applied to any primal and dual step direction ΔZ and Δx , such as

$$\begin{aligned} F(x + \Delta x) &= F(x) + \Delta F \geq 0 \\ F_i \bullet \Delta Z &= 0 \quad \forall i = 1, \dots, m \end{aligned} \quad (40)$$

so that one has scaled search directions

$$\begin{aligned} D_F &= D^{-\frac{1}{2}} \Delta F D^{\frac{1}{2}} \\ D_Z &= D^{-\frac{1}{2}} \Delta Z D^{\frac{1}{2}} \end{aligned} \quad (41)$$

and a scaled primal-dual step direction

$$D_V = D_X + D_S \quad (42)$$

which is computed as

$$\begin{aligned} D_V^* &= \arg \min_{D_V} V \bullet (V + D_V) \\ &\text{subject to } \|V^{-\frac{1}{2}} D_V V^{-\frac{1}{2}}\| \leq 1. \end{aligned} \quad (43)$$

The step that is computed using the Frobenius-Norm is called the *Dikin – Step*. If the spectral norm is used, the resulting step is called *Primal – Dual Affine – Scaling Direction*. It is worth mentioning that these step directions are also known as predictor directions in the predictor-corrector methods discussed in the next subsection. For a complete layout of the algorithms, the reader can refer to [13].

2.4.4 Primal-Dual Path-Following Methods

As the name suggests, these types of algorithms seek to follow the central path to its limit point, i.e. the optimal solution. This is achieved by solving the optimality conditions of (24) (which are basically given by [13] where the last equality is replaced by [25]) for different (decreasing) values of η . Thus at every iteration one looks for step directions such that

$$\begin{aligned} F_i \bullet \Delta Z &= 0 \\ Z + \Delta Z &\geq 0 \\ \sum_{i=1}^m \Delta x_i F_i &= \Delta F \\ F(x) + \Delta F(x) &\geq 0 \\ (Z + \Delta Z)(F(x) + \Delta F) &= \left(\frac{\eta}{n}\right) I. \end{aligned} \quad (44)$$

This is an overdetermined nonlinear system, where the Nonlinearity stems from the last equation. Using the scaling strategy as in the previous subsection, one can

symmetrize and linearize the equation such that a unique solution is obtained. This solution is given by

$$D_V = \left(\frac{\eta}{n}\right) V^{-1} - V \quad (45)$$

and a suitable step direction can be obtained from

$$\begin{aligned} \Delta Z + D\Delta F D &= \left(\frac{\eta}{n}\right) F^{-1} - Z \\ F_i \bullet \Delta Z &= 0 \\ \sum_{i=1}^m \Delta x_i F_i &= \Delta F \end{aligned} \quad (46)$$

A popular way to implement these ideas is in combination with affine scaling methods, which results in algorithms known as *predictor – corrector* algorithms. The predictor steps arise from the affine scaling methods, and are carried out first, but not fully. The step is then corrected using the path following step directions. For a detailed discussion see [29].

2.5 Some useful remarks on formulating SDPs

Semidefinite Programming contains many instances of convex optimization as special cases. However, it is not always trivial to find the respective SDP-formulation. In the following, a quick survey of useful formulation strategies is presented, some of which will be also helpful in later chapters. This is by far not a complete survey, and even in later chapters of this thesis, more complicated procedures of problem formulation will be needed.

2.5.1 Linear Programs

A linear program is of the form

$$\begin{aligned} \min_x & c^T x \\ \text{subject to} & Ax \leq b \end{aligned} \quad (47)$$

where $A \in \mathbb{R}^{m \times n}$, $c, x \in \mathbb{R}^n$ and $b \in \mathbb{R}^m$. Note that the inequality in (47) is meant component-wise, such that we have m linear inequalities of the form $A_k x \leq b_k$ where A_k is the k th row of A and b_k k th component of b . Thus, the set of feasible solutions is the intersection of m half-spaces, which gives a polyhedral unless the set is empty. Recall that the diagonal matrix is *p.s.d.* only if all its diagonal elements are greater

or equal zero. This is because the eigenvalues of that matrix are given by these very elements. It then becomes clear that the constraint above can be rewritten as

$$\begin{pmatrix} b_1 - A_1x & \dots & 0 \\ \vdots & \ddots & \vdots \\ 0 & \dots & b_m - A_mx \end{pmatrix} \geq 0. \quad (48)$$

Replacing the inequality in (47) with (48) gives a semidefinite program. Whenever we further encounter a diagonal matrix that comprises linear functions like in (48), it will be denoted by **Diag**($b - Ax$), albeit with slightly different arguments. The somewhat reverse operation **diag**(A) gives the vector that is comprised of the diagonal elements of the argument matrix.

2.5.2 Schur Complements and other Nonlinear Problems

First, there shall be a quick review of the idea behind the *Schur Complement* (as provided in [44] Appendix A.5.5). Given a linear system of equalities of the form

$$\begin{pmatrix} A & B \\ C & D \end{pmatrix} \begin{pmatrix} x \\ y \end{pmatrix} = \begin{pmatrix} f \\ g \end{pmatrix} \iff \begin{array}{ll} I & Ax + By = f \\ II & Cx + Dy = g \end{array} \quad (49)$$

where all capital letters denote matrices and all lower case letter denote vectors of an appropriate size. One can perform iterations of the Gaussian algorithm such that $II^{new} = II - C(A^{-1})I = D - C(A^{-1})By = g - C(A^{-1})f$ and the system becomes

$$\begin{pmatrix} A & B \\ 0 & D - C(A^{-1})B \end{pmatrix} \begin{pmatrix} x \\ y \end{pmatrix} = \begin{pmatrix} f \\ g - C(A^{-1})f \end{pmatrix}. \quad (50)$$

The lower right entry of the new matrix, $D - C(A^{-1})B$, is called the *Schur Complement* of the original matrix. It has many interesting properties of which the following is of particular interest for SDPs. A symmetric matrix of the form like in (49) is *p.s.d* only if its Schur complement and A are *p.s.d* (given that A is square symmetric and nonsingular). Now consider the following nonlinear program (an example taken from [45]):

$$\begin{aligned} & \min_x \frac{(c^T x)^2}{d^T x} \\ & \text{subject to } Ax + b \geq 0 \end{aligned} \quad (51)$$

We assume that the objective function exists and is always positive for all feasible x , i.e $d^T x > 0$ whenever $Ax + b \geq 0$. Again we can linearize the target function by introducing an extra variable as well as an extra constraint so that we have

$$\begin{aligned} & \min_{t,x} t \\ & \text{subject to } \frac{(c^T x)^2}{d^T x} \leq t \\ & \quad Ax + b \geq 0. \end{aligned} \tag{52}$$

The new constraint can be written as $t - (c^T x)^2 (d^T x)^{-1} \geq 0$, where the left hand side is a Schur complement. So this constraint is implied by

$$\begin{pmatrix} d^T x & c^T x \\ c^T x & t \end{pmatrix} \geq 0 \tag{53}$$

as long as $d^T x$ is positive, which holds for feasible x by assumption. Thus we see that (57) can be written as an SDP given by

$$\begin{aligned} & \min_{t,x} t \\ & \text{subject to } \begin{pmatrix} \text{diag}(Ax + b) & 0 & 0 \\ 0 & d^T x & c^T x \\ 0 & c^T x & t \end{pmatrix} \geq 0. \end{aligned} \tag{54}$$

2.5.3 Quadratic Programms and Quadratic constraints

Consider the problem

$$\begin{aligned} & \min_x (A_0 x + b_0)^T (A_0 x + b_0) - c_0^T x - d_0 \\ & \text{subject to } (A_i x + b_i)^T (A_i x + b_i) - c_i^T x - d_i \leq 0 \quad \forall i = 1, \dots, m \end{aligned} \tag{55}$$

where $A_i \in \mathbb{R}^{k_i \times n}$, $b_i \in \mathbb{R}^{k_i}$, $x \in \mathbb{R}^n$ and $d_i \in \mathbb{R}$. The target function and the left hand sides of the constraints are convex quadratic functions. First, note that the m constraints can be written as linear matrix inequalities of the form

$$\begin{pmatrix} I & A_i x + b_i \\ (A_i x + b_i)^T & c_i^T x + d_i \end{pmatrix} \geq 0. \tag{56}$$

In order to linearize the target function, it is simply replaced by t and introduced as an additional constraint of the form

$$\text{subject to } (A_0x + b)^T(A_0x + b) - c_0^T x - d_0 \leq t \quad (57)$$

and finally reformulated as

$$\begin{pmatrix} I & A_0x + b_i \\ (A_0x + b_i)^T & c_0^T x + d_0 + t \end{pmatrix} \succeq 0. \quad (58)$$

We then obtain an SDP of the form

$$\begin{aligned} \min_{t,x} t \quad (59) \\ \begin{pmatrix} I & A_0x + b_0 \\ (A_0x + b_0)^T & c_0^T x + d_0 + t \end{pmatrix} \succeq 0 \\ \begin{pmatrix} I & A_i x + b_i \\ (A_i x + b_i)^T & c_i^T x + d_i \end{pmatrix} \succeq 0. \end{aligned}$$

Note that this is of the standard form, since the two linear matrix inequalities can be summarized into a single one where the respective matrix has a block diagonal structure. The ideas presented above are a review from [45].

2.5.4 Second Order Cone Programms (SOCP)

The Material covered here is mainly due to [18]. The unit second order cone (SOC) is given by the epigraph of the 2-norm function.

$$SOC^n = \left\{ \begin{pmatrix} x_0 \\ x \end{pmatrix} \middle| x_0 \in \mathfrak{R}, x \in \mathfrak{R}^{n-1}, x_0 \geq \|x\|_2 \right\} \quad (60)$$

where $\|x\|_2 = (x^T x)^{\frac{1}{2}}$. In $\mathfrak{R}^3$ the graph has the shape of an upright ice-cream cone that emerges from the origin and is symmetric around the z-axis. Note that a SOC is a closed convex cone. In a SOCP, the task is to minimize a linear function over a SOC^n :

$$\begin{aligned} \min_x c^T x \quad (61) \\ \text{subject to } \|A_i x + b_i\| \leq d_i^T x + e_i \quad \forall i = 1, \dots, m \end{aligned}$$

where $e_i \in \mathfrak{R}$, $c, d_i, x \in \mathfrak{R}^n$, $b_i \in \mathfrak{R}^{(k_i-1)}$ and $A_i \in \mathfrak{R}^{(k_i-1) \times n}$. The constraint is called a Second-order cone constraint. All points that satisfy the constraint belong to a set, which is given by the inverse image of the SOC under an affine mapping, i.e.

$$\|A_i x + b_i\| \leq d_i^T x + e_i \iff \begin{pmatrix} A_i \\ d_i^T \end{pmatrix} x + \begin{pmatrix} b_i \\ e_i \end{pmatrix} \in \text{SOC}^{k_i}. \quad (62)$$

Many different instances of convex programming can be cast as a SOCP, but here the focus is on the relation between SOCPs and SDPs. Note that

$$\begin{aligned} \|x\| &\leq t \\ (x^T x)^{\frac{1}{2}} &\leq t \\ 0 &\leq t - x^T t^{-1} x \end{aligned} \quad (63)$$

gives a formulation of the second-order cone constraint that involves a Schur complement. Thus it can be rewritten as a linear matrix inequality:

$$\begin{pmatrix} tI & x \\ x^T & t \end{pmatrix} \geq 0 \quad (64)$$

Thus the SOCP can be cast as an SDP of the form

$$\begin{aligned} &\min_x c^T x \\ &\text{subject to } \begin{pmatrix} (d_i^T x + e_i)I & A_i x + b_i \\ (A_i x + b_i)^T & d_i^T x + e_i \end{pmatrix} \geq 0 \quad \forall i = 1, \dots, m. \end{aligned} \quad (65)$$

It is worth mentioning that from an SDP point of view this problem can be seen as a special case of restricting a spectral- or maximum singular value norm of an affine linear combination of matrices. Such a restriction is given by

$$\|A(x)\| \leq t \quad (66)$$

where $A(x) = A_0 + \sum_{i=1}^m x_i A_i \in \mathfrak{R}^{k \times n}$ and $t \in \mathfrak{R}^n$. By a similar argument as in the SOC-case, it can be written as a Linear Matrix Inequality:

$$\begin{pmatrix} tI & A(x) \\ A(x)^T & tI \end{pmatrix} \geq 0 \quad (67)$$

Considering t a decision variable, one can formulate the problem of finding the

smallest spectral norm of $A(x)$ as

$$\begin{aligned} & \min_{t,x} t \\ & \text{subject to } \begin{pmatrix} tI & A(x) \\ A(x)^T & tI \end{pmatrix} \succeq 0. \end{aligned} \quad (68)$$

2.5.5 Relaxations of some NP-Hard Problems

Consider the nonconvex quadratic problem

$$\begin{aligned} & \min_x x^T A_0 x + 2b_0^T x + c_0 \\ & \text{subject to } x^T A_i x + 2b_i^T x + c_i \leq 0 \quad \forall i = 1, \dots, m. \end{aligned} \quad (69)$$

A_i are not confined to S_+ , so neither the objective function nor the constraints are convex, which makes (69) a hard problem. However, one can derive the dual and introduce an extra variable t , which results in an SDP. The Lagrangian of (69) is given by

$$\begin{aligned} L(x, z, t) &= t + (x^T A_0 x + 2b_0^T x + c_0 - t) + \sum_{i=1}^m z_i (x^T A_i x + 2b_i^T x + c_i) \\ &= t + \begin{pmatrix} x \\ 1 \end{pmatrix}^T \begin{pmatrix} A_0 & b_0 \\ b_0 & c_0 - t \end{pmatrix} \begin{pmatrix} x \\ 1 \end{pmatrix} + \sum_{i=1}^m z_i \begin{pmatrix} x \\ 1 \end{pmatrix}^T \begin{pmatrix} A_i & b_i \\ b_i & c_i \end{pmatrix} \begin{pmatrix} x \\ 1 \end{pmatrix} \end{aligned} \quad (70)$$

which for a given $(z, t) \in \mathfrak{R}^{m+1}$ is bounded below if

$$\begin{pmatrix} A_0 & b_0 \\ b_0 & c_0 - t \end{pmatrix} + \sum_{i=1}^m z_i \begin{pmatrix} A_i & b_i \\ b_i & c_i \end{pmatrix} \succeq 0. \quad (71)$$

Thus, the dual can be cast as an SDP:

$$\begin{aligned} & \max_{z,t} t \\ & \text{subject to } \begin{pmatrix} A_0 & b_0 \\ b_0 & c_0 - t \end{pmatrix} + \sum_{i=1}^m z_i \begin{pmatrix} A_i & b_i \\ b_i & c_i \end{pmatrix} \succeq 0 \\ & \quad z_i \geq 0 \quad \forall i = 1, \dots, m. \end{aligned} \quad (72)$$

As usual, weak duality holds, i.e. the dual yields a lower bound on the primal problem. However, since (69) is not a convex problem, strong duality is not guaranteed. But having a good and easily computed bound on a problem is still useful, e.g. in

order to evaluate the performance of heuristics or for the implementation of branch and bound type algorithms.

The MAX-CUT problem is a famous instance, where the relaxation is a special case of (72). It aims to separate the nodes of a graph into two groups such that the weight of the arcs across groups is maximized. Formally it is given by

$$\begin{aligned} \max_x \quad & \frac{1}{4} \sum_{i=1}^n \sum_{j=1}^n w_{ij}(1 - x_i x_j) \\ \text{subject to } & x_i \in \{-1, 1\} \quad \forall j = 1, \dots, m. \end{aligned} \quad (73)$$

Here w_{ij} denotes the weight on the arc between nodes i and j . x_i is 1 if node i belongs to group one, and -1 otherwise. The objective function can be reformulated into

$$\frac{1}{4} \sum_{i=1}^n \sum_{j=1}^n w_{ij}(1 - x_i x_j) = \frac{1}{4} \sum_{i=1}^n \sum_{j=1}^n w_{ij} - x^T W x. \quad (74)$$

The first term is now a constant and W is a matrix whose j th entry is given by w_{ji} . Further the constraint is equivalent to $x_i^2 = 1$ which can be represented by two quadratic inequalities. Thus (73) is equivalent to

$$\begin{aligned} \min_x \quad & x^T W x \\ \text{subject to } & x_i^2 \leq 1 \quad \forall j = 1, \dots, n \\ & -x_i^2 \leq -1 \quad \forall j = 1, \dots, n \end{aligned} \quad (75)$$

which clearly is a special case of (69) and might be relaxed as an SDP accordingly. The above example is taken from [40].

2.5.6 S-Lemma and S-Procedure

The above result is a special case of a wider class of problems that are not as easily dealt with. The basic motivation behind these is the question of when a quadratic inequality is a consequence of other quadratic inequalities. As an illustration one could consider the problem of finding the smallest ball that contains a set of k other ellipsoids (this is an example from [45]). Since ellipsoids may be described as the sub level sets of quadratic functions, we may define the respective sets as

$$\mathcal{E}_i = \{x | f_i(x) = x^T A_i x + 2b_i^T x + c_i \leq 0\}. \quad (76)$$

If the ellipsoid that is to be minimized is described by

$$\tilde{\mathcal{E}} = \{x | \tilde{f}_i(x) = x^T x - 2x_c^T x + \gamma \leq 0\}, \quad (77)$$

then we want to find a condition that is equivalent to

$$x \in \mathcal{E}_i \quad \forall i = 1, \dots, k \Rightarrow x \in \tilde{\mathcal{E}}. \quad (78)$$

Such a condition can be obtained from the so called *S-lemma*, a detailed discussion of which can be found in [47]. We present here a corollary which can be found in [37] (chapter 9.3)

Lemma 2.5.6: Let

$$f_i(x) = x^T A_i x + 2b_i^T x + c_i \quad i = 0, 1, \dots, p \quad (79)$$

be quadratic function of $x \in \mathfrak{R}^n$. Then,

$$f_i(x) \geq 0, i = 1, \dots, p \Rightarrow f_0(x) \geq 0 \quad (80)$$

if there exist $\tau_i \geq 0$ such that

$$\begin{pmatrix} A_0 & b_0 \\ b_0^T & c_0 \end{pmatrix} - \sum_{i=1}^p \tau_i \begin{pmatrix} A_i & b_i \\ b_i^T & c_i \end{pmatrix} \geq 0. \quad (81)$$

If $p = 1$, then the converse also holds as long as $\exists x_0$ s.t. $f_1(x_0) > 0$.

This way of obtaining constraints is often referred to as *S-Procedure*. For the above example, we can thus formulate constraints, which guarantee that the minimized ball contains all the ellipsoids by applying that very procedure, thus obtaining

$$\begin{pmatrix} I & -x_c \\ -x_c^T & \gamma \end{pmatrix} - \sum_{i=1}^p \tau_i \begin{pmatrix} A_i & b_i \\ b_i^T & c_i \end{pmatrix} \leq 0. \quad (82)$$

2.5.7 Sufficient Conditions for Positive Polynomials

For any polynomial of the form $g(x) = \sum_{r=0}^n y_r x^r$, one could ask for the conditions on the coefficients $y_r \in \mathfrak{R}$ such that the polynomial stays non negative on a given domain. It turns out that such conditions can be stated in terms of linear and positive semidefinite constraints. These are due to [20] and since they will be very useful in later chapters, their results will be presented in the following. The interested reader may refer to the source literature for the proofs. [16] found some minor mistakes in

the original version the paper from 1999, which was later revised. What follows is the corrected version:

Proposition 2.5.7

(a) The polynomial $g(x) = \sum_{r=0}^{2k} y_r x^r$ satisfies $g(x) \geq 0$ for all $x \in \mathfrak{R}$ if and only if there exists a p.s.d. matrix $X \in \mathfrak{R}^{(k+1) \times (k+1)}$, $X = [x_{ij}]_{i,j=0,\dots,k}$, such that

$$\begin{aligned} y_r &= \sum_{i,j:i+j=r} x_{i,j} \quad \forall r = 0, \dots, 2k \\ X &\geq 0. \end{aligned} \tag{83}$$

(b) The polynomial $g(x) = \sum_{r=0}^k y_r x^r$ satisfies $g(x) \geq 0$ for all $x \geq 0$ if and only if there exists a p.s.d. matrix $X \in \mathfrak{R}^{(k+1) \times (k+1)}$, $X = [x_{ij}]_{i,j=0,\dots,k}$, such that

$$\begin{aligned} 0 &= \sum_{i,j:i+j=2l-1} x_{i,j} \quad \forall l = 1, \dots, k \\ y_l &= \sum_{i,j:i+j=2l} x_{i,j} \quad \forall l = 0, \dots, k \\ X &\geq 0. \end{aligned} \tag{84}$$

(c) The polynomial $g(x) = \sum_{r=0}^k y_r x^r$ satisfies $g(x) \geq 0$ for all $x \in [0, a]$ if and only if there exists a p.s.d. matrix $X \in \mathfrak{R}^{(k+1) \times (k+1)}$, $X = [x_{ij}]_{i,j=0,\dots,k}$, such that

$$\begin{aligned} 0 &= \sum_{i,j:i+j=2l-1} x_{i,j} \quad \forall l = 1, \dots, k \\ \sum_{r=0}^l y_r \binom{k-r}{l-r} a^r &= \sum_{i,j:i+j=2l} x_{i,j} \quad \forall l = 0, \dots, k \\ X &\geq 0. \end{aligned} \tag{85}$$

(d) The polynomial $g(x) = \sum_{r=0}^k y_r x^r$ satisfies $g(x) \geq 0$ for all $x \in [a, \infty)$ if and only if there exists a p.s.d. matrix $X \in \mathfrak{R}^{(k+1) \times (k+1)}$, $X = [x_{ij}]_{i,j=0,\dots,k}$, such that

$$\begin{aligned} 0 &= \sum_{i,j:i+j=2l-1} x_{i,j} \quad \forall l = 1, \dots, k \\ \sum_{r=l}^k y_r \binom{r}{l} a^{r-l} &= \sum_{i,j:i+j=2l} x_{i,j} \quad \forall l = 0, \dots, k \\ X &\geq 0. \end{aligned} \tag{86}$$

(e) The polynomial $g(x) = \sum_{r=0}^k y_r x^r$ satisfies $g(x) \geq 0$ for all $x \in (-\infty, a]$ if and only

if there exists a p.s.d. matrix $X \in \Re^{(k+1) \times (k+1)}$, $X = [x_{ij}]_{i,j=0,\dots,k}$, such that

$$\begin{aligned} 0 &= \sum_{i,j:i+j=2l-1} x_{i,j} \quad \forall l = 1, \dots, k \\ (-1)^l \sum_{r=0}^k y_r \binom{r}{l} a^{r-l} &= \sum_{i,j:i+j=2l} x_{i,j} \quad \forall l = 0, \dots, k \\ X &\geq 0. \end{aligned} \tag{87}$$

(f) The polynomial $g(x) = \sum_{r=0}^k y_r x^r$ satisfies $g(x) \geq 0$ for all $x \in [a, b]$ if and only if there exists a p.s.d. matrix $X \in \Re^{(k+1) \times (k+1)}$, $X = [x_{ij}]_{i,j=0,\dots,k}$, such that

$$\begin{aligned} 0 &= \sum_{i,j:i+j=2l-1} x_{i,j} \quad \forall l = 1, \dots, k \\ \sum_{m=0}^l \sum_{r=m}^{k+m-l} y_r \binom{r}{l} \binom{k-r}{l-m} a^{r-m} b^m &= \sum_{i,j:i+j=2l} x_{i,j} \quad \forall l = 0, \dots, k \\ X &\geq 0. \end{aligned} \tag{88}$$

3 Asset Pricing Approaches using Conic and Semidefinite Optimization

One of the most important areas of research in finance is the problem of pricing a security, whose payoff lies in the future, which is a especially hard task if that payoff is uncertain. The following chapter will provide some insight in how SDPs are helpful in such an endeavor.

3.1 Arbitrage Pricing Theory: Recovering Risk Neutral Probability Measures ("RNPM") from Option Prices

In this chapter we will review the results by [25]. Before the actual problem is introduced, the basic ideas behind the arbitrage pricing theory will be briefly revisited, albeit in a sketchy way. The interested reader may refer to [49] or [42] for a more detailed discussion.

3.1.1 Arbitrage Pricing : Basic Theory

Consider a one period model of an economy (i.e. the set of points in time is given by $t \in T = \{0, 1\}$) endowed with a risk free rate r . Let $s \in S = \{1, \dots, m\}$ denote a set of

possible states of the economy at $t=1$. Further assume that there are m fundamental assets with indices $i = 1, \dots, m$, each costing q_i at $t = 0$ and paying X_i^s at $t = 1$ if the economy is in state s . The respective vector of payoffs of asset i is denoted by X_i , and the vector of prices as q . The matrix whose entry in the i -th column and the j -th row gives X_i^j is denoted as $X \in \mathfrak{R}^{m \times m}$. Further assume all X_i to be linearly independent, i.e. X has full column rank. Then one can always find a vector $b_j \in \mathfrak{R}^m$ such that

$$Xb_j = X_{m+j} \quad (89)$$

with $j > 0$ where X_{m+j} is the payoff-vector of an arbitrary, non-fundamental asset. Thus any non-fundamental assets payoff can be generated as a linear combination of the payoffs of the fundamental ones, which is called a replicating portfolio. The price of such an asset can then be inferred via the price of the replicating portfolio $q_{m+j} = b_j^T q$. However, another way to price such an asset is as follows. Find a vector $\pi \in \mathfrak{R}^m$ such that

$$X^T \pi = q(1 + r). \quad (90)$$

The price of any asset is then given by

$$X_j^T \pi = q_j(1 + r). \quad (91)$$

This already bears some resemblance to a discounted expected value, where the entries π_s of the vector π play the role of quasi probabilities for the states of the economy. But under the conditions introduced so far, there is no guarantee that $0 \leq \pi_s < 1 \ \forall s \in S$. Thus, we are looking for conditions for the existence of a probability measure under which the price of an asset is equal to its expected payoff. Such a probability measure is called a *Martingale Measure* or *Risk Neutral Probability Measure*. It exists and is uniquely implied by (90) if there is no b_j such that $b_j^T q < 0$ and $Xb_j \geq 0$, which is the condition that the market is free of arbitrage opportunities. This gives rise to the first fundamental theorem of asset pricing.

Theorem 3.1.3.: Given a discrete and complete market that is arbitrage free, there exists a unique probability measure $f(x)$ on a probability space $(S, \mathcal{F}, f(x))$ such that $q_j = \frac{1}{1+r} \sum_{s=1}^m f(s) X_j^s$, i.e. that is a martingale measure.

These basic ideas can be transferred to the continuous case where $S = [a, b]$ and $f(x)$ is a density function. The prices are then given by

$$q_j = \frac{1}{1+r} \int_S f(s) X_j^s ds. \quad (92)$$

In this case the no-arbitrage-condition has to be tightened, such that market returns of assets follow a Brownian motion.

3.1.2 Model for Recovering the RNPM

The question potentially raised is how one can recreate the RNPM by observing market data. Of course this would be useful since we could price any asset accordingly, as soon as we know its value in any given state. But at first glance this seems like a hopeless endeavor, since the number of possible states is now infinite, and we can only observe a finite number of prices. However, [25] gave a sensible solution to this problem using cubic splines.

The idea is to recreate the density function $f : [a, b] \rightarrow \mathfrak{R}$ first using estimates for $f(x_i)$ on a set of points $\{x_i\}, i = 1, \dots, m + 1$, where $x_1 = a$ and $x_{m+1} = b$, and then filling the gaps with appropriate polynomial functions. A spline in this regard is a piecewise polynomial approximation $\hat{f}(x)$ to f , such that they are identical at the nodes x_i , i.e. $\hat{f}(x_i) = f(x_i)$. Here cubic splines are of specific interest, i.e. splines where the respective polynomials are of the form $\hat{f}_i(x) = \alpha_i x^3 + \beta_i x^2 + \gamma_i x + \delta_i$. Further the spline function as a whole has to fulfill the requirement that its second derivative exists at each node, as to ensure the smoothness of $\hat{f}(x)$. Given $m + 1$ nodes, we have m intervals for which the respective polynomials have to meet the following $4m$ conditions:

$$\begin{aligned} \hat{f}_i(x_i) &= f(x_i) \quad i = 1, \dots, m \quad \text{and} \quad \hat{f}_m(x_{m+1}) = f(x_{m+1}) \\ \hat{f}_{i-1}(x_i) &= \hat{f}_i(x_i) \quad i = 2, \dots, m \\ \hat{f}'_{i-1}(x_i) &= \hat{f}'_i(x_i) \quad i = 2, \dots, m \\ \hat{f}''_{i-1}(x_i) &= \hat{f}''_i(x_i) \quad i = 2, \dots, m \\ \hat{f}''_1(x_1) &= \hat{f}''_m(x_{m+1}) = 0. \end{aligned} \tag{93}$$

The last condition leads to a so called natural spline, and ensures that there is no curvature at the tails of the density function. Each of the m partial functions $\hat{f}_i(x)$ bring along unknown parameters $\alpha_i, \beta_i, \gamma_i, \delta_i$. The aim is to set their values in a way such that the resulting density function $\hat{f}(x)$ implies option prices for a chosen security that are as close as possible to observed option prices. To this end we chose an arbitrary security and identify the state space S in terms of possible future prices of that security, i.e. $s \in S = [a, b]$ contain the future prices of the stock. Option prices

on the chosen security in the arbitrage pricing framework are given by

$$\begin{aligned} C_K &= \frac{1}{1+r} \int_a^b f(s)(s-K)^+ ds \\ P_K &= \frac{1}{1+r} \int_a^b f(s)(K-s)^+ ds \end{aligned} \quad (94)$$

where $C_K, P_K \in \mathfrak{R}$ are the call and put prices, respectively, and $K \in \mathfrak{R}$ is the strike price. If $\hat{C}_K$ and $\hat{P}_K$ are the put and call prices, respectively, according to $\hat{f}(x)$, then the objective is to minimize

$$\sum_{K \in \mathcal{K}} (C_K - \hat{C}_K)^2 + (P_K - \hat{P}_K)^2, \quad (95)$$

i.e. the squared difference between the actual option price and its theoretical counterpart ($\mathcal{K}$ is the set of strike prices which is chosen conveniently). The former are easily obtained from market data and their respective strike prices will serve as nodes such that we have $x_i \in \mathcal{K}$ with $x_1 = a, x_{m+1} = b$ and $x_i < x_{i+1}$. The latter can be computed as follows:

$$\begin{aligned} (1+r)\hat{C}_K &= \int_a^b \hat{f}(s)(s-K)^+ ds \\ &= \sum_{i=1}^m \int_{x_i}^{x_{i+1}} \hat{f}(s)(s-K)^+ ds \\ &= \sum_{i=j+1}^m \int_{x_i}^{x_{i+1}} \hat{f}(s)(s-K) ds \\ &= \sum_{i=j+1}^m \int_{x_i}^{x_{i+1}} (\alpha_i s^3 + \beta_i s^2 + \gamma_i s + \delta_i)(s-K) ds \end{aligned} \quad (96)$$

$$\begin{aligned}
(1+r)\hat{P}_K &= \int_a^b \hat{f}(s)(K-s)^+ ds \\
&= \sum_{i=1}^m \int_{x_i}^{x_{i+1}} \hat{f}(s)(K-s)^+ ds \\
&= \sum_{i=1}^{j-1} \int_{x_i}^{x_{i+1}} \hat{f}(s)(K-s) ds \\
&= \sum_{i=1}^{j-1} \int_{x_i}^{x_{i+1}} (\alpha_i s^3 + \beta_i s^2 + \gamma_i s + \delta_i)(K-s) ds
\end{aligned} \tag{97}$$

where $x_j = \min_i \{x_i | x_i \geq K\}$. Note that these expressions are linear in the parameters of the spline function which are our decision variables. So far we have ensured that $\hat{f}(x)$ is a natural cubic spline, with the desired properties of matching the real density function at the nodes, and implying the observed option prices as close as possible. For it to be a proper density function its integral over S has to equal one and it must be non-negative. The first requirement is easily enforced via

$$\sum_{i=1}^m \int_{x_i}^{x_{i+1}} \hat{f}(s) ds = 1. \tag{98}$$

The second is much harder and requires some additional theory, which is already discussed in chapter 2.5.7. However, for the sake of self-containment of this chapter, the result in question is stated again one more.

Theorem 3.1.3.b: The polynomial $g(x) = \sum_{r=0}^k y_r x^r$ satisfies $g(x) \geq 0$ for all $x \in [a, b]$ if and only if there exists a positive semidefinite matrix $X = [x_{ij}]_{i,j=0,\dots,k}$ such that

$$\begin{aligned}
\sum_{i,j:i+j=2l-1} x_{ij} &= 0, \quad \forall l = 1, \dots, k \\
\sum_{i,j:i+j=2l} x_{ij} &= \sum_{m=0}^l \sum_{r=m}^{k+m+l} y_r \binom{r}{m} \binom{k-r}{l-m} a^{r-m} b^m \quad \forall l = 0, \dots, k \\
X &\geq 0.
\end{aligned} \tag{99}$$

For the case of a third degree polynomial, the following corollary can be stated:

Corollary 3.1.3.b: The polynomial $\hat{f}_q(x) = \alpha_q x^3 + \beta_q x^2 + \gamma_q x + \delta_q$ satisfies $\hat{f}_q(x) \geq 0$

for all $x \in [x_q, x_{q+1}]$ if and only if there exists a 4×4 matrix $X^q = [x_{ij}^q]_{i,j=0,\dots,3}$ such that

$$\begin{aligned}
x_{ij}^q &= 0, \quad \text{if } i+j \in \{1, 5\} \\
x_{03}^q + x_{12}^q + x_{21}^q + x_{30}^q &= 0 \\
x_{00}^q &= \alpha_q x_q^3 + \beta_q x_q^2 + \gamma_q x_q + \delta_q \\
x_{02}^q + x_{11}^q + x_{20}^q &= 3\alpha_q x_q^2 x_{q+1} + \beta_q (2x_q x_{q+1} + x_q^2) + \gamma_q (x_{q+1} + 2x_q) + 3\delta_q \\
x_{13}^q + x_{22}^q + x_{31}^q &= 3\alpha_q x_q x_{q+1}^2 + \beta_q (2x_q x_{q+1} + x_{q+1}^2) + \gamma_q (x_q + 2x_{q+1}) + 3\delta_q \\
x_{33}^q &= \alpha_q x_{q+1}^3 + \beta_q x_{q+1}^2 + \gamma_q x_{q+1} + \delta_q \\
X^q &\geq 0.
\end{aligned} \tag{100}$$

Now that we have a found sufficient condition for a polynomial to be positive we can state the whole problem. Note that if we choose to represent the $4m$ decision variables as a $4m$ dimensional vector y , we can find appropriate vectors $c, f_i, g_k^q \in \mathfrak{R}^{4m}$, and matrices $Q \in \mathfrak{R}^{4m \times 4m}$ and $H_k^q \in \mathfrak{R}^{4 \times 4}$ such that the model can be represented as

$$\begin{aligned}
\min_{y, X^1, \dots, X^m} \quad & c^T y + \frac{1}{2} y^T Q y \\
\text{subject to} \quad & f_i^T y = b_i \quad \forall i = 1, \dots, 4m \\
& H_k^q \bullet X^q = 0 \quad \forall k = 1, 2 \quad q = 1, \dots, m \\
& (g_k^q)^T y + H_k^q \bullet X^q = 0 \quad \forall k = 3, 4, 5, 6 \quad q = 1, \dots, m \\
& X^q \geq 0 \quad \forall q = 1, \dots, m.
\end{aligned} \tag{101}$$

Since the objective function is quadratic and convex, it can be turned into a linear matrix inequality constraint such that the problem becomes an SDP.

3.2 Option Pricing

In this section we start out by briefly discussing the famous result of [27] and [7] as well as its shortcomings. Next we review the work of [16] as an alternative framework for finding tight upper and lower bounds for the price of options based on SDPs. Additionally, an outlook on other applications of similar sort is provided at the end of the second part of this chapter.

3.2.1 Merton-Black-Scholes Model (MBSM)

The MBSM in its basic form is a closed formula solution to the problem of finding the price of a European call or put. What follows is an abbreviated version of the very comprehensive derivation of the formula given in [42]. The basic idea is to find the

replicating portfolio, comprised of a stock and the risk free asset that yields the same payoff as the option that is to be priced, and to price the option accordingly, given that the market is arbitrage free. The concept has been presented above (chapter 3.1.1), where the generic case of an arbitrary asset in a one period economy was considered. The idea similarly extends into the multi period case. When pricing an option which pays off only in the last period, the replicating portfolio has to be adjusted in every period such that its payoff mimics the payoff of the option in the last period, but is zero in every other period. Let's further specify the model in a way where at every point in time $t \in \mathcal{T} = \{1, \dots, T\}$ the economy is in one of two states $s \in \mathcal{S} = \{up, down\}$ and there are only two assets, one of which earns the risk free rate r_f in every period, while the other's value at time t is $S_t = S_{t-1} * u$ if at that point in time the economy is in the up state and $S_t = S_{t-1} * d$ otherwise. Coefficients u and d are set such that $u > 1$ and $0 < d < 1$. Then at any point in time the risk neutral probability of the "up"-state is given by

$$\phi = \frac{1 + r_f - d}{u - d}. \quad (102)$$

For a call option which is exercised at $t = T$ at which point it pays off $(S_T - K, 0)_+$ (K being its strike price), we can find its price by calculating the expected value of its replicating portfolio consisting of the riskless and the risky asset as

$$C_0 = \frac{1}{(1 + r_f)^T} \sum_{k=0}^T \binom{T}{k} \phi^k (1 - \phi)^{T-k} (S_0 u^k d^{T-k} - K, 0)_+ \quad (103)$$

where k is the number of states where the payoff of the option is non zero. Note that for $K = 0$ (103) gives a valuation of the risky asset as the discounted expected payoff in period T under the martingale measure which we can denote as

$$S_0 = e^{-r_f^* \tau} E[\tilde{S}_T] \quad \text{where} \quad \tilde{S}_T = S_0 u^{\tilde{k}} d^{T-\tilde{k}} \quad (104)$$

where we now introduce continuous compounded discounting ($r_f^* = \ln(1 + r_f)$ and $\tau = T - t$) which allows us to identify the continuous rate of the risky assets price development as

$$\begin{aligned} \tilde{r}_s &= \ln \left(\frac{\tilde{S}_T}{S_0} \right) = \tilde{k} \ln(u) + T \ln(d) - \tilde{k} \ln(d) \\ &= T \ln(d) + \ln \left(\frac{u}{d} \right) \tilde{k}. \end{aligned} \quad (105)$$

Note that $\tilde{k}$ is a binomially distributed random variable which converges to a normally distributed one as $T \rightarrow \infty$. Further, since $\tilde{r}_s$ is a linear combination of $\tilde{k}$, it also becomes normally distributed, where the mean is given by $\tau\mu_s$ and the variance given by $\tau\sigma_s^2$. Here μ_s and σ_s denote the one period moments of $\tilde{r}_s$. We can thus state

$$\tilde{S}_T = S_0 e^{\tilde{r}_s} \quad \text{thus} \quad \mathbb{E}[\tilde{S}_T] = S_0 e^{\mu_s \tau + \frac{1}{2} \sigma_s^2 \tau}. \quad (106)$$

Since in the arbitrage free framework every asset is expected to earn the risk free rate, we μ_s has to fulfill

$$r_f^* \tau = \mu_s^* \tau + \frac{1}{2} \sigma_s^2 \tau, \quad (107)$$

and the respective mean and the resulting rate of return of the risky asset shall be marked with an asterisk. If $f(r_s^*)$ denotes the density of the normally distributed return, then the price of a call is given by

$$C_0 = e^{-r_f^* \tau} \int_{-\infty}^{\infty} f(r_s^*) (S_0 e^{r_s^*} - K, 0)_+ dr_s^*. \quad (108)$$

The maximum function at the end is only greater than zero if

$$\begin{aligned} S_0 e^{r_s^*} - K &\geq 0 \\ r_s^* &\geq \ln\left(\frac{K}{S_0}\right) \end{aligned} \quad (109)$$

such that (108) can be written as

$$C_0 = S_0 e^{-r_f^* \tau} \int_{\ln\left(\frac{K}{S_0}\right)}^{\infty} e^{r_s^*} f(r_s^*) dr_s^* - K e^{-r_f^* \tau} \int_{\ln\left(\frac{K}{S_0}\right)}^{\infty} f(r_s^*) dr_s^* \quad (110)$$

Solving the integral (a difficult task which is spared here) finally yields the famous formula

$$C_0 = S_0 e^{\mu_s^* \tau + \frac{1}{2} \sigma_s^2 \tau - r_f^* \tau} N\left(\frac{\ln\left(\frac{S_0}{K}\right) + \mu_y^* \tau + \sigma_s^2 \tau}{\sigma_s \tau}\right) - K e^{-r_f^* \tau} N\left(\frac{\ln\left(\frac{S_0}{K}\right) + \mu_y^* \tau}{\sigma_s \tau}\right) \quad (111)$$

where $N(\cdot)$ is the c.d.f of the standard normal distribution.

3.2.2 Tight Bounds for Option Prices via an SDP-Approach

As any model, the MBSM rests on assumptions, some of which might be too restrictive, in order for it to accurately describe reality. One of the most restrictive assumptions of the model is that the underlying stock prices follow geometric Brownian Motion. Of course this hampers the reliability of the model, leading to what is called a pricing bias. As a consequence many alternative approaches were pursued, some of which rely on mathematical programming (see for example [35] [43] [33]). We will focus here on the work of [20] (the original version dates back to 1999) which was later expanded upon in [16]. Their approach aims to find tight upper and lower bounds on option prices, given only the first n moments of the RNPM. We start out by stating the problem in its basic form. Let $T \in \mathcal{N}$ be the exercise date and $K \in \mathfrak{R}$ the strike price of an option. The price of the underlying at any point in time $t \in \mathcal{N}$ is given by $S_t \in \mathfrak{R}$. Further r_f is the risk free rate and we use continuous compounding for discounting. Finally the RNPM is denoted by $f(S_T)$. Hence the price of a call at time t is given by

$$C_t = e^{-r(T-t)} \int_0^\infty (S_T - K)_+ f(S_T) dS_T. \quad (112)$$

The idea now is to find the density function $f(S_T)$ with known moments $q_j \in \mathfrak{R}$, $j = 1, \dots, n$ which imply the highest and the lowest option price, respectively, i.e.

$$\begin{aligned} \max_{f(S_T)} \quad & e^{-r(T-t)} \int_0^\infty (S_T - K)_+ f(S_T) dS_T \\ \text{subject to} \quad & e^{-jr(T-t)} \int_0^\infty S_T^j f(S_T) dS_T = q_j \quad \forall j = 0, 1, \dots, n \\ & f(S_T) \geq 0 \end{aligned} \quad (113)$$

and

$$\begin{aligned} \min_{f(S_T)} \quad & e^{-r(T-t)} \int_0^\infty (S_T - K)_+ f(S_T) dS_T \\ \text{subject to} \quad & e^{-jr(T-t)} \int_0^\infty S_T^j f(S_T) dS_T = q_j \quad \forall j = 0, 1, \dots, n \\ & f(S_T) \geq 0. \end{aligned} \quad (114)$$

Both of these problems can be viewed as linear programs with an infinite number of variables $f(S_T)$. Therefore their duals may be obtained via the usual procedure. We will contract notation in that we define $x = e^{r(T-t)} S_T$, which allows us to write the

duals as

$$\begin{aligned} & \min_y \sum_{i=0}^n q_i y_i \\ & \text{subject to } \sum_{i=0}^n x^i y_i \geq (x - e^{r(T-t)}K)_+ \quad \forall x \in \mathfrak{R}_+ \end{aligned} \quad (115)$$

and

$$\begin{aligned} & \max_y \sum_{i=0}^n q_i y_i \\ & \text{subject to } \sum_{i=0}^n x^i y_i \leq (x - e^{r(T-t)}K)_+ \quad \forall x \in \mathfrak{R}_+. \end{aligned} \quad (116)$$

In what follows we will use the results from [20] as discussed in chapter 2.5.7. These propositions allow us to formulate two problems that are equivalent to the upper bound problem. The first is given by

$$\begin{aligned} & \min \sum_{i=0}^n y_i q_i \\ & \text{subject to } \sum_{i,j:i+j=2l-1} x_{ij} = 0 \quad \forall l = 1, \dots, n \\ & \quad y_l - \sum_{i,j:i+j=2l} x_{ij} = 0 \quad \forall l = 0, \dots, n \\ & \quad \sum_{i,j:i+j=2l-1} z_{ij} = 0 \quad \forall l = 1, \dots, n \\ & \quad y_0 - z_{00} = -e^{r(T-t)}K \\ & \quad y_1 - \sum_{i,j:i+j=2} z_{ij} = 1 \\ & \quad y_l - \sum_{i,j:i+j=2l} z_{ij} = 0 \quad \forall l = 2, \dots, n \\ & \quad Z, X \geq 0. \end{aligned} \quad (117)$$

To see that these constraints are equivalent to those in (115), observe that the feasible

region described in that problem can be represented as

$$\begin{aligned} \sum_{r=0}^n y_r x^r &\geq 0, \quad \forall x \geq 0 \\ (y_0 + e^{-r(T-t)}K) + (y_1 - 1)x + \sum_{r=2}^n y_r x^r &\geq 0, \quad \forall x \geq 0. \end{aligned} \quad (118)$$

Following proposition 2.5.7 b), these constraints can be replaced by those stated above. Another way to describe the feasible set in (115) is given by

$$\begin{aligned} \sum_{r=0}^n y_r x^r &\geq 0, \quad \forall x \in [0, e^{-r(T-t)}K) \\ (y_0 + e^{-r(T-t)}K) + (y_1 - 1)x + \sum_{r=2}^n y_r x^r &\geq 0, \quad \forall x \in [e^{-r(T-t)}K, \infty). \end{aligned} \quad (119)$$

According to Proposition 2.5.7 c) and d), these can also be represented by semidefinite constraints such that another version of the problem is given by

$$\begin{aligned} \min \quad & \sum_{h=0}^n q_h y_h \\ \text{subject to} \quad & \sum_{i,j:i+j=2l-1} x_{ij} = 0 \quad \forall l = 1, \dots, n \\ & \sum_{i,j:i+j=2l} x_{ij} = \sum_{h=0}^l y_h \binom{n-h}{l-h} (e^{-r(T-t)}K)^h \quad \forall l = 0, 1, \dots, n \\ & \sum_{i,j:i+j=2l-1} z_{ij} = 0 \quad \forall l = 1, \dots, n \\ & z_{00} = \sum_{h=0}^n (e^{-r(T-t)}K)^h y_h \\ & \sum_{i,j:i+j=2} z_{ij} = -1 + \sum_{h=1}^n h (e^{-r(T-t)}K)^{(h-1)} y_h \\ & \sum_{i,j:i+j=2l} z_{ij} = \sum_{h=l}^n \binom{h}{l} (e^{-r(T-t)}K)^{(h-l)} y_h \quad \forall l = 2, \dots, n \\ & X, Z \geq 0. \end{aligned} \quad (120)$$

Solving any of these problems yields the optimal bound, yet there might be computational differences due to the formulations. However, the source literature

states nothing about which of the two formulations is advantageous. We proceed in deriving a similar reformulation for the lower bound problem (117). Again, we can observe that the respective feasible region is given by

$$\begin{aligned} -\sum_{r=0}^n y_r x^r &\geq 0, \quad \forall x \in [0, e^{-r(T-t)}K] \\ -(y_0 + e^{-r(T-t)}K) - (y_1 - 1)x - \sum_{r=2}^n y_r x^r &\geq 0, \quad \forall x \in [e^{-r(T-t)}K, \infty). \end{aligned} \quad (121)$$

Using Proposition 2.5.7 c) and d), we can thus formulate the lower bound problem as

$$\begin{aligned} &\max \sum_{h=0}^n q_h y_h \\ &\text{subject to} \quad \sum_{i,j:i+j=2l-1} x_{ij} = 0 \quad \forall l = 1, \dots, n \\ &\quad \sum_{h=0}^l y_h \binom{n-h}{l-h} (e^{-r(T-t)}K)^h + \sum_{i,j:i+j=2l} x_{ij} = 0 \quad \forall l = 0, \dots, n \\ &\quad \sum_{i,j:i+j=2l-1} z_{ij} = 0 \quad \forall l = 1, \dots, n \\ &\quad \sum_{h=0}^n (e^{-r(T-t)}K)^h y_h + z_{00} = 0 \\ &\quad \sum_{h=1}^n h(e^{-r(T-t)}K)^{(h-1)} y_h + \sum_{i,j:i+j=2} z_{ij} = 1 \\ &\quad \sum_{h=l}^n \binom{h}{l} (e^{-r(T-t)}K)^{(h-l)} y_h + \sum_{i,j:i+j=2l} z_{ij} = 0 \quad \forall l = 2, \dots, n \\ &\quad X, Z \geq 0. \end{aligned} \quad (122)$$

All of the problems given above are SDPs, and might be solved by standard algorithms. The authors of [16] argue that their approach could potentially also be used to price more complicated exotic options. But despite some research, we have not found any publication that deals with this particular problem of pricing exotic options using the framework presented above. This might be due to the fact that other authors have proposed more advanced techniques to address this issue. [28] for example have developed SDP-Relaxations for pricing a class of exotic options such as Asian and barrier options modeling the underlying assets price by a Brownian motion and a number of mean reverting processes. However, [21] generalized the

ideas above in order to formulate SDP-Relaxations for the Multi-Asset Options (i.e. options where there is more than one underlying asset).

4 Sparse Principal Component Analysis: SDP-Relaxations and Application

Assume one has a data set in matrix form, where the rows represent observations and the columns represent different variables. In case the number of variables is very extensive, one might want to combine those that are highly correlated into artificial variables, which still account for most of the variance in the data while being easily interpreted. This is achieved via Principal Component Analysis (PCA), which was further developed into Sparse Principal Component Analysis (SPCA), which will be the centre of attention in this section. Specifically we will review the results by [12] and [10]. Still there will be a brief discussion on the former, again for the sake of the self containment of this text. The main source for this discussion within this paper is [52].

4.1 Standard Principal Component Analysis

Again, now more formally, assume a data matrix $X \in \mathbb{R}^{n \times k}$ that is comprised of n observations for which k variables were recorded. The seminal idea is to find vectors $a_j \in \mathbb{R}^k$ such that the vectors $z_j = Xa_j \in \mathbb{R}^n$ give a new matrix $Z = (z_1, z_2, \dots, z_m)$ with $m < k$. This has to happen in way, such that Z can be seen as a new data matrix with fewer variables that might be more easily interpreted, and that conveys a similar amount of information as did X . In other words, we want to project X onto z_j , such that as much of the variance as possible is preserved. z_j is accordingly called a *Principal Component* of X . Let's assume for now that every value in X has been standardized and centered and that X_j are ordered by decreasing variances. Next we want to chose a_j such that

$$\text{Var}(z_j) = \frac{1}{n} \sum_{i=1}^n z_{ji}^2 = \frac{1}{n} \|z_j\|^2 = \frac{1}{n} a_j^T X^T X a_j \quad (123)$$

are maximized. Note that $\Sigma = \frac{1}{n} X^T X$ is the empirical Covariance matrix of X . The maximization of course can be simply achieved by making $|a_j|$ arbitrarily big. This option has to be ruled out with an appropriate constraint. The problem can thus be

stated as

$$\begin{aligned} & \max a_j^T \Sigma a_j \\ & \text{subject to } \|a_j\|^2 = 1. \end{aligned} \tag{124}$$

i.e. to find the vector of unit length that maximizes the variance of the projection on the component. By stating the Lagrangian function $L(x, \lambda_j) = a_j^T \Sigma a_j - \lambda_j(a_j^T a_j - 1)$, taking the derivative and setting it to zero, one finds the optimality condition to be

$$\Sigma a_j = \lambda_j a_j. \tag{125}$$

Thus a_j are eigenvectors of Σ , and λ_j are the respective eigenvalues. Under the conditions introduced so far, λ_j is the biggest eigenvalue, and we further express this fact by giving it as well as the respective eigenvector and the resulting principal component the index $j = 1$. So far we have found that the optimal coefficients for generating z_1 from X are given by the elements of a_1 , which are also called loadings. Yet one can capture the remaining variance in a second principal component z_2 which is obtained in a similar procedure, with the addition of an extra constraint in the optimization program.

$$\max a_j^T \Sigma a_j \tag{126}$$

$$\text{subject to } \|a_j\|^2 = 1 \tag{127}$$

$$a_1^T a_j = 0 \tag{128}$$

The orthogonality condition provides that the correlation between the two principal components is zero, which implies that they cover different portions of the overall variance in the data. Again, we find a_2 to be the eigenvector of Σ which has the second largest eigenvalue λ_2 . Applying this procedure consecutively yields all eigenvectors of Σ in order of their eigenvalues, where their respective principal components capture decreasing amounts of the variance in the data. Of course, all principal components together result in $\sum_{j=1}^k \text{Var}(z_j) = \sum_{j=1}^k \text{Var}(x_j)$, but if the leading principal components have a significant contribution, one can ignore the other and thus gain a reduction of complexity. Nicely, the size of the eigenvalue gives an indication of the importance of the respective principal component. To see this,

observe that

$$\text{Var}(z_j) = a_j^T \Sigma a_j = \lambda_j \quad (129)$$

thus

$$\frac{\lambda_j}{\lambda_1 + \dots + \lambda_k} \quad (130)$$

gives a measure of the percentage of variance, captured by the j -th principal component, and one simply has to choose $m < k$ such that

$$\frac{\sum_{j=1}^m \lambda_j}{\sum_{j=1}^k \lambda_j} \quad (131)$$

becomes reasonably big, which might vary depending on the specific application. We skip a further discussion on the topic of determining the appropriate number of components and turn to another issue that, as we will see, is a domain where SDPs come in as very helpful.

4.2 Sparse Principal Component Analysis: Two SDP-Relaxation Approaches

In general, factor loadings (i.e. the entries of a_j) will be nonzero, making the principal components linear combinations of all original variables, which might be an undesirable effect, since it hampers interpretability of the principal components. This is the very problem tackled in the so called Sparse Principal Component Analysis. The aim is simply to find sparse factor loadings, while still keeping the goal of maximizing the variance of the principal component. Obviously it is not enough to simply find the ordered eigenvalues along with the associated eigenvectors, since a choice has to be made as to which original variable should be excluded from a certain component. Rather, the task becomes a combinatorial problem, which can be shown as NP-hard. [12] and [10] however have shown that there are powerful relaxations of the problem, which we will revisit in the following.

4.2.1 SDP-Relaxation based on the cardinality constrained Eigenvalue Problem

In the following we discuss the results due to [12]. The problem of SPCA can be stated as

$$\begin{aligned} & \max_{a_j} \quad a_j \Sigma a_j - \rho \text{Card}(a_j) \\ & \text{subject to} \quad \|a_j\| \leq 1 \end{aligned} \quad (132)$$

where $\rho > 0$ is a parameter that gives a penalty for an increase of $\text{Card}(a_j)$ which is a function that returns the cardinality of a vector, i.e. the number of its nonzero elements. This problem is related to

$$\begin{aligned} & \max_{a_j} \quad a_j \Sigma a_j \\ & \text{subject to} \quad \|a_j\| = 1 \\ & \quad \text{Card}(a_j) \leq k \end{aligned} \quad (133)$$

which is the Cardinality constrained maximum eigenvalue problem. We find that a dual of (133) is given by

$$\inf_{\rho \in P} \sup_{\|a_j\| \leq 1} (a_j^T \Sigma a_j - \rho \text{Card}(a_j) + \rho k) \quad (134)$$

which by weak duality, gives an upper bound on its optimal value. This in turn implies that an optimizer a^* for (132) is also globally optimal for (133) with $k = \text{Card}(a^*)$, given that P is the set of all values for ρ , which leads to sparse loadings of cardinality k . Thus we turn our attention to (133) and find the semidefinite relaxation to be

$$\begin{aligned} & \max \Sigma \bullet X \\ & \text{subject to} \quad \text{Tr}(X) = 1 \\ & \quad \text{Card}(X) \leq k^2 \\ & \quad \text{Rank}(X) = 1 \\ & \quad X \succeq 0. \end{aligned} \quad (135)$$

This is achieved by noting that $a_j^T \Sigma a_j = \text{Tr}(\Sigma a_j a_j^T) = \Sigma \bullet X$ and $\|a_j\| = \sqrt{a_j^T a_j} = \sqrt{\text{Tr}(X)}$. To insure that X is of the form $X = a_j a_j^T$ it has to be *p.s.d.* and of rank one. Further $\text{Card}(xx^T) = \text{Card}(x)^2$. (135) is still a hard problem due to the constraints on the rank and the cardinality, which are non-convex. To resolve this problem, note that

$x^t \mathbf{1} = x^t \mathbf{1}_k$ if $\mathbf{1}$ is a vector of ones and $\mathbf{1}_k$ is suitable vector of k ones and $n - k$ zeros. Obviously, $\text{Card}(x) = \text{Card}(\mathbf{1}_k)$ which we set equal to k for convenience. The Cauchy-Schwarz Inequality then implies that $x^t \mathbf{1} = x^t \mathbf{1}_k \leq \sqrt{x^T x} \sqrt{\mathbf{1}_k^T \mathbf{1}_k} = \sqrt{x^T x} \sqrt{k}$. Thus if $\text{Card}(x) = k$, $X = xx^t$ and $\text{Tr}(X) = 1$ we can show that it is implied that

$$\begin{aligned} x^t \mathbf{1} &\leq \sqrt{x^T x} \sqrt{k} = \sqrt{\text{Tr}(X)} \sqrt{k} = \sqrt{k} \\ (x^t \mathbf{1})(x^t \mathbf{1}) &\leq k \\ \mathbf{1}^T X \mathbf{1} &\leq k. \end{aligned} \tag{136}$$

Thus we can substitute the cardinality constraint with the weaker but linear constraint obtained above. The rank constraint is simply dropped to finally formulate the relaxed problem

$$\begin{aligned} \max \quad & \Sigma \bullet X \\ \text{subject to} \quad & \text{Tr}(X) = 1 \\ & \mathbf{1}^T X \mathbf{1} \leq k \\ & X \geq 0 \end{aligned} \tag{137}$$

which is an SDP. In the original paper by [12], an even stronger constraint is introduced by turning the third constraint into $\mathbf{1}^T |X| \mathbf{1} \leq k$, where $|X|$ denotes the matrix whose entries are the absolute values of X . The weaker version was chosen for the purposes of this text, since it yields an SDP in standard form, which is handled easily with the algorithms presented in previous chapters. The stronger version requires a more specific algorithm, which the interested reader may study from the primary source. In any case, it is obvious that the optimizer will not always be of rank one. To obtain an approximate solution, one has to truncate the optimizer and keep only its dominant eigenvector (i.e. whose corresponding eigenvalue is the largest in magnitude). [12] argue that their approach usually yields solutions X^* that have a rank near one, and that the respective dominant eigenvector a_j^* is indeed sparse. In order to derive further sparse loadings, one has to deflate Σ as to obtain

$$\Sigma_{j+1} = \Sigma_j - (a_j^T \Sigma_j a_j) a_j a_j^T, \tag{138}$$

thus removing the respective eigenvector and eigenvalue pair from the current iterate. The so generated matrix is the new iterate for further computation. The natural question again is when to stop the decomposition process. Of course, as in the standard case there might be qualitative reasoning depending on the application

in question, but unlike the standard case there is no natural barrier in the form of $\text{Rank}(\Sigma_j)$ becoming zero eventually. There is no guarantee that the decomposition procedure above will ever stop finding eigenvectors whose corresponding eigenvalues are unequal to zero. [12] found a meaningful interpretation of their procedure that may guide a reasonable stopping criterion. One can transform (137) into a model where the relaxed cardinality constrained is enforced via a penalty

$$\begin{aligned} \max \quad & \Sigma \bullet X - \rho \mathbf{1}^T X \mathbf{1} \\ \text{subject to} \quad & \text{Tr}(X) = 1 \\ & X \succeq 0 \end{aligned} \tag{139}$$

where again ρ is a parameter controlling sparsity. This problem can be rewritten as

$$\begin{aligned} \max_{\Sigma} \min_U \quad & \text{Tr}(X(\Sigma + U)) \\ \text{subject to} \quad & \text{Tr}(X) = 1 \\ & |U_{ij}| \leq \rho \\ & X \succeq 0 \end{aligned} \tag{140}$$

which is equivalent to

$$\begin{aligned} \min_U \quad & \lambda^{\max}(\Sigma + U) \\ & |U_{ij}| \leq \rho. \end{aligned} \tag{141}$$

This can be interpreted as a worst-case maximum eigenvalue computation of the matrix Σ of which each component is subject to noise bounded by ρ . Note that once problem (137) has been solved, ρ can easily be obtained from the Lagrange multiplier associated with the relaxed cardinality constraint. This interpretation however gives rise to the following stopping criterion for the decomposition process. If the absolute values of all coefficients at an iterate Σ_j are smaller than the noise level ρ (obtained from the last problem calculated), then the iterate is indistinguishable from the zero matrix and the process has to come to an end.

4.2.2 SDP-Relaxation based on a convex Restriction

Another approach is attributable to [10]. They again start with (132), and show that the optimal solution is $a_j = 0$ if $\rho \geq \Sigma_{11}$. To see this, recall that $a_j^T \Sigma a_j \leq \Sigma_{11} (\sum_{i=1}^n |(a_j)_i|)^2$

and $(\sum_{i=1}^n |(a_j)_i|)^2 \leq \|a_j\|^2 \mathbf{Card}(a_j)$ which implies that

$$\max_{\|a_j\| \leq 1} a_j^T \Sigma a_j - \rho \mathbf{Card}(a_j) \leq (\Sigma - \rho) \mathbf{Card}(a_j) \leq 0. \quad (142)$$

Thus we turn to the opposite case. Here the inequality $a_j \geq 1$ is tight and one can represent the sparsity pattern of a_j by a vector $u \in \{0, 1\}^n$. (132) can thus be cast as

$$\begin{aligned} & \max_{\|a_j\| \leq 1} a_j^T \Sigma a_j - \rho \mathbf{Card}(a_j) \\ &= \max_{u \in \{0,1\}^n} \lambda_{\max}(\mathbf{Diag}(u) \Sigma \mathbf{Diag}(u)) - \rho \mathbf{1}^T u \\ &= \max_{u \in \{0,1\}^n} \lambda_{\max}(\mathbf{Diag}(u) S^T S \mathbf{Diag}(u)) - \rho \mathbf{1}^T u \\ &= \max_{u \in \{0,1\}^n} \lambda_{\max}(S \mathbf{Diag}(u) S^T) - \rho \mathbf{1}^T u \\ &= \max_{\|x\|=1} \max_{u \in \{0,1\}^n} x^T S \mathbf{Diag}(u) S^T x - \rho \mathbf{1}^T u \\ &= \max_{\|x\|=1} \max_{u \in \{0,1\}^n} \sum_{i=1}^n u_i ((s_i^T x) - \rho) \\ &= \max_{\|x\|=1} \sum_{i=1}^n ((s_i^T x) - \rho)_+. \end{aligned} \quad (143)$$

Here we used the facts that $\mathbf{Diag}(u)^2 = \mathbf{Diag}(u)$ for $u \in \{0, 1\}^n$, that $\lambda_{\max}(B^T B) = \lambda_{\max}(B B^T)$ and that $\max_{v \in \{0,1\}} \beta v = \beta_+$ where $\beta_+ = \max(\beta, 0)$. By the same argument as before we can restate this above problem in terms of $X = x x^T$, i.e.

$$\begin{aligned} & \max \sum_{i=1}^n ((s_i^T X s_i) - \rho)_+ \\ & \text{subject to } \text{Tr}(X) = 1 \\ & \mathbf{Rank}(X) = 1 \\ & X \geq 0. \end{aligned} \quad (144)$$

Note that since both the objective function and the elliptope $\Delta_n = \{X \geq 0 : \text{Tr}(X) = 1\}$ are convex, the solution must be an extreme point of Δ_n , which are of rank one anyway. Thus the rank constraint can be dropped. The problem is still hard since the convex function is maximized. [10] however show that on the rank one elements of Δ_n , the objective function equals a concave function of X which gives rise to semidefinite relaxation.

Theorem 4.2.2.: Let $S \in \mathfrak{R}^{n \times n}$, $\rho \geq 0$ and denote by $s_1, \dots, s_n \in \mathfrak{R}^n$ the columns of A ,

then an upper bound on (144) can be computed by solving

$$\begin{aligned} \max \quad & \sum_{i=1}^n ((X^{\frac{1}{2}} B_i X^{\frac{1}{2}}))_+ \\ \text{subject to} \quad & \text{Tr}(X) = 1 \\ & X \geq 0 \end{aligned} \tag{145}$$

in the variables $X \in S_+^n$, where $B_i = s_i s_i^t - \rho \mathbf{I}$, or also

$$\begin{aligned} \max \quad & \sum_{i=1}^n \text{Tr}(P_i B_i) \\ \text{subject to} \quad & \text{Tr}(X) = 1 \\ & X \geq P_i \\ & P_i \geq 0 \end{aligned} \tag{146}$$

which is an SDP in the variables $X, P_i \in S_+^n$.

The proof can be found in [10]. The optimal value of (146) is of course always greater or equal to its counterpart for the un-relaxed problem. But if the optimizer X is of rank one, then the optimal values of both are the same.

4.3 Applications of SPCA in Finance

PCA is useful for analysts who are trying to find out how a plethora of underlying fundamentals drives asset prices. An example for such an analysis is given by [8]. The author shows that local and global variables that explain returns in emerging stock markets are summarized by GDP, inflation, money and interest rates (local factors) as well as by world industrial production and world inflation (global factors). These categories were obtained using PCA on a wide variety of variables. SPCA could potentially facilitate these kinds of analysis, allowing for even more variables that could be summarized in a meaningful way. Another example for such an endeavor is given in [11], where the historical daily returns of S&P stocks ($N = 472$) over a five year period were investigated. The authors showed that if one wishes to represent 90% of the overall variance with the first sparse principal component, one needs almost 300 out of 472 stocks to do so. When they inspected the first principal components for increasing cardinality, they found that initially only stocks from "Financials" are incorporated, while stocks from other sectors start to appear only later at cardinality 32. They thus set the cardinality of the first principal component to 32 and proceeded to the second one, using the same guideline, which eventually

leads them to a cardinality of 26 for the second principal component. They observe that it mainly incorporates stocks from the "Energy" and the "Materials" sector. Further they find that there is a positive correlation between the first and the second principal component. Void of any theoretical background, it is hard to give an interpretation of such results. But nonetheless it illustrates the usefulness of SPCA for data analysis, especially considering that none of the phenomena described above were observed when standard PCA was applied.

Another way to utilize the SPCA in a finance context is portfolio hedging. Given a portfolio of assets with known covariance matrix, instead of hedging across all components of the portfolio, one could hedge only with respect to the most important factors. If this is done using SPCA, one could greatly reduce the number of assets necessary for the hedge, thus minimizing transaction costs.

5 Portfolio selection under data certainty

5.1 Classical Portfolio Theory

The Problem of selecting financial securities from a given set of existing titles was first famously explored in [19] (a very detailed introduction and discussion of the theory can be found in [50]). The decisive measure the theory relies on is the one-period-return of a security which is modeled as a random variable. The random vector of returns is jointly normal and given by $r \sim N_n(\mu, \Sigma)$, where first and second moments are assumed to be known. Σ in this regard is also referred to as the risk model. The variance of the return of a title or a portfolio is the only measure of riskiness in this framework. By combining securities into a portfolio, one can diversify this risk due to the covariance structure of the risk model. Portfolios are selected according to the $\mu - \sigma$ -criteria, where portfolios with high expected returns and low risk dominate portfolios with low expected returns and high risk. The decision variable then is the vector of weights $x \in \mathbb{R}^n$ of the n securities in the optimized portfolio. Investors wish either to minimize the portfolio risk $\sigma_p = x^T \Sigma x$ while targeting a certain expected portfolio return $\mu_p = x^T \mu$, or to maximize that return for a given target risk. The respective optimization problems are given by

$$\begin{aligned} \min_x \frac{1}{2} x^T \Sigma x & \quad \max_x x^T \mu \\ \text{subject to } x^T \mu = \mu_{\text{target}} & \quad \text{and} \quad \text{subject to } x^T \Sigma x = \sigma_{\text{target}}^2 \\ x^T \mathbf{1} = 1 & \quad x^T \mathbf{1} = 1 \end{aligned} \tag{147}$$

where $\mathbf{1} \in \mathbb{R}^n$ is a vector of ones. Since both are quadratic or quadratically constrained, respectively, they could in principle be cast as SDPs, while in practice there would be no benefit for doing so.

The model was later expanded upon in [22], where a risk-less asset was introduced. It yields the risk free return r_f , and can be combined with any risky portfolio, one of which has it that its combinations with r_f dominate combinations of r_f with any other risky portfolio. Thus the choice of the optimal bundle of assets is separable into identifying this dominant risky portfolio, and afterwards choosing a combination of that risky portfolio with r_f . The respective optimization problem can be cast as

$$\begin{aligned} \min_x \frac{1}{2} x^T \Sigma x & \qquad \max x^T \mu + (1 - x^T \mathbf{1}) r_f \\ \text{subject to } x^T \mu + (1 - x^T \mathbf{1}) r_f = \mu_{\text{target}} & \quad \text{and} \quad \text{subject to } x^T \Sigma x = \sigma_{\text{target}}^2 \end{aligned} \quad (148)$$

where $(1 - x^T \mathbf{1})$ gives the share of the risk free asset in the portfolio. Obviously, any of the cases discussed so far give instances of quadratic programming, since either the objective function is quadratic while the constraints are linear or vice versa. Thus one could in principle apply SDP-tools in order to solve these problems, but this is hardly of practical significance. However, classical portfolio theory gives the foundation for more advanced methods that will be of greater interest for the purpose of this text.

5.2 Tracking error

In practice, portfolio managers are often confronted with the task of "beating a benchmark". This might be a stock index, or another funds portfolio, whose return should be outperformed by the return of the actively managed portfolio. That is to maximize

$$r^T x_{BM} - r^T x = r^T (x_{BM} - x) \quad (149)$$

where $x_{BM} \in \mathbb{R}^n$ is the vector of weights of the securities in the benchmark portfolio, while $x \in \mathbb{R}^n$ is the vector of weights in the portfolio that is to be optimized. In order to assess whether the achieved excess return is in line with the risk undertaken, it is common use to measure the Tracking Error Volatility (TEV), i.e. the volatility of excess returns,

$$TE := \sqrt{(x - x_{BM})^T \Sigma (x - x_{BM})}, \quad (150)$$

and to demand that a certain level of TE not be exceeded. Thus the overall problem can be stated as

$$\begin{aligned} & \max_x \mu^T(x_{BM} - x) \\ & \text{subject to } (x - x_{BM})^T \Sigma (x - x_{BM}) \leq (TE)^2. \end{aligned} \quad (151)$$

The inequality is quadratic, and by $\Sigma \geq 0$ also convex, thus it can be stated in terms of linear matrix inequality. To this end, note that $\exists R : \Sigma = R^T R$ thus

$$\begin{aligned} & (x - x_{BM})^T R^T R (x - x_{BM}) \leq (TE)^2 \\ & 0 \leq (TE) - (x - x_{BM})^T R^T (TE)^{-1} R (x - x_{BM}) \\ & \begin{pmatrix} (TE)I & R(x - x_{BM}) \\ (x - x_{BM})^T R^T & TE \end{pmatrix} \geq 0. \end{aligned}$$

The problem was first considered by [39] and further analyzed by [34], where the author points out that a portfolio strategy that is purely driven by the emphasis of beating the benchmark is likely to result in the creation of portfolios that have systematically high risk. This is especially the case if the manager to whom the task is delegated is exclusively rewarded for outperforming the benchmark. In light of classical portfolio theory, the target even seems nonsensical, since there is no guarantee that the benchmark portfolio is efficient (with respect to the $\mu - \sigma$ -criteria), or if it is even possible to create a tracking portfolio that is itself efficient. However, due to its relevance in practice, it is a model worth mentioning.

5.3 SDP-Relaxation for Portfolio Selection with Higher Order Moments

In this chapter we review the results by [55]. They seek to overcome one of the shortcomings of the classical portfolio selection framework in that higher moments of the returns distribution are incorporated in the model. Hence, besides mean and variance, skewness and the kurtosis are taken into account. However, transaction costs are not accounted for, and shortselling is forbidden. The vector of returns is given by $R = (r_1, \dots, r_n)^T$ and there is a risk-less asset that yields return r_f . The vector of the respective means is denoted $\mu = (\mu_1, \dots, \mu_n)^T$. The covariance, coskewness

and cokurtosis matrices are given in that order by

$$\begin{aligned} M_2 &= \mathbb{E}[(R - \mu)(R - \mu)^T] \\ M_3 &= \mathbb{E}[(R - \mu)(R - \mu)^T \otimes (R - \mu)^T] \\ M_4 &= \mathbb{E}[(R - \mu)(R - \mu)^T \otimes (R - \mu)^T \otimes (R - \mu)^T] \end{aligned} \quad (152)$$

where $\otimes$ denotes the Kronecker product, i.e.

$$A \otimes B = \begin{pmatrix} A_{11}B & \dots & A_{1n}B \\ \dots & \ddots & \dots \\ A_{m1}B & \dots & A_{mn}B \end{pmatrix}. \quad (153)$$

Hence, $M_2 \in \mathfrak{R}^{n \times n}$, $M_3 \in \mathfrak{R}^{n \times n^2}$ and $M_4 \in \mathfrak{R}^{n \times n^3}$. Given vector weights $x = (x_1, \dots, x_n)^T$ in a portfolio, the expected return, variance, skewness and kurtosis are calculated from

$$\begin{aligned} R_p &= x^T \mu + (1 - x^T \mathbf{1})r_f \\ V_p &= x^T M_2 x \\ S_p &= x^T M_3 (x \otimes x) \\ K_p &= x^T M_4 (x \otimes x \otimes x). \end{aligned} \quad (154)$$

The model considered in [55] is a constrained kurtosis minimization model:

$$\begin{aligned} \min x^T M_4 (x \otimes x \otimes x) \\ x^T M_2 x &\leq \bar{V}_p \\ x^T M_3 (x \otimes x) &\geq \bar{S}_p \\ x^T \mu + (1 - x^T \mathbf{1})r_f &\geq \bar{R}_p \\ x_i &\geq 0 \quad \forall i = 1, \dots, n \end{aligned} \quad (155)$$

where $\bar{R}_p, \bar{V}_p, \bar{S}_p$ are critical levels for the expected return, variance and skewness that are acceptable to decision makers. The problem stated above is hard due to its objective nature and being that the second constraint is non convex. In order to

derive the SPD-Relaxation, additional notation has to be introduced. Let

$$p(x) = x^T M_4(x \otimes x \otimes x) = \sum_{i=1}^n \sum_{j=1}^n \sum_{m=1}^n \sum_{l=1}^n x_i x_j x_m x_l k_{ijml} = \sum_{\alpha} p_{\alpha} x^{\alpha} \quad (156)$$

where

$$x^{\alpha} = x_1^{\alpha_1} x_2^{\alpha_2} \dots x_n^{\alpha_n} \quad \text{with} \quad \sum_{i=1}^n \alpha_i \leq 4$$

and p_{α} be the respective coefficient. The vector of $(1, x_1, x_2, \dots, x_1^2, x_1 x_2, \dots, x_1 x_n, \dots, x_n^2, x_1^4, \dots, x_n^4)$ is called the basis of $p(x)$. In order to be identifiable, α is a n -digit number where the i -th digit equals α_i . Further, let

$$\begin{aligned} K &= \{x \in \mathfrak{R}^n | x^T M_3(x \otimes x) = \bar{S}_P, x^T M_2 x = \bar{V}_P, \\ &\quad x^T \mu + (1 - x^T \mathbf{1}) r_f = \bar{R}_P, x_i \geq 0 \quad \forall i = 1, \dots, n\}. \end{aligned} \quad (157)$$

Now one can show that the following proposition holds (see [55] for a proof)

Proposition 5.3.1: If $p_K^* = \min_{x \in K} p(x)$ then $p_K^* = \min_{f \in P(K)} \int_K p(x) df$ where f is an arbitrary measure over the semi-algebraic set K , and $P(K)$ is the finite space of Borel probability measures over K .

The objective of (155) can thus be converted in $\min_{f \in P(K)} \int_K p(x) df$. We can then reformulate as follows:

$$\int_K p(x) df = \int_K \sum_{\alpha} p_{\alpha} x^{\alpha} df = \sum_{\alpha} p_{\alpha} \int_K x^{\alpha} df = \sum_{\alpha} p_{\alpha} y_{\alpha} \quad (158)$$

where $\{y_{\alpha}\}$ is the sequence of moments associated with $f \in P(K)$. Let Γ be the set of sequences of moments associated with all probability measures over K . Then the objective of (155) can be cast as

$$\min_{\{y_{\alpha}\}_{\alpha \in \Gamma}} \sum_{\alpha} p_{\alpha} y_{\alpha} \quad (159)$$

which achieves a linearization of the original objective. We now turn to further characterization of $\{y_{\alpha}\}$ in order to obtain the SDP-Relaxation of (155). We define the

moment matrix as

$$M(y) = \begin{pmatrix} y_{000\dots 0} & y_{100\dots 0} & y_{010\dots 0} & \dots & y_a & \dots \\ y_{100\dots 0} & y_{200\dots 0} & y_{110\dots 0} & \dots & \dots & \dots \\ y_{010\dots 0} & y_{110\dots 0} & y_{020\dots 0} & \dots & \dots & \dots \\ \dots & \dots & \dots & \dots & \dots & \dots \\ y_b & \dots & \dots & \dots & y_{a+b} & \dots \\ \dots & \dots & \dots & \dots & \dots & \dots \end{pmatrix}, \quad (160)$$

the truncated moment matrix $M_t(y)$ as the leading principal sub-matrix of dimension $s(t)$, where t denotes the degree of the objective function, and $s(t)$ as the dimension of its basis. We further define $M_t(hy)$ as the matrix whose components $M_t(hy)_{ij} = \sum_{\alpha} h_{\alpha} y_{\{\beta(i,j)+\alpha\}}$ where $\beta(i, j)$ is the subscript of the component in $M_t(y)$ in the i -th row and the j -th column, while h_{α} is the coefficient of the corresponding component in the basis of the function $h(x)$ which is a polynomial function. We then can state the following proposition:

Proposition 5.3.2: If $h(x)$ is a polynomial of degree $2d$ or $2d - 1$ and $H = \{x|h(x)\} \geq 0$ then $M_t(y) \geq 0$ and M_{t-d} .

We can now state the relaxed problem as

$$\begin{aligned} & \min p^T y \\ & \text{subject to } M_t(y) \geq 0 \\ & \quad M_{t-d_j}(h_j y) \geq 0 \end{aligned} \quad (161)$$

where $t \geq \max(d_0, d_1, \dots, d_6)$, $d_0 = \frac{\deg[p(x)]}{2}$, $d_j = \frac{\deg[h_j]}{2}$, for $j = 1, \dots, 6$.

5.4 The Knapsack Problem: SDP-Relaxations and Applications in Finance

The classical Knapsack problem deals with the task of assigning n items to be packed into a knapsack that has capacity restrictions, in a way that maximizes the utility gained from the items packed. The knapsack has restriction with respect to more than just one dimension (space, weight ect.), such that multiple constraints have to be considered. In finance on the other hand, one often has to choose which investments out of a set of positive NPV (Net Present Value) projects should be undertaken, while maintaining a budget constraint or constraints on the number of investments in a

certain industrial sector. In either case the problem can be formulated as

$$\begin{aligned}
& \max_x m^T x \\
& \text{subject to } w_j^T x \leq d_j \quad j = 1, \dots, J \\
& x \in \{0, 1\}^n
\end{aligned} \tag{162}$$

where $m, w_j \in \mathfrak{R}_+^n$, $j = 1, \dots, J$ are vectors of added value (i.e. NPV in the finance context) and resource consumption respectively and d_j is the amount of the j -th resource available. The binary decision variables $x_i \in \{0, 1\}$ encode whether an item is packed ($x_i = 1$) or unpacked ($x_i = 0$). Sometimes there might be interactions between two items with respect to the utility they deliver. For example, the utility of packing a screw increases, if the screwdriver is packed as well. In a finance context, their might be synergies or cannibalization effects between projects that effect their NPV if both are realized. In this case the problem becomes what is called the quadratic knapsack problem

$$\begin{aligned}
& \max_x x^T M x \\
& \text{subject to } w_j^T x \leq d_j \quad j = 1, \dots, J \\
& x \in \{0, 1\}^n
\end{aligned} \tag{163}$$

where the diagonal elements of the matrix $M \in \mathcal{S}_n$ capture the marginal utility of the items while the off-diagonal elements m_{ij} give half of the interaction effect on utility of items j and i are packed simultaneously. The quadratic Knapsack problem was first introduced by [6] and is known to be NP-hard. A recent survey of the most important ways of solving this problems is given in [9]. In what follows we will briefly introduce the SDP based relaxations due to [38] (as presented in [9]).

We start out by introducing variables y_{ij} which are defined as the products $x_i x_j$ for $i \leq j$, i.e. the product xx^T is replaced by as symmetric matrix $Y \in \mathcal{S}^n$. The diagonal of the this matrix is denoted as y . In order for a solution for Y to satisfy $Y = xx^T$ with $x \in \{0, 1\}^n$ it must fulfill $Y \geq 0$ and $\text{rank}(Y) = 1$ and $Y_{ii} \in \{0, 1\}$. An SDP formulation of (163) is given by

$$\begin{aligned}
& \max_Y Y \bullet M \\
& \text{subject to } \text{Diag}(w_j) \bullet Y \leq d_j \\
& Y \geq 0 \\
& \text{Rank}(Y) = 1 \\
& Y_{ii} \in \{0, 1\}
\end{aligned} \tag{164}$$

Dropping the rank constraint and relaxing the last one to $0 \leq Y_{ii} \leq 1$, the problem becomes an SDP, which unfortunately yields a very poor upper bound for the original problem (see [38]). However, the bound can be tightened by the following strategy. First observe, that if $Y \geq 0$, $\mathbf{Rank}(Y) = 1$ and $Y_{ii} \in \{0, 1\}$, then $Y - yy^T \geq 0$. This can be seen by observing that $(x + v)(x + v)^T \geq 0$ for any $v \in \mathfrak{R}^n$. Thus

$$xx^T + vx^T + xv^T + vv^T \geq 0 \quad (165)$$

Using the fact that $\mathbf{diag}(xx^T) = x$ when $x \in \{0, 1\}^n$ yields

$$\begin{aligned} Y + vy^T + yv^T + vv^T &\geq 0 \\ Y + (v + y)(v + y)^T - yy^T &\geq 0 \end{aligned} \quad (166)$$

we can now choose $v = -y$ to obtain the statement. Further, using Schur's complement we can reformulate $Y - yy^T \geq 0$ as

$$\begin{pmatrix} 1 & y^T \\ y & Y \end{pmatrix} \geq 0 \quad (167)$$

which binds the diagonal elements y_i between zero and one, which is not the case for the off-diagonal elements. This gives rise to an SDP-relaxation given by

$$\begin{aligned} &\max_Y M \bullet Y \\ &\text{subject to } \mathbf{Diag}(w_j) \bullet Y \leq d_j \\ &\quad \begin{pmatrix} 1 & y^T \\ y & Y \end{pmatrix} \geq 0 \\ &\quad y = \mathbf{diag}(Y) \\ &\quad Y_{ij} \geq 0 \quad \forall i \neq j \end{aligned} \quad (168)$$

The model can be further enhanced if we reformulate the constraints by squaring

both sides in order to obtain $w_j^T x x^T w_j \leq d_j^2$. Again setting $Y = x x^T$ we obtain

$$\begin{aligned}
& \max_Y M \bullet Y \\
& \text{subject to } w_j w_j^T \bullet Y \leq d_j^2 \\
& \begin{pmatrix} 1 & y^T \\ y & Y \end{pmatrix} \succeq 0 \\
& y = \mathbf{diag}(Y) \\
& Y_{ij} \geq 0 \quad \forall i \neq j
\end{aligned} \tag{169}$$

An alternative relaxation is achieved by multiplying both sides of a constraint by the respective right hand side thus obtaining

$$\begin{aligned}
(w_j^T x)^2 & \leq d_j (w_j^T x) \\
0 & \leq w_j^T x (d_j - x^T w_j) \\
& = w_j^T \begin{pmatrix} 1 & x^T \\ 0 & -w_j \end{pmatrix} \begin{pmatrix} d_j \\ -w_j \end{pmatrix} \\
& = \begin{pmatrix} 0 & w_j^T \end{pmatrix} \begin{pmatrix} 1 \\ x \end{pmatrix} \begin{pmatrix} 1 & x^T \\ 0 & -w_j \end{pmatrix} \begin{pmatrix} d_j \\ -w_j \end{pmatrix}
\end{aligned} \tag{170}$$

Setting $Y' = \begin{pmatrix} 1 \\ x \end{pmatrix} \begin{pmatrix} 1 & x^T \end{pmatrix} = x' x'^T$ the last expression can be written as

$$\begin{pmatrix} d_j \\ -w_j \end{pmatrix} \begin{pmatrix} 0 & w_j^T \end{pmatrix} \bullet Y' \tag{171}$$

thus obtaining the relaxed problem in the form of

$$\begin{aligned}
& \max_Y M \bullet Y \\
& \text{subject to } \begin{pmatrix} d_j \\ -w_j \end{pmatrix} \begin{pmatrix} 0 & w_j^T \end{pmatrix} \bullet Y' \geq 0 \\
& \begin{pmatrix} 1 & y^T \\ y & Y \end{pmatrix} \succeq 0 \\
& y = \mathbf{diag}(Y) \\
& Y_{ij} \geq 0 \quad \forall i \neq j \\
& Y' = \begin{pmatrix} 1 \\ x \end{pmatrix} \begin{pmatrix} 1 & x^T \end{pmatrix}
\end{aligned} \tag{172}$$

$$Y' = \begin{pmatrix} 1 \\ x \end{pmatrix} \begin{pmatrix} 1 & x^T \end{pmatrix} \tag{173}$$

A final tightening of the relaxation is obtained by multiplying each capacity constraint with each of the variables x_i as to get $x_i w_j^T x \leq x_i c$, which can be written as

$$\sum_{k=1}^n (w_j)_k x_k x_i \leq x_i c \quad (174)$$

or by substituting Y_{ik} for $x_i x_k$ and Y_{ii} for x_i as

$$\sum_{k=1}^n (w_j)_k Y_{ik} \leq Y_{ii} c \quad (175)$$

The final formulation is thus given by

$$\begin{aligned} & \max_Y M \bullet Y \\ & \text{subject to } \sum_{k=1}^n (w_j)_k Y_{ik} - Y_{ii} d_j \leq 0 \\ & \begin{pmatrix} 1 & y^T \\ y & Y \end{pmatrix} \succeq 0 \\ & y = \mathbf{diag}(Y) \\ & Y_{ij} \geq 0 \quad \forall i \neq j \end{aligned} \quad (176)$$

Each of the relaxations above yields an upper bound to the original problem, but with varying quality. The problems are presented here in order of decreasing upper bound, such that the last formulation is to be preferred (for a proof see [38] or [9]).

6 Portfolio Selection with uncertain data

The parameters in any application of optimization are often estimates based on past events or on genuine models of the future. As such, they are subject to error and uncertainty. It is therefor preferable to find solutions to a given problem that are not perfect in a certain case, but rather perform well in any, or at least a lot of cases. In the following we discuss some instances of uncertainty, how they are modeled, and the relevance of such models in finance.

6.1 Modeling Uncertainty using Ellipsoidal Sets

One of the most popular ways to account for uncertainty in the input data is to assume that the set of possible values of a given datum is given by ellipsoidal sets,

i.e. sets that are described using ellipsoids, which are themselves simply given as subsets of second order cones. The later is easily handled in convex optimization as discussed above, making these types of models particularly useful. In the following, we will discuss how to model ellipsoidal uncertainty in the case of linear and quadratic constraints. All theory presented below is due to [1], but the notation and presentation was mainly adapted from [37] since it is more compact than the one from the original source. Further, after reading [1], one might notice that there is an entire chapter dealing with uncertainty in the case of SDP-constraints, which raises the obvious question of why this was not included in this text. The main conclusion that can be drawn from the aforementioned chapter is that such problems are intractable, with only a few exceptions for which an approximation can be found. For all we know, no finance relevant models of that type have been thus far developed.

6.1.1 Ellipsoidal uncertainty for linear constraints

We now consider the case where the constraint coefficients of a linear constraint are subject to ellipsoidal uncertainty, i.e.

$$\begin{aligned} & \min c^T x \\ & \text{subject to } a^T x + b \geq 0 \quad \forall (a, b) \in \mathcal{U} \\ & \text{where} \\ & \mathcal{U} = \{(a^0, b^0) + \sum_{j=1}^k u_j (a^j, b^j) : \|u\| \leq 1\}. \end{aligned} \tag{177}$$

Note that $a^j, a, x \in \mathbb{R}^n$, $b^j, b \in \mathbb{R}$ for $j = 0, \dots, k$ and $u \in \mathbb{R}^k$, so that $(a^j, b^j), (a, b) \in \mathbb{R}^{n+1}$. In the context of the minimization problem, this constraint will force the optimal solution to be feasible even in the worst case scenario that is captured by the restriction (which is not necessarily the worst case scenario in the model from which the set is inferred). Thus for any given x , the robust constraint is fulfilled if

$$0 \leq \min_{(a,b) \in \mathcal{U}} [a^T x + b] = \min_{u: \|u\| \leq 1} [\alpha + u^T \beta] \tag{178}$$

where $\alpha = (a^0)^T x + b^0$ and $\beta = (\beta_1, \dots, \beta_k)$ with $\beta_j = (a^j)^T x + b^j$. Note that here α is constant such that the optimal value of the second optimization problem in (178) is obtained by minimizing $u^T \beta$ constrained by $\|u\| \leq 1$. From trigonometry we have that $u^T \beta = \|u\| \|\beta\| \cos(\theta)$ which is minimized when $\|u\| = 1$ and $\cos(\theta) = -1$ in which case u points in the opposite direction of β . Thus $u^* = \frac{-\beta}{\|\beta\|}$, i.e. the normalized vector,

points in the opposite direction of β . We can therefore substitute

$$0 \leq \min_{(a,b) \in \mathcal{U}} a^T x + b = \alpha - \|\beta\| = a^0 x + b^0 - \sqrt{\sum_{j=1}^k ((a^j)^T x + b^j)^2} \quad (179)$$

and finally formulate the robust version of the linear constraint as SOC one:

$$0 \leq a^0 x + b^0 - \sqrt{\sum_{j=1}^k ((a^j)^T x + b^j)^2} = a^0 x + b^0 - \|Ax + b\| \quad (180)$$

or equivalently

$$\begin{pmatrix} ((a^0)^T x + b^0)I & Ax + b \\ (Ax + b)^T & (a^0)^T x + b^0 \end{pmatrix} \succeq 0 \quad (181)$$

where A is the matrix comprised of all a^j and b is the vector comprised of all b^j , where $j = 1, \dots, k$.

6.1.2 Ellipsoidal uncertainty for quadratic constraints

In a similar fashion as above, one can address convex quadratic constraints, where the coefficients as well as the coefficient matrices are subject to uncertainty described by an ellipsoidal set, albeit of slightly different sort. Consider the quadratically constrained minimization problem

$$\begin{aligned} & \min d^T x \\ & \text{subject to } -x^T(A^T A)x + 2b^T x + c \geq 0 \quad \forall (A, b, c) \in \mathcal{U} \end{aligned} \quad (182)$$

where

$$\mathcal{U} = \{(A^0, b^0, c^0) + \sum_{j=1}^k u_j(A^j, b^j, c^j) : \|u\| \leq 1\}$$

where $A, A^j \in \mathbb{R}^{m \times n}$, $b, b^j, x \in \mathbb{R}^n$ and $c, c^j \in \mathbb{R}$ for $j = 0, \dots, k$. Thus we want to choose x such that

$$(A, b, c) \in \mathcal{U} \Rightarrow -x^T(A^T A)x + 2b^T x + c \geq 0. \quad (183)$$

Substituting the definition of (A, b, c) that is in $\mathcal{U}$ we find that this is equivalent to

$$\begin{aligned} \|u\| \leq 1 \Rightarrow & -x^T \left(A^0 + \sum_{j=1}^k A^j u_j \right) \left(A^0 + \sum_{j=1}^k A^j u_j \right)^T x \\ & + 2 \left(b^0 + \sum_{j=1}^k b^j u_j \right)^T x + \left(c^0 + \sum_{j=1}^k c^j u_j \right) \geq 0. \end{aligned} \quad (184)$$

Now we define

$$\begin{aligned} A(x) &= (A^1 x, A^2 x, \dots, A^k x), \\ b(x) &= (x^T b^1, x^T b^2, \dots, x^T b^k)^T + \frac{1}{2}(c^1, c^2, \dots, c^k) - A(x)^T A^0 x \\ &\quad \text{and} \\ c(x) &= c^0 + 2(b^0)^T x - x^T (A^0)^T A^0 x. \end{aligned}$$

Further we rewrite $\|u\| \leq 1$ as $-u^T I u + 1 \geq 0$. We can then reformulate (184) as

$$-u^T I u + 1 \geq 0 \Rightarrow -u^T A(x)^T A(x) u + 2b(x)^T u + c(x) \geq 0. \quad (185)$$

Thus we want to find a condition for the unit ball to be engulfed by the ellipsoid described by the second part of the constraint. Such a condition is given by the S-lemma:

$$\begin{aligned} & \exists \tau \geq 0 \\ \text{subject to } & \begin{pmatrix} c(x) - \tau & b(x)^T \\ b(x) & A(x)^T A(x) - \tau I \end{pmatrix} \geq 0 \end{aligned} \quad (186)$$

The last entry of the above matrix is non linear. However, the condition is equivalent to the following:

$$\begin{aligned} & \tau \geq 0 \\ & \begin{pmatrix} c'(x) - \tau & b'(x)^T & (A^0 x)^T \\ b'(x) & \tau I & A(x)^T \\ A^0 x & A(x) & I \end{pmatrix} \geq 0 \end{aligned} \quad (187)$$

where $b'(x) = (x^T b^1, x^T b^2, \dots, x^T b^k)^T + \frac{1}{2}(c^1, c^2, \dots, c^k)$ and $c'(x) = c^0 + 2(b^0)^T x$. With this formulation we have that $A(x)$, $b'(x)$ and $c'(x)$ are linear in x , so that (187) is a linear matrix inequality.

6.1.3 A Robust Multi-Period Portfolio Selection Model

In this section, we discuss a model that uses ellipsoidal uncertainty sets in order to illustrate the theory discussed above. The model is attributed to [5], and will be presented in the notation of [37]. However, the model will be presented in a different way than it is presented in the original source, in order to make the connection to the above theory more apparent. We consider an existing portfolio where the shares of asset $i = 0, \dots, n$ are given by $x_i^0 \in \mathfrak{X}$ and the corresponding vector of shares by $x^0 \in \mathfrak{X}^n$. By convention, asset 0 is cash. In the next L investment periods the composition of the portfolio is to be readjusted, pursuing the goal of maximizing total wealth at the end of period L . In each period $l = 1, \dots, L$ a decision is made on b_i^l , i.e. how many shares of asset i should be sold, and on the number s_i^l of shares of asset i that should be sold such that x_i^l shares of asset i are in the portfolio in period l . In any period the number of shares is thus given by

$$x_i^l = x_i^{l-1} - s_i^l + b_i^l \quad \forall i, l. \quad (188)$$

In period l , asset i is traded at price P_i^l . Cash does not earn interest, so that its price is $P_0^l = 1 \quad \forall l$, which for real world applications can always be achieved by a simple change of numeraire. We further assume proportional transaction costs for buying denoted by α_i^l , and for selling, β_i^l . Cash holdings in each period are thus given:

$$x_0^l \leq x_0^{l-1} + \sum_{i=1}^n (1 - \alpha_i) P_i^l s_i^l - \sum_{i=1}^n (1 + \beta_i) P_i^l b_i^l \quad \forall l \quad (189)$$

The inequality is due to technical reasons. Adding the objective function to the picture, the deterministic model is given by

$$\begin{aligned} \max \quad & \sum_{i=0}^n P_i^L x_i^L \\ x_0^l \leq & x_0^{l-1} + \sum_{i=1}^n (1 - \alpha_i) P_i^l s_i^l - \sum_{i=1}^n (1 + \beta_i) P_i^l b_i^l \quad \forall l \\ x_i^l = & x_i^{l-1} - s_i^l + b_i^l \quad \forall i, l \\ s_i^l \geq & 0 \quad \forall i, l \\ b_i^l \geq & 0 \quad \forall i, l \\ x_i^l \geq & 0 \quad \forall i, l \\ x_0^l \geq & 0 \end{aligned} \quad (190)$$

which is easily solved as it is an LP. Note that the non-negativity constraints which forbid short selling and borrowing imposed in [5] are entirely optional. The obvious restriction of the model is due to the fact that asset prices cannot be known in advance, which motivates the formulation of a model that is robust with respect to that very uncertainty. In the previous two chapters we only considered uncertainty in constraints, while the model requires us to take care of the objective function as well. This problem is easily dealt with by simply moving the objective function into a constraint such that we have

$$\begin{aligned} & \max t \\ & t \leq \sum_{i=0}^n P_i^L x_i^L \end{aligned} \quad (191)$$

in place of the original objective. In order to account for the uncertainty of future prices we introduce ellipsoidal uncertainty sets

$$\mathcal{U}^l = \{\mu^l + (\Sigma^l)^{\frac{1}{2}} u : \|u\| \leq 3\} \quad (192)$$

where μ^l is the vector of expected prices in period l , and Σ^l the respective covariance matrix. Thus (192) gives the set of prices, which are no further than three standard deviations away from their mean. Assuming that prices are normally distributed, we can guarantee that constraints relying on these uncertainty sets will be fulfilled in 99,9% of cases. Applying this uncertainty set to the constraint in (191) we have

$$(P^L)^T x^L - t \geq 0 \quad \forall P^L \in \mathcal{U}^L \quad (193)$$

which is of similar form as in (177), the differences being that there is no constant and no uncertainty in the coefficient of t and that u is restricted to the unit ball times three. We thus can use a similar strategy for reformulation in order to obtain

$$t \leq (\mu^L)^T x^L - 3 \sqrt{(x^L)^T V^L x^L} \quad (194)$$

as the robust version of the respective constraint. Now we use the same procedure on the second inequality of the deterministic problem to see that

$$x_0^l \leq x_0^{l-1} + \sum_{i=1}^n (1 - \alpha_i) P_i^l s_i^l - \sum_{i=1}^n (1 + \beta_i) P_i^l b_i^l \quad \forall P^l \in \mathcal{U}^l \quad \forall l \quad (195)$$

is equivalent to

$$x_0^l - x_0^{l-1} \geq (\mu^l)^T (D_{\alpha'}^l - D_{\beta}^l) \begin{pmatrix} s^l \\ b^l \end{pmatrix} - 3 \sqrt{\begin{pmatrix} s^l \\ b^l \end{pmatrix}^T \begin{pmatrix} D_{\alpha}^l \\ -D_{\beta}^l \end{pmatrix} \Sigma^l (D_{\alpha'}^l - D_{\beta}^l) \begin{pmatrix} s^l \\ b^l \end{pmatrix}} \quad (196)$$

where D_{α}^l and D_{β}^l are diagonal matrices whose i -th diagonal entry is given by $(1 - \alpha_i)$ and $(1 + \beta_i)$, respectively. Since both of these new constraints are conic, the robust problem can be cast as an SDP.

6.2 Distributional Robust Modelling

The example in the previous section illustrates an approach where robustness rests on the idea that the uncertainty in the data is subject to a certain known distribution. This is viable if there is a lot of information, e.g. from past events, that provides valid estimates or supports one or the other hypothesis about the true distribution. In many situations however, the distribution itself is subject to uncertainty, either because information is scarce or not valid for future events, as is the case if the true distribution changes over time. Thus it is desirable to have a model that is not only robust in the face of uncertainty by obeying a single specified distribution but also governed by a distribution which is itself uncertain. This problem was first addressed by [17], who explored the problem of purchasing an amount of an item given future uncertain demand. At that time (1958) the approach was not very fruitful due to computational difficulties. More recently, and substantially due to the work of [51] and [30], the topic has become more popular as it turned out that SDPs provide an efficient way for solving many instances of distributionally robust optimization. In the following the framework presented in [30] will be summarized, as it is very comprehensive and rich in its possible applications. We then turn to the models presented in [2] and [51]. Both deal with the knapsack problem and the portfolio selection problem, respectively. We will illustrate the power of the general framework discussed in the following by applying it to these problems.

6.2.1 A General Framework

The framework presented in [30] deals with tractable representations of robust counterparts of deterministic constraints. Let $z \in \mathfrak{R}^S$ be a vector of uncertain problem parameters and $x \in \mathfrak{R}^N$ a vector of decision variables that have to be set before the realization of z . Further $C \subseteq \mathfrak{R}^S$ is the uncertainty set and $v : \mathfrak{R}^N \times \mathfrak{R}^P \rightarrow \mathfrak{R}$ a given constraint function with a prescribed threshold $w \in \mathfrak{R}$. A robust counterpart of a

constraint is then given by

$$v(x, z) \leq w \quad \forall z \in C \quad (197)$$

or equivalently by

$$\sup_{z \in C} v(x, z) \leq w.$$

Such a constraint provides that the solution for x will not violate the threshold in any case of z . This is the very idea that guided the approach in the previous chapter. But we can address the problem of uncertainty also from another angle, specifically *Stochastic Programming*, where distributional information is incorporated via expectation constraints, i.e.

$$\mathbb{E}_{Q^0}[v(x, \tilde{z})] \leq w \quad (198)$$

where $\tilde{z} \in \mathfrak{R}^S$ is now a random vector. The expectation is taken with respect to the distribution Q^0 which is the true distribution of $\tilde{z}$. The final step is now to account for uncertainty in that distribution, again using the robustness approach and thus obtaining the *distributionally robust counterpart* given by

$$\mathbb{E}_P[v(x, \tilde{z})] \leq w \quad \forall P \in \mathcal{P} \quad (199)$$

or equivalently

$$\sup_{P \in \mathcal{P}} \mathbb{E}_P[v(x, \tilde{z})] \leq w$$

where Q^0 is known to be in the so called *ambiguity set* $\mathcal{P}$ which is a set of joint distributions of $\tilde{z}$ that are generally denoted as P . The characterization of this set is critical, and we will further assume that it is in any case representable in the standard form

$$\mathcal{P} = \left\{ P \in \mathcal{P}_0(\mathfrak{R}^S \times \mathfrak{R}^T) : \begin{array}{l} \mathbb{E}_P[A\tilde{z} + B\tilde{u}] = b, \\ P(\tilde{z}, \tilde{u} \in C_i) \in [\underline{p}_i, \bar{p}_i] \quad \forall i \in \mathcal{I} \end{array} \right\}. \quad (200)$$

Here $\tilde{u} \in \mathfrak{R}^S$ is an auxiliary random vector, which might be helpful in structuring the marginal distribution of $\tilde{z}$. Further $A \in \mathfrak{R}^{K \times S}$, $B \in \mathfrak{R}^{K \times T}$ and $b \in \mathfrak{R}^K$. The confidence

sets C_i , with $i \in \mathcal{I} = 1, \dots, I$ are defined as

$$C_i = \{(z, u) \in \mathbb{R}^S \times \mathbb{R}^T : C_i z + D_i u \preceq_{\mathcal{K}_i} c_i\} \quad (201)$$

where $C_i \in \mathbb{R}^{L_i \times Q}$, $D_i \in \mathbb{R}^{L_i \times Q}$, $c_i \in \mathbb{R}^{L_i}$, and $\mathcal{K}_i$ are proper cones (i.e. closed, convex, pointed and nonempty cones, with the respective dual denoted as $\mathcal{K}_i^*$). In the fashion of the Löwner partial order we note $x \preceq_{\mathcal{K}} y$ whenever $y - x \in \mathcal{K}$. The Löwner Partial order is a special case where $\mathcal{K} = \mathcal{S}_+$. K and T are allowed to be zero in which case the expectation condition or $\tilde{u}$ vanishes. Also let $\underline{p}_i, \bar{p}_i \in [0, 1]$ and $\underline{p}_i \leq \bar{p}_i \forall i \in \mathcal{I}$. In order for later theorems to hold it is necessary that $\mathcal{P}$ satisfies the following two regularity conditions:

(C1) The last confidence set C_I is bounded and has a probability of one, i.e. $\underline{p}_I = \bar{p}_I = 1$

(C2). There is a $P \in \mathcal{P}$ such that $P(\tilde{z}, \tilde{u} \in C_i) \in (\underline{p}_i, \bar{p}_i)$, whenever $\underline{p}_i < \bar{p}_i$, $i \in \mathcal{I}$.

The first condition ensures that the confidence set with the highest index contains the support of the joint random vector $(\tilde{z}, \tilde{u})$. The second requires the existence of a distribution in the ambiguity set that satisfies the probability bounds with strict inequality whenever $[\underline{p}_i, \bar{p}_i]$ is non-degenerate. Further the following requirement on the constraint function v has to be met

(C3). The constraints function $v(x, z)$ can be written as

$$v(x, z) = \max_{l \in \mathcal{L}} v_l(x, z) \quad (202)$$

where $\mathcal{L} = 1, \dots, L$ and the functions $v_l : \mathbb{R}^N \times \mathbb{R}^P \rightarrow \mathbb{R}$ are of the form

$$v_l(x, z) = s_l(z)^T x + t_l(z) \quad (203)$$

with $s_l(z) = S_l z + s_l$, $S_l \in \mathbb{R}^{N \times S}$ and $s_l \in \mathbb{R}^N$, $t_l(z) = \mathbf{t}_l^T z + t_l$, $\mathbf{t}_l \in \mathbb{R}^S$ and $t_l \in \mathbb{R}$.

This renders v a convex, piecewise bilinear function in (x, z) . However, (C3) is relatively restrictive but in can be suspended if the following two conditions hold.

(C3a) The constraint function $v(x, z)$ is convex in x for all z and can be evaluated in polynomial time.

(C3b) For $i \in \mathcal{I}$, $x \in \chi$ (which is an arbitrary set of feasible solutions) and $\theta \in \mathbb{R}$,

it can be decided in polynomial time whether

$$\max_{(z,u) \in C_i} v(x, z) \leq \theta \quad (204)$$

Moreover, if (204) does not hold, then a separating hyperplane $(\pi, \phi) \in \mathcal{R}^N \times \mathcal{R}$ can be constructed in polynomial time such that $\pi^T x > \phi$ and $\pi^T \hat{x} \leq \phi$ for all $\hat{x} \in \chi$ satisfying (204).

(C3) implies (C3a) as well as (C3b). The former conditions ensures that the distributionally robust expectation constraint (199) has a convex feasible set. The latter enables us to efficiently separate the semi-infinite constraints, which we will see, arise from the dual formulation of (199). Finally the confidence sets have to fulfill the following nesting condition

(N) For all $i, i' \in \mathcal{I}, i \neq i'$ we have either $C_i \subseteq \text{int}(C_{i'})$, $C_{i'} \subseteq \text{int}(C_i)$ or $C_i \cap C_{i'} = \emptyset$

Under the conditions (C1)-(C3) and (N) we can formulate equivalent deterministic counterparts of (199). Since it is a nice exercise in duality theory and also illustrates the necessity of the conditions, the derivation of those counterparts will be provided in the following.

We start by observing that if (C3) is satisfied, then the constraint

$$v(x, z) + f^T z + g^T u \leq h \quad \forall (z, u) \in C_i \quad (205)$$

is satisfied if and only if there is $\phi_l \in \mathcal{K}_i^*, l \in \mathcal{L}$, such that $c_i^T + \phi_l + s_l^T x + t_l \leq h$, $C_i^T \phi_l = S_l^T x + \mathbf{t}_l + f$ and $D_i^T \phi_l = g$ for all $l \in \mathcal{L}$. This can be seen by observing that (205) is equivalent to

$$(S_l^T x + \mathbf{t}_l + f)^T z + g^T u \leq h - s_l^T x - t_l \quad \forall l \in \mathcal{L}, \forall (z, u) \in \mathcal{R}^S \times \mathcal{R}^T : C_i z + D_i u \leq_{\mathcal{K}_i} c_i \quad (206)$$

which is the same as requiring that no optimal solution of the L strictly feasible problems

$$\begin{aligned} & \max (S_l^T x + \mathbf{t}_l + f)^T z + g^T u \\ & \text{subject to } C_i z + D_i u \leq_{\mathcal{K}_i} c_i \\ & \quad z \in \mathcal{R}^S \quad u \in \mathcal{R}^T \end{aligned} \quad (207)$$

is smaller than the respective values $h - s_l^T x - t_l$. This is the case only if the optimal values of the L dual problems

$$\begin{aligned} & \min c_i^T \phi_l \\ & \text{subject to } C_i^T \phi_l = S_l^T x + \mathbf{t}_l + f \\ & D_i^T \phi_l = g \\ & \phi_l \in \mathcal{K}_i^* \end{aligned} \tag{208}$$

do not exceed the respective values $h - s_l^T x - t_l$ as well. This implies the above statement. Now further observe that the left hand side of distributionally robust constraint

$$\sup_{P \in \mathcal{P}} E_P[v(x, \tilde{z})] \leq w$$

is equivalent to the optimal value of the following moment problem:

$$\begin{aligned} & \max \int_{C_I} v(x, z) d\mu(z, u) \\ & \text{subject to } \int_{C_I} [Az + Bu] d\mu(z, u) = b \\ & \int_{C_I} \mathbf{1}_{[(z, u) \in C_i]} d\mu(z, u) \geq \underline{p}_i \quad \forall i \in \mathcal{I} \\ & \int_{C_I} \mathbf{1}_{[(z, u) \in C_i]} d\mu(z, u) \leq \bar{p}_i \quad \forall i \in \mathcal{I} \\ & \mu \in \mathcal{M}_+(\mathbb{R}^S \times \mathbb{R}^T) \end{aligned} \tag{209}$$

of which the dual is given by

$$\begin{aligned} & \min b^T \beta + \sum_{i \in \mathcal{I}} [\bar{p}_i \kappa_i - \underline{p}_i \lambda_i] \\ & \text{subject to } [Az + Bu]^T \beta + \sum_{i \in \mathcal{I}} \mathbf{1}_{[(z, u) \in C_i]} [\kappa_i - \lambda_i] \geq v(x, z) \quad \forall (z, u) \in C_I \\ & \beta \in \mathbb{R}^K, \quad \kappa_i, \lambda_i \in \mathbb{R}_+ \quad \forall i \in \mathcal{I}. \end{aligned} \tag{210}$$

Due to (C2), strong duality holds (for a proof see [4]). Further, the nesting condition allows us to partition C_I into I nonempty and disjoint sets $\tilde{C}_i = C_i \setminus \bigcup_{i' \in \mathcal{D}(i)} C_{i'}$, where $\mathcal{D}(i)$ is the index set of strict subsets of C_i , i.e $\mathcal{D}(i) = \{i' \in \mathcal{I} : C_{i'} \subset \text{ri}(C_i)\}$, while the respective index set of all supersets is given by $\mathcal{A}(i) = \{i\} \cup \{i' \in \mathcal{I} : C_i \subset C_{i'}\}$. This

allows us to reformulate the duals constraint into a set of constraints given by

$$[Az + Bu]^T \beta + \sum_{i' \in \mathcal{A}(i)} [\kappa_{i'} - \lambda_{i'}] \geq v(x, z) \quad \forall (z, u) \in \bar{C}_i, \forall i \in \mathcal{I}$$

of which the i -th constraint is equivalently given by

$$\max_{(z, u) \in \bar{C}_i} \left\{ v(x, z) - [Az + Bu]^T \beta - \sum_{i' \in \mathcal{A}(i)} [\kappa_{i'} - \lambda_{i'}] \right\} \leq 0. \quad (211)$$

Due to the convexity of v , the expression that is maximized is also convex, implying that the solution to the maximization lies on the boundary of $\bar{C}_i$ which due to the nesting condition coincides with the boundary of C_i . The distributionally robust constraint is thus satisfied if and only if there exists a $\beta \in \mathfrak{R}^K$ and $\kappa_i, \lambda_i \in \mathfrak{R}, i = 1, \dots, I$ such that

$$b^T \beta + \sum_{i \in \mathcal{I}} [\bar{p}_i \kappa_i - \underline{p}_i \lambda_i] \leq w \quad (212)$$

$$[Az + Bu]^T \beta - \sum_{i' \in \mathcal{A}(i)} [\kappa_{i'} - \lambda_{i'}] \geq v(x, z) \quad \forall (z, u) \in C_i, \forall i \in \mathcal{I}.$$

Now we can use the initially obtained result with $f = -A^T \beta$, $g = -B^T \beta$ and $h = \sum_{i' \in \mathcal{A}(i)} [\kappa_{i'} - \lambda_{i'}]$ in order to finally obtain the following theorem:

Theorem 6.2.1a: Assume that conditions (C1)-(C3) and (N) hold, then the distributionally robust constraint (199) is satisfied for the ambiguity set (200) if and only if there are $\beta \in \mathfrak{R}^K$, $\kappa_i, \lambda_i \in \mathfrak{R}$ and $\phi_{il} \in \mathcal{K}_i^*, i \in \mathcal{I}$ and $l \in \mathcal{L}$ that satisfy the constraint system

$$\begin{aligned} b^T \beta + \sum_{i \in \mathcal{I}} [\bar{p}_i \kappa_i - \underline{p}_i \lambda_i] &\leq w \\ \left. \begin{aligned} c_i^T \phi_{il} + s_l^T x + t_l &\leq \sum_{i' \in \mathcal{A}(i)} [\kappa_{i'} - \lambda_{i'}] \\ C_i^T \phi_{il} + A^T \beta &= S_l^T x + \mathbf{t}_l \\ D_i^T \phi_{il} + B^T \beta &= 0 \end{aligned} \right\} \forall i \in \mathcal{I}, \forall l \in \mathcal{L}. \end{aligned} \quad (213)$$

If we choose $\mathcal{K}_i = S_+^{n_i} = (S_+^{n_i})^*$ (since it is self dual), then the above system gives the constraints of an SDP. Besides the above result, [30] derive a conservative approximation for the generic case of N being violated. However, as this case will not be relevant to the models provided later in the text, we will not discuss these approximations here and instead turn to a useful remark on the representability of the ambiguity set, which will allow us to incorporate information on moments into

the model.

Before proceeding in this direction, further notation has to be introduced. A $\mathcal{K}$ -epigraph of a function $f : \mathfrak{X}^M \rightarrow \mathfrak{X}^N$ and proper cone $\mathcal{K}$ given by the set $\{(x, y) \in \mathfrak{X}^M \times \mathfrak{X}^N : f(x) \leq_{\mathcal{K}} y\}$ is said to be *conic representable* if that very set can be expressed in terms of conic inequalities, allowing for cones other than $\mathcal{K}$ and auxiliary variables to be involved. Now the following "Lifting Theorem" can be stated as follows:

Theorem 6.2.1b: Let $f \in \mathfrak{X}^M$ and $g : \mathfrak{X}^S \rightarrow \mathfrak{X}^M$ be a function with a conic representable $\mathcal{K}$ -epigraph and consider the ambiguity set

$$\mathcal{P}' = \left\{ P \in \mathcal{P}_0(\mathfrak{X}^S) : \begin{array}{l} \mathbb{E}_P[g(\tilde{z})] \leq_{\mathcal{K}} f \\ P[\tilde{z} \in C_i] \in [\underline{p}_i, \bar{p}_i] \quad \forall i \in \mathcal{I} \end{array} \right\} \quad (214)$$

as well as the lifted ambiguity set

$$\mathcal{P} = \left\{ P \in \mathcal{P}_0(\mathfrak{X}^S \times \mathfrak{X}^M) : \begin{array}{l} \mathbb{E}_P[\tilde{u}] = f \\ P[g(\tilde{z}) \leq_{\mathcal{K}} \tilde{u}] = 1 \\ P[\tilde{z} \in C_i] \in [\underline{p}_i, \bar{p}_i] \quad \forall i \in \mathcal{I} \end{array} \right\} \quad (215)$$

which involves the auxiliary random vector $\tilde{u} \in \mathfrak{X}^M$. We then have that (i) $\mathcal{P}'$ is the set of marginal distributions of $\tilde{z}$ contained in $\mathcal{P}$, and (ii) $\mathcal{P}$ can be reformulated as an instance of the standardized ambiguity set.

Finally, one might observe that the distributionally robust constraint requires the stochastic constraint to be satisfied in the expected sense. This of course is problematic, e.g. assume that the worst distribution in $\mathcal{P}$ is symmetric and that it is tight for the optimal solution, then the constraint would be satisfied in only half of the cases. However, the risk aversion of a decision maker can be accounted for by introducing a disutility function of the form $U(y) = \max_{u \in \mathcal{U}} \{\gamma_u y + \delta_u\}$ where $\mathcal{U}$ is a finite index set and $\gamma > 0$. The distributionally robust constraint is then formulated as

$$\sup_{P \in \mathcal{P}} \mathbb{E}_P[U(v(x, \tilde{z}))] \leq w.$$

The decision maker thus becomes more and more cautious with consuming an extra

unit of the capacity. The deterministic counterpart for such a constraint is given by

$$b^T \beta + \sum_{i \in \mathcal{I}} [\bar{p}_i \kappa_i - \underline{p}_i \lambda_i] \leq w \quad (216)$$

$$\left. \begin{aligned} c_i^T \phi_{il} + \gamma_u s_l^T x + \gamma_u t_l + \delta_u &\leq \sum_{i' \in \mathcal{A}(i)} [\kappa_{i'} - \lambda_{i'}] r \\ C_i^T \phi_{il} + A^T \beta &= \gamma_u S_l^T x + \gamma_u \mathbf{t}_l \\ D_i^T \phi_{il} + B^T \beta &= 0 \end{aligned} \right\} \forall i \in \mathcal{I}, \forall l \in \mathcal{L}, \forall u \in \mathcal{U}$$

Similarly, one can formulate safeguarding constraints based on the shortfall risk measure introduced by [3] which is defined through

$$\rho^{SR}[v(x, \tilde{z})] = \inf_{\eta \in \mathbb{R}} \left\{ \eta : \sup_{P \in \mathcal{P}} E_P[U(v(x, \tilde{z}) - \eta)] \leq 0 \right\} \quad (217)$$

where U is a non-decreasing convex disutility function with $U(0) = 0$ and the sub-differential map of $\partial U(0) = \{1\}$. Shortfall risk can be interpreted as the amount of cash one would have to add to position in order to make it an acceptable position, i.e. to have non negative disutility. If U is of the form introduced above, then $\rho^{SR}[v(x, \tilde{z})] \leq w$ is equivalent to

$$\sup_{P \in \mathcal{P}} E_P[U(v(x, \tilde{z}) - w)] \leq 0 \quad (218)$$

and the deterministic counterpart is given by

$$b^T \beta + \sum_{i \in \mathcal{I}} [\bar{p}_i \kappa_i - \underline{p}_i \lambda_i] \leq w \quad (219)$$

$$\begin{aligned} c_i^T \phi_{il} + \gamma_u s_l^T x + \gamma_u t_l - \gamma_u w + \delta_u &\leq \sum_{i' \in \mathcal{A}(i)} [\kappa_{i'} - \lambda_{i'}] \quad \forall i \in \mathcal{I}, \forall l \in \mathcal{L}, \forall u \in \mathcal{U} \\ C_i^T \phi_{il} + A^T \beta &= \gamma_u S_l^T x + \gamma_u \mathbf{t}_l \quad \forall i \in \mathcal{I}, \forall l \in \mathcal{L}, \forall u \in \mathcal{U} \\ D_i^T \phi_{il} + B^T \beta &= 0 \quad \forall i \in \mathcal{I}, \forall l \in \mathcal{L}, \forall u \in \mathcal{U}. \end{aligned}$$

6.2.2 Applying the theory to the Knapsack Problem

Recently [2] and [51] have proposed distributionally robust models for the Knapsack and the Portfolio selection problem respectively. Their approaches differ from each other in the way they deal with the semi-infinite constraints arising from dual of the moment problem, and in the ambiguity sets they assume. What follows is an attempt to apply the framework reviewed above to those very same problems, in order to demonstrate its strengths and restrictions. We start out with simple models

and extend them step by step as to capture more challenging cases. The notation from Chapter 5.4 will be continued here. For the most part we will assume the following ambiguity set

$$\mathcal{P}' = \left\{ P \left| \begin{array}{l} P(\tilde{z} \in \mathcal{S}) = 1 \\ \mathbb{E}_P[\tilde{z}] = \mu \\ \mathbb{E}_P[(\tilde{z} - \mu)(\tilde{z} - \mu)^T] \leq \Sigma \end{array} \right. \right\} \quad (220)$$

which by the lifting Theorem 6.2.1b can be stated as

$$\mathcal{P}' = \left\{ P \left| \begin{array}{l} \mathbb{E}_P[\tilde{z}] = \mu \\ \mathbb{E}_P[\tilde{U}] = \Sigma \end{array} \right. , P \left[\begin{pmatrix} 1 & (\tilde{z} - \mu)^T \\ (\tilde{z} - \mu) & \tilde{U} \end{pmatrix} \geq 0 \right] = 1, P(\tilde{z} \in \mathcal{S}) = 1 \right\} \quad (221)$$

where $\mathcal{S}$ is the support of $\tilde{z}$. If we choose $\mathcal{S} = \mathfrak{R}^m$, the last condition is implied by the preceding one, so it can be omitted. Thus there is only one confidence set, which will be denoted by C_I so that

$$C_I = \left\{ (\tilde{z}, \tilde{U}) : \begin{pmatrix} 1 & (\tilde{z} - \mu)^T \\ (\tilde{z} - \mu) & \tilde{U} \end{pmatrix} \geq 0 \right\} \quad (222)$$

is of the standard form. Further, we will assume, that the data of the problems will linearly depend on the random vector $\tilde{z} \in \mathfrak{R}^m$ i.e. for the classical and the quadratical knapsack problem respectively

$$m = \sum_{i=1}^m m^i \tilde{z}_i \quad M = \sum_{i=1}^m M^i \tilde{z}_i \quad w_j = \sum_{i=1}^m w_j^i \tilde{z}_i \quad (223)$$

where $m^i, w^i \in \mathfrak{R}^n$, $M^i \in \mathfrak{R}^{n \times n}$. The robust counterpart of the classical knapsack problem is then obtained by a straightforward application of the theory discussed above. We still provide the step by step derivation in order to maintain transparency of the results. Let's start out with the stochastic representation

$$\begin{aligned} & \max_x \inf_{P \in \mathcal{P}'} \mathbb{E}_P[m^T x] \\ & \text{subject to } \sup_{P \in \mathcal{P}'} \mathbb{E}_P[w_j^T x] \leq d_j \quad j = 1, \dots, J \\ & \quad x \in \{0, 1\}^n \end{aligned} \quad (224)$$

or equivalently

$$\begin{aligned}
& \max_{x, d_0} d_0 \\
& \text{subject to } \sup_{P \in \mathcal{P}'} \mathbb{E}_P[-m^T x] \leq -d_0 \\
& \quad \sup_{P \in \mathcal{P}'} \mathbb{E}_P[w_j^T x] \leq d_j \quad j = 1, \dots, J \\
& \quad x \in \{0, 1\}^n
\end{aligned} \tag{225}$$

Now the former objective and the constraints are of the same form, with the argument of the expectation function is given by

$$v_j(x, \tilde{z}) = v_j^T x = \left(\sum_{i=1}^m v_j^i \tilde{z}_i \right)^T x = \sum_{i=1}^m x^T v_j^i \tilde{z}_i = x^T \mathcal{V}_j \tilde{z} \tag{226}$$

where $\mathcal{V}_j$ is the matrix composed of the vectors $v_j \in \mathbb{R}^n$, $j = 1, \dots, m$. Thus the constraints are fulfilled, if the respective right hand side is not exceeded by the optimal value of the moment problem given by

$$\begin{aligned}
& \max_{\mathcal{F}} \int_{C_I} v_j(x, z) d\mathcal{F} \\
& \quad \int_{C_I} z d\mathcal{F} = \mu \\
& \quad \int_{C_I} U d\mathcal{F} = \Sigma \\
& \quad \int_{C_I} \mathbf{1}_{[(z, u) \in C_I]} d\mathcal{F} = 1 \\
& \quad \mathcal{F} \in \mathcal{M}_+(\mathbb{R}^m \times \mathbb{R}^{m \times m})
\end{aligned} \tag{227}$$

the dual of which is given by

$$\begin{aligned}
& \min_{\beta_j, \Gamma_j, \kappa_j} \beta_j^T \mu + \Gamma_j \bullet \Sigma + \kappa_j \\
& \text{subject to } \beta_j^T z + \Gamma_j \bullet U + \kappa_j \mathbf{1}_{[(z, u) \in C_I]} \geq v_j(x, z) \quad \forall (z, U) \in C_I \\
& \quad \beta_j \in \mathbb{R}^m \\
& \quad \Gamma_j \in S_+^n \cup S_-^n \\
& \quad \kappa_j \in \mathbb{R}
\end{aligned} \tag{228}$$

Since there is only one confidence set, we don't need the subdivision argument that was used above. We rather observe, that the indicator function always gives one, and we can thus reformulate the dual as

$$\begin{aligned}
& \min_{\beta_j, \Gamma_j, \kappa_j} && \beta_j^T \mu + \Gamma_j \bullet \Sigma + \kappa_j && (229) \\
& \text{subject to} && \beta_j^T z + \Gamma_j \bullet U + \kappa_j \geq v_j(x, z) && \forall (z, U) \in C_I \\
& && \beta_j \in \mathfrak{R}^m && \\
& && \Gamma_j \in S_+^n \cup S_-^n && \\
& && \kappa_j \in \mathfrak{R} &&
\end{aligned}$$

where the constraint is of the form

$$x^T \mathcal{V}_j z + f^T z + G \bullet U \leq h \quad \forall (z, U) \in C_I \quad (230)$$

which is equivalent to stating that the optimum of the problem

$$\begin{aligned}
& \max_{z, U} && x^T \mathcal{V}_j z + f^T z + G \bullet U && (231) \\
& \text{subject to} && \begin{pmatrix} 1 & (z - \mu)^T \\ (z - \mu) & U \end{pmatrix} \succeq 0 &&
\end{aligned}$$

should not exceed h , which in turn is implied if the optimal value of its dual given by

$$\begin{aligned}
& \min && \bar{r} - 2r^T \mu && (232) \\
& && \mathcal{V}_j^T x + f + 2r = 0 && \\
& && G + R = 0 && \\
& && \begin{pmatrix} \bar{r} & r^T \\ r & R \end{pmatrix} \succeq 0 &&
\end{aligned}$$

does neither. Setting $f = -\beta_j$, $G = -\Gamma_j$, $h = \bar{k}_j$, and indexing all of the dual variables with j , the deterministic counterpart of the distributionally robust constraint is given

by

$$\begin{aligned}
\beta_j^T \mu + \Gamma_j \bullet \Sigma + \kappa_j &\leq d_j \\
\bar{r}_j - 2r_j^T \mu &\leq \kappa_j \\
\mathcal{V}_j^T x - \beta_j + 2r_j &= 0 \\
-\Gamma_j + R_j &= 0 \\
\begin{pmatrix} \bar{r}_j & r_j^T \\ r_j & R_j \end{pmatrix} &\geq 0
\end{aligned}
\Rightarrow
\begin{aligned}
x^T \mathcal{V}_j \mu + 2r_j^T \mu + \Gamma_j \bullet \Sigma + \kappa_j &\leq d_j \\
\bar{r}_j - 2r_j^T \mu &\leq \kappa_j \\
\begin{pmatrix} \bar{r}_j & r_j^T \\ r_j & R_j \end{pmatrix} &\geq 0
\end{aligned}$$

where the second step is achieved by solving $\Gamma_j = R$ and $\beta_j = \mathcal{V}^T x + 2r$ which can be simplified further into

$$\begin{aligned}
x^T \mathcal{V}_j \mu + \bar{r}_j + \Gamma_j \bullet \Sigma &\leq d_j \\
\bar{r}_j &\geq 0 \\
\Gamma_j &\geq 0
\end{aligned} \tag{233}$$

We can now insert $\mathcal{V}_0 = -\mathcal{M} = -(m^1, \dots, m^m)$, $\mathcal{V}_j = \mathcal{W}_j = (w^1, \dots, w^m)$ $j = 1, \dots, J$, and remove the auxiliary variable d_0 to obtain

$$\begin{aligned}
\max x^T \mathcal{M} \mu - \bar{r}_0 - \Gamma_0 \bullet \Sigma \\
x^T \mathcal{W}_j \mu + \bar{r}_j + \Gamma_j \bullet \Sigma &\leq d_j \quad j = 1, \dots, J \\
\left. \begin{aligned} \bar{r}_j &\geq 0 \\ \Gamma_j &\geq 0 \end{aligned} \right\} j = 0, \dots, J \\
x &\in \{0, 1\}
\end{aligned} \tag{234}$$

One immediately observes, that the variance is irrelevant for solving this problem, which is the case, because we assumed a linear (i.e. risk neutral) utility function for the objective as well as for the constraints. Of course this model is useless for any practical purpose, so we need to extend it. We start out by modeling risk aversion regarding the constraints, by introducing a disutility function of the form proposed in [30], i.e.

$$D(y) = \max_{u \in \mathcal{U}} [\gamma_u y + \delta_u] \tag{235}$$

where $\mathcal{U} = \{1, \dots, \bar{U}\}$ is an index set and all $\gamma_u \geq 0$. The constraints are thus given by

$$\sup_{P \in \mathcal{P}'} \mathbb{E}_P[D(v_j^T x)] \leq d_j \quad j = 1, \dots, J \tag{236}$$

and . Such a constraint can be interpreted as an increasing reluctance of the decision maker to further consume a recourse, thus seeking to avoid situations where this is likely to occur. Since now $v_j = U(x^T \mathcal{V}_j \tilde{z})$ the dual of the moment problem is unaffected, given the notation above, but when we insert

$$\beta_j^T z + \Gamma_j \bullet U + \kappa_j \geq \max_{u \in \mathcal{U}} [\gamma_u x^T \mathcal{V}_j z + \delta_u] \quad \forall (z, U) \in C_I$$

we see that in order for this to be fulfilled, the optimal values of the $\bar{U}$ problems of the form (231) must not exceed the respective right hand side. Again, by formulating the $\bar{U}$ duals, we find the equivalent conditions to be

$$\begin{aligned} \bar{r}_{ju} - 2r_{ju}^T \mu &\leq \kappa_j - \delta_u \\ \gamma_u \mathcal{V}_j^T x - \beta_j + 2r_{ju} &= 0 \\ -\Gamma_j + R_{ju} &= 0 \end{aligned} \Rightarrow \begin{aligned} \bar{r}_{ju} + \gamma_u x^T \mathcal{V}_j \mu - \beta_j^T \mu &\leq \kappa_j - \delta_u \\ \begin{pmatrix} \bar{r}_{ju} & r_{ju}^T \\ \frac{1}{2}(\beta_j - \gamma_u \mathcal{V}_j^T x) & \Gamma_j \end{pmatrix} &\geq 0 \end{aligned}$$

Note that we do not assume that each constraint is treated with the same disutility function, but for ease of Notation we forgo an accordant subindex for the u indices. Thus the distributionally robust counterpart is given by

$$\begin{aligned} \beta_j^T \mu + \Gamma_j \bullet \Sigma + \kappa_j &\leq d_j \\ \left. \begin{aligned} \bar{r}_{ju} + \gamma_u x^T \mathcal{V}_j \mu - \beta_j^T \mu &\leq \kappa_j - \delta_u \\ \begin{pmatrix} \bar{r}_{ju} & r_{ju}^T \\ \frac{1}{2}(\beta_j - \gamma_u \mathcal{V}_j^T x) & \Gamma_j \end{pmatrix} &\geq 0 \end{aligned} \right\} \forall u \in \mathcal{U} \end{aligned} \quad (237)$$

For the objective function, there are two strategies that will allow us to model risk aversion. The first one is to define a utility function $U_0(y) = \min_{u \in \mathcal{U}} [\gamma_u y + \delta_u]$ and change the objective function accordingly

$$\max_x \inf_{P \in \mathcal{P}'} \mathbb{E}_P[U_0(m^T x)] \Rightarrow \max_{x, d_0} d_0 \quad \text{subject to } \sup_{P \in \mathcal{P}'} \mathbb{E}_P[\max_{u \in \mathcal{U}} (\gamma_u (-m^T x) - \delta_u)] \leq -d_0 \quad (238)$$

The distributionally robust counterpart is then of similar form as it was the case for

the constraints. The model can thus be stated as

$$\begin{aligned}
& \max -\beta_0^T \mu - \Gamma_0 \bullet \Sigma - \kappa_0 \quad (239) \\
& \text{subject to} \\
& \left(\begin{array}{cc} \bar{r}_{0u} - \gamma_u x^T \mathcal{M} \mu - \beta_0^T \mu \leq \kappa_0 + \delta_u & \\ \bar{r}_{0u} & \frac{1}{2}(\beta_0 + \gamma_u \mathcal{M}^T x)^T \\ \frac{1}{2}(\beta_0 + \gamma_u \mathcal{M}^T x) & \Gamma_0 \end{array} \right) \geq 0 \quad \left. \vphantom{\left(\begin{array}{cc} \bar{r}_{0u} - \gamma_u x^T \mathcal{M} \mu - \beta_0^T \mu \leq \kappa_0 + \delta_u & \\ \bar{r}_{0u} & \frac{1}{2}(\beta_0 + \gamma_u \mathcal{M}^T x)^T \\ \frac{1}{2}(\beta_0 + \gamma_u \mathcal{M}^T x) & \Gamma_0 \end{array} \right)} \right\} \forall u \in \mathcal{U} \\
& \left(\begin{array}{cc} \beta_j^T \mu + \Gamma_j \bullet \Sigma + \kappa_j \leq d_j & \\ \bar{r}_{ju} + \gamma_u x^T \mathcal{W}_j \mu - \beta_j^T \mu \leq \kappa_j + \delta_u & \\ \bar{r}_{ju} & \frac{1}{2}(\beta_j - \gamma_u \mathcal{W}_j^T x)^T \\ \frac{1}{2}(\beta_j - \gamma_u \mathcal{W}_j^T x) & \Gamma_j \end{array} \right) \geq 0 \quad \left. \vphantom{\left(\begin{array}{cc} \beta_j^T \mu + \Gamma_j \bullet \Sigma + \kappa_j \leq d_j & \\ \bar{r}_{ju} + \gamma_u x^T \mathcal{W}_j \mu - \beta_j^T \mu \leq \kappa_j + \delta_u & \\ \bar{r}_{ju} & \frac{1}{2}(\beta_j - \gamma_u \mathcal{W}_j^T x)^T \\ \frac{1}{2}(\beta_j - \gamma_u \mathcal{W}_j^T x) & \Gamma_j \end{array} \right)} \right\} \forall u \in \mathcal{U} \quad \forall j = 1, \dots, J \\
& x \in \{0, 1\}^n \\
& \beta_j \in \mathfrak{R}^m \quad \Gamma_j \in S_+^n \quad \kappa_j \in \mathfrak{R} \quad \forall j = 0, \dots, J \\
& \bar{r}_{ju} \in \mathfrak{R} \quad \forall j = 0, \dots, J \quad \forall u \in \mathcal{U} \quad (240)
\end{aligned}$$

Alternatively we can minimize the shortfall risk as presented in the previous chapter i.e.

$$\begin{aligned}
\min \rho^{SR}[m^T x] & \Rightarrow \min_{x, d_0} d_0 \\
& \text{subject to } \sup_{P \in \mathcal{P}} \mathbb{E}_P[U(m^T x - d_0)] \leq 0 \quad (241)
\end{aligned}$$

The model then becomes

$$\begin{aligned}
& \min d_0 \quad (242) \\
& \text{subject to } \beta_0^T \mu + \Gamma_0 \bullet \Sigma + \kappa_0 \leq 0 \\
& \left(\begin{array}{cc} \bar{r}_{0u} + \gamma_u x^T \mathcal{M} \mu - \beta_0^T \mu \leq \kappa_0 - \delta_u + \gamma_u d_0 & \\ \bar{r}_{0u} & \frac{1}{2}(\beta_0 - \gamma_u \mathcal{M}^T x)^T \\ \frac{1}{2}(\beta_0 - \gamma_u \mathcal{M}^T x) & \Gamma_0 \end{array} \right) \geq 0 \quad \left. \vphantom{\left(\begin{array}{cc} \bar{r}_{0u} + \gamma_u x^T \mathcal{M} \mu - \beta_0^T \mu \leq \kappa_0 - \delta_u + \gamma_u d_0 & \\ \bar{r}_{0u} & \frac{1}{2}(\beta_0 - \gamma_u \mathcal{M}^T x)^T \\ \frac{1}{2}(\beta_0 - \gamma_u \mathcal{M}^T x) & \Gamma_0 \end{array} \right)} \right\} \forall u \in \mathcal{U} \\
& \left(\begin{array}{cc} \beta_j^T \mu + \Gamma_j \bullet \Sigma + \kappa_j \leq d_j & \\ \bar{r}_{ju} + \gamma_u x^T \mathcal{W}_j \mu - \beta_j^T \mu \leq \kappa_j - \delta_u & \\ \bar{r}_{ju} & \frac{1}{2}(\beta_j - \gamma_u \mathcal{W}_j^T x)^T \\ \frac{1}{2}(\beta_j - \gamma_u \mathcal{W}_j^T x) & \Gamma_j \end{array} \right) \geq 0 \quad \left. \vphantom{\left(\begin{array}{cc} \beta_j^T \mu + \Gamma_j \bullet \Sigma + \kappa_j \leq d_j & \\ \bar{r}_{ju} + \gamma_u x^T \mathcal{W}_j \mu - \beta_j^T \mu \leq \kappa_j - \delta_u & \\ \bar{r}_{ju} & \frac{1}{2}(\beta_j - \gamma_u \mathcal{W}_j^T x)^T \\ \frac{1}{2}(\beta_j - \gamma_u \mathcal{W}_j^T x) & \Gamma_j \end{array} \right)} \right\} \forall u \in \mathcal{U} \quad \forall j = 1, \dots, J \\
& x \in \{0, 1\}^n, \quad d_0 \in \mathfrak{R} \\
& \beta_j \in \mathfrak{R}^m \quad \Gamma_j \in S_+^n \quad \kappa_j \in \mathfrak{R} \quad \forall j = 0, \dots, J \\
& \bar{r}_{ju} \in \mathfrak{R} \quad \forall j = 0, \dots, J \quad \forall u \in \mathcal{U} \quad (243)
\end{aligned}$$

In each of the models the relaxations of the binary constraint can be achieved by substituting it with $0 \leq x_i \leq 1, \forall i = 1, \dots, n$. Further, by replacing all the capacity constraints with $\mathbf{1}^T x = 1$ the models become instances of the portfolio selection problem, where m is a vector of uncertain revenues and x_i gives the share of the i th stock in the portfolio. This illustrates, that the framework presented above is suitable to reproduce the results achieved in [51]. Of course, the ambiguity set the authors assumed there is much richer than the one we are working with here, but it is still well within the scope of the general framework. Further remarks on the above models will be made in the conclusion of this chapter. We meanwhile proceed with tackling the quadratic knapsack problem.

What has been written so far, is a straight forward application of the theory by [30]. If we now want to consider the quadratic knapsack problem we unfortunately cannot proceed with the same strategy, which becomes apparent when we observe that for the rearranged problem

$$\begin{aligned}
& \min_{x, d_0} d_0 \tag{244} \\
& \text{subject to } \sup_{P \in \mathcal{P}'} \mathbb{E}_P[U_0(x^T M x - d_0)] \leq 0 \\
& \sup_{P \in \mathcal{P}'} \mathbb{E}_P[U(w_j^T x)] \leq d_j \quad j = 1, \dots, J \\
& x \in \{0, 1\}^n
\end{aligned} \tag{245}$$

the left hand side of the constraint is not necessarily convex, and thus (C3a) is not always fulfilled unless $M \geq 0$. This has a couple of consequences. Set aside for a moment, that (C3b) might be problematic as well, requiring M to be p.s.d. implies, that we can only ever model situations, where even in the worst case, we can only ever end up with an positive overall NPV. One could argue, that such a comfortable situation rarely occurs, and if it did, one could raise the question whether applying sophisticated optimization techniques is worth the effort. We will come back to this issue in the conclusion of this chapter.

Now, we have to adopt the model in order ensure that $M = \sum_{i=1}^m M^i \tilde{z}_i$ stays positive semidefinite for all possible values of $\tilde{z}$. To this end we have to specify the support $\mathcal{S}$ of $\tilde{z}$ accordingly. A simple choice is to restrict $\tilde{z}$ so that $\mathbf{1}^T \tilde{z} = 1$ and $z_i \geq 0, i = 1, \dots, m$, such that M is a convex combination of all M^i , which of course have to

be p.s.d. for all i . The respective set can be represented as

$$\mathcal{S} = \left\{ \tilde{z} \left| \begin{pmatrix} \mathbf{Diag}(\tilde{z}) & 0 & 0 \\ 0 & \mathbf{1}^T \tilde{z} - 1 & 0 \\ 0 & 0 & -\mathbf{1}^T \tilde{z} + 1 \end{pmatrix} \geq 0 \right. \right\} \quad (246)$$

Unfortunately C_I and $\mathcal{S}$ as stated above do not fulfill the the nesting condition (N), but since $\tilde{z}$ is in either of those with certainty we can combine them into one confidence set

$$C_S = \left\{ (\tilde{z}, \tilde{U}) \left| \begin{pmatrix} 1 & (\tilde{z} - \mu)^T & 0 & 0 & 0 \\ (\tilde{z} - \mu) & \tilde{U} & 0 & 0 & 0 \\ 0 & 0 & \mathbf{Diag}(\tilde{z}) & 0 & 0 \\ 0 & 0 & 0 & \mathbf{1}^T \tilde{z} - 1 & 0 \\ 0 & 0 & 0 & 0 & -\mathbf{1}^T \tilde{z} + 1 \end{pmatrix} \geq 0 \right. \right\} \quad (247)$$

where μ has to be chosen, such that the set is nonempty. The function $v(x, \tilde{z})$ is now of the form

$$v(x, \tilde{z}) = x^T \left(\sum_{i=1}^m V_i \tilde{z}_i \right) x = \sum_{i=1}^m x^T V_i x \tilde{z}_i = \begin{pmatrix} x^T V_1 x \\ \dots \\ x^T V_m x \end{pmatrix}^T \tilde{z} \quad (248)$$

The moment problem is similar as before, with the only difference being, that the integral is now over C_S . If we again have the objective to minimize shortfall risk, the dual of the moment problem is given by

$$\begin{aligned} & \min_{\beta_j, \Gamma_j, \kappa_j} \quad \beta_j^T \mu + \Gamma_j \bullet \Sigma + \kappa_j \\ & \text{subject to} \quad \beta_j^T z + \Gamma_j \bullet U + \kappa_j \geq \max_{u \in \mathcal{U}} \gamma_u[(\cdot)^T z - d_j] + \delta_u \quad \forall (z, U) \in C_S \\ & \quad \beta_j \in \mathfrak{R}^m \\ & \quad \Gamma_j \in S_+^n \cup S_-^n \\ & \quad \kappa_j \in \mathfrak{R} \end{aligned} \quad (249)$$

where the semi-infinite constraint is fulfilled if the $\bar{U}$ equation systems

$$\begin{aligned} \bar{r}_{ju} + c_u - 2r_{ju}^T \mu &\leq \kappa_j - \delta_u + \gamma_u d_j \\ \gamma_u \begin{pmatrix} x^T V_1 x \\ \cdots \\ x^T V_m x \end{pmatrix} - \beta_j + 2r_{ju} + b_u - c_u \mathbf{1} &= 0 \\ -\Gamma_j + R_{ju} &= 0 \\ \begin{pmatrix} \bar{r}_{ju} & r_{ju}^T \\ r_{ju} & R_{ju} \end{pmatrix} &\geq 0 \end{aligned}$$

$\Downarrow$

$$\begin{aligned} \bar{r}_{ju} + c_u + \delta_u + \gamma_u d_j \begin{pmatrix} x^T V_1 x \\ \cdots \\ x^T V_m x \end{pmatrix}^T \mu - \beta_j^T \mu + b_u^T \mu - c_u \mathbf{1}^T \mu &\leq \kappa_j - \delta_u + \gamma_u d_j \\ \begin{pmatrix} \bar{r}_{ju} & \frac{1}{2} \left(\beta - \gamma_u \begin{pmatrix} x^T V_1 x \\ \cdots \\ x^T V_m x \end{pmatrix} - b_u + c_u \mathbf{1} \right)^T \\ \frac{1}{2} \left(\beta - \gamma_u \begin{pmatrix} x^T V_1 x \\ \cdots \\ x^T V_m x \end{pmatrix} - b_u + c_u \mathbf{1} \right) & \Gamma_j \end{pmatrix} &\geq 0 \end{pmatrix} \quad (250)$$

are fulfilled.

We further impose safeguarding constraints in order to model risk aversion with regard to the distributionally robust constraints. The moment problems and the respective duals as well as the finite reformulation of the semi-infinite constraints are very similar to the ones obtained above. We therefore skip the step by step derivation. The deterministic counter part of the distributionally robust model is then given by

$$\begin{aligned}
& \min d_0 \\
& \text{subject to} \\
& \beta_0^T \mu + \Gamma_0 \bullet \Sigma + \kappa_j \leq 0 \\
& \left(\begin{array}{cc} \bar{r}_{0u} + c_u + \gamma_u \begin{pmatrix} x^T M_1 x \\ \cdots \\ x^T M_m x \end{pmatrix}^T & \mu - \beta_0^T \mu + b_u^T \mu - c_u \mathbf{1}^T \mu \leq \kappa_0 - \delta_u + \gamma_u d_0 \\ \bar{r}_{0u} & \frac{1}{2} \left(\beta_0 - \gamma_u \begin{pmatrix} x^T M_1 x \\ \cdots \\ x^T M_m x \end{pmatrix} - b_u + c_u \mathbf{1} \right)^T \\ \frac{1}{2} \left(\beta_0 - \gamma_u \begin{pmatrix} x^T M_1 x \\ \cdots \\ x^T M_m x \end{pmatrix} - b_u + c_u \mathbf{1} \right) & \Gamma_0 \end{array} \right) \geq 0 \quad \forall u \in \mathcal{U} \\
& \left(\begin{array}{cc} \bar{r}_{ju} + c_u + \gamma_u x^T \mathcal{W}_j \mu - \beta_j^T \mu + b^T \mu - c_u \mathbf{1}^T \mu \leq \kappa_j - \delta_u & \beta_j^T \mu + \Gamma_j \bullet \Sigma + \kappa_j \leq d_j \\ \bar{r}_{ju} & \frac{1}{2} (\beta_j - \gamma_u \mathcal{W}_j^T x - b_u + c_u \mathbf{1})^T \\ \frac{1}{2} (\beta_j - \gamma_u \mathcal{W}_j^T x - b_u + c_u \mathbf{1}) & \Gamma_j \end{array} \right) \geq 0 \quad \forall u \in \mathcal{U} \quad \forall j = 1, \dots, J \\
& x \in \{0, 1\}^n \quad \beta_j \in \mathfrak{K}^m \quad \Gamma_j \in S_+^n \quad \kappa_j \in \mathfrak{K} \quad \forall j = 0, \dots, J \\
& d_0 \in \mathfrak{K} \quad \bar{r}_{ju} \in \mathfrak{K} \quad \forall j = 0, \dots, J \quad \forall u \in \mathcal{U} \\
& \quad \quad \quad c_u \in \mathfrak{K} \quad b_u \in \mathfrak{K}^m \quad \forall u \in \mathcal{U}
\end{aligned} \tag{251}$$

This model gives an exact distributionally robust counterpart for the quadratic knapsack problem with safeguarding constraints that account for the risk aversion of the decision maker. As with the deterministic model, there might be many ways to find relaxations for the binary constraint. Here we will restrict ourselves to the approach presented in Chapter 5.4 for obtaining (168). To this end we substitute $Y = yy^T$ and relax $Y \geq yy^T$ which can be represented as

$$\begin{pmatrix} 1 & y^T \\ y & Y \end{pmatrix} \geq 0 \tag{252}$$

The relaxed model is given by

$$\begin{aligned}
& \min d_0 \\
& \text{subject to} \\
& \beta_0^T \mu + \Gamma_0 \bullet \Sigma + \kappa_j \leq 0 \\
& \left(\begin{array}{cc} \bar{r}_{0u} + c_u + \gamma_u \begin{pmatrix} Y \bullet M_1 \\ \cdots \\ Y \bullet M_m \end{pmatrix}^T & \mu - \beta_0^T \mu + b_u^T \mu - c_u \mathbf{1}^T \mu \leq \kappa_0 - \delta_u + \gamma_u d_0 \\ \bar{r}_{0u} & \frac{1}{2} \left(\beta - \gamma_u \begin{pmatrix} Y \bullet M_1 \\ \cdots \\ Y \bullet M_m \end{pmatrix} - b_u + c_u \mathbf{1} \right)^T \\ \frac{1}{2} \left(\beta - \gamma_u \begin{pmatrix} Y \bullet M_1 \\ \cdots \\ Y \bullet M_m \end{pmatrix} - b_u + c_u \mathbf{1} \right) & \Gamma_0 \end{array} \right) \geq 0 \quad \forall u \in \mathcal{U}
\end{aligned} \tag{253}$$

$$\begin{aligned}
& \left(\begin{array}{cc} \bar{r}_{ju} + c_u + \gamma_u x^T \mathcal{W}_j \mu - \beta_j^T \mu + b_u^T \mu - c_u \mathbf{1}^T \mu \leq \kappa_j - \delta_u & \beta_j^T \mu + \Gamma_j \bullet \Sigma + \kappa_j \leq d_j \\ \bar{r}_{ju} & \frac{1}{2} (\beta_j - \gamma_u \mathcal{W}_j^T x - b_u + c_u \mathbf{1})^T \\ \frac{1}{2} (\beta_j - \gamma_u \mathcal{W}_j^T x - b_u + c_u \mathbf{1}) & \Gamma_j \end{array} \right) \geq 0 \quad \forall u \in \mathcal{U} \quad \forall j = 1, \dots, J \\
& \begin{pmatrix} 1 & y^T \\ y & Y \end{pmatrix} \geq 0, \quad y = \mathbf{diag}(Y), \quad Y_{ij} \geq 0 \quad \forall i \neq j, \\
& x_i = y_i \quad \forall i = 1, \dots, n \\
& d_0 \in \mathfrak{R} \\
& \beta_j \in \mathfrak{R}^m \quad \Gamma_j \in S_+^n \quad \kappa_j \in \mathfrak{R} \quad \forall j = 0, \dots, J \\
& \bar{r}_{ju} \in \mathfrak{R} \quad \forall j = 0, \dots, J \quad \forall u \in \mathcal{U} \\
& c_u \in \mathfrak{R} \quad b_u \in \mathfrak{R}^m \quad \forall u \in \mathcal{U}
\end{aligned}$$

While the classical knapsack problem was handled with relative ease using the theory by [30], the quadratic knapsack was a little more resilient. In any case the approach chosen above is questionable for relying on safeguarding constraints. First of all, it is hard to quantify the probability of each individual constraint being fulfilled. But even if we could formulate safeguarding constraints, that guarantee a high probability for the constraints to hold individually, the chance of them to hold collectively would deteriorate exponentially as more constraints are added. Without further analytical investigation and numerical experiments we cannot asses

the impact of this issue yet. The alternative approach proposed by [2] is to use chance constraints, that enforce a certain probability for the constraints to be fulfilled jointly. However, their approach only gives conservative approximations if the chance condition is imposed in the case of multiple constraints.

Another obvious disadvantage of the model is, that the NPV-Matrix M needs to be positive semidefinite in order to ensure convexity of the of the constraint function. The consequence is, that we can only model situations, where it is sure that even in the worst case it would be impossible to achieve negative total NPVs. Needless to say, that such a restriction renders the models quite useless for practical application. The convexity of the constraint function comes into play first when deriving the dual of the moment problem, where in absence of convexity there would be no guarantee, that the duality gap is zero. Let's denote the optimal value of the primal moment problem as p_{max} and the optimal value of the dual d_{min} . The Argument is that the left-hand side of the $\sup_{p \in \mathcal{P}} E_P[v(x, \tilde{z})] \leq w$ coincides with p_{max} , and thus the constraint is fulfilled if $p_{max} \leq w$ which is in turn implied by $d_{min} \leq w$. This is certainly true even in absence of strong duality since $p_{max} \leq d_{min}$ holds by weak duality. However if there is a duality gap such that $d_{min} = p_{max} + d_{Gap}$, where d_{Gap} is some positive duality gap, we have that $d_{min} \leq w \Rightarrow p_{max} \leq w - d_{Gap}$, which is a tightening of the original distributionally robust constraint, which in the worst case could lead to an infeasible problem. The second time convexity is important, is when the semi-infinite constraints are reformulated by partitioning the confidence set and then using an optimization argument to find equivalent conditions. However, in the above model there is only one confidence set, in which case this reformulation does not depend on such an argument, so non-convexity does not harm the model here. Finally, convexity of the constraint function comes into play when the reformulated semi-infinite constraints are again transformed into finite constraints using a duality argument. In our case, the respective function is linear in $\tilde{z}$, such that there is no issue with convexity here. Thus the requirement of convexity might be restricted to the objective function of the moment problem, whose convexity rests not only on the definiteness of the NPV-Matrix but also on the probability measure. Therefore, even if not all M_i are positive definite, convexity of the moment problems objective could still be ensured if the set of feasible measures is restricted accordingly. Finding out whether such constraints exist and how they could be formulated could be a worthwhile task for the future.

Finally, the ambiguity set that was necessary in order to achieve the regularization of M excludes a lot of distributions that are common in optimization modeling such as the normal distribution. While the condition, that the random vector gives weights

of a convex combination, can be omitted easily, this is not the case with the non-negativity constraints since they have to be imposed in order to guarantee that the NPV matrix is p.s.d. for all realizations of the random vector. This of course also hampers the applicability of the model if e.g. the normal distribution is known to belong to the "real" ambiguity set.

Conclusion

The field of semidefinite programming has grown extensively in the last two decades, as was demonstrated in this thesis. Many of the ideas discussed are relatively young, but due to the vast attention the field lately received from researchers, the above discussion is by far not exhaustive even within the finance world. That does not mean however, that there is no room for further research. The conclusion of the final chapter already addressed many unanswered questions regarding the distributionally robust knapsack problem with respect to the way of obtaining the robust model. Shading further light on that area will require the combination of ideas from stochastic and convex analysis. But beyond that, there is also much to be discovered concerning the quality of the bounds obtained by the relaxed model, and whether or not one can achieve improvements in that regard by applying alternative relaxation strategies. Besides the strategies presented for the deterministic model, copositive programming could potentially yield new approaches, which deserve thorough studying. Finally, we have not addressed at all the issue of the performance of each of the models when applied in a real world setting. This as well remains a challenging task for the future.

References

- [1] Ben-Tal A.; Nemirovski A. Robust convex optimization. *Mathematics of Operations Research*, 23(4):769–805, 1998.
- [2] Cheng J.; Delage E.; Lisser A. Distributionally robust stochastic knapsack problem. *SIAM Journal of Optimization*, 24(3), 2014.
- [3] Föllmer H.; Schied A. Convex measures of risk and trading constraints. *Finance and Stochastics*, 6(4):429–447, 2002.
- [4] Shapiro A. On duality theory of conic linear problems. *Semi-Infinite Programming*, pages 135–165, 2001. Kluwer Academic Publishers.

- [5] Ben-Tal A.; Margalit T.; Nemirovski A.N. Robust modeling of multi-stage portfolio problems. pages 303–328, 2002. Dordrecht: Kluwer.
- [6] Gallo G.; Hammer P.L.; Simeone B. Quadratic knapsack problems. *Mathematical Programming*, (12):132–149, 1980.
- [7] Merton R. C. Theory of rational option pricing. *The Bell Journal of Economics and Management Science*, 4(1), 1973.
- [8] Fifield S. G. M.; Power D. M.; Sinclair C.D. Macroeconomic factors and share returns: An analysis using emerging markets data. *International Journal of Finance and Economics*, 7:51–62, 2002.
- [9] Pisinger D. The quadratic knapsack problem - a survey. *Discrete Applied Mathematics*, 155:623–648, 2006.
- [10] d’Aspremont A.; Bach F.; El Ghaoui L. Optimal solutions for sparse principal component analysis. *Journal of Machine Learning Research*, (9):1269–1294, 2008.
- [11] Zhang Y.; d’Aspremont A.; El Ghaoui L. Sparse pca: Convex relaxations, algorithms and applications. pages 915–940, 2012. Springer.
- [12] d’Aspremont A.; El Ghaoui L.; Jordan M. I.; Lanckriet G. R. G. A direct formulation for sparse pca using semidefinite programming. *SIAM Review*, 49(3):434–448, 2007.
- [13] de Klerk E. *Aspects of Semidefinite Programming*. 2002. Kulwar Academic Publishers, Dordrecht, The Netherlands.
- [14] Alizadeh F. Combinatorial optimization with interior point methods and semidefinite matrices. *Proc. 2nd Annual Integer Programming and Combinatorial Optimization Conference*, 1992. Carnegie-Mellon University, Pittsburgh.
- [15] Jarre F. Convex analysis on symmetric matrices. *Handbook of Semidefinite Programming*, pages 13–28, 2003.
- [16] Gotoh J.; Konno H. Bounding option prices by semidefinite programming: A cutting plane algorithm. *Management Science*, 48(5):665–678, 2002. INFORMS.
- [17] Scarf H. A min-max solution of an inventory problem, studies in the mathematical theory of inventory and production. pages 201–209, 1958. <http://dido.econ.yale.edu/hespubstudies-12.pdf>, Stanford University Press.

- [18] Lobo M.S.; Vandenberghe L.; Boyd S.; Lebret Hervé. Applications of second-order cone programming. *Linear Algebra and its Applications*, 284:193–228, 1998.
- [19] Markowitz H.M. Portfolio selection. *The Journal of Finance*, 1952.
- [20] Bertsimas D.; Popescu I. Optimal inequalities in probability theory: A convex optimization approach. *SIAM Journal on Optimization*, 15(3):780–804, 2005.
- [21] Han D.; X. Li; Sun D.; Sun J. Bounding option prices of multiassets: a semidefinite programming approach. *Pacific Journal of Optimization*, 1(1):59–79, 2005.
- [22] Tobin J. Liquidity preference as behavior towards risk. *Review of Economic Studies*, pages 65–86, 1958.
- [23] Tanabe K. Centered newton method for mathematical programming. *System Modelling and Optimization*, 113:197–206, August/September 1987.
- [24] Wolkowicz H.; Saigal R.; Vandenberghe L. *Handbook of Semidefinite Programming*. 2000. Kluwer Academic Publishers.
- [25] Monteiro A.M.; R.H Tütüncü; Vicente L.N. Recovering risk-neutral probability density functions from options prices using cubic splines and ensuring non-negativity. *European Journal of Operational Research*, 187:525–542, 2008.
- [26] Alizadeh F.; Haeberly J.P.; Overton M. Primal-dual interior-point methods for semidefinite programming. *SIAM Journal of Optimization*, 8(3), 2006.
- [27] Black F.; Scholes M. The pricing of options and corporate liabilities. *The Journal of Political Economy*, 81:637–654, 1973.
- [28] Lasserre J.B.; Prieto-Rumeau T.; Zervos M. Pricing a class of exotic options via moments and SDP relaxations. *Mathematical Finance*, 16(3):469–494, 2006. Blackwell Publishing Inc.
- [29] Monteiro R.; Todd M. Path-following methods. *Handbook of Semidefinite Programming*, pages 267–306, 2003.
- [30] Wiesemann W.; Kuhn D.; Sim M. Distributionally robust convex optimization. *Operations Research*, 62(6):1358–1376, 2014.
- [31] Monteiro R.D.C.; Adler I.; Resende M.G.C. A polynomial-time primal-dual affine scaling algorithm for linear and convex quadratic programming and its power series extension. *Mathematics of Operations Research*, 15:191–214, 1997.

- [32] Nesterov Y.E.; Todd M.J. Self-scaled cones and interior-point methods in non-linear programming. *Tech. Report*, 1994.
- [33] Basso A.; Pianca P. Decreasing absolute risk aversion and option pricing bounds. *Management Science*, 43:206–216, 1997.
- [34] Jorion P. Portfolio optimization with constraints an tracking error. *Financial Analysts Journal*, 5(59):70–82, 2003.
- [35] Ritchken P. Option pricing bounds. *Journal of Finance*, 40:1219–1233, 1985.
- [36] Roos C.; Vial J. Ph. A polynomial mehod of approximate centres for linear programming. *Mathematical Programming*, 54:295–305, 1992.
- [37] Cornuejols G.; Tütüncü R. *Optimization Methods in Finance*. Mathematics, Finance and Risk. January 2006. Cambrige.
- [38] Helmberg C.; Rendl F.; Weismantel R. A semidefinite programming approach to the quadratic knapsack problem. *Journal of Combinatorial Optimization*, 4:197–215, 2000.
- [39] Roll R. A mean-variance analysis of tracking error. *Journal of Portfolio Management*, 18(4):13–22, 1992.
- [40] Freund R.M. Introduction to semidefinite programming. *unpublished Paper available at <http://ocw.mit.edu/courses/electrical-engineering-and-computer-science>*.
- [41] Alizadeh F.; Schmieta S. Symmetric cones, potential reduction methods. *Handbook of Semidefinite Programming*, pages 195–134, 2003.
- [42] Kruschwitz L.; Husmann S. *Finanzierung und Investition*, volume 7. 2012. Oldenbourg.
- [43] Ritchken P.; Kuo S. On stochastic dominance and decreasing absolute risk averse option pricing bounds. *Management Science*, (35):51–59, 1989.
- [44] Vandenberghe L.; Boyd S. Convex optimization. 25:601–615, 2000. Cambridge University Press.
- [45] Vandenberghe L.; Boyd S. Semidefinite programming. *SIAM Review*, 38:49–95, March, 1996.
- [46] Jansen B.; Roos C.; Terlaky T. A family of polynomial affine scaling algorithms for positiv semi-definite linear complementary problems. *SIAM Journal on Optimization*, 7:126–140, 1997.

- [47] Polik I.; Terlaky T. A survey of the s-lemma. *SIAM Review*, 49(3):371–418, 2007.
- [48] Rockafellar R. T. *Convex Analysis*. 1972. Princeton University Press, New Jersey.
- [49] Delbaen F.; Schachermayer W. *The Mathematics of Arbitrage*. 2000. Springer Finance.
- [50] Sharpe W. *Portfolio Theory and Capital Markets*. 1964. McGraw-Hill, New York.
- [51] Delage E.; Ye Y. Distributionally robust optimization under moment uncertainty with application to data-driven problems. *Operations Research*, 58:596–612, 2010.
- [52] Keho Y. The basics of linear principal components analysis. *Principal Component Analysis*, pages 181–206, 2012. InTech.
- [53] Nesterov Y. Interior-point polynomial methods in convex programming. *Studies in Applied Mathematics*, 13:117–204, 1994. Society for Industrial and Applied Mathematics.
- [54] Todd M.J.; Ye Y. A centered algorithm for linear programming. *Mathematics of Operations Research*, 15:508–529, 1990.
- [55] Sheng zhi P.; Fu-sheng W. Semidefinite programming relaxation for portfolio selection with higher order moments. In *International Conference on Management Science and Engineering*, pages 99–104, Sept. 13-15,2011.