

Approximation of the Whole Pareto Optimal Set for the Vector Optimization Problem

In Memoriam András Prékopa (1929-2016)

Tibor Illés¹, Gábor Lovics²

¹ Optimization Research Group, Department of Differential Equations, Budapest University of Technology and Economics, Eötvös József u. 1., 1111 Budapest, Hungary, illes@math.bme.hu

² Hungarian Central Statistical Office, Keleti Károly u. 5-7., 1024 Budapest, Hungary, Gabor.Lovics@ksh.hu

In multi-objective optimization problems several objective functions have to be minimized simultaneously. In this work, we present a new computational method for the linearly constrained, convex multi-objective optimization problem. We propose some techniques to find joint decreasing directions for both the unconstrained and the linearly constrained case as well. Based on these results, we introduce a method using a subdivision technique to approximate the whole Pareto optimal set of the linearly constrained, convex multi-objective optimization problem. Finally, we illustrate our algorithm by solving the Markowitz model on real data.

Keywords: multi-objective optimization; joint decreasing direction; approximation of Pareto optimal set; Markowitz model

2000 MSC: 90C29

1 Introduction

In the literature of economics and finance the measure of risk has always been a very interesting topic, and nowadays it may be even more important than ever. One of the first idea to be taken into consideration is the risk in financial activities coming from Markowitz [26], who developed his famous model, where the investors make portfolios from different securities, and try to maximize their profit and minimize their risk at the same time. In this model, the profit was linear and the risk was defined as the variance of the securities. The Markowitz model can be formulated as a linearly constrained optimization problem with two objective (linear profit and quadratic risk) functions.

In the general case, the least risky portfolio is not the most profitable one; thus we could not optimize the two objectives at the same time. Therefore, we need to find portfolios, where one of the goals cannot be improved without worsening the

other. This kind of solutions are called *Pareto optimal* or *Pareto efficient* solutions discussed by Pareto [30].

Single Pareto optimal solution of the Markowitz model can be computed by scalarization methods, see Luc [23] and Miettinen [28], where the weighted sum of the objective functions, serving as a new objective function, defines a single quadratic objective function. Therefore, after the scalarization of the Markowitz model, the optimization problem simplifies to a quadratic optimization problem over linear constraints [25]. The simplified problem's optimal solution is a single Pareto optimal solution of the original problem. The effect of the weights of the objective functions determine the computed Pareto efficient solution of the original problem. The weights might have unpredictable effects on the computed Pareto efficient solution in general. Weakness of this approach is that it restricts the Pareto efficient solution set to a single element and its local neighborhood. In this way, we lose some information, like how much extra profit can be gained by accepting a larger risk. Finding - or at least approximating - the whole Pareto efficient solution set of the original, multi-objective problem, may lead to a better understanding of the modeled practical problem [27].

For some unconstrained multi-objective optimization problems there are research papers [7, 8, 13, 34] discussing algorithms applicable for approximating the Pareto efficient solution set. However, many multi-objective optimization problems - naturally - have constraints [13, 14]. A simple example for a constrained multi-objective optimization problem is the earlier mentioned Markowitz model. In this paper, we extend and generalize the algorithm of Dellnitz *et al.* [7] for approximating the Pareto efficient set of a linearly constrained convex multi-objective problem. Fliege and Svaiter [13] obtained some theoretical results, that are similar to our approach for finding joint decreasing directions.

In the next section, most important definitions and results of vector optimization problems are summarized. In the third section, we discuss some results about the unconstrained vector optimization problem. The method called subdivision technique introduced by Dellnitz *et al.* [7, 8] was developed to approximate the Pareto efficient solution set of an unconstrained vector optimization problems. The subdivision method uses some results described in [34]. An important ingredient of all methods, that can approximate the Pareto optimal set of a convex vector optimization problem, is the computation of a joint decreasing direction for all objective functions. We show that - using results from linear optimization - a joint decreasing direction for an unconstrained vector optimization problem can be computed.

In the fourth section, the computation of a feasible joint decreasing direction for linearly constrained convex vector optimization problem is discussed. The set of feasible joint decreasing directions forms a finitely generated cone and can be computed, as shown in Section 4. Interesting optimality conditions of Eichfelder and Ha [10] for multi-objective optimization problems show some similarities to those that are used in this paper during the computations of joint decreasing directions; however their results do not have practical, algorithmic applications yet.

Section 5 contains an algorithm, which is a generalization of the subdivision method

for the linearly constrained convex vector optimization problem. In Section 6, we show some numerical results obtained on a real data set (securities from Budapest Stock Exchange) for the Markowitz model.

Comparing our method presented in Section 5, with the subdivision algorithm of Dellnitz et al. [7, 8], clearly, our method works for unconstrained vector optimization problems (UVOP), like that of Dellnitz et al. [7, 8], but we generate different joint decreasing directions. Furthermore, our method is applicable to vector optimization problems (VOP) with convex objective functions and linear constraints, keeping all advantageous properties of subdivision algorithm of Dellnitz et al. [7, 8].

For the (VOP) there are some sophisticated scalarization methods reported in [15]. This method, as scalarization methods in general, finds a single Pareto optimal solution. The scalarization method introduced by [15] defines a weighted optimization problem (WOP) of (VOP) with fixed weights and a feasible solution set, that depends on the current feasible vector $\mathbf{x}$. Thus, in each iteration of the algorithm the actual feasible solution set is restricted to such a subset of the original feasible solution set, where some Pareto optimal solutions are located.

Subdivision techniques - including ours - find a cover set of the Pareto optimal solution set, formed by boxes used in the subdivision procedure (see Sections 5 and 6). The approximation of the whole Pareto optimal solution set is controlled by the diameter of the covering boxes computed in the subdivision procedure of the algorithm.

Although both the scalarization algorithm of Gianessi et al. [15] and our subdivision algorithm have some similarities (i.e., in each iteration the feasible solution set decreases), there are significant differences as well. The scalarization algorithm finds a single Pareto optimal solution under some quite general assumptions, while the subdivision algorithm approximates the whole Pareto optimal solution set for the (VOP) with convex objective functions and linear constraints.

We use the following notations throughout the paper: scalars and indices are denoted by lowercase Latin letters, column vectors by lowercase boldface Latin letters, matrices by capital Latin letters, and finally sets by capital calligraphic letters.

The vector, with all 1 elements is denoted by $\mathbf{e}$, i.e.

$$\mathbf{e}^T := (1, 1, \dots, 1) \in \mathbb{R}^n,$$

for some $n \in \mathbb{N}$, where T stands for the transpose of a (column) vector (or a matrix). Vector $\mathbf{e}_i \in \mathbb{R}^n$ is the i th unit vector of the n dimensional Euclidean space.

2 Basic Definitions and Results in Vector Optimization

In this section, we discuss some notations, define the vector (or multi-objective) optimization problem and the concept of Pareto optimal solutions. Furthermore, we state two well known results of vector optimization, which are important for our approach.

We define the *unit simplex set*, as,

Definition 1. Let $\mathcal{S}_k$ denote the unit simplex in the k dimensional vector space, and define it as follows:

$$\mathcal{S}_k := \{\mathbf{w} \in \mathbb{R}^k : \mathbf{e}^T \mathbf{w} = 1, \mathbf{w} \geq \mathbf{0}\}.$$

Let $\mathcal{F} \subseteq \mathbb{R}^n$ be a set and $F : \mathcal{F} \rightarrow \mathbb{R}^k$ is a function defined as

$$F(\mathbf{x}) = [f_1(\mathbf{x}), f_2(\mathbf{x}), \dots, f_k(\mathbf{x})]^T,$$

where $f_i : \mathcal{F} \rightarrow \mathbb{R}$ is a coordinate function for all i .

The *general vector optimization problem* (GVOP) can be formulated as

$$(GVOP) \quad \text{MIN } F(\mathbf{x}), \text{ subject to } \mathbf{x} \in \mathcal{F}.$$

If the set $\mathcal{F}$ and the function F are convex, then (GVOP) is a *convex vector optimization problem*. We assume that F is differentiable.

Usually, different objective functions of (GVOP) describe conflicting goals, therefore such $\mathbf{x} \in \mathcal{F}$, that minimize all objective functions at the same time, is unlikely to exist. For this reason, the following definitions naturally extend the concept of an optimal solution for (GVOP) settings.

Definition 2. Let (GVOP) be given. We say that $\mathbf{x}^* \in \mathcal{F}$ is a

1. weakly Pareto optimal solution of problem (GVOP) if there does not exist a feasible solution $\mathbf{x} \in \mathcal{F}$ which satisfies the vector inequality $F(\mathbf{x}) < F(\mathbf{x}^*)$;
2. Pareto optimal solution if does not exist feasible solution $\mathbf{x} \in \mathcal{F}$ which satisfies the vector inequality $F(\mathbf{x}) \leq F(\mathbf{x}^*)$ and $F(\mathbf{x}) \neq F(\mathbf{x}^*)$.

Furthermore, we call the set $\mathcal{F}^* \subseteq \mathcal{F}$ a weakly Pareto optimal set if every $x^* \in \mathcal{F}^*$ is a weakly Pareto optimal solution of the (GVOP).

Our goal is to approximate the whole Pareto optimal or weakly Pareto optimal solution sets for different vector optimization problems. During the approximation procedure of the whole Pareto optimal solution set, we compute many Pareto optimal solutions and produce an outer approximation of the whole Pareto optimal solution set.

The literature contains several methods that, find one of the Pareto optimal solutions, see [9, 23, 28], but sometimes it is interesting to compute all of them, or at least as much as we can.

One of the frequently used method to compute a Pareto optimal solution uses a weighted sum of the objective functions to obtain a single objective optimization problem. Let $\mathbf{w} \in \mathcal{S}_k$ be a given vector of weights. From a vector optimization problem, using a vector of weights, we can define the weighted optimization problem as follows

$$(WOP) \quad \min \mathbf{w}^T F(\mathbf{x}), \text{ subject to } \mathbf{x} \in \mathcal{F}.$$

We state without proof two well-known theorems, that describe the relationship between (GVOP) and (WOP). The first theorem shows that the (WOP) can be used to find a Pareto optimal solution, see for instance [9, 23, 28].

Theorem 1. *Let a (GVOP) and the corresponding (WOP) for a $\mathbf{w} \in \mathcal{S}_k$ be given. Assume that $\mathbf{x}^* \in \mathcal{F}$ is an optimal solution of the (WOP) problem; then $\mathbf{x}^*$ is a weak Pareto optimal solution for the (GVOP).*

Next theorem needs a bit more complicated reasoning, but for the convex case each Pareto optimal solution of the (GVOP) can be found through a (WOP) using the proper weights [9, 23, 28].

Theorem 2. *Let (GVOP) be a convex vector optimization problem, and assume that $\mathbf{x}^* \in \mathcal{F}$ is a Pareto optimal solution of the (GVOP); then there is a $\mathbf{w} \in \mathcal{S}_k$ weight vector, and a (WOP) problem, for which $\mathbf{x}^*$ is an optimal solution.*

The method, that will be described in section 5, decreases every coordinate function of F at the same time and always moves from a feasible solution to another feasible solution; hence we introduce the following useful definition.

Definition 3. *Let problem (GVOP) and feasible point $\mathbf{x} \in \mathcal{F}$ be given. Vector $\mathbf{v} \in \mathbb{R}^n$, $\mathbf{v} \neq \mathbf{0}$ is called a*

1. *joint decreasing direction at point $\mathbf{x}$ iff there exists $h_0 > 0$ for every $h \in]0, h_0]$ satisfying that*

$$F(\mathbf{x} + h\mathbf{v}) < F(\mathbf{x});$$

2. *feasible joint decreasing direction iff it is a joint decreasing direction and there exists $h_1 > 0$ such that, for every $h \in]0, h_1]$ we have $\mathbf{x} + h\mathbf{v} \in \mathcal{F}$.*

Example. Let the following unconstrained vector optimization problem

$$(GVOP_1) \quad \text{MIN } F(x_1, x_2) = \begin{pmatrix} f_1(x_1, x_2) = x_1^2 + x_2^2 \\ f_2(x_1, x_2) = (x_1 - 1)^2 + (x_2 - 1)^2 \end{pmatrix},$$

be given alongside a point $\mathbf{x}^T = (x_1, x_2) = (0, 1)$ and direction $\mathbf{v}^T = (1, -1)$. Now we show that $\mathbf{v}$ is a joint decreasing direction for the objective function F at point $\mathbf{x}$.

It is easy to show that

$$f_1(\mathbf{x} + h\mathbf{v}) = f_2(\mathbf{x} + h\mathbf{v}) = h^2 + (1 - h)^2 = 2 \left(h - \frac{1}{2} \right)^2 + \frac{1}{2}$$

From the last form of the coordinate functions, it is easy to see that the coordinate functions are decreasing on the $[0; \frac{1}{2}]$ interval; therefore $\mathbf{v}$ is a joint decreasing direction with $h_0 = \frac{1}{2}$.

If we add a single constraint to our example, then we obtain a new problem

$$(GVOP_2) \quad \text{MIN } F(x_1, x_2) = \begin{pmatrix} f_1(x_1, x_2) = x_1^2 + x_2^2 \\ f_2(x_1, x_2) = (x_1 - 1)^2 + (x_2 - 1)^2 \end{pmatrix},$$

$$\text{subject to } x_1 \leq \frac{1}{3}.$$

It is easy to see that $\mathbf{v}$ is a feasible joint decreasing direction for problem $(GVOP_2)$ too, with $h_1 = \frac{1}{3}$.

From now on, let us consider $(GVOP)$ with a convex, differentiable objective function F and let us denote the Jacobian-matrix of F at point $\mathbf{x}$ by $J(\mathbf{x})$. Then, $\mathbf{v} \in \mathbb{R}^n$ is a joint decreasing direction of function F at point $\mathbf{x}$, if and only if

$$[J(\mathbf{x})]\mathbf{v} < 0, \tag{1}$$

as $\mathbf{v}$ is a decreasing direction for the i th coordinate function f_i at point $\mathbf{x}$, if and only if $[\nabla f_i(\mathbf{x})]^T \mathbf{v} < 0$.

3 Results for Unconstrained Vector Optimization

In this section, we review some results of unconstrained vector optimization, namely for $\mathcal{F} = \mathbb{R}^n$. We assume that F is a differentiable function. The unconstrained vector optimization problem is denoted by $(UVOP)$.

Before we show how a joint decreasing direction can be computed, we need a criterion to decide whether an $\mathbf{x}$ is a Pareto optimal solution or not, see Schäffler et al. [34].

Definition 4. Let $J(\mathbf{x}) \in \mathbb{R}^{k \times n}$ be the Jacobian matrix of a differentiable function $F : \mathbb{R}^n \rightarrow \mathbb{R}^k$ at a point $\mathbf{x} \in \mathbb{R}^n$. An $\mathbf{x}^*$ is called substationary point of F iff there exist a $\mathbf{w} \in \mathcal{S}_k$, which fulfills the following equation:

$$\mathbf{w}^T [J(\mathbf{x}^*)] = \mathbf{0}.$$

In the unconstrained case, point $\mathbf{x}^*$ is a substationary point of the objective function F of $(GVOP)$, if it is a stationery point of the weighted objective function of (WOP) . From Theorem 1 and Theorem 2 we can see that in convex case, substationary points are weak Pareto optimal solutions of the unconstrained vector optimization problem $(GVOP)$.

We are ready to discuss two models to find joint decreasing directions. The first model has been discussed in [34] and uses a quadratic programming problem formulation to compute joint decreasing directions. Later we show that a joint decreasing direction can be computed in a simpler way by using a special linear programming problem.

Let us define the following quadratic programming problem ($QOP(\mathbf{x})$) for any $\mathbf{x} \in \mathbb{R}^n$, with variable $\mathbf{w}$

$$(QOP(\mathbf{x})) \quad \min \mathbf{w}^T \left(J(\mathbf{x}) [J(\mathbf{x})]^T \right) \mathbf{w}, \text{ subject to } \mathbf{w} \in \mathcal{S}_k.$$

From the well known Weierstarss Theorem it follows that this problem always has an optimal solution, since the feasible set is compact and the function

$$g : \mathcal{S}_k \rightarrow \mathbb{R}, \quad g(\mathbf{w}) = \mathbf{w}^T \left(J(\mathbf{x}) [J(\mathbf{x})]^T \right) \mathbf{w}$$

is a convex, quadratic, continuous function for any given $\mathbf{x} \in \mathbb{R}^n$.

Next theorem is an already known statement [34, Theorem 2.1], for which we give a new and shorter proof. This shows that using the ($QOP(\mathbf{x})$) problem we can find a joint decreasing direction of F or a certificate that $\mathbf{x}$ is a Pareto optimal solution of problem ($UVOP$).

Theorem 3. *Let a problem ($UVOP$), a point $\mathbf{x} \in \mathbb{R}^n$ and the associated ($QOP(\mathbf{x})$) be given. Let $\mathbf{w}^* \in \mathbb{R}^k$ denote the optimal solution of ($QOP(\mathbf{x})$). We define vector $\mathbf{q} \in \mathbb{R}^n$ as $\mathbf{q} = [J(\mathbf{x})]^T \mathbf{w}^*$. If $\mathbf{q} = \mathbf{0}$, then $\mathbf{x}$ is a substationary point, otherwise $-\mathbf{q}$ is a joint decreasing direction for F at point $\mathbf{x}$.*

Proof. When $\mathbf{q} = \mathbf{0}$ then Definition 4 shows that $\mathbf{x}$ is substationary point. When $\mathbf{q} \neq \mathbf{0}$, we indirectly assume that $-\mathbf{q}$ is not a decreasing direction for the i th coordinate function, f_i of F and $[\nabla f_i(\mathbf{x})]^T \mathbf{q} \neq \mathbf{0}$. It means that $[\nabla f_i(\mathbf{x})]^T \mathbf{q} < 0$. Since $[\nabla f_i(\mathbf{x})]^T = \mathbf{e}_i^T J(\mathbf{x})$, so our indirect assumption means

$$[\nabla f_i(\mathbf{x})]^T \mathbf{q} = \mathbf{e}_i^T [J(\mathbf{x})] [J(\mathbf{x})]^T \mathbf{w}^* < 0.$$

We show that $\mathbf{e}_i - \mathbf{w}^* \neq \mathbf{0}$ is a feasible decreasing direction of $g(\mathbf{w}^*)$ which contradicts the optimality of $\mathbf{w}^*$. The $\mathbf{e}_i = \mathbf{w}^*$ can not be fulfilled because it contradicts the indirect assumption, and it is easy to see, that $\mathbf{e}_i$ is a feasible solution of ($QOP(\mathbf{x})$) so $\mathbf{e}_i - \mathbf{w}^*$ is a feasible direction at point $\mathbf{w}^*$.

Since

$$\nabla g(\mathbf{w}) = 2[J(\mathbf{x})][J(\mathbf{x})]^T \mathbf{w}$$

thus

$$\begin{aligned} [\nabla g(\mathbf{w}^*)]^T (\mathbf{e}_i - \mathbf{w}^*) &= 2\mathbf{w}^{*T} [J(\mathbf{x})] [J(\mathbf{x})]^T (\mathbf{e}_i - \mathbf{w}^*) \\ &= 2\mathbf{w}^{*T} [J(\mathbf{x})] [J(\mathbf{x})]^T \mathbf{e}_i - 2\mathbf{w}^{*T} [J(\mathbf{x})] [J(\mathbf{x})]^T \mathbf{w}^* \\ &= 2\mathbf{q}^T [\nabla f_i(\mathbf{x})] - 2\mathbf{w}^{*T} [J(\mathbf{x})] [J(\mathbf{x})]^T \mathbf{w}^* < 0, \end{aligned}$$

where the first term of the sum is negative because of the indirect assumption, and the second term is not positive, because $[J(\mathbf{x})][J(\mathbf{x})]^T$ is a positive semidefinite matrix. $\square$

The previous result underline the importance of solving ($QOP(\mathbf{x})$) problem efficiently. For solving smaller size linearly constrained convex quadratic problems

pivot algorithms [1, 4–6, 16, 22] can be used. In case of larger size linearly constrained, convex quadratic problems, interior point algorithms can be used to solve the problem (see for instance [18, 20]).

Theorem 3 shows that a joint decreasing direction can be computed as the convex combination of the gradient vectors of coordinate functions of F . Following the ideas discussed above, we can formulate a linear programming problem such that any optimal solution of the linear program defines a joint decreasing direction of problem (UVOP). Some similar results can be found in [13].

Let us define the linear optimization problem ($LP(\mathbf{x})$) in the following way:

$$(LP(\mathbf{x})) \quad \max q_0, \quad \text{subject to} \quad [J(\mathbf{x})]\mathbf{q} + q_0\mathbf{e} \leq \mathbf{0}, \quad 0 \leq q_0 \leq 1,$$

where $\mathbf{q} \in \mathbb{R}^n$ and $q_0 \in \mathbb{R}$ are the decision variables of the problem $LP(\mathbf{x})$. Now we are ready to state and prove a theorem that discusses a connection between (UVOP) and ($LP(\mathbf{x})$).

Theorem 4. *Let a point $\mathbf{x} \in \mathbb{R}^n$, an (UVOP) and an associated ($LP(\mathbf{x})$) be given. Then the ($LP(\mathbf{x})$) always has an optimal solution $(\mathbf{q}^*, q_0^*)$. There are two cases for the optimal value of the ($LP(\mathbf{x})$), either $q_0 = 0$ thus $\mathbf{x}$ is a substationary point of the (UVOP), or $q_0 = 1$ thus $\mathbf{q}^*$ is a joint decreasing direction for the function F at point $\mathbf{x}$.*

Proof. It is easy to see that $\mathbf{q} = \mathbf{0}$, $q_0 = 0$ is a feasible solution of problem ($LP(\mathbf{x})$) and 1 is an upper bound of the objective function, which means ($LP(\mathbf{x})$) should have an optimal solution.

Let us examine the case

$$[J(\mathbf{x})]\mathbf{q} + q_0\mathbf{e} \leq \mathbf{0}, \quad q_0 > 0. \quad (2)$$

If system (2) has a solution, then $\left(\frac{1}{q_0}\mathbf{q}, 1\right)$ is a solution of the system, so the optimal value of the objective function is 1. This means that

$$[J(\mathbf{x})]\mathbf{q} \leq -\mathbf{e}$$

so the $\mathbf{q}$ is a joint decreasing direction of function F .

If the system (2) has no solution then the optimal value of the objective function is 0, and from the Farkas lemma ([11, 12, 17, 22, 29, 31, 32, 35]) we know that there exists a $\mathbf{w}$ which satisfies the following:

$$\mathbf{w}^T [J(\mathbf{x})] = \mathbf{0}, \quad \mathbf{e}^T \mathbf{w} = 1, \quad \mathbf{w} \geq \mathbf{0}. \quad (3)$$

It means that if the optimal value of the problem ($LP(\mathbf{x})$) is 0, then $\mathbf{x}$ is a substationary point. $\square$

A linear programming problem ($LP(\mathbf{x})$) (and later on ($LPS(\mathbf{x})$)) can be solved by either pivot or interior point algorithms, see [21]. In case of applying pivot methods to solve linear programming problem, simplex algorithm is a natural choice, see [24,

29, 35]. A recent study on anti-cycling pivot rules for linear programming problem contains a numerical study on different pivot algorithms, see [6]. Sometimes, if the problem is well structured and small, criss-cross algorithm of T. Terlaky can be used for solving the linear programming problem as well, see [17, 36]. More about interior point algorithms for linear programming problems can be learnt from [19, 24, 33].

In this section two techniques were introduced to decide whether a point $\mathbf{x}$ is a weak Pareto optimal solution of problem (*UVOP*) or to find a joint decreasing direction. Before we generalize this result to the linear constrained case let us compare this technique with some known procedures. The classical scalarization technique based on (*WOP*) finds a Pareto optimal solution $\mathbf{x}^*$ of (*UVOP*). Due to the requirements defined by the concept of the weak Pareto optimal solution, it may happen that in some iterations of the scalarization algorithm such feasible solutions are computed for which some objective function's value increases. In our method this phenomena can not happen, because in each iteration we select a joint decreasing direction.

4 Vector Optimization with Linear Constraints

In this section, we show how we can find a feasible joint decreasing direction for linearly constrained vector optimization problems. First we find a feasible joint decreasing direction for a special problem, where we only have sign constraints on the variables. After that we generalize our results to general linearly constrained vector optimization problems. Our method can be considered as the generalization of the well known reduced gradient method to vector optimization problems. Some similar result can be found in [13], for the feasible direction method of Zoutendijk.

First we define the vector optimization problem with sign constraints (*SVOP*):

$$(\text{SVOP}) \quad \min F(\mathbf{x}), \quad \text{subject to } \mathbf{x} \geq \mathbf{0},$$

where F is a convex function. From Theorem 1 we know that $\mathbf{x}^* \geq \mathbf{0}$ is a Pareto optimal solution if there exists a $\mathbf{w} \in \mathcal{S}_k$ vector such that $\mathbf{x}^*$ is an optimal solution of

$$(\text{SWOP}) \quad \min \mathbf{w}^T F(\mathbf{x}), \quad \text{subject to } \mathbf{x} \geq \mathbf{0}.$$

Since Slater regularity and convexity conditions hold, from the KKT theorem [24] we know that $\mathbf{x}^* \geq \mathbf{0}$ is an optimal solution of (*SWOP*) iff it satisfies the following system:

$$\mathbf{w}^T [J(\mathbf{x}^*)] \geq \mathbf{0}, \quad \mathbf{w}^T [J(\mathbf{x}^*)] \mathbf{x}^* = 0. \quad (4)$$

Let the vector $\mathbf{x} \geq \mathbf{0}$ be given and we would like to decide whether it is an optimal solution of the (*SWOP*) problem or not. Let us define the index sets

$$I_+ = \{i : x_i > 0\}, \quad \text{and} \quad I_0 = \{i : x_i = 0\},$$

that depend on the selected vector $\mathbf{x}$. Using the index sets I_0, I_+ , we partition the column vectors of matrix $J(\mathbf{x})$ into two parts. The two parts are denoted by $J(\mathbf{x})_{I_0}$ and $J(\mathbf{x})_{I_+}$. Taking into consideration the partition, the KKT conditions can be written in an equivalent form as

$$\mathbf{w}^T [J(\mathbf{x})]_{I_+} = \mathbf{0}, \quad \mathbf{w}^T [J(\mathbf{x})]_{I_0} \geq \mathbf{0}, \quad \mathbf{w} \in \mathcal{S}_k. \quad (5)$$

The inequality system (5) plays the same role for (SVOP) as (3) for (UVOP), namely $\mathbf{x}$ is a Pareto-optimal solution if (5) has a solution.

Now we can define a linear programming problem corresponding to (SVOP) such that an optimal solution of the linear programming problem either defines a joint decreasing direction or gives a certificate that the solution $\mathbf{x}$ is a Pareto optimal solution of (SVOP).

$$\begin{aligned} (LPS(\mathbf{x})) \quad & \max z, \\ \text{subject to} \quad & [J(\mathbf{x})]_{I_+} \mathbf{u} + [J(\mathbf{x})]_{I_0} \mathbf{v} + z \mathbf{e} \leq \mathbf{0}, \\ & \mathbf{v} \geq \mathbf{0}, \quad 0 \leq z \leq 1, \end{aligned}$$

where $\mathbf{u}, \mathbf{v}$ and z are the decision variables of problem (LPS($\mathbf{x}$)). Now we are ready to prove the following theorem.

Theorem 5. *Let a (SVOP) and an associated (LPS($\mathbf{x}$)) be given, where $\mathbf{x} \in \mathcal{F}$ is a feasible point. The problem (LPS($\mathbf{x}$)) always has an optimal solution $(\mathbf{u}^*, \mathbf{v}^*, z^*)$. There are two cases for the optimal value of problem (LPS($\mathbf{x}$)), $z^* = 0$ which means that $\mathbf{x}$ is a Pareto optimal solution of the (SVOP), or $z^* = 1$ which means that $\mathbf{q}^T = (\mathbf{u}^*, \mathbf{v}^*)$ is a feasible joint decreasing direction of function F .*

Proof. It is easy to see that $\mathbf{u} = \mathbf{0}, \mathbf{v} = \mathbf{0}, z = 0$ is a feasible solution of the problem (LPS($\mathbf{x}$)) and 1 is an upper bound of the objective function, therefore (LPS($\mathbf{x}$)) has an optimal solution.

Let us examine the following system

$$[J(\mathbf{x})]_{I_+} \mathbf{u} + [J(\mathbf{x})]_{I_0} \mathbf{v} + z \mathbf{e} \leq \mathbf{0}, \quad \mathbf{v} \geq \mathbf{0}, \quad z > 0. \quad (6)$$

If system (6) has a solution, then $(\frac{1}{z} \mathbf{u}, \frac{1}{z} \mathbf{v}, 1)$ is an optimal solution of the problem (LPS($\mathbf{x}$)) with optimal value 1. Thus the vector $\mathbf{q}^T = (\mathbf{u}, \mathbf{v})$ satisfies

$$[J(\mathbf{x})] \mathbf{q} \leq -\mathbf{e} < \mathbf{0},$$

so the $\mathbf{q}$ is a feasible joint decreasing direction for function F at $\mathbf{x} \in \mathcal{F}$. Vector $\mathbf{q}$ is feasible because $\mathbf{q}_{I_0} = \mathbf{v} \geq \mathbf{0}$.

If the system (6) has no solution then the optimal value of the objective function is 0, and from a variant of the Farkas lemma, see [11, 12, 17, 22, 31, 32, 35] we know that there exists a $\mathbf{w}$ which satisfies the following system of inequalities:

$$[J(\mathbf{x})]_{I_+}^T \mathbf{w} = \mathbf{0}, \quad [J(\mathbf{x})]_{I_0}^T \mathbf{w} \geq \mathbf{0}, \quad \mathbf{e}^T \mathbf{w} = 1, \quad \mathbf{w} \geq \mathbf{0}. \quad (7)$$

It means that if the optimal value of problem $(LPS(\mathbf{x}))$ is 0, then point $\mathbf{x}$ is a Pareto optimal solution of $(SVOP)$. $\square$

We are ready to find feasible joint decreasing direction to a linearly constrained vector optimization problem at a feasible solution $\tilde{\mathbf{x}}$. Let the matrix $A \in \mathbb{R}^{m \times n}$ and vector $\mathbf{b} \in \mathbb{R}^m$ be given. Without loss of generality we may assume that $\text{rank}(A) = m$. Furthermore, let us assume the following non degeneracy assumption (for details see [2]): any m columns of A are linearly independent and every basic solution is non degenerate. We have a vector optimization problem with linear constraint in the following form

$$(LVOP) \quad \text{MIN } F(\mathbf{x}), \quad \text{subject to } A\mathbf{x} = \mathbf{b}, \quad \mathbf{x} \geq \mathbf{0}.$$

Like in the reduced gradient method, see [2], we can partition the matrix A into two parts $A = [B, N]$, where B is a basic and N the non-basic part of the matrix. Similarly every $\mathbf{v} \in \mathbb{R}^n$ vector can be partitioned as $\mathbf{v} = [\mathbf{v}_B, \mathbf{v}_N]$. We call $\mathbf{v}_B$ basic and $\mathbf{v}_N$ a nonbasic vector. We can chose the matrix B such that the $\tilde{\mathbf{x}}_B > 0$ is fulfilled. While $A\mathbf{x} = \mathbf{b}$ holds, we know that

$$B\mathbf{x}_B + N\mathbf{x}_N = \mathbf{b}, \quad \text{and} \quad \mathbf{x}_B = B^{-1}(\mathbf{b} - N\mathbf{x}_N).$$

We can redefine function F in a reduced form as

$$F_N(\mathbf{x}_N) = F(\mathbf{x}_B, \mathbf{x}_N) = F(B^{-1}(\mathbf{b} - N\mathbf{x}_N), \mathbf{x}_N).$$

Let us define using the partition (B, N) and the following sign constraint optimization problem

$$(SVOP_B(\tilde{\mathbf{x}})) \quad \text{MIN } F_N(\mathbf{x}_N), \quad \text{subject to } \mathbf{x}_N \geq \mathbf{0}.$$

Let $\mathbf{q}_N$ denote a feasible joint decreasing direction for $(SVOP_B(\tilde{\mathbf{x}}))$ at point $\tilde{\mathbf{x}}_N$, which can be found by applying Theorem 5. Let $\mathbf{q}_B = -B^{-1}N\mathbf{q}_N$; then we show that $\mathbf{q} = [\mathbf{q}_B, \mathbf{q}_N]$ is feasible joint decreasing direction for $(LVOP)$ at point $\tilde{\mathbf{x}}$. Let us notice that

$$A(\tilde{\mathbf{x}} + h\mathbf{q}) = A\tilde{\mathbf{x}} + h(B\mathbf{q}_B + N\mathbf{q}_N) = \mathbf{b} + h(-B(B^{-1}N\mathbf{q}_N) + N\mathbf{q}_N) = \mathbf{b},$$

for every $h \in \mathbb{R}$. So $F_N(\tilde{\mathbf{x}} + h\mathbf{q}_N) = F(\tilde{\mathbf{x}} + h\mathbf{q})$ and while $\mathbf{q}_N$ is a feasible joint decreasing direction with $h_1 > 0$ stepsize for $(SVOP_B(\tilde{\mathbf{x}}))$. We can show that $\mathbf{q}$ is a joint decreasing direction for problem $(LVOP)$ with the same $h_1 > 0$ step size. While $\tilde{\mathbf{x}}_B > 0$, there exists $h_2 > 0$, such that $\tilde{\mathbf{x}}_B + h_2\mathbf{q}_B \geq 0$, therefore $\mathbf{q}$ is a feasible joint decreasing direction for problem $(LVOP)$ with a step-size

$$h_3 = \min(h_1, h_2) > 0. \quad (8)$$

We can compute h_1 and h_2 using a ratio-test, since $\tilde{x}_i + hq_i \geq 0$ for all i is required, therefore

$$h_1 = \min \left\{ -\frac{\tilde{x}_i}{q_i} : q_i < 0, \quad i \in \mathcal{N} \right\} > 0, \text{ and}$$

$$h_2 = \min \left\{ -\frac{\tilde{x}_i}{q_i} : q_i < 0, \quad i \in \mathcal{B} \right\} > 0.$$

5 The Subdivision Algorithm for Linearly Constrained Vector Optimization Problem

In this section, we show how can we build a subdivision method to approximate the Pareto optimal set of a linearly constrained vector optimization problem. Our method is a generalization of the algorithm discussed in [7], where you can find some results about convergence of the subdivision technique. The original method can not handle linear constraints.

Our algorithm approximates the Pareto optimal solution set $\mathcal{F}^*$, using small boxes that each contains at least one computed Pareto optimal solution. The smaller the sets, the better approximation of the $\mathcal{F}^*$, therefore we define the following measure of sets involved in the approximation of $\mathcal{F}^*$.

Definition 5. Let $\mathcal{H} \subseteq \mathbb{R}^n$ be given; the diameter of $\mathcal{H}$ is defined as

$$\text{diam}(\mathcal{H}) := \sup_{x, y \in \mathcal{H}} \|x - y\|.$$

Let $\mathbb{H}$ be a family of sets, which contains a finite number of sets from $\mathbb{R}^n$; then the diameter of $\mathbb{H}$ is

$$\text{diam}(\mathbb{H}) = \max_{\mathcal{H} \in \mathbb{H}} \text{diam}(\mathcal{H}).$$

Let us assume that the feasible set of our problem is nonempty, closed and bounded. In this case Pareto optimal solution set of the problem (LVOP) is a nonempty set, defined as,

$$\mathcal{F}_L^* = \{\mathbf{x} \in \mathcal{F} : \mathbf{x} \text{ is Pareto optimal solution of (LVOP)}\}.$$

Based on our assumption, that $\mathcal{F}$ is bounded set, there exists

$$\mathcal{H}_0 = \{\mathbf{x} \in \mathbb{R}^n | \mathbf{l} \leq \mathbf{x} \leq \mathbf{u}\},$$

where $\mathbf{l}, \mathbf{u} \in \mathbb{R}^n$ are given vectors and

$$\mathcal{F}_L^* \subseteq \mathcal{F} \subseteq \mathcal{H}_0 \cap \{\mathbf{x} \in \mathbb{R}^n : A\mathbf{x} = \mathbf{b}\}.$$

Our goal is to introduce such a subdivision algorithm for (LVOP) problem, that iteratively defines better and better inner and outer approximation of the Pareto optimal solution set of a problem (LVOP). The inner approximation will be an increasing sequence of sets $\mathcal{F}_i$, containing finitely many Pareto optimal solutions produced by the algorithm. The outer approximation will be a family of sets $\mathbb{H}_i$ produced in each iteration of the algorithm, that covers $\mathcal{F}_L^*$ with decreasing diameter, $\text{diam}(\mathbb{H}_i)$.

The input of our method are data of the (LVOP) problem, namely matrix $A \in \mathbb{R}^{m \times n}$, vector $\mathbf{b} \in \mathbb{R}^m$, function $F: \mathbb{R}^n \rightarrow \mathbb{R}^k$, set $\mathcal{H}_0$ and a constant $\varepsilon > 0$.

The output of our algorithm is a family of sets $\mathbb{H}$, such that $\text{diam}(\mathbb{H}) < \varepsilon$ and each $\mathcal{H} \in \mathbb{H}$ contains at least one Pareto optimal solution.

The algorithm uses some variables and subroutines, too. The $\mathcal{SP}$, $\mathcal{FP}$ are finite-element sets of points from $\mathbb{R}^n$, $\mathcal{H}, \mathcal{G} \subseteq \mathcal{F}$, $\mathbb{H}', \mathbb{K}, \mathbb{K}'$ and $\mathbb{A}$ are family of sets like $\mathbb{H}$.

Our algorithm in the first step defines the family of sets $\mathbb{H}$, which contains only the $\mathcal{H}_0$ set. The cycle in step 2 runs while the diameter of $\mathbb{H}$ is not small enough. The algorithm reaches this goal in a finite number of iteration, because as you will see in subroutine Newset($\mathbb{H}$) the diameter of $\mathbb{H}$ converges to zero. Nevertheless we show that after every execution of the cycle the family of sets $\mathbb{H}$ contains sets $\mathcal{H}$ which have Pareto optimal solutions. At the beginning it is trivial, because $\mathbb{H}$ contain the whole feasible set.

Subdivision algorithm for (LVOP)

1. $\mathbb{H} = \{\mathcal{H}_0\}$
 2. **While** $diam(\mathbb{H}) \geq \varepsilon$ **do**
 - (a) $\mathbb{H}' = \text{Newsets}(\mathbb{H})$
 - (b) $\mathcal{SP} = \emptyset$
 - (c) **While** $\mathbb{H} \neq \emptyset$ **do**
 - i. $\mathcal{H} \in \mathbb{H}$
 - ii. $\mathcal{SP} = \mathcal{SP} \cup \text{Startpoint}(\mathcal{H})$
 - iii. $\mathbb{H} = \mathbb{H} \setminus \{\mathcal{H}\}$
 - End While**
 - (d) $\mathcal{FP} = \text{Points}(\mathcal{SP}, A, \mathbf{b}, F)$
 - (e) **While** $\mathbb{H}' \neq \emptyset$ **do**
 - i. $\mathcal{H} \in \mathbb{H}'$
 - ii. **If** $\mathcal{H} \cap \mathcal{FP} \neq \emptyset$ **then** $\mathbb{H} = \mathbb{H} \cup \{\mathcal{H}\}$ **End If**
 - iii. $\mathbb{H}' = \mathbb{H}' \setminus \{\mathcal{H}\}$
 - End While**
 3. **Output**($\mathbb{H}$)
-

In step 2(a) we define a family of sets $\mathbb{H}'$ using the subroutine Newset($\mathbb{H}$). The sets from $\mathbb{H}'$ are smaller than sets from $\mathbb{H}$ and cover the same set. Therefore the result of this subroutine has two important properties:

1. $\bigcup_{\mathcal{H} \in \mathbb{H}'} (\mathcal{H} \cap \mathcal{F}) = \bigcup_{\mathcal{H} \in \mathbb{H}} (\mathcal{H} \cap \mathcal{F})$,
 2. $diam(\mathbb{H}') = \frac{1}{K} diam(\mathbb{H})$,
-

where $K > 1$ is a constant.

The steps in cycle 2, from step 2(b), delete the sets from $\mathbb{H}'$ which do not contain any Pareto optimal points. Step 2(b) makes set $\mathcal{SP}$ empty. The cycle in step 2(c) produces a finite number of random starting points in set $\mathcal{H} \cap \{\mathbf{x} \in \mathbb{R}^n : A\mathbf{x} = \mathbf{b}\}$ for every $\mathcal{H} \in \mathbb{H}$ using subroutine $\text{Startpoint}(\mathcal{H})$, and puts the generated points into the set SP .

The main step of our algorithm 2(d) is the subroutine $\text{Points}(\mathcal{SP}, A, \mathbf{b}, F)$ that produce a set $\mathcal{FP}$ which contains Pareto optimal points. This subroutine uses our results from Section 4.

In cycle 2(e) we keep every set from $\mathbb{H}'$ which contains Pareto optimal solutions and add those to $\mathbb{H}$. Finally, we check the length of the diameter of $\mathbb{H}$ and repeat the cycle until the diameter is larger than the accuracy parameter ε .

Subroutine $\text{Points}(\mathcal{SP}, A, \mathbf{b}, F)$

1. **While** $\mathcal{SP} \neq \emptyset$ **do**
 - (a) $\mathbf{s} \in \mathcal{SP}$
 - (b) $\mathbf{x} = \mathbf{s}, z = 1$
 - (c) **While** $z = 1$ **do**
 - i. $(B, N) = A$
 - ii. $(\mathbf{x}_B, \mathbf{x}_N) = \mathbf{x}$
 - iii. $(\mathbf{q}, z) = \text{Solve}(LPS(\mathbf{x}_N))$
 - iv. **If** $z = 1$ **then**
 - A. $h_3 = \text{stepsize}(F, B, N, \mathbf{b}, \mathbf{x}_N, \mathbf{q})$
 - B. $\mathbf{x}_N = \mathbf{x}_N + h_3 \mathbf{q}$
 - C. $\mathbf{x}_B = \mathbf{b} - B^{-1} N \mathbf{x}_N$
 - D. $\mathbf{x} = (\mathbf{x}_B, \mathbf{x}_N)$
 - End If**
 - End While**
 - (d) $\mathcal{SP} = \mathcal{SP} \setminus \{\mathbf{s}\}$
 - (e) $\mathcal{FP} = \mathcal{FP} \cup \{\mathbf{x}\}$
 - End While**
 2. $\text{Output}(\mathcal{FP})$
-

Subroutine Points uses a version of the reduced gradient method for computing Pareto optimal solutions or joint decreasing directions, as discussed in Section 4.

This subroutine works until from all points in $\mathcal{SP}$ search for a Pareto optimal point has been executed. The cycle 1(c) runs until it finds a Pareto optimal point. As we discussed in section 4 the cycle finishes when $z = 0$. In line 1(c)i the matrix A is partitioned into a basic and a non basic parts, denoted by B and N , respectively. The same partition is made with vector $\mathbf{x}$ according to 1(c)ii, and we choose the basis such that $\mathbf{x}_B > 0$ should be satisfied. The $LCP(\mathbf{x}_N)$ is solved in step 1(c)iii. If the variable $z = 0$ then $\mathbf{x}$ is a Pareto optimal solution and we select a new starting point from $\mathcal{SP}$, unless $\mathcal{SP}$ is empty. Otherwise $\mathbf{q}$ is a feasible joint decreasing direction for the reduced function F_N . In step 1(c)ivB we compute step-size h_3 which was defined in (8), and a new feasible solution $\mathbf{x}$ is computed.

Let us summarize the properties of our subdivision algorithm at the end of this section. In each iteration of the algorithm, a set of Pareto optimal solutions, $\mathcal{SP}_i$, as inner approximation of $\mathcal{F}_L^*$ and a family of box sets $\mathbb{H}_i$, with diameter $diam(\mathbb{H}_i)$, as outer approximation has been produced. The input data for inner and outer approximations of $\mathcal{F}_L^*$ are as follows

$$\mathcal{SP}_0 := \emptyset \quad \text{and} \quad \mathbb{H}_0 := \{\mathcal{H}_0\}.$$

Proposition 1. *Let a problem (LVOP) with nonempty, polytope $\mathcal{F}$ and differentiable, convex objective function F be given. Furthermore, let us assume that $\mathcal{F} \subseteq \mathcal{H}_0$ holds, where $\mathcal{H}_0 \subset \mathbb{R}^n$ is a box set (i.e. generalized interval). Our subdivision algorithm in iteration k produces two outputs:*

- a) a subset of Pareto optimal solutions $\mathcal{SP}_k$, and
- b) a family of box sets $\mathbb{H}_k$,

with the following properties

1. $\mathcal{SP}_i \subseteq \mathcal{SP}_{i+1}$ for all $i = 0, 1, \dots, k-1$,
2. $\cup_{\mathcal{H} \in \mathbb{H}_{i+1}} \mathcal{H} \subseteq \cup_{\mathcal{H} \in \mathbb{H}_i} \mathcal{H}$ for all $i = 0, 1, \dots, k-1$,
3. $diam(\mathbb{H}_{i+1}) = \frac{1}{K} diam(\mathbb{H}_i)$, where $K > 1$ is a constant,
4. $\mathcal{SP}_i \subset \mathcal{F}_L^*$, for all $i = 0, 1, \dots, k$,
5. $\mathcal{F}_L^* \subset \cup_{\mathcal{H} \in \mathbb{H}_i} \mathcal{H}$, for all $i = 0, 1, \dots, k$.

Dellnitz et al. [7] discussed two important issues related to subdivision methods: (i) convergence (see Section 3), and (ii) possibility of deleting box that contains Pareto optimal solution (see paragraph 4.2).

Convergence of subdivision methods according to Dellnitz and his coauthors follows under mild smoothness assumptions of the objective functions and compactness of their domains together with some useful properties of the iteration scheme. All these necessary properties of the objective functions and iteration scheme are satisfied in our case, too. Convergence of our subdivision method is based on similar arguments as in case discussed by Dellnitz et al. [7].

It may occur that a box containing Pareto optimal solution is deleted during the subdivision algorithm, as stated in [7]. This phenomenon is related to the discretization

induced by the iteration scheme. Decision whether keep or delete a box during the course of the algorithm depends on whether we found Pareto optimal solution in that box or not. From each box, finite number of points are selected and tested whether those are Pareto optimal solutions or joint decreasing direction corresponds to them. From those test points that are not Pareto optimal solutions, using joint decreasing direction an iterative process is started that stops with founding a Pareto optimal solution. If all Pareto optimal solutions computed in this way lay out of the box, we may conclude that the box under consideration does not contain Pareto optimal solution. However, this conclusion is based only on finitely many test points thus there is a chance to delete a box even if it contains Pareto optimal solution. This situation, in practice, can be handled by applying different strategies. All these strategies decrease the opportunity of deleting a box containing Pareto optimal solution.

In [7] discussed a recovering algorithm that could be used after the diameter of boxes reached the prescribed $\varepsilon > 0$ accuracy. Their recovering algorithm finds all those boxes with the current diameter that are necessary to ensure that a cover set of the Pareto optimal solutions is obtained. For details, see [7], recovering algorithm (paragraph 4.2).

6 The Markowitz Model and Computational Results

Let us illustrate our method by solving the Markowitz model to find the most profitable and less risky portfolios. The standard way of solving the model is to find one of the Pareto optimal solution with an associated (*WOP*) see [26, 28]. The question is whether such single Pareto optimal solution is what we really need for decision making. Naturally, if we would like to make extra profit, we should accept larger risk. Therefore, a single Pareto optimal solution does not contain enough information for making a practical decision. If we produce or approximate the Pareto optimal solution set then we can make use of the additional information for making more established decision.

The analytical description of the whole Pareto optimal set for the Markowitz model is known [37]. Thus as a test problem, the Markowitz model has the following advantage: it is possible to derive its Pareto optimal solution set in an analytical way [37], therefore the result of our subdivision algorithm can be compared with the analytical description of the Pareto optimal solution set.

We are now ready to formulate the original Markowitz model. Our goal is to make a selection from n different securities. Let x_i denote how much percentage we spend from our budget on security i ($i = 1, 2, \dots, n$), based on more approximate information than a single Pareto optimal solution. Therefore, our decision space is the n -dimensional unit simplex, $\mathcal{S}_n$.

Let $\mathbf{a} \in \mathbb{R}^n$ denote the expected return of the securities, while $C \in \mathbb{R}^{n \times n}$ denotes the covariant matrix of the securities return. It is known that the expected return of our portfolio is equal to $\mathbf{a}^T \mathbf{x}$. One of our goal is to maximize the expected return.

It is much harder to measure the risk of the portfolio, but in this model it is equal to the variance of the securities return, namely $\mathbf{x}^T C \mathbf{x}$. Our second goal is to minimize this value. Now we are ready to formulate our model

$$(MM) \quad \text{MIN} \left(\begin{array}{c} -\mathbf{a}^T \mathbf{x} \\ \mathbf{x}^T C \mathbf{x} \end{array} \right), \quad \text{subject to} \quad \mathbf{x} \in \mathcal{S}_n.$$

For computational purposes we used data from the spot market [3] and daily prices of A category shares has been collected for a one year period from 01. 09. 2010. to 01. 09. 2011. Let $P_{i,d}$ denote the daily price of the i -th share on date d , then the i -th coordinate of the vector $\mathbf{a}$ is equal to $(P_{i,01.09.2011} - P_{i,01.09.2010})/P_{i,01.09.2010}$. Thus we only work with the relative returns from the price change and do not deal with shares dividend. We compute the daily return of the shares for every day (d) from 01. 09. 2010. to 31. 08. 2011. as $(P_{i,d} - P_{i,d+1})/P_{i,d}$, and C is the co-variant matrix of this daily return. To illustrate our method we use three shares ($i = \text{MOL}$, MTELEKOM, OTP) that are usually selected into portfolios because these shares correspond to large and stable Hungarian companies. We used the following data:

$$\mathbf{a} = \begin{pmatrix} -0,1906 \\ -0,2556 \\ -0,1665 \end{pmatrix}$$

$$C = 10^{-5} \begin{pmatrix} 27,1024 & 7,5655 & 17,1768 \\ 7,5655 & 16,4816 & 8,1816 \\ 17,1768 & 8,1816 & 34,2139 \end{pmatrix}$$

The input data for the *Subdivision algorithm for (LVOP)* are: matrix $A = \mathbf{e} \in \mathbb{R}^3$, $\mathbf{b} = 1$ since we have a single constraint in our model, and the objective function $F(\mathbf{x}) = \begin{pmatrix} -\mathbf{a}^T \mathbf{x} \\ \mathbf{x}^T C \mathbf{x} \end{pmatrix}$. Let $\mathcal{H}_0 = \{\mathbf{0} \leq \mathbf{x} \leq \mathbf{e}\}$, $\varepsilon = \frac{1}{2^6}$ and $K = 2$.

At the beginning of the algorithm in step 1 the family of set $\mathbb{H}$ has been defined (see Figure 1).

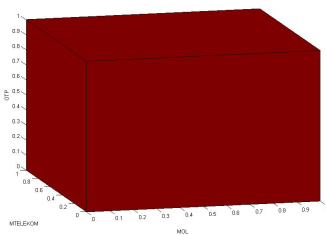


Figure 1

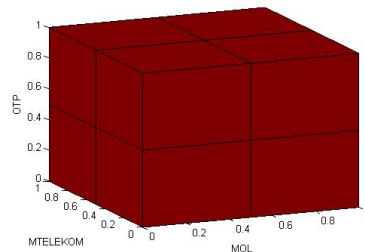


Figure 2

At the first iteration the procedure *Newsets* in step 2(a), defines $\mathbb{H}'$ in the following way. First, it cuts the set $\mathcal{H}_0$ into eight equal pieces as you see in Figure 2. After that

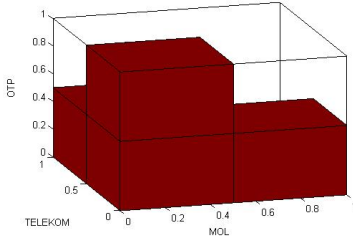


Figure 3

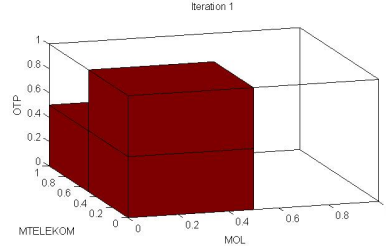


Figure 4

all those sets are deleted from $\mathbb{H}'$ that does not contain any point from the feasible solution set of the problem. Thus the result of the procedure *Newsets*, the family of $\mathbb{H}'$ covering the feasible solution set of the given problem has been shown in Figure 3.

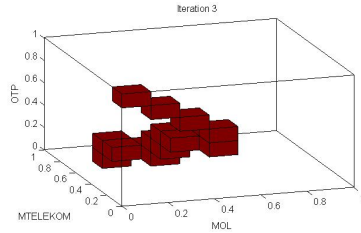


Figure 5

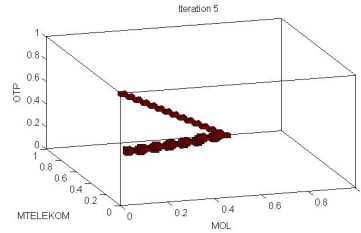


Figure 6

The main part of the algorithm starts at step 2(c). Two hundred random points are generated from the unit simplex (set $\mathcal{SP}$). For each generated point either a joint decreasing direction is computed and after that a corresponding Pareto optimal solution has been identified through some iteration or it has been shown that the generated point itself is a Pareto optimal solution of the problem. After we obtained 200 Pareto optimal solutions in set $\mathcal{FP}$ at step 2(d) we delete those boxes that does not contain any point from $\mathcal{FP}$ at step 2(e). The result of the first iteration can be seen in Figure 4.

From the original eight boxes remains three. For these three boxes the procedure has been repeated in the second iteration. The results of iteration 3, 5 and 7 are illustrated in Figures 5, 6, and 7, respectively.

These figures illustrate the flow of our computations. Finally to illustrate the convergence of our method the whole Pareto optimal set was determined based on the result of [37], and compared to the result computed in the fifth iteration, see in Figure 8.

The summary of our computations are shown in the 1 where I stands for the iteration number; B_{in} and B_{out} denotes the number of boxes at the beginning and at the end

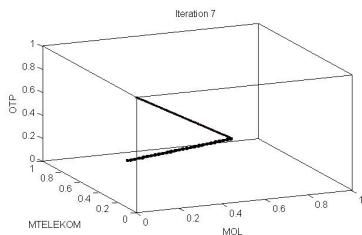


Figure 7

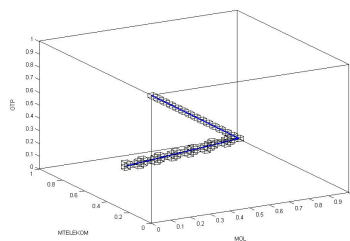


Figure 8

of iteration, respectively. Furthermore, $T(s)$ is the computational time of the I^{th} iteration in seconds, while d is the diameter of the family of sets, $\mathbb{H}$.

I	B_{in}	$T(s)$	B_{out}	d
1	1	8	3	2^{-3}
2	24	29	7	2^{-6}
3	56	70	15	2^{-9}
4	120	166	29	2^{-12}
5	232	312	56	2^{-15}
6	448	696	110	2^{-18}
7	880	1429	228	2^{-21}

Table 1

Computational results for Markowitz model using subdivision method.

The total computational time for our MATLAB implementation using a laptop with the following characteristics (processor: Intel(r) Core(TM) i3 [3.3 GHZ], RAM Memory: 4096 MB), took 2710 seconds for the subdivision algorithm for the given Markowitz model to approximate the whole Pareto optimal solution set with the accuracy $\varepsilon = 2.4 \cdot 10^{-8}$.

Analyzing our approximation of the Pareto optimal solution set, we can conclude that our first option is to buy OTP shares only. From the data it can be understood that this share has the biggest return (smallest loss in the financial crisis), so this solution represents the strategy when someone does not care about the risk but only about the return. From that point a line starts which represents strategies related to portfolios based on OTP and MOL shares. Clearly, there exists a breaking point where a new line segment starts. From the braking point the line lies in the interior of the simplex suggesting a portfolio based on all three selected shares.

Acknowledgement

This research has been supported by the TÁMOP-4.2.2/B-10/1-2010-0009, Hungarian National Office of Research and Technology with the financial support of the European Union from the European Social Fund.

Tibor Illés acknowledges the research support obtained as a part time *John Anderson Research Lecturer* from the Management Science Department, Strathclyde University, Glasgow, UK. This research has been partially supported by the UK Engineering and Physical Sciences Research Council (grant no. EP/P005268/1).

References

- [1] A.A. Akkeleş, L. Balogh and T. Illés, New variants of the criss-cross method for linearly constrained convex quadratic programming. *European Journal of Operational Research*, 157:74–86, 2004.
- [2] M.S. Bazaraa H.D. Sherali and M.C. Shetty, *Nonlinear Programing Theory and Algorithms*. 3-rd edition, Jhon Wiley & Son Inc., New Jersey, 2006.
- [3] Budapest Stock Exchange, http://client.bse.hu/topmenu/trading_data/stat_hist_download/trading_summ_pages_dir.html
Accessed: 2012. 06. 10.
- [4] Zs. Csizmadia, *New pivot based methods in linear optimization, and an application in petroleum industry*, PhD Thesis. Eötvös Loránd University of Sciences, Budapest, Hungary, 2007.
- [5] Zs. Csizmadia and T. Illés, New criss-cross type algorithms for linear complementarity problems with sufficient matrices. *Optimization Methods and Software*, 21(2):247–266, 2006.
- [6] Zs. Csizmadia, T. Illés and A. Nagy, The s-monotone index selection rules for pivot algorithms of linear programming. *European Journal of Operation Research*, 221(3):491–500, 2012.
- [7] M. Dellnitz, O. Schütze and T. Hestermeyer, Covering Pareto Sets by Multi-level Subdivision Techniques. *Journal of Optimization Theory and Applications*, 124(1):113–136, 2005.
- [8] M. Dellnitz, O. Schütze and Q. Zheng, Locating all the zeros of an analytic function in one complex variable, *Journal of Computational and Applied Mathematics*, 138(2):325–333, 2002.
- [9] G. Eichfelder, *Adaptive Scalarization Methods in Multiobjective Optimization*, Springer-Verlag, Berlin, Heidelberg, 2008.
- [10] G. Eichfelder and T.X.D. Ha, Optimality conditions for vector optimization problems with variable ordering structures. *Optimization: A Journal of Mathematical Programming and Operations Research*, 62(5):597–627, 2013.
- [11] Farkas Gy., A Fourier-féle mechanikai elv alkalmazásai, (The applications of the mechanical principle of Fourier [in Hungarian]). *Mathematikai és Természettudományi Értesítő*, 12:457–472, 1894.
- [12] Gy. Farkas, Theorie der einfachen Ungleichungen. *Journal für die Reine und Angewandte Mathematik*, 124:1–27, 1902.

- [13] J. Fliege, B.F. Svaiter, Steepest Descent Methods for Multicriteria Optimization, *Mathematical Methods of Operations Research. Mathematical Methods of Operations Research*, 51:479–494, 2000.
- [14] J. Fliege, An Efficient Interior-Point Method for Convex Multicriteria Optimization Problems, *Mathematical Methods of Operations Research*, 31(4):825–845, 2006.
- [15] F. Giannessi, G. Mastroeni and L. Pellegrini, On the Theory of Vector Optimization and Variational Inequalities. Image Space Analysis and Separation. In *Vector Variational Inequalities and Vector Equilibria: Mathematical Theories*, editor: F. Giannessi. Series on Nonconvex Optimization and its Applications, Vol.38, Kluwer, Dordrecht, 2000.
- [16] D. den Hertog, C. Roos and T. Terlaky, The linear complementarity problem, sufficient matrices and the criss-cross method. *Linear Algebra and Its Applications*, 187:1–14, 1993.
- [17] T. Illés and K. Mészáros, A New and Constructive Proof of Two Basic Results of Linear Programming. *Yugoslav Journal of Operations Research*, 11:15–30, 2001.
- [18] T. Illés and M. Nagy, A new variant of the Mizuno-Todd-Ye predictor-corrector algorithm for sufficient matrix linear complementarity problem. *European Journal of Operational Research*, 181(3):1097–1111, 2007.
- [19] T. Illés, M. Nagy and T. Terlaky, Interior Point Algorithms [in Hungarian] (Belsőpontos algoritmusok), pp. 1230–1297. Iványi Antal (alkotó szerkesztő), *Informatikai Algoritmusok 2*, ELTE Eötvös Kiadó, Budapest, 2005.
- [20] T. Illés, C. Roos and T. Terlaky, Polynomial Affine–Scaling Algorithms for $P_*(\kappa)$ Linear Complementarity Problems. In P. Gritzmann, R. Horst, E. Sachs, R. Tichatschke, editors, *Recent Advances in Optimization*, Proceedings of the 8th French-German Conference on Optimization, Trier, July 21–26, 1996, *Lecture Notes in Economics and Mathematical Systems* 452, pp. 119–137, Springer Verlag, 1997.
- [21] T. Illés and T. Terlaky, Pivot Versus Interior Point Methods: Pros and Cons. *European Journal of Operations Research* 140:6–26, 2002.
- [22] E. Klafszky and T. Terlaky, The role of pivoting in proving some fundamental theorems of linear algebra. *Linear Algebra and its Application* 151:97–118, 1991.
- [23] D.T. Luc, *Theory of Vector Optimization*. Lecture Notes in Economics and Mathematical Systems, No. 319, Springer-Verlag, Berlin, 1989.
- [24] G D. Luenberger and Y Ye, *Linear and Nonlinear Programming*. International Series in Operations Research and Management Science, Springer-Verlag, New York, 2008.

- [25] H. Markowitz, The Optimalization of Quadratic Function Subject to Linear Constarint. *Naval Reserch Logistic Quartely*, 3(1-2):111–133, 1956.
- [26] H. Markowitz, Portfolio Selection. *The Journal of Finance*, 7(1):77–91, 1952.
- [27] H. M. Markowitz, *Portfolio Selection: Efficient Diversification of Investment*, Jhon Wiley and Son Inc., New York, 1959.
- [28] K. Miettinen, *Nonlinear Multiobjective Optimization*. International Series in Operations Research and Management Science, Springer US, 1998.
- [29] K. G. Murty, *Linear Complementarity, Linear and Nonlinear Programming*. Heldermann Verlag, Berlin, 1988.
- [30] V. Pareto, *Cours D'Économie Politique*. F. Rouge, Lausanne, 1896.
- [31] Prékopa A., *Lineáris programozás I.* Bolyai János Matematikai Társulat, Budapest, 1968.
- [32] A. Prékopa, A Very Short Introduction to Linear Programming. *RUTCOR Lecture Notes* 2-92, 1992.
- [33] C. Roos, T. Terlaky and J-Ph. Vial, *Theory and Algorithms for Linear Optimization: An Interior Point Approach*. Wiley-Interscience Series in Discrete Mathematics and Optimization, John Wiley & Sons., 1997.
- [34] S. Schäßler, R. Schultz and K. Weinzierl, A Stochastic Method for the Solution of Unconstrained Vector Optimization Problems. *Journal of Optimalization Theory and Application*, 114(1):112–128, 2002.
- [35] A. Schrijver, *Theory of Linear and Integer Programming*, John Wiley and Sons Inc., New York, 1986.
- [36] T. Terlaky, A convergent criss-cross method. *Math. Oper. und Stat. Ser. Optimization* 16(5):683–690, 1985.
- [37] J. Vörös, Portfolio analysis – An analytic derivation of the efficient portfolio frontier. *European Journal of Operational Reserch*, 23(3):294–300, 1986.