

Optimisation en finance

Clément W. Royer

Notes de cours - M2 MIAGE - 2022/2023

- La version la plus récente de ces notes se trouve à l'adresse :
<https://www.lamsade.dauphine.fr/~croyer/ensdocs/OF/PolyOF.pdf>.
- Pour tout commentaire, merci d'envoyer un mail à clement.royer@lamsade.dauphine.fr.
- **Historique du document**
 - 2022.12.13 : Correction d'une ambiguïté dans la partie 2.2.2.
 - 2022.12.04 : Version complète.
 - 2022.11.15 : Première version.
- **Objectifs d'apprentissage**
 - Comprendre les caractéristiques des modèles d'optimisation les plus classiques en finance.
 - Connaître les grands principes derrière la résolution de ces modèles et les informations à en tirer.
 - Étudier des exemples de problèmes en finance et leur modélisation via des outils d'optimisation.

Sommaire

1	Introduction à l'optimisation	6
1.1	Définition et contexte	6
1.1.1	Problème et processus d'optimisation	6
1.1.2	Le contexte de la finance	7
1.2	Logiciels et optimisation	7
1.3	Bases de l'optimisation mathématique	8
1.3.1	Formulation générale et premières définitions	8
1.4	Conclusion du chapitre 1	9
2	Optimisation linéaire en finance	10
2.1	Bases de la programmation linéaire	10
2.2	Résolution d'un programme linéaire	11
2.2.1	Dualité et conditions de KKT	11
2.2.2	Analyse de sensibilité	13
2.3	Algorithmes et logiciels de programmation linéaire	14
2.3.1	Algorithmes du simplexe	14
2.3.2	Algorithmes de points intérieurs	16
2.4	Application : arbitrage et programmation linéaire	17
2.4.1	Formulation en programme linéaire	17
2.4.2	Théorème fondamental de l'arbitrage	18
2.5	Conclusion du chapitre 2	19
2.6	Exercices	19
2.6.1	Solution de l'exercice 1	20
3	Optimisation quadratique en finance	23
3.1	Bases de la programmation quadratique	23
3.2	Résolution d'un programme quadratique	24
3.2.1	Résolution théorique	24
3.2.2	Résolution algorithmique	24
3.3	Application : modèle de Markowitz	25
3.3.1	Principe	25
3.3.2	Problèmes d'optimisation	26
3.3.3	Analyse des modèles et de leurs solutions	27
3.4	Conclusion du chapitre 3	28
3.5	Exercices	28

4	Optimisation stochastique en finance	32
4.1	Formulation générale	32
4.1.1	Programmes stochastiques sans et avec recours	32
4.1.2	Programme linéaire stochastique en deux temps	33
4.2	Approche par scénarios	34
4.2.1	Principe	34
4.2.2	Décomposition de Benders pour programmes linéaires stochastiques	35
4.3	Mesures de risque	36
4.3.1	Motivation et définition générique	36
4.3.2	VaR et CVaR	36
4.4	Conclusion du chapitre 4	38
4.5	Exercices	38
Annexe A	Bases mathématiques	43
A.1	Éléments d'algèbre linéaire	43
A.1.1	Algèbre linéaire vectorielle	43
A.1.2	Algèbre linéaire matricielle	45

Avant-propos

Le but de ce cours est de présenter les classes de problème d'optimisation les plus utilisés en finance, et les algorithmes associés à leur résolution. Il se veut découpé en plusieurs parties suivant le déroulement des séances de cours.

Ce cours n'est pas un cours de mathématiques, mais il repose sur plusieurs concepts élémentaires d'algèbre linéaire et de calcul différentiel. Ceux-ci ne seront pas rappelés en détail en cours, mais sont en revanche détaillés en annexe de ce polycopié. De même, les notations cruciales du polycopié sont décrites ci-après.

Notations

- Les scalaires seront représentés par des lettres minuscules : $a, b, c, \alpha, \beta, \gamma$.
- Les vecteurs seront représentés par des lettres minuscules **en gras** : $\mathbf{a}, \mathbf{b}, \mathbf{c}, \boldsymbol{\alpha}, \boldsymbol{\beta}, \boldsymbol{\gamma}$.
- Les lettres majuscules en gras seront utilisées pour les matrices : $\mathbf{A}, \mathbf{B}, \mathbf{C}$.
- Les lettres majuscules cursives seront utilisées pour les ensembles : $\mathcal{A}, \mathcal{B}, \mathcal{C}$.
- La définition d'une nouvelle quantité ou d'un nouvel opérateur sera indiquée par $:=$.
- L'ensemble des entiers naturels sera noté $\mathbb{N}$, l'ensemble des entiers relatifs sera noté $\mathbb{Z}$.
- L'ensemble des réels sera noté $\mathbb{R}$. L'ensemble des réels positifs sera noté $\mathbb{R}_+$ et l'ensemble des réels strictement positifs sera noté $\mathbb{R}_{++}$.
- On notera $\mathbb{R}^n$ l'ensemble des vecteurs à n composantes réelles, et on considèrera toujours que n est un entier supérieur ou égal à 1.
- Un vecteur $\mathbf{x} \in \mathbb{R}^n$ sera pensé (par convention) comme un vecteur colonne. On notera $x_i \in \mathbb{R}$ sa i -ème coordonnée dans la base canonique de $\mathbb{R}^n$. On aura ainsi $\mathbf{x} = \begin{bmatrix} x_1 \\ \vdots \\ x_n \end{bmatrix}$. Étant donné un vecteur (colonne) $\mathbf{x} \in \mathbb{R}^n$, le vecteur ligne correspondant sera noté $\mathbf{x}^T$. On aura donc $\mathbf{x}^T = [x_1 \ \cdots \ x_n]$ et $[\mathbf{x}^T]^T = \mathbf{x}$.
- Le produit scalaire entre deux vecteurs de $\mathbb{R}^n$ est défini par $\mathbf{x}^T \mathbf{v} = \mathbf{v}^T \mathbf{x} = \sum_{i=1}^n x_i v_i$. On définit ensuite la norme d'un vecteur $\mathbf{x} \in \mathbb{R}^n$ par $\|\mathbf{x}\| = \sqrt{\mathbf{x}^T \mathbf{x}}$.
- On notera $\mathbb{R}^{m \times n}$ l'ensemble des matrices à m lignes et n colonnes à coefficients réels, où m et n seront des entiers supérieurs ou égaux à 1. Une matrice $\mathbf{A} \in \mathbb{R}^{n \times n}$ est dite carrée (dans le cas général, on parlera de matrice rectangulaire).
- On identifiera les vecteurs de $\mathbb{R}^n$ avec les matrices de $\mathbb{R}^{n \times 1}$.
- Étant donnée une matrice $\mathbf{A} \in \mathbb{R}^{m \times n}$, on notera $\mathbf{A}_{ij}$ le coefficient en ligne i et colonne j de la matrice. La diagonale de $\mathbf{A}$ est formée par l'ensemble des coefficients $\mathbf{A}_{ii}$ pour $i = 1, \dots, \min\{m, n\}$. La notation $[\mathbf{A}_{ij}]_{\substack{1 \leq i \leq m \\ 1 \leq j \leq n}}$ sera donc équivalente à $\mathbf{A}$. Sans ambiguïté sur la taille de la matrice, on notera simplement $[\mathbf{A}_{ij}]$.
- Selon les besoins, on utilisera $\mathbf{a}_i^T$ pour la i -ème ligne de $\mathbf{A}$ ou $\mathbf{a}_j$ pour la j -ème colonne de $\mathbf{A}$.
Selon le cas, on aura donc $\mathbf{A} = \begin{bmatrix} \mathbf{a}_1^T \\ \vdots \\ \mathbf{a}_m^T \end{bmatrix}$ ou $\mathbf{A} = [\mathbf{a}_1 \ \cdots \ \mathbf{a}_n]$.
- Pour une matrice $\mathbf{A} = [\mathbf{A}_{ij}] \in \mathbb{R}^{m \times n}$, la matrice transposée de $\mathbf{A}$, notée $\mathbf{A}^T$, est la matrice de $\mathbb{R}^{n \times m}$ telle que
$$\forall i = 1 \dots m, \forall j = 1 \dots n, \quad \mathbf{A}_{ji}^T = \mathbf{A}_{ij}.$$
- Pour tout $n \geq 1$, $\mathbf{I}_n$ correspond à la matrice identité $\mathbb{R}^{n \times n}$ (avec 1 sur la diagonale et 0 ailleurs).

Chapitre 1

Introduction à l'optimisation

Le but de ce cours est d'étudier les problèmes d'optimisation issus de la science des données, ainsi que les algorithmes qui leur sont appliqués. Le concept d'optimisation existe bien au-delà de l'apprentissage, mais il est particulièrement important pour ce dernier. De nombreuses tâches d'apprentissage peuvent en effet être formulées comme des problèmes d'optimisation, et ceux-ci peuvent être résolus efficacement par des algorithmes modernes. Ce chapitre présente les principes de base derrière la formulation et l'étude (à la fois mathématique et algorithmique) d'un problème d'optimisation.

1.1 Définition et contexte

1.1.1 Problème et processus d'optimisation

Avant d'être un concept mathématique, un problème d'optimisation peut être verbalisé : on en donne une décomposition ci-dessous.

Définition 1.1.1 *Un problème d'optimisation mathématique est formé par les trois éléments suivants.*

- Une **fonction objectif** : il s'agit du critère qui permet de quantifier la qualité d'une décision. Selon le cas, une meilleure décision peut signifier une valeur de l'objectif plus faible (auquel cas on aura affaire à un problème de minimisation) ou plus élevée (il s'agira alors d'un problème de maximisation).
- Des **variables de décision** : ce sont les différents choix qu'il est possible d'effectuer, dont la valeur influence celle de la fonction objectif. On cherche les meilleures valeurs de ces variables relativement à l'objectif.
- Des **contraintes** : celles-ci spécifient des conditions qui doivent être satisfaites par les variables de décision pour que les valeurs correspondantes soient acceptables.

Tout processus d'optimisation passe nécessairement par une phase de **modélisation**, durant laquelle on transforme un problème concret en objet mathématique que l'on peut essayer de résoudre. Pour cela, on se base sur des **algorithmes d'optimisation** qui, bien que pouvant être formulés à la main, ont généralement vocation à être implémentés sur un ordinateur. Lorsque l'algorithme appliqué

au problème renvoie une réponse, celle-ci doit être examinée pour savoir si elle fournit une solution satisfaisante au problème : au besoin, le modèle pourra être modifié, et l'algorithme ré-exécuté. En pratique, le processus d'optimisation peut donc boucler après une phase d'**interprétation** ou de **discussion** avec des experts du domaine d'application dont le problème est issu.

Remarque 1.1.1 *L'optimisation numérique obéit généralement aux principes suivants :*

- *Il n'y a pas d'algorithme universel : certaines méthodes seront très efficaces sur certains types de problèmes, et peu efficaces sur d'autres. L'étude de la structure du problème facilite le choix d'un algorithme.*
- *Il peut y avoir un fort écart entre la théorie et la pratique : le calcul en précision finie peut induire des erreurs d'arrondis qui conduisent à une performance en-deçà de celle prévue par la théorie mathématique.*
- *La théorie guide la pratique, et vice-versa : pour la plupart des problèmes, il est possible de définir des expressions mathématiques qui permettent de vérifier si l'on a trouvé une solution du problème. Celles-ci sont à la base de nombreux algorithmes que nous étudierons. De manière générale, l'étude théorique d'un problème permet des avancées garanties qui surpassent souvent les heuristiques. Réciproquement, la performance de certaines méthodes a amené les chercheurs en optimisation à en proposer une étude théorique, permettant ainsi d'expliquer le comportement pratique.*

1.1.2 Le contexte de la finance

L'optimisation est utilisée en finance comme un outil de modélisation et de résolution efficace de problèmes.

Gestion des risques Durant la dernière décennie et suite à la crise dite des *subprimes* de 2007-2008, la réglementation sur les niveaux de risques autorisés dans les milieux financiers a considérablement évolué. Les opérations financières sont aujourd'hui exécutées. En termes d'optimisation, on considèrera ainsi des problèmes sous contraintes de niveau de risque, tels que la maximisation du retour sur investissement.

Gestion des actifs (assets) et des dettes (liabilities) La modélisation de valeurs d'actifs ou de dettes prend en compte l'aléatoire sous-jacent au marché. En termes d'optimisation, on obtient ainsi des problèmes d'optimisation stochastique, où le caractère aléatoire est modélisé par l'introduction de variables aléatoires dont la valeur évolue sur une ou plusieurs périodes.

1.2 Logiciels et optimisation

L'optimisation a connu un essor majeur au cours des années 1980 avec le développement des ordinateurs. Une composante majeure

Les langages de programmation les plus utilisés pour développer des algorithmes d'optimisation étaient historiquement C/C++/Fortran, notamment pour développer des codes performants. Le développement de prototypes de recherche se fait lui plutôt via des langages interprétés tels que

Matlab/Octave/Python/Julia, avec Python prenant de plus en plus d'importance de par son utilisation dans les problèmes de sciences des données et apprentissage.

En plus des langages de programmation standard, de nombreux langages ou bibliothèques ont été développés pour modéliser les problèmes d'optimisation. On peut citer notamment GAMS, AMPL, CVX, Pyomo, qui ont un champ d'application génériques, ou MATPOWER et PyTorch, qui sont dédiés à un domaine particulier (respectivement les réseaux d'énergie et l'apprentissage machine).

En finance Dans le contexte de la finance, le recours aux tableurs est prépondérant, ce qui a conduit au développement de solveurs en tant qu'outil interne aux tableurs. C'est ainsi que Microsoft Excel, LibreOffice Calc et Google Sheets possèdent des solveurs, notamment de programmation linéaire. Si certains sont disponibles gratuitement (*open source*), le solveur le plus abouti (Microsoft Excel Solver) requiert une licence d'exploitation¹

D'autres outils plus bas niveau ont également été développés, parfois avec un spectre plus large d'applications. Le logiciel MATLAB (propriétaire) dispose ainsi d'une boîte à outils de finance très fournie, qui comporte notamment des algorithmes d'optimisation. On s'intéressera dans ce cours à la bibliothèque *cvxpy*, développée à l'origine à l'université de Stanford (États-Unis), et qui dispose de nombreux outils pour résoudre les problèmes de programmation convexe.

1.3 Bases de l'optimisation mathématique

1.3.1 Formulation générale et premières définitions

La transformation d'une description formelle d'un problème d'optimisation en objet mathématique conduit à l'écriture suivante :

$$\underset{x \in \mathbb{R}^n}{\text{minimiser}} f(x) \quad \text{s. c.} \quad x \in \mathcal{X}, \quad (1.3.1)$$

où minimiser représente ce que l'on cherche à faire (on peut vouloir minimiser ou maximiser), $x \in \mathbb{R}^n$ est un vecteur regroupant les variables de décision du problème, $f : \mathbb{R}^n \rightarrow \mathbb{R}$ est la fonction objectif qui mesure la qualité des décisions, et $\mathcal{X} \subset \mathbb{R}^n$ est l'ensemble des points réalisables ou admissibles, qui vérifient les contraintes posées sur les variables de décision.²

Définition 1.3.1 i) Un point $x \in \mathbb{R}^n$ tel que $x \in \mathcal{X}$ est dit **admissible**, ou **réalisable**.

ii) Un point $x \in \mathbb{R}^n$ tel que $x \notin \mathcal{X}$ est dit **irréalisable**.

iii) Si $\mathcal{X} = \emptyset$, alors le problème (1.3.1) est dit irréalisable.

Définition 1.3.2 i) La **valeur minimale** (ou le **minimum**) du problème (1.3.1) est notée :

$$\min_{x \in \mathbb{R}^n} \{f(x) \text{ s. c. } x \in \mathcal{X}\}. \quad (1.3.2)$$

Lorsque le problème possède une solution, la valeur minimale est la valeur de f en toute solution, et cette valeur est finie. Lorsque le problème est irréalisable, on pose par convention

$$\min_{x \in \mathbb{R}^n} \{f(x) \text{ s. c. } x \in \mathcal{X}\} = +\infty.$$

¹Il est à noter que si un compte passeport Dauphine permet d'accéder à Microsoft Excel, il ne donne pas l'autorisation d'installer des extensions, et donc ne permet pas l'utilisation de Solver.

²L'abréviation s.c. signifie "sous la/les contrainte(s)"; en anglais, on utilise s.t., pour *subject to*.

Enfin, si le problème est réalisable mais qu'il n'existe pas de solution, on aura

$$\min_{x \in \mathbb{R}^n} \{f(x) \text{ s. c. } x \in \mathcal{X}\} = -\infty,$$

et on dit alors que le problème est **non borné**.

ii) L'ensemble des solutions du problème (1.3.1) est noté

$$\operatorname{argmin}_{x \in \mathbb{R}^n} \{f(x) \text{ s. c. } x \in \mathcal{X}\}, \quad (1.3.3)$$

qu'on appelle argument minimal (ou *argmin*) du problème. C'est un sous-ensemble de $\mathbb{R}^n$ (et de $\mathcal{X}$) : il s'agit de l'ensemble des points de $\mathcal{X}$ qui donnent la valeur minimale de f . Il est vide si le problème est irréalisable ou réalisable et non borné.

Remarque 1.3.1 Le problème (1.3.1) n'admet pas forcément de solution (prendre par exemple $d = 1$, $\mathcal{F} = \mathbb{R}$ et $f(w) = w$) : dans ce cas, l'argument minimal est l'ensemble vide, et la valeur minimale est $-\infty$.

On définit de la même manière que précédemment les problèmes de maximisation ainsi que les opérateurs argmax et $\max$. Dans ce cours, on se concentrera sur les problèmes de minimisation, utilisant en cela la propriété ci-dessous.

Proposition 1.3.1 Pour toute fonction $f : \mathbb{R}^n \rightarrow \mathbb{R}$ et tout ensemble $\mathcal{F} \subset \mathbb{R}^n$, résoudre le problème de maximisation

$$\operatorname{maximiser}_{x \in \mathbb{R}^n} f(x) \quad \text{s. c.} \quad x \in \mathcal{F}$$

équivalent à résoudre le problème de minimisation

$$\operatorname{minimiser}_{x \in \mathbb{R}^n} -f(x) \quad \text{s. c.} \quad x \in \mathcal{F},$$

dans la mesure où

$$\operatorname{argmax}_{x \in \mathbb{R}^n} \{f(x) \text{ s. c. } x \in \mathcal{F}\} = \operatorname{argmin}_{x \in \mathbb{R}^d} \{-f(x) \text{ s. c. } x \in \mathcal{F}\}$$

et

$$\max_{x \in \mathbb{R}^n} \{f(x) \text{ s. c. } x \in \mathcal{F}\} = - \min_{x \in \mathbb{R}^n} \{-f(x) \text{ s. c. } x \in \mathcal{F}\}.$$

Tout problème de maximisation se ramène donc à un problème de minimisation via une reformulation. Nous verrons d'autres exemples de reformulation, qui est une technique très utilisée pour transformer un problème en un autre problème équivalent mais que l'on saura mieux traiter théoriquement et/ou algorithmiquement.

Dans la majorité de ce cours, et notamment dans le reste de ce chapitre, on se concentrera sur des problèmes d'optimisation sans contraintes pour en caractériser plus précisément les solutions.

1.4 Conclusion du chapitre 1

L'optimisation est un outil mathématique qui permet de modéliser de nombreux problèmes en sciences des données et au-delà. L'optimisation comporte toujours une phase de modélisation, qui consiste à formuler mathématiquement le problème en définissant une fonction objectif, des variables de décisions et des contraintes. Cela permet de placer le problème dans une classe, et ainsi de caractériser ce qu'est une solution de ce problème.

Chapitre 2

Optimisation linéaire en finance

Dans ce chapitre, nous nous intéressons aux problèmes de programmation linéaire. Le développement algorithmique lié à la résolution de ces problèmes, dû notamment à George Dantzig dans les années 1960, a conduit à des avancées majeures dans de nombreux domaines d'application, dont la finance. Les programmes linéaires forment une classe bien comprise de problèmes d'optimisation qui permettent à la fois de modéliser de nombreux problèmes concrets et également de les résoudre avec des algorithmes efficaces, même en présence d'un très grand nombre de variables.

2.1 Bases de la programmation linéaire

Formellement, un programme linéaire consiste à minimiser ou maximiser une fonction linéaire des variables de décision sous la contrainte que ces variables appartiennent à un ensemble décrit par des équations et inéquations linéaires.

Définition 2.1.1 *Un programme linéaire est un problème d'optimisation de la forme*

$$\begin{aligned} & \text{minimiser}_{x \in \mathbb{R}^n} && c^T x \\ & \text{s.c.} && Ax = b \\ & && Dx \geq f. \end{aligned}$$

avec $A \in \mathbb{R}^{m \times n}$, $b \in \mathbb{R}^m$, $D \in \mathbb{R}^{\ell \times n}$, $f \in \mathbb{R}^\ell$, $c \in \mathbb{R}^n$.

Définition 2.1.2 *Un programme linéaire en **forme standard** est de la forme*

$$\begin{aligned} & \text{minimiser}_{x \in \mathbb{R}^n} && c^T x \\ & \text{s.c.} && Ax = b \\ & && x \geq 0, \end{aligned}$$

où les contraintes d'intervalle $x \geq 0$ peuvent porter seulement sur un sous-ensemble des coordonnées de x .

Tout programme linéaire peut se mettre sous forme standard en introduisant des variables supplémentaires. Par exemple, la contrainte $x_1 \geq 1$ peut se réécrire comme

$$x_1 - s_1 = 1, s_1 \geq 0.$$

Exemple 2.1.1 (Allocation de fonds (Cornuéjols et al. 2018)) On considère un ensemble de quatre fonds d'investissement ayant les propriétés suivantes:

Fonds	Fonds 1	Fonds 2	Fonds 3	Fonds 4
Part capitalisation haute	50%	30%	25%	60%
Part capitalisation moyenne	30%	10%	40%	20%
Part capitalisation basse	20%	60%	35%	20%
Retour sur investissement	10%	15%	16%	8%

On souhaite allouer 80000€ parmi cet ensemble, de sorte à avoir

- Au moins 35% de l'allocation en capitalisation haute;
- Au moins 30% de l'allocation en capitalisation moyenne;
- Au moins 15% de l'allocation en capitalisation basse.

On suppose que l'on ne peut investir que positivement dans un fonds (long-only positions), et on cherche l'allocation qui maximise le retour sur investissement.

Pour modéliser cet exemple via un programme linéaire, on commence par identifier les variables : on définit ainsi $x \in \mathbb{R}^4$, où x_i correspond à l'investissement dans le fond i (exprimé en k€). On obtient ainsi le problème de maximisation suivant

$$\begin{aligned}
 & \text{maximiser}_{x \in \mathbb{R}^4} && 0.1x_1 + 0.15x_2 + 0.16x_3 + 0.08x_4 \\
 & \text{s.c.} && 0.5x_1 + 0.3x_2 + 0.25x_3 + 0.6x_4 \geq 28 \\
 & && 0.3x_1 + 0.1x_2 + 0.4x_3 + 0.2x_4 \geq 24 \\
 & && 0.2x_1 + 0.6x_2 + 0.35x_3 + 0.2x_4 \geq 12 \\
 & && x_1 + x_2 + x_3 + x_4 = 80 \\
 & && x \geq 0.
 \end{aligned} \tag{2.1.1}$$

Le problème (2.1.1) n'est pas en forme standard, mais peut être converti sous cette forme en utilisant la procédure décrite plus haut. Il possède trois contraintes d'inégalité linéaires, une contrainte d'égalité linéaire et quatre contraintes de positivité concaténées sous la forme $x \geq 0$.

2.2 Résolution d'un programme linéaire

La théorie de la programmation linéaire permet à la fois de déterminer si un programme linéaire possède une solution et de donner des conditions qui caractérisent ces solutions si elles existent. Un outil majeur

2.2.1 Dualité et conditions de KKT

Dans cette partie, on se concentre sur le problème suivant, appelé **problème primal**

$$\begin{cases} \text{minimiser}_{x \in \mathbb{R}^n} & c^T x \\ \text{s.c.} & Ax = b \\ & x \geq 0, \end{cases} \tag{P}$$

On fera les hypothèses simplificatrices suivantes.

Hypothèse 2.2.1 • $A \in \mathbb{R}^{m \times n}$ est de rang plein.

- Les contraintes de positivité s'appliquent à tous les x_i .

Théorème 2.2.1 (Conditions de Karush-Kuhn-Tucker (KKT)) On considère le problème $(\mathcal{P})$ sous l'hypothèse (2.2.1). Un point $x^* \in \mathbb{R}^n$ est une solution de $(\mathcal{P})$ si et seulement si il existe $y^* \in \mathbb{R}^m$ et $s^* \in \mathbb{R}^n$ tels que (x^*, y^*, s^*) vérifie le système

$$\begin{cases} A^T y^* + s^* = c \\ Ax^* = b \\ x^* \geq 0 \\ s^* \geq 0 \\ x_i^* s_i^* = 0 \quad \forall i = 1, \dots, n. \end{cases} \quad (\text{KKT-Lin})$$

Les vecteurs y^* et s^* sont les multiplicateurs de Lagrange associés aux contraintes d'égalité et de positivité. En optimisation, on les appelle variables duales.

Remarque 2.2.1 En finance, les variables duales associées à une contrainte linéaire sont appelées des shadow prices, ou **prix fictifs**. Dans le cas d'une contrainte de positivité, le multiplicateur associé sera appelé un **coût réduit**.

Définition 2.2.1 (Problème dual) Le problème dual de $(\mathcal{P})$ est donné par

$$\begin{cases} \text{maximiser}_{y \in \mathbb{R}^m} & b^T y \\ \text{s.c.} & A^T y + s = c \\ & s \geq 0. \end{cases} \quad (\mathcal{D})$$

Proposition 2.2.1 (Conditions d'optimalité) Si (y^*, s^*) est une solution de $(\mathcal{D})$, il existe $x^* \in \mathbb{R}^n$ tel que (x^*, y^*, s^*) vérifie les conditions de KKT.

L'intérêt de formuler le problème dual réside dans le lien intrinsèque entre ces deux problèmes comme suit.

Théorème 2.2.2 On considère les problèmes $(\mathcal{P})$ et $(\mathcal{D})$.

1. Soit les deux problèmes sont irréalisables;
2. Soit $(\mathcal{P})$ est irréalisable et $(\mathcal{D})$ est non borné;
3. Soit $(\mathcal{D})$ est irréalisable et $(\mathcal{P})$ est non borné;
4. Soit les deux problèmes sont réalisables.

La preuve de ce résultat repose sur le résultat suivant d'algèbre linéaire.

Théorème 2.2.3 (Théorème des alternatives) Soient $A \in \mathbb{R}^{m \times n}$ et $b \in \mathbb{R}^m$. Alors, les alternatives suivantes sont mutuellement exclusives :

- Soit $\exists x, Ax = b, x \geq 0$, soit $\exists y, A^T y \leq 0, b^T y > 0$.

- Soit $\exists x, Ax = 0, x \geq 0$, soit $\exists y, A^T y > 0$.
- Soit $\exists x, Ax = 0, x > 0$, soit $\exists y, A^T y \geq 0$.

En plus de résultats sur le caractère réalisable ou non borné des problèmes, on peut aussi lier les solutions de chaque problème.

Théorème 2.2.4 (Dualité faible) Pour tout (y, s) admissible pour $(\mathcal{D})$ et tout x admissible pour $(\mathcal{P})$, on a $b^T y \leq c^T x$.

Théorème 2.2.5 (Dualité forte) Si les deux problèmes sont réalisables, alors il existe (x^*, y^*, s^*) tel que :

- x^* est solution de $(\mathcal{P})$;
- (y^*, s^*) est solution de $(\mathcal{D})$;
- (x^*, y^*, s^*) vérifie (KKT);
- $b^T y^* = c^T x^*$.

Pour résoudre un problème linéaire, on peut donc chercher à la fois une solution primale de $(\mathcal{P})$ ainsi qu'une solution duale de $(\mathcal{D})$ en résolvant (KKT-Lin).

2.2.2 Analyse de sensibilité

On s'intéresse maintenant à un usage *a posteriori* des variables duales. Celles-ci permettent en effet de décrire la sensibilité de la valeur optimale d'un problème relativement à une contrainte, en décrivant l'effet d'une perturbation du second membre de ces contraintes.

Théorème 2.2.6 Soit y une variable dual associée à une contrainte d'un programme linéaire. Alors, pour tout Δ suffisamment petit en valeur absolue, un changement de second membre dans la contrainte de Δ induit un changement de valeur optimale de $y \cdot \Delta$.

Ainsi, la valeur absolue d'un multiplicateur permet de juger de la sensibilité du problème par rapport à une contrainte. Dans le cas où $y = 0$, par exemple, on sait que la solution ne changera pas si l'on perturbe un peu la contrainte : cela signifie en général que cette contrainte n'est pas critique pour calculer la solution.

Remarque 2.2.2 En règle générale, l'intervalle de valeurs Δ pour lesquelles le résultat du théorème 2.2.6 s'applique n'est pas connu. Cependant, dans le cadre de la programmation linéaire et notamment des algorithmes de type simplexe, il est possible de quantifier ce Δ .

Une interprétation plus générale de la sensibilité est fournie ci-dessous. On se place comme dans la section précédente dans le cadre d'un programme linéaire en forme standard.

Définition 2.2.2 (Problèmes perturbés) Partant du problème $(\mathcal{P})$ (sous l'hypothèse 2.2.1), on définit pour tous $u \in \mathbb{R}^m$ et $v \in \mathbb{R}^n$:

$$p(u, v) := \min_{x \in \mathbb{R}^n} \{c^T x \mid Ax = b + u, x \geq v\}. \quad (2.2.1)$$

On notera que la valeur $p(0, 0)$ correspond à la valeur optimale du problème de départ.

Théorème 2.2.7 Si $(\mathbf{y}^*, \mathbf{s}^*)$ sont les variables duales du problème de départ, alors pour tous $\mathbf{u}$ et $\mathbf{v}$, on a

$$p(\mathbf{u}, \mathbf{v}) \geq p(0, 0) + (\mathbf{y}^*)^T \mathbf{u} + (\mathbf{s}^*)^T \mathbf{v}. \quad (2.2.2)$$

En finance, le résultat du théorème 2.2.7 s'interprète comme le fait que chaque variable duale (ou multiplicateur) représente un prix d'équilibre pour la ressource représentée par la contrainte. Cela permet d'identifier les effets positifs ou négatifs d'agir sur la contrainte. A titre d'exemple, notons que si $s_i^* \gg 1$ et $v_i > 0$, alors l'équation 2.2.2 indique que la valeur optimale du problème augmentera si l'on requiert $x_i \geq v_i$ au lieu de $x_i \geq 0$, ce qui semble logique car on réduit ici l'ensemble des contraintes.

2.3 Algorithmes et logiciels de programmation linéaire

Comme on l'a vu dans la partie précédente, les conditions d'optimalité (dites de KKT) d'un programme linéaire permettent de caractériser ses solutions si elles existent. Cependant, déterminer les solutions du problème directement depuis les conditions de KKT s'avère délicat, car celles-ci forment un système d'équations non linéaires.

On décrit ci-après le fonctionnement de ces méthodes dans le cas d'un programme linéaire en forme standard, et on supposera que ce problème et son dual sont tous deux réalisables. Les conditions d'optimalité pour ce programme, ou de KKT s'écrivent

$$\begin{cases} \mathbf{A}^T \mathbf{y} + \mathbf{s} &= \mathbf{c} \\ \mathbf{A} \mathbf{x} &= \mathbf{b} \\ \mathbf{x} &\geq \mathbf{0} \\ \mathbf{s} &\geq \mathbf{0} \\ x_i s_i &= 0 \quad \forall i = 1, \dots, n. \end{cases} \quad (2.3.1)$$

On décrit ci-après les deux grandes catégories de méthodes qui visent à résoudre ces conditions d'optimalité de manière itérative.

2.3.1 Algorithmes du simplexe

La méthode du simplexe classique revient à une suite $\{\mathbf{x}_k, \mathbf{y}_k, \mathbf{s}_k\}$ vérifiant les conditions de KKT sauf $\mathbf{s} \geq \mathbf{0}$, et vise à satisfaire cette condition asymptotiquement. La variante duale du simplexe cherche à satisfaire $\mathbf{x} \geq \mathbf{0}$.

L'idée principale de l'algorithme du simplexe est d'exploiter le fait que l'ensemble réalisable d'un programme linéaire (réalisable) est un polyèdre, et que la solution du problème se trouve nécessairement sur un sommet de ce polyèdre. En ce sommet, toutes les contraintes d'inégalité ne sont généralement pas actives : dans le cadre d'un problème en forme standard ($\mathcal{P}$), cela signifie que certaines coordonnées de la solution seront nulles (contrainte $x_i \geq 0$ active) tandis que d'autres ne le seront pas. Le théorème ci-dessous garantit qu'identifier ces contraintes actives est suffisant pour déterminer la solution.

Théorème 2.3.1 On considère le problème ($\mathcal{P}$) et son dual sous l'hypothèse 2.2.1. Alors, il existe une partition $\mathcal{B} \cup \mathcal{N} = \{1, \dots, n\}$ telle que

$$i) |\mathcal{B}| = m;$$

ii) la matrice $\mathbf{A}_B := [A_{ij}]_{\substack{1 \leq i \leq m \\ j \in B}}$ est inversible;

iii) la solution du problème $\mathbf{x}^*$ se décompose selon la partition en

$$\mathbf{x}_B^* = [x_j]_{j \in B} = -\mathbf{A}_B^{-1} \mathbf{b}$$

et $\mathbf{x}_N^* = \mathbf{0}$.

Dans ce cas, une solution duale du problème est donnée par $\mathbf{y}^* = [\mathbf{A}_B^T]^{-1} \mathbf{c}_B$.

L'ensemble B , qui caractérise un ensemble de contraintes, est appelé une **base** du problème.

L'algorithme du simplexe décrit une manière algorithmique de parcourir l'ensemble des sommets du polyèdre en recherchant la base dans laquelle la solution optimale doit être exprimée. Ce processus termine nécessairement après un nombre fini d'itérations, mais celui-ci peut être très grand si le nombre de contraintes est important.

Principe général L'algorithme commence par déterminer une base B telle que les propriétés i) et ii) du théorème 2.3.1 soient vérifiées. Si le vecteur $\mathbf{x}$ associé vérifie $\mathbf{x} \geq \mathbf{0}$, on parle alors de solution basique (même s'il ne s'agit pas forcément de la solution optimale, il s'agit d'un point réalisable). Il s'agit ensuite de regarder le vecteur

$$\bar{\mathbf{c}} = \mathbf{c} - \mathbf{A}^T \mathbf{A}_B^{-T} \mathbf{c}_B,$$

qui correspond aux variables duales $\mathbf{s}$ dans les conditions de KKT. Si les coordonnées de ce vecteur sont positives ou nulles, alors on a obtenu une solution optimale. Sinon, cela signifie qu'il existe un indice de colonne $j \notin B$ tel que $\bar{c}_j < 0$ et

$$\mathbf{c}^T (\mathbf{x} - \alpha \mathbf{A}_B^{-1} \mathbf{a}_j) < \mathbf{c}^T \mathbf{x}$$

pour tout $\alpha > 0$. Cette propriété signifie que l'on ne dispose pas de la bonne base B , et qu'il en existe une meilleure contenant j : l'algorithme du simplexe explique comment échanger j avec un élément de la base.

Variante duale Le principe de méthode du simplexe décrit plus haut est un principe **primal**, qui se base essentiellement sur le calcul d'une solution primale $\mathbf{x} \in \mathbb{R}^n$. Il existe également des variantes duales, dans lesquelles on calcule $(\mathbf{y}, \mathbf{s})$ tels que $\mathbf{A}^T \mathbf{y} + \mathbf{s} = \mathbf{0}$, $\mathbf{s} \geq \mathbf{0}$, $\mathbf{A}\mathbf{x} = \mathbf{b}$ et $x_i s_i = 0$ pour tout $i = 1, \dots, n$.

Implémentation L'algorithme du simplexe est notamment utilisé dans l'extension *Solver* de Microsoft Excel. Ce solveur fournit une analyse de sensibilité détaillée, avec notamment un intervalle (plus ou moins interprétable) de valeurs dans lequel l'approximation via les multiplicateurs décrite dans le théorème 2.2.6 est valide. Le solveur de LibreOffice/OpenOffice Calc est également un solveur de type simplexe, ce qui se reconnaît notamment par rapport à la saturation des contraintes en la valeur renvoyée. De même, *Google Solve* et *OpenSolver* (accessibles dans Google Sheets) se basent sur le même genre de techniques, sans pour autant fournir de marges de validité des multiplicateurs. Enfin, le solveur *HIGHS*, qui est maintenant l'algorithme par défaut de simplexe dans la bibliothèque SciPy, est lui aussi basé sur une approche de simplexe.

2.3.2 Algorithmes de points intérieurs

La méthode du simplexe repose sur le fait de parcourir les sommets de l'ensemble des contraintes d'un programme linéaire, qui (lorsqu'il est non vide) est un polyèdre. De par la structure même du programme linéaire, on sait que la solution se trouve sur un sommet, et donc qu'il existe un certain nombre de contraintes d'inégalité satisfaites avec égalité en la solution (c'est le sens du Théorème 2.3.1).

Le principe des méthodes de points intérieurs, comme leur nom l'indique, consiste à passer par l'intérieur de l'ensemble réalisable de sorte à arriver plus vite au sommet solution qu'en parcourant l'ensemble des sommets.

Principe Un algorithme de points intérieurs part des conditions de KKT (KKT-Lin) et considère des points dans l'intérieur de l'ensemble réalisable, c'est-à-dire des triplets de la forme (x, y, s) tels que

$$Ax = b, \quad A^T y + s = c, \quad x > 0, y > 0.$$

Un tel triplet est dit **primal-dual admissible**, car il vérifie les contraintes du problème primal et celles du problème dual. En revanche, il ne vérifie pas la dernière condition d'optimalité de (KKT-Lin), car $x_i s_i > 0$ pour tout i . On pose alors

$$\mu = \sum_{i=1}^n \frac{x_i s_i}{n}, \quad (2.3.2)$$

qui représente une certaine distance entre le point calculé et une solution des conditions de KKT. Le but d'une méthode de points intérieurs est de calculer une suite de triplets telle que la valeur de μ tende vers 0, et ainsi de devenir de plus en plus proche d'un point optimal. En pratique, on peut calculer pour tout $\mu > 0$ une solution vérifiant les conditions souhaitées en trouvant (x, y, s) tels que $x_i > 0$ et $s_i > 0$ pour tout i , et le triplet est solution du système linéaire

$$\begin{aligned} Ax &= b \\ A^T y + s &= c \\ XSe &= \mu e, \end{aligned} \quad (2.3.3)$$

où X et S sont deux matrices diagonales donc les coefficients diagonaux sont ceux des vecteurs x et s , respectivement, et $e \in \mathbb{R}^n$ est le vecteur dont toutes les coordonnées valent 1. On peut ainsi calculer ce type de triplet pour des valeurs décroissantes de μ : c'est ce que font les approches de points intérieurs dites primales-duales.

La théorie des points intérieurs est extrêmement riche en résultats, en particulier pour des programmes linéaires. Le succès originel des points intérieurs réside dans l'obtention de garanties de convergence en temps polynomial vers une solution approchée du problème (qui dépend de la valeur du paramètre μ défini en (2.3.2)). Ce type de résultat contraste avec ceux existants pour le simplexe, où le temps de calcul peut être exponentiel pour atteindre la solution.

Implémentation CVX (MATLAB) et sa variante CVXPY en Python est à la fois un langage de modélisation et un outil de résolution basé sur les points intérieurs. Il permet de toujours renvoyer les variables duales d'un problème. Cette librairie permet notamment l'interface avec les meilleurs solveurs commerciaux du marché que sont Gurobi, CPLEX et Mosek. Ceux-ci sont basés sur une approche mixte simplexe/points intérieurs en ce qui concerne la programmation linéaire.

2.4 Application : arbitrage et programmation linéaire

Les problèmes d'arbitrage se posent lorsque l'on cherche à échanger différentes devises. On parle d'**opportunités d'arbitrage** lorsqu'il est possible de réaliser un bénéfice via un cycle d'échanges de devises.

Exemple 2.4.1 *Supposons que l'euro s'échange à 0,9 dollar, et que le dollar s'échange à 1,22 euro. Dans ce cas, pour tout euro échangé en dollar et à nouveau en euro, on obtiendrait 1,098 euro, soit un bénéfice.*

Dans cette section, on s'intéresse à la détection d'opportunités d'arbitrage lorsque l'on considère un ensemble arbitraire de devises.

2.4.1 Formulation en programme linéaire

On considère que l'on dispose d'un tableau $A = [a_{ij}]_{\substack{1 \leq i \leq n \\ 1 \leq j \leq n}}$ représentant des taux d'échange entre n monnaies. Une unité de monnaie i s'échange ainsi contre $a_{ij} > 0$ unités de monnaie j . On posera $a_{ii} = 1$ pour tout $i = 1, \dots, n$ de sorte à éviter les arbitrages triviaux, et on fait l'hypothèse simplificatrice que le coût de conversion entre devises est nul.

On choisit les quantités suivantes comme variables du problème :

- x_{ij} : quantité de devise i échangée en devise j , que l'on souhaitera toujours positive ou nulle (on s'attend notamment à ce que $x_{ii} = 0$);
- y_k : quantité de devise k obtenue après tous les échanges, définie par

$$y_k = \sum_{i=1}^n a_{ik} x_{ik} - \sum_{j=1}^n x_{kj},$$

où la première somme correspond à la quantité de devise k obtenue en échangeant d'autres monnaies, et la seconde correspond à la quantité de devise k échangée contre d'autres monnaies. On souhaitera toujours avoir $y_k \geq 0$.

Lorsque $y_k > 0$ pour une devise k donnée, on a une opportunité d'arbitrage. Notre objectif consiste donc à définir un ensemble de transactions (c'est-à-dire de valeurs pour $\{x_{ij}\}_{ij}$ et, par là-même, de valeurs pour $\{y_k\}_k$) telles que l'une des quantités de devise après transaction soit positive. Dans ce qui suit, on se concentrera sur l'existence d'un arbitrage pour la devise 1, et on cherchera donc $\{x_{ij}\}_{ij}$ tels que $y_1 > 0$.

Une première modélisation consiste à maximiser la valeur de y_1 , ce qui conduit au problème

$$\left\{ \begin{array}{ll} \text{maximiser}_{\{x_{ij}\}, \{y_k\}} & y_1 \\ \text{s.c.} & y_k = \sum_{i=1}^n a_{ik} x_{ik} - \sum_{j=1}^n x_{kj} \quad k = 1, \dots, n \\ & y_k \geq 0 \quad k = 1, \dots, n \\ & x_{ij} \geq 0 \quad i, j = 1, \dots, n. \end{array} \right. \quad (2.4.1)$$

Si ce problème est un programme linéaire presque en forme standard (il faudrait minimiser plutôt que maximiser, ce que l'on peut faire en minimisant $-y_1$), il n'est pas pour autant bien posé. En effet, s'il existe effectivement une opportunité d'arbitrage pour la devise 1, alors il existe un point

admissible à ce problème avec $y_1 > 0$. Il en existe même une infinité, et pour tout choix des $\{x_{ij}\}$ qui conduit à un bénéfice de y_1 , alors le choix $\{\lambda x_{ij}\}$ où $\lambda > 1$ conduit à $\lambda y_1 > y_1$. Par conséquent, le problème est non borné. Un remède à cela consiste simplement à imposer une borne sur y_1 , qui sera une valeur cible pour l'arbitrage. En choisissant cette valeur cible comme étant égale à 1, on obtient

$$\left\{ \begin{array}{ll} \text{maximiser}_{\{x_{ij}\}, \{y_k\}} & y_1 \\ \text{s.c.} & y_k = \sum_{i=1}^n a_{ik} x_{ik} - \sum_{j=1}^n x_{kj} \quad k = 1, \dots, n \\ & y_1 \leq 1 \\ & y_k \geq 0 \quad k = 1, \dots, n \\ & x_{ij} \geq 0 \quad i, j = 1, \dots, n. \end{array} \right. \quad (2.4.2)$$

La résolution du programme linéaire (2.4.2) conduit aux deux cas suivants : soit il n'y a pas d'arbitrage possible, auquel cas la valeur optimale du problème est 0, et $x_{ij} = y_k = 0$ pour tous i, j, k , soit il y a un arbitrage possible, et on aura alors $y_1 = 1$.

2.4.2 Théorème fondamental de l'arbitrage

La notion d'arbitrage peut s'exprimer de manière plus abstraite, mais toujours en lien avec la programmation linéaire. Soit un système composé de m actifs durant une période donnée de $T = 0$ à $T = 1$. On définit le vecteur

$$s_0 = \begin{bmatrix} s_0^1 \\ \vdots \\ s_0^m \end{bmatrix} \in \mathbb{R}^m,$$

qui représente les valeurs des actifs à l'instant $T = 0$ (on considère les valeurs comme pouvant être positives ou négatives). À l'instant $T = 1$, le système peut être dans un état $\omega \in \{\omega_1, \dots, \omega_n\}$, qui définit un nouveau vecteur de valeurs pour les actifs. On pose ainsi

$$\forall j = 1, \dots, n, \quad s_1(\omega_j) = \begin{bmatrix} s_1^1(\omega_j) \\ \vdots \\ s_1^m(\omega_j) \end{bmatrix} \in \mathbb{R}^m.$$

Une **opportunité d'arbitrage** dans ce contexte est un vecteur $y \in \mathbb{R}^m$ tel que

$$\left\{ \begin{array}{ll} s_0^T y \leq 0 \\ s_1(\omega_j)^T y \geq 0 & j = 1, \dots, n \\ s_0^T y \neq 0 & \text{ou } s_1(\omega_j)^T y > 0 \quad j = 1, \dots, n. \end{array} \right. \quad (2.4.3)$$

L'existence ou l'absence d'un tel arbitrage est établie par le résultat suivant, appelé en anglais *asset pricing*.

Théorème 2.4.1 (Théorème fondamental de l'arbitrage) On considère les vecteurs $s_0, s_1(\omega_1), \dots, s_1(\omega_n)$ définis ci-dessus. Alors,

i) Soit il existe un arbitrage, c'est-à-dire un vecteur $y \in \mathbb{R}^m$ vérifiant les conditions (2.4.3);

ii) Soit il existe $x \in \mathbb{R}^n$ tel que

$$s_0 = \sum_{j=1}^n x_j s_1(\omega_j) \quad \text{avec} \quad x_j > 0 \quad \forall j = 1, \dots, n. \quad (2.4.4)$$

Le résultat du théorème 2.4.1 peut s'écrire comme un théorème des alternatives similaire au théorème 2.2.3 de la manière suivante :

i) Soit il existe $\mathbf{x} \in \mathbb{R}^n$ solution du système

$$\mathbf{S}\mathbf{x} = \mathbf{s}_0, \quad \mathbf{S} = [\mathbf{s}_1(\omega_1) \cdots \mathbf{s}_1(\omega_n)],$$

avec $\mathbf{x} > \mathbf{0}$;

ii) Soit il existe $\mathbf{y} \in \mathbb{R}^m$ tel que

$$\mathbf{0} \neq [\mathbf{S} - \mathbf{s}_0]^T \mathbf{y} \geq \mathbf{0}.$$

Ce type de condition n'est pas un problème d'optimisation linéaire, mais est en revanche considéré comme un programme linéaire (résoudre un système d'inégalités et d'égalités linéaires). De fait, en introduisant une fonction objectif constante, il serait possible d'écrire les conditions ci-dessus comme contraintes d'un programme linéaire.

2.5 Conclusion du chapitre 2

La programmation linéaire permet de modéliser un grand nombre de problèmes, parfois grâce à des hypothèses simplificatrices (que l'on cherchera à relâcher dans la suite grâce à des formulations plus complexes). Un des grands avantages des programmes linéaires réside dans leur facilité de résolution, que ce soit par l'algorithme du simplexe ou les approches de points intérieurs. Par ailleurs, l'analyse de ces problèmes au moyen de la théorie de la dualité et des multiplicateurs de Lagrange (ou coûts fictifs/réduits) permet à la fois de caractériser les solutions et d'analyser la sensibilité d'une solution aux différentes contraintes. Ce procédé se révèle essentiel pour ne pas avoir à résoudre un problème de manière répétée.

2.6 Exercices

Exercice 1 : Gestion actifs-passifs

Dans cet exercice, on s'intéresse aux problèmes dits de gestion actifs-passifs¹, que l'on peut modéliser par des programmes linéaires en supposant que le seul risque inhérent au problème réside dans les taux d'intérêt. On souhaite construire un portefeuille qui génère suffisamment de liquidités pour couvrir des besoins (passifs) sur une période de m années consécutives. On note $\ell_i > 0$ le besoin en liquidités pour l'année i .

On suppose que le portefeuille peut être construit à partir de n obligations ayant chacune un prix unitaire et qui génèrent un flot de liquidités sur la période des m années. On notera F_{ij} le flot de liquidités généré par l'obligation j en l'an i .

Année	1	2	...	m	
Obligation					Prix unitaire
B_1	F_{11}	F_{21}	...	F_{m1}	p_1
$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$
B_n	F_{1n}	F_{2n}	...	F_{mn}	p_n

¹Ou *asset and liability management (ALM)* en anglais.

On souhaite déterminer le portefeuille de coût minimal parmi ceux pouvant couvrir les besoins en liquidités sur les m années. On supposera qu'on construit le portefeuille à partir de quantités d'obligations non négatives.

Formuler le problème comme un programme linéaire selon les trois manières suivantes :

- en supposant que le flot de liquidités à l'année i doit a minima couvrir les besoins sur l'année ℓ_i , et en utilisant des contraintes d'inégalité linéaires. Justifier le choix des inégalités linéaires par rapport aux égalités linéaires.
- en supposant que le flot de liquidités à l'année i doit a minima couvrir les besoins sur l'année ℓ_i , et en mettant le problème sous forme standard;
- en supposant qu'un surplus de liquidités peut être reportés d'une année sur l'autre, et en restant en forme standard.

2.6.1 Solution de l'exercice 1

Les variables du problème, concaténées dans un vecteur $\mathbf{x} \in \mathbb{R}^n$, représentent les quantités des différentes obligations. L'objectif consiste à minimiser le coût du portefeuille qui est donné par

$$\mathbf{p}^T \mathbf{x} = \sum_{i=1}^n p_i x_i,$$

où $\mathbf{p}$ est le vecteur de prix qui concatène $p_1, \dots, p_n$. En termes de contrainte, on sait d'après l'énoncé que l'on cherche à avoir $x_i \geq 0$ pour tout $i = 1, \dots, n$.

a) Pour cette première modélisation, on cherche à couvrir les besoins en liquidité sur chaque année uniquement avec les liquidités obtenues durant cette période. Pour chaque année i , cela s'exprime par la contrainte

$$\sum_{j=1}^n F_{ij} x_j \geq \ell_i.$$

Il est à noter que formuler cette contrainte comme une contrainte d'égalité linéaire n'est pas souhaitable, car cela pourrait conduire à ce que le problème soit irréalisable. Au final, on obtient le problème

$$\begin{cases} \text{minimiser}_{\mathbf{x} \in \mathbb{R}^n} & \mathbf{p}^T \mathbf{x} \\ \text{s.c.} & \mathbf{F}^T \mathbf{x} \geq \boldsymbol{\ell}, \quad i = 1, \dots, m \\ & \mathbf{x} \geq \mathbf{0}, \end{cases} \quad (2.6.1)$$

où $\mathbf{F} \in \mathbb{R}^{n \times m}$ est la matrice des F_{ij} (flots de liquidités) et $\boldsymbol{\ell} = [\ell_i] \in \mathbb{R}^m$ est le vecteur des besoins.

b) Pour mettre le problème en forme standard, on transforme les contraintes d'inégalité en contraintes d'égalité en rajoutant des variables d'écart. On obtient ainsi

$$\begin{cases} \text{minimiser}_{\substack{\mathbf{x} \in \mathbb{R}^n \\ \mathbf{s} \in \mathbb{R}^m}} & \mathbf{p}^T \mathbf{x} \\ \text{s.c.} & \mathbf{F}^T \mathbf{x} - \mathbf{s} = \boldsymbol{\ell} \\ & \mathbf{x} \geq \mathbf{0} \\ & \mathbf{s} \geq \mathbf{0}. \end{cases} \quad (2.6.2)$$

c) Les contraintes $F^T x - s = \ell$ formulées en question b) modélisent de fait la présence éventuelle d'un surplus. En effet, le surplus de liquidités en l'an j est donné par $\sum_{i=1}^n F_{ij}x_i - \ell_j = s_j$. Par conséquent, pour modéliser le fait que l'on puisse reporter un surplus d'une année sur l'autre, on remplace la contrainte

$$\sum_{i=1}^n F_{ij}x_i - s_j = \ell_j$$

par

$$\sum_{i=1}^n F_{ij}x_i - s_j + s_{j-1} = \ell_j$$

lorsque $j \geq 2$ (on ne peut pas effectuer de report la première année).

Exercice 2 : Contraintes de levier

Dans les exemples considérés dans ce chapitre, on a considéré pour simplifier que l'on n'effectue que des investissements positifs (*long-only positions*). Cependant, en pratique, on autorise un certain niveau d'investissements négatifs (ou *short positions*), qui représentent une réinjection d'un investissement existant dans un autre fonds. Mathématiquement, on modélise cette liberté par l'introduction d'une contrainte dite **de levier** dans la modélisation, de la forme :

$$\sum_{j=1}^n \min\{x_j, 0\} \geq -L, \quad (2.6.3)$$

où $L > 0$ est un niveau de levier fixé. Cette contrainte est **non linéaire** en les variables x_j du fait de l'utilisation de l'opérateur $\min$.

a) En programmation linéaire, l'inégalité (2.6.3) est en général remplacée par le système

$$\begin{cases} \sum_{j=1}^n z_j & \geq -L \\ z_j \leq x_j & j = 1, \dots, n \\ z_j \leq 0 & j = 1, \dots, n, \end{cases} \quad (2.6.4)$$

où $z \in \mathbb{R}^n$ est un vecteur de variables auxiliaires. Pour tout $x \in \mathbb{R}^n$ vérifiant (2.6.3), donner un vecteur z vérifiant (2.6.4).

b) Convertir le système d'inégalités et d'égalités obtenu en question a) en forme standard ² s'il ne l'est pas déjà.

Solution de l'exercice 2

a) On introduit le vecteur $z \in \mathbb{R}^n$ dont les composantes sont données par $z_j = \min\{x_j, 0\}$. Ce vecteur vérifie bien le système (2.6.4) dès lors que x vérifie (2.6.3).

²Par abus de langage, on considérera qu'un tel système est en forme standard si un programme linéaire avec de telles contraintes le serait.

b) Pour reformuler le système en forme standard, on introduit $2n+1$ variables auxiliaires $\mathbf{s} \in \mathbb{R}^{2n+1}$ qui seront contraintes d'être positives ou nulles. On obtient ainsi le système suivant :

$$\begin{cases} \sum_{j=1}^n z_j - s_1 &= -L \\ z_j + s_{j+1} &= x_j \quad j = 1, \dots, n \\ z_j + s_{j+n+1} &= 0 \quad j = 1, \dots, n \\ \mathbf{s} &\geq \mathbf{0}, \end{cases}$$

qui est bien en forme standard.

Chapitre 3

Optimisation quadratique en finance

Dans ce chapitre, on cherche à étendre les résultats du chapitre précédent au cas où la fonction objectif n'est plus linéaire mais quadratique en les variables de décision. Comme on le verra, il est possible de développer une théorie de la dualité ainsi que des algorithmes similaires à ceux utilisés pour la programmation linéaire. On s'intéressera tout particulièrement au modèle de Markowitz, qui conduit à des programmes quadratiques spécifiques.

3.1 Bases de la programmation quadratique

Définition 3.1.1 Un programme quadratique, ou problème d'optimisation quadratique, est un problème d'optimisation de la forme

$$\begin{aligned} \text{minimiser}_{x \in \mathbb{R}^n} \quad & \frac{1}{2} x^T H x + g^T x \\ \text{s.c.} \quad & A x = b \\ & C x \geq d. \end{aligned}$$

avec $H \in \mathbb{R}^{n \times n}$, $g \in \mathbb{R}^n$, $A \in \mathbb{R}^{m \times n}$, $b \in \mathbb{R}^m$, $C \in \mathbb{R}^{\ell \times n}$, $d \in \mathbb{R}^\ell$.

Définition 3.1.2 Un programme linéaire en **forme standard** est de la forme

$$\begin{aligned} \text{minimiser}_{x \in \mathbb{R}^n} \quad & \frac{1}{2} x^T H x + g^T x \\ \text{s.c.} \quad & A x = b \\ & x \geq 0, \end{aligned} \tag{3.1.1}$$

où les contraintes d'intervalle $x \geq 0$ peuvent porter seulement sur un sous-ensemble des coordonnées de x .

En général, une modélisation en programme quadratique est faite de sorte à satisfaire l'hypothèse ci-dessous.

Hypothèse 3.1.1 La matrice H est symétrique et semi-définie positive.

Un cas particulier de programme quadratique, très important dans la pratique, est décrit dans l'exemple ci-dessous.

Exemple 3.1.1 (Moindres carrés linéaires) Soient $A \in \mathbb{R}^{m \times n}$ et $b \in \mathbb{R}^m$. On cherche un vecteur $x \in \mathbb{R}^n$ tel que $Ax = b$, ce qui peut ou non être faisable. On considère donc les deux cas suivants.

a) S'il n'existe pas de solution à l'équation $Ax = b$, on résout

$$\underset{x \in \mathbb{R}^n}{\text{minimiser}} \frac{1}{2} \|Ax - b\|^2 = \frac{1}{2} x^T A^T A x - b^T A x + \frac{1}{2} b^T b.$$

b) S'il existe plusieurs solutions, on cherche celle de norme minimale en résolvant :

$$\underset{x \in \mathbb{R}^n}{\text{minimiser}} \frac{1}{2} \|x\|^2 = \frac{1}{2} x^T I_n x \quad \text{s. c.} \quad Ax = b,$$

où I_n est la matrice identité dans $\mathbb{R}^{n \times n}$.

3.2 Résolution d'un programme quadratique

3.2.1 Résolution théorique

La théorie de la programmation quadratique est très similaire à celle de la programmation linéaire. On peut notamment formuler le dual d'un programme quadratique, et en déduire des conditions d'existence de solutions pour le problème de départ et son dual. On se contente ici de donner les conditions d'optimalité, dites de KKT (Karush-Kuhn-Tucker).

Théorème 3.2.1 On considère le programme quadratique (3.1.1) sous l'hypothèse 3.1.1. Si $x^* \in \mathbb{R}^n$ est solution du problème, alors il existe $y^* \in \mathbb{R}^m$ et $s^* \in \mathbb{R}^n$ tels que

$$\begin{cases} -Hx^* + A^T y^* + s^* &= g \\ Ax^* &= b \\ x^* &\geq 0 \\ s^* &\geq 0 \\ x_i^* s_i^* &= 0 \quad \forall i = 1, \dots, n. \end{cases} \quad (\text{KKT-Quad})$$

On notera que les conditions (KKT-Quad) sont quasiment identiques à celles pour la programmation linéaire (KKT-Lin). La seule différence réside dans la présence d'un terme en $-Hx^*$ dans la première condition, dû au caractère quadratique de la fonction.

3.2.2 Résolution algorithmique

Comme dans le cas des programmes linéaires, la résolution des programmes quadratiques se base sur les conditions de KKT. Les techniques principales de résolution se divisent en deux catégories.

Méthodes de contraintes actives (active set) Ces méthodes peuvent être vues comme une extension de l'algorithme du simplexe à la programmation quadratique. Ces méthodes cherchent à déterminer l'ensemble des contraintes actives en la solution, c'est-à-dire les contraintes d'inégalité qui sont satisfaites avec égalité en la solution. Dans le cas d'un programme quadratique en forme standard pour lequel la solution est x^* , on dira que la contrainte $x_i \geq 0$ est active en x^* si $x_i^* = 0$. Une contrainte inactive (pour laquelle $x_i^* > 0$) peut être ignorée pour déterminer la solution.

La méthode des contraintes actives utilise à chaque itération un ensemble d'indices $\mathcal{I}_k \subset \{1, \dots, n\}$, et résout le problème

$$\begin{aligned} &\underset{x \in \mathbb{R}^n}{\text{minimiser}} \quad \frac{1}{2} x^T H x + g^T x \\ &\text{s.c.} \quad Ax = b \\ &\quad x_i = 0 \quad i \in \mathcal{I}_k. \end{aligned} \quad (3.2.1)$$

En fonction de la solution, voire de la réalisabilité du problème, le point calculé est conservé, et l'ensemble $\mathcal{I}_k$ est mis à jour en ajoutant ou supprimant un élément, comme pour le simplexe.

Remarque 3.2.1 *Les approches par contraintes actives sont notamment utilisées dans le solveur d'Excel, ainsi que dans des routines de résolution de programmation quadratique présentes dans la bibliothèque SciPy du langage Python.*

Points intérieurs Les approches de points intérieurs pour la programmation quadratique procèdent de manière identique à celles pour la programmation linéaire. Pour un problème en forme standard tel que (3.1.1), il s'agit toujours de calculer une suite $(\mathbf{x}_k, \mathbf{y}_k, \mathbf{s}_k)$ telle que $\mathbf{x}_k > \mathbf{0}$ et $\mathbf{s}_k > \mathbf{0}$, et de réduire la valeur de $\mathbf{x}_k^T \mathbf{s}_k$ progressivement. La valeur des prochains itérés est toujours calculée via les deux premières conditions de KKT, exprimées comme un système linéaire en les variables.

Remarque 3.2.2 *Les solveurs de points intérieurs pour la programmation quadratique sont utilisés dans différentes bibliothèques Python, dont SciPy et cvxpy.*

Pour terminer cette section, on notera que de nombreux solveurs sont capables d'exploiter une structure particulière d'un programme quadratique. Par exemple, les problèmes aux moindres carrés linéaires (cf l'exemple 3.1.1) bénéficient souvent de solveurs dédiés. La librairie cvxpy choisit automatiquement l'algorithme adapté à la structure du problème considéré.

3.3 Application : modèle de Markowitz

3.3.1 Principe

Le modèle de Markowitz, qui valut le prix Nobel d'économie à son auteur en 1990, consiste en la construction d'un portefeuille qui réalise un compromis entre un risque réduit et un rendement moyen élevé.

Le portefeuille peut être défini parmi un ensemble de n actifs qui présentent chacun une part d'aléatoire, et donc un risque, sur une période d'investissement allant d'un instant t_0 à un instant

$t > t_0$. On part ainsi de prix fixés à l'instant t_0 , que l'on écrit sous la forme d'un vecteur $\mathbf{v}_0 = \begin{bmatrix} v_{1,0} \\ \vdots \\ v_{n,0} \end{bmatrix}$

et on cherche à définir un portefeuille en prévision de l'instant t , pour lequel les prix seront les composantes d'un vecteur aléatoire $\mathbf{v}_t = \begin{bmatrix} v_{1,t} \\ \vdots \\ v_{n,t} \end{bmatrix}$. La décision de la construction du portefeuille est

prise à l'instant t_0 , alors que $\mathbf{v}_t$ est inconnu. On suppose que l'on connaît néanmoins des statistiques sur le vecteur $\mathbf{v}_t$, de sorte à pouvoir optimiser une certaine mesure de qualité du portefeuille.

Mathématiquement, on note $\mathbf{h} \in \mathbb{R}^n$ les différentes quantités d'actifs qui composent le portefeuille choisi. Les valeurs du portefeuille à t_0 et à t correspondent à

$$w_0 = \mathbf{v}_0^T \mathbf{h} \quad \text{et} \quad w = \mathbf{v}_t^T \mathbf{h}.$$

En pratique, plutôt que de raisonner sur w et w_0 , on raisonne sur les retours sur investissements, définis via la normalisation suivante :

$$r_p = \frac{w - w_0}{w_0}, \quad \mathbf{r} = \left[r_i = \frac{v_{i,t} - v_{i,0}}{v_{i,0}} \right]_{i=1}^n.$$

Par ailleurs, plutôt que de considérer les quantités d'actifs $\mathbf{h}$, on raisonne sur les pourcentages d'actifs présents dans le portefeuille, que l'on concatène sous la forme d'un vecteur $\mathbf{x} \in \mathbb{R}^n$ avec

$$x_i = \frac{h_i v_{i,0}}{w_0} \in [0, 1].$$

Par construction, on a alors $\sum_{i=1}^n x_i = 1$. De plus, lorsque l'on souhaite maximiser la valeur $w = \mathbf{v}^T \mathbf{h}$ du portefeuille, cela revient à maximiser

$$r_p = \frac{w - w_0}{w_0} = \frac{(\mathbf{v} - \mathbf{v}_0)^T \mathbf{h}}{w_0} = \mathbf{r}^T \mathbf{x},$$

où la dernière égalité s'obtient par

$$\frac{(\mathbf{v} - \mathbf{v}_0)^T \mathbf{h}}{w_0} = \frac{\sum_{i=1}^n (v_i - v_{i,0}) h_i}{w_0} = \sum_{i=1}^n \frac{(v_i - v_{i,0})}{v_{i,0}} \frac{h_i v_{i,0}}{w_0} = \sum_{i=1}^n r_i^T x_i = \mathbf{r}^T \mathbf{x}.$$

On se demande alors quels problèmes d'optimisation peuvent être définis si l'on ne connaît que des statistiques sur le vecteur $\mathbf{r}$ (et donc sur les valeurs à venir des actifs), et à quel genre de portefeuille on aboutit en résolvant ces problèmes

3.3.2 Problèmes d'optimisation

Les problèmes d'optimisation que nous allons définir se basent sur les statistiques suivantes :

- les moyennes $\{\mu_i = \mathbb{E}[r_i]\}_{i=1,\dots,n}$, ou retours sur investissement moyens (par rapport à $\mathbf{v}_0$);
- les variances $\{\sigma_i^2 = \mathbb{E}[(r_i - \mathbb{E}[r_i])^2]\}_{i=1,\dots,n}$;
- les covariances définies pour tous $(i, j) \in \{1, \dots, n\}$ par

$$\sigma_{ij} = \sigma_{ji} = \frac{\mathbb{E}[(r_i - \mathbb{E}[r_i])(r_j - \mathbb{E}[r_j])]}{\sigma_i \sigma_j}.$$

On notera $\boldsymbol{\mu} = [\mu_i]$ et $\mathbf{V} = [\sigma_{ij}]$; on a alors

$$\boldsymbol{\mu} = \mathbb{E}[\mathbf{r}^T \mathbf{x}] \quad \text{et} \quad \mathbf{V} = \mathbb{E}[(\mathbf{r} - \mathbb{E}[\mathbf{r}])(\mathbf{r} - \mathbb{E}[\mathbf{r}])^T] = \mathbb{E}[(\mathbf{r} - \boldsymbol{\mu})(\mathbf{r} - \boldsymbol{\mu})^T].$$

On peut alors formuler les trois problèmes d'optimisation suivants :

a) **Calcul de portefeuille de risque minimal avec retour moyen sur investissement garanti :**

$$\begin{aligned} & \text{minimiser}_{\mathbf{x} \in \mathbb{R}^n} \quad \frac{1}{2} \mathbf{x}^T \mathbf{V} \mathbf{x} \\ & \text{s.c.} \quad \boldsymbol{\mu}^T \mathbf{x} \geq \bar{\mu} \\ & \quad \mathbf{e}^T \mathbf{x} = 1, \end{aligned} \tag{3.3.1}$$

où $\mathbf{e} = [1 \cdots 1]^T$ et $\bar{\mu} \in \mathbb{R}$ représente un retour sur investissement espéré. Ce problème est un programme quadratique.

- b) **Calcul de portefeuille à retour moyen sur investissement maximal pour un niveau de risque contrôlé :**

$$\begin{aligned} & \text{maximiser}_{x \in \mathbb{R}^n} \quad \mu^T x \\ & \text{s.c.} \quad \frac{1}{2} x^T V x \leq \frac{\sigma^2}{2} \\ & \quad e^T x = 1, \end{aligned} \quad (3.3.2)$$

où σ^2 borne le niveau de risque souhaité. Ce problème possède un objectif linéaire et des contraintes dont l'une est quadratique. Il ne s'agit pas d'un programme quadratique¹

- c) **Approche de Markowitz**

$$\begin{aligned} & \text{maximiser}_{x \in \mathbb{R}^n} \quad \mu^T x - \frac{\gamma}{2} x^T V x \\ & \quad e^T x = 1, \end{aligned} \quad (3.3.3)$$

où $\gamma > 0$ est un paramètre de pénalité qui permet de contrôler le poids du terme en $-x^T V x$ dans la maximisation en l'introduisant dans l'objectif et non via une contrainte.

On notera que la dernière formulation (3.3.3) est équivalente au programme quadratique en forme standard :

$$\begin{aligned} & \text{minimiser}_{x \in \mathbb{R}^n} \quad \frac{\gamma}{2} x^T V x - \mu^T x \\ & \quad e^T x = 1. \end{aligned} \quad (3.3.4)$$

La formulation (3.3.4) est un outil de base de calcul de portefeuille qui prend en compte un risque. La valeur de γ représente un compromis entre le retour moyen sur investissement $\mu^T x$ et la volatilité $x^T V x$. Lorsque γ tend vers 0, on tend à minimiser $-\mu^T x$, et donc à maximiser uniquement le retour moyen sur investissement; à l'inverse, lorsque γ tend vers ∞ , on tend à minimiser $x^T V x$ uniquement, et donc à rechercher un portefeuille de risque minimal (cf Section 3.5).

Plus globalement, les modèles définis plus haut peuvent être liés, comme le montre le résultat suivant.

Proposition 3.3.1 *Sous les définitions de cette section, on suppose que V est inversible et que la quantité $\bar{\mu}$ utilisée dans la formulation (3.3.4) vérifie $\bar{\mu} > \frac{\mu^T V^{-1} e}{e^T V^{-1} e}$. Alors, si x^* est une solution de la formulation (3.3.4), il existe σ^2 et γ (qui dépendent de $\bar{\mu}$) tels que x^* est également une solution des formulations (3.3.3) et (3.3.3).*

Ce raisonnement s'applique aux trois formulations considérées, ainsi qu'à la formulation (3.3.4).

3.3.3 Analyse des modèles et de leurs solutions

Dans cette section, on fait l'hypothèse que la matrice V représentant les covariances des actifs est inversible. Il s'agit alors d'une matrice symétrique définie positive, c'est-à-dire telle que $x^T V x > 0$ si $x \neq 0$. Sous cette hypothèse, on peut montrer que le problème

$$\begin{aligned} & \text{minimiser}_{x \in \mathbb{R}^n} \quad \frac{1}{2} x^T V x \\ & \quad e^T x = 1 \end{aligned} \quad (3.3.5)$$

¹Il s'agit cependant d'un programme conique, qui peut lui aussi être résolu par des approches de points intérieurs. En revanche, cette résolution est plus complexe que celle d'un programme quadratique.

possède une unique solution $x^* = \frac{V^{-1}e}{e^T V^{-1}e}$ qui vérifie les conditions de KKT

$$\begin{cases} Vx^* + ey^* &= 0 \\ e^T x^* &= 1 \end{cases}$$

pour un prix fictif $y^* = -\frac{1}{e^T V^{-1}e}$. Le portefeuille ainsi obtenu est noté

$$x_R = \frac{V^{-1}e}{e^T V^{-1}e}$$

et s'appelle le **portefeuille de risque minimum**.

Dans le cadre du problème (3.3.4), on montre de manière similaire qu'il existe deux solutions du problème particulières :

- le vecteur x_R est solution, et minimise le risque parmi les solutions;
- le vecteur $x_T = \frac{V^{-1}\mu}{e^T V^{-1}\mu}$, appelé **portefeuille tangent** (*tangency portfolio*), maximise le retour moyen sur investissement parmi les solutions.

Le résultat suivant permet de lier ces portefeuilles particuliers aux autres solutions.

Théorème 3.3.1 (Two-fund theorem) *Toute solution du problème (3.3.4) avec V inversible est une combinaison des portefeuilles de risque minimum et tangent. Pour toute solution x^* , il existe ainsi $\alpha \in \mathbb{R}$ tel que $x^* = \alpha x_R + (1 - \alpha)x_T$.*

3.4 Conclusion du chapitre 3

La programmation quadratique offre des possibilités de modélisation plus riches que la programmation linéaire, en pouvant prendre en compte des concepts tels que les corrélations entre couples d'actifs. Des algorithmes efficaces existent encore pour résoudre ces problèmes, notamment en se basant sur leur structure.

L'application majeure de l'optimisation quadratique dans le cadre de la finance est due aux problèmes d'optimisation de portefeuille impliquant un modèle moyenne-variance dit de Markowitz. Dans ce cadre, la programmation quadratique est un outil indispensable pour prendre en compte des aspects de risque via un terme quadratique dans l'objectif.

3.5 Exercices

Exercice 3 : Allocation à volatilité minimale

On considère un problème d'allocation de portefeuille sur trois actifs désignés par 1, 2, 3 dont le rendement est une variable aléatoire. On définit les quantités suivantes :

- σ_{ij} est la corrélation entre les actifs i et j pour tout couple $(i, j) \in \{1, 2, 3\}^2$, avec $\sigma_{ii} = 1$ pour tout $i = 1, 2, 3$;
- $C = [\sigma_{ij}] \in \mathbb{R}^{3 \times 3}$ est la matrice des corrélations;
- σ_i^2 est la variance de l'actif i , et σ_i son écart-type;

- $\mathbf{s} \in \mathbb{R}^3$ est le vecteur des écarts-types;
- $\mathbf{V} \in \mathbb{R}^{3 \times 3}$ est la matrice de covariance des actifs, donnée par $\mathbf{V} = \mathbf{C}\mathbf{s}\mathbf{s}^T$.

Le risque ou la volatilité d'un portefeuille $\mathbf{x} \in \mathbb{R}^n$ sera mesuré via la quantité

$$\frac{1}{2} \mathbf{x}^T \mathbf{V} \mathbf{x} = \frac{1}{2} \sum_{i=1}^3 V_{ii} x_i^2 + \sum_{i=1}^3 \sum_{j \neq i} V_{ij} x_i x_j.$$

Le premier terme de la somme combine les variances relatives aux investissements dans un actif, tandis que le second représente le risque dû aux corrélations entre actifs. Le but de cet exercice est de trouver un portefeuille de risque minimal.

- Formuler le problème en tant que programme quadratique, en supposant que le portefeuille est exprimé en pourcentages (sommant à 1) et que ses composantes ne peuvent être que positives ou nulles (*long positions*).
- La résolution numérique de ce problème vue en séance² donne comme solution

$$\mathbf{x}^* = \begin{bmatrix} 0.089 \\ 0 \\ 0.91 \end{bmatrix} \text{ et comme valeur optimale } f^* = 0.0012. \text{ Lorsque les contraintes}$$

$$\mathbf{x} \geq \mathbf{0} \text{ sont supprimées de la modélisation, on obtient en revanche } \mathbf{x}^* = \begin{bmatrix} 0.19 \\ -0.14 \\ 0.94 \end{bmatrix}$$

et $f^* = 0.0009$, qui est une meilleure valeur de fonction. Pourquoi le fait d'enlever une contrainte permet-il d'obtenir une meilleure valeur de l'objectif ?

Solution de l'exercice 3

- Le problème s'écrit

$$\begin{aligned} & \text{minimiser}_{\mathbf{x} \in \mathbb{R}^3} \quad \frac{1}{2} \mathbf{x}^T \mathbf{V} \mathbf{x} \\ & \text{s.c.} \quad x_1 + x_2 + x_3 = 1 \\ & \quad \mathbf{x} \geq \mathbf{0}. \end{aligned}$$

Note : ce problème est en forme standard.

- Lorsque l'on supprime la contrainte $\mathbf{x} \geq \mathbf{0}$, on obtient

$$\begin{aligned} & \text{minimiser}_{\mathbf{x} \in \mathbb{R}^3} \quad \frac{1}{2} \mathbf{x}^T \mathbf{V} \mathbf{x} \\ & \text{s.c.} \quad x_1 + x_2 + x_3 = 1. \end{aligned}$$

L'ensemble admissible de ce problème contient (strictement) celui du problème écrit en question a). Par conséquent, la valeur optimale de ce second problème est au moins aussi bonne que celle du premier problème. Comme on s'autorise plus de liberté dans le choix des x_i , on s'attend à ce que la valeur optimale soit strictement meilleure, et c'est ce que l'on observe ici numériquement.

²Cf les sources associées.

Exercice 4 : Ratio de Sharpe

Le ratio de Sharpe est une alternative au modèle de Markowitz qui combine le retour moyen et la volatilité en une seule quantité, sous réserve que la matrice de covariance V soit inversible. Avec les mêmes notations que celles de la section 3.3, on considère donc le problème :

$$\begin{aligned} & \text{maximiser}_{x \in \mathbb{R}^n} \frac{\mu^T x}{\sqrt{x^T V x}} \\ & \text{s.c. } e^T x = 1 \end{aligned} \quad (3.5.1)$$

Ce problème n'est ni un programme linéaire, ni un programme quadratique. Le but de cet exercice est néanmoins de trouver une reformulation équivalente à ce problème en un programme quadratique.

- a) On considère le changement de variable $z = \alpha x$ avec $\alpha > 0$ et $\mu^T z = 1$. Reformuler le problème (3.5.1) en un problème d'optimisation sur α et z où l'objectif est de la forme $\frac{1}{f(z)}$, où f est une fonction que l'on précisera.
- b) Pour toute fonction strictement positive f , minimiser f revient à maximiser $\frac{1}{\sqrt{f}}$. Par ailleurs, le cas $\alpha = 0$ peut être incorporé au problème. En déduire alors que le problème (3.5.1) se reformule comme le programme quadratique :

$$\begin{aligned} & \text{maximiser}_{\substack{z \in \mathbb{R}^n \\ \alpha \in \mathbb{R}}} z^T V z \\ & \text{s.c. } \begin{aligned} \mu^T z &= 1 \\ e^T z - \alpha &= 0 \\ \alpha &\geq 0. \end{aligned} \end{aligned} \quad (3.5.2)$$

- c) D'après les conditions de KKT, on sait que (z^*, α^*) est solution du problème (3.5.2) s'il existe $y^* \in \mathbb{R}^2$ et $s^* \in \mathbb{R}$ tels que

$$\begin{aligned} -2Vz^* + y_1^* \mu + y_2^* e &= 0 \\ -y_2^* + s^* &= 0 \\ \mu^T z^* &= 1 \\ e^T z^* - \alpha^* &= 0 \\ \alpha^* &\geq 0 \\ \alpha^* s^* &= 0. \end{aligned}$$

Montrer alors que le portefeuille tangent $x^* = \frac{V^{-1}\mu}{e^T V^{-1}\mu}$ définit z^* et α^* qui forment une solution du problème (3.5.2).

Solution de l'exercice 4

- a) On utilise le changement de variable $z = \alpha x$, qui donne $x = \frac{1}{\alpha} z$. En remplaçant les x dans l'expression de (3.5.1), on obtient :

$$\begin{aligned} & \text{maximiser}_{\substack{z \in \mathbb{R}^n \\ \alpha \in \mathbb{R}}} \frac{\frac{1}{\alpha} \mu^T z}{\sqrt{\frac{1}{\alpha} z^T V \left(\frac{1}{\alpha} z\right)}} \\ & \text{s.c. } \begin{aligned} \frac{1}{\alpha} e^T z &= 1 \\ \mu^T z &= 1 \\ \alpha &> 0, \end{aligned} \end{aligned}$$

où les deux dernières contraintes proviennent de la définition même de z et α . En simplifiant les α dans la valeur de l'objectif, et en exploitant la contrainte $\mu^T z = 1$, on aboutit à

$$\begin{aligned} & \underset{\substack{z \in \mathbb{R}^n \\ \alpha \in \mathbb{R}}}{\text{maximiser}} && \frac{1}{\sqrt{z^T V z}} \\ & \text{s.c.} && e^T z - \alpha = 0 \\ & && \mu^T z = 1 \\ & && \alpha > 0, \end{aligned}$$

b) D'après l'indication donnée dans la question, on sait que le problème obtenu en fin de question a) est équivalent à

$$\begin{aligned} & \underset{\substack{z \in \mathbb{R}^n \\ \alpha \in \mathbb{R}}}{\text{minimiser}} && z^T V z \\ & \text{s.c.} && e^T z - \alpha = 0 \\ & && \mu^T z = 1 \\ & && \alpha > 0, \end{aligned}$$

Il suffit ensuite de transformer la contrainte $\alpha > 0$ en $\alpha \geq 0$ comme demandé par l'énoncé pour arriver à (3.5.2).

c) Il s'agit d'utiliser les conditions de KKT pour obtenir une formule explicite pour α^* et z^* en fonction de celle de x^* . On tire tout d'abord de la contrainte $\mu^T z^* = 1$ le fait que

$$\alpha^* = \frac{1}{\mu^T x^*} = \frac{e^T V^{-1} \mu}{\mu^T V^{-1} \mu}.$$

En utilisant cette formule dans la définition de $z^* = \alpha^* x^*$, il vient

$$z^* = \frac{e^T V^{-1} \mu}{\mu^T V^{-1} \mu} x^* = \frac{e^T V^{-1} \mu}{\mu^T V^{-1} \mu} \frac{V^{-1} \mu}{e^T V^{-1} \mu} = \frac{V^{-1} \mu}{\mu^T V^{-1} \mu}.$$

Il est ensuite aisé de vérifier algébriquement que les conditions de KKT sont vérifiées pour ces valeurs.

Chapitre 4

Optimisation stochastique en finance

Dans les chapitres précédents, nous avons soit négligé l'aléatoire en nous concentrant sur le rendement en moyenne (cf chapitre 2), soit pris en compte cet aléatoire au travers de quantités statistiques (rendement moyen et volatilité, cf chapitre 3 et plus particulièrement la section 3.3).

Dans ce chapitre, on va s'intéresser à des formulations dans lesquelles l'aléatoire est présent. On présente tout d'abord un cadre général, en restant sur les formulations en deux temps, puis on décrit une technique de résolution dans le cadre de l'approche par scénarios. On parle enfin de mesures de risque, et de la façon dont elles généralisent les concepts déjà entrevus.

4.1 Formulation générale

4.1.1 Programmes stochastiques sans et avec recours

On considère une fenêtre de temps $[0, 1]$. On suppose que l'on doit prendre une décision au temps $t = 0$ sous la forme d'un vecteur $\mathbf{x} \in \mathbb{R}^n$, mais que la qualité de cette décision est évaluée en fonction d'une variable aléatoire ω , dont la valeur ne sera révélée qu'au temps $t = 1$. Un programme stochastique classique vise à prendre la meilleure solution en moyenne.

Définition 4.1.1 (Programme stochastique) Un **programme stochastique** est un problème d'optimisation de la forme

$$\begin{aligned} \text{minimiser}_{\mathbf{x} \in \mathbb{R}^n} \quad & \mathbb{E}_{\omega} [F(\mathbf{x}, \omega)] \\ \text{s.c.} \quad & \mathbf{x} \in \mathcal{X}, \end{aligned} \tag{4.1.1}$$

où $\omega \in \Omega$ est une variable aléatoire, $F : \mathbb{R}^n \times \Omega \rightarrow \mathbb{R}$ est une fonction et $\mathcal{X} \subset \mathbb{R}^n$.

Un exemple de fonction F est $\mathbf{r}^T \mathbf{x}$, où $\omega = \mathbf{r}$ est un vecteur de rendement aléatoire. Dans ce cas, l'objectif devient de maximiser le rendement moyen.

Il est fréquent qu'une partie de la décision prise à $t = 0$ ne dépende pas de la variable aléatoire. Cela donne lieu à une modélisation particulière du problème d'optimisation, décrite ci-dessous.

Définition 4.1.2 Un **programme stochastique avec recours** est de la forme

$$\begin{aligned} \text{minimiser}_{\mathbf{x} \in \mathbb{R}^n} \quad & f(\mathbf{x}) + \mathbb{E}_{\omega} [Q(\mathbf{x}, \omega)] \\ \text{s.c.} \quad & \mathbf{x} \in \mathcal{X}, \end{aligned} \tag{4.1.2}$$

où $f : \mathbb{R}^n \rightarrow \mathbb{R}$ est une fonction déterministe et $Q : \mathbb{R}^n \times \Omega \rightarrow \mathbb{R}$ est une fonction dite de recours donnée par la solution d'un problème d'optimisation :

$$Q(x, \omega) := \min_{\mathbf{y}(\omega) \in \mathbb{R}^m} g(\mathbf{y}(\omega), \omega) \quad \text{s. c.} \quad \mathbf{y}(\omega) \in \mathcal{Y}(x, \omega). \quad (4.1.3)$$

Le vecteur $\mathbf{y}(\omega)$ représente la décision prise au temps $t = 1$, qui dépend de la valeur de la variable aléatoire ω . Cette décision vise à minimiser une fonction de coût $g : \mathbb{R}^m \times \Omega \rightarrow \mathbb{R}$ sur un ensemble $\mathcal{Y}(x, \omega) \subset \mathbb{R}^m$. Cet ensemble contient toute dépendance de $\mathbf{y}(\omega)$ vis-à-vis de x .

Dans le langage de la programmation stochastique, le problème (4.1.2) s'appelle également un **programme stochastique en deux temps** (*two-phase stochastic programming*). La décision de recours (ou dans un second temps) $\mathbf{y}(\omega)$ est fixe dès lors que la variable ω a été observée et que x est fixé.

Remarque 4.1.1 Les notions présentées dans ce chapitre se généralisent sur des temps multiples (on parle en anglais de multistage stochastic programming). Dans ce cours, on se contentera du cas plus simple de la programmation en deux temps, qui illustre déjà les défis de modélisation.

4.1.2 Programme linéaire stochastique en deux temps

Définition 4.1.3 Un **programme linéaire stochastique en deux temps** est donné par

$$\begin{aligned} \text{minimiser}_{x \in \mathbb{R}^n} \quad & \mathbf{c}^T x + \mathbb{E}_\omega [Q(x, \omega)] \\ \text{s.c.} \quad & \mathbf{A}x = \mathbf{b} \\ & x \geq \mathbf{0}, \end{aligned} \quad (4.1.4)$$

où les contraintes d'intervalle $x \geq \mathbf{0}$ peuvent porter seulement sur un sous-ensemble des coordonnées de x . La fonction Q est définie comme

$$Q(x, \omega) := \min_{\mathbf{y}(\omega) \in \mathbb{R}^m} \begin{aligned} & \mathbf{q}(\omega)^T \mathbf{y}(\omega) \\ \text{s.c.} \quad & \mathbf{T}(\omega)x + \mathbf{N}(\omega)\mathbf{y}(\omega) = \mathbf{h}(\omega) \\ & \mathbf{y}(\omega) \geq \mathbf{0}, \end{aligned} \quad (4.1.5)$$

où $\mathbf{q}(\omega) \in \mathbb{R}^m$, $\mathbf{T}(\omega) \in \mathbb{R}^{\ell \times n}$, $\mathbf{N}(\omega) \in \mathbb{R}^{\ell \times m}$, $\mathbf{h}(\omega) \in \mathbb{R}^\ell$, et les contraintes d'intervalle $\mathbf{y}(\omega) \geq \mathbf{0}$ peuvent porter seulement sur un sous-ensemble des coordonnées de $\mathbf{y}(\omega)$.

On notera qu'en l'absence de recours, le problème serait un programme linéaire en forme standard. Ici, la fonction objectif dépend d'un second programme linéaire (lui aussi en forme standard). Cependant, on observera que les matrices $\mathbf{T}(\omega)$ et $\mathbf{N}(\omega)$, ainsi que les vecteurs $\mathbf{q}(\omega)$ et $\mathbf{h}(\omega)$, peuvent être des fonctions non linéaires de ω .

La définition 4.1.3 met en avant le fait que la décision $\mathbf{y}(\omega)$ est prise relativement à la valeur de la variable ω (tandis que la décision x dite "dans un premier temps" n'en dépend pas). Il est courant d'omettre cette dépendance pour obtenir une formulation simplifiée, de la forme suivante :

$$\begin{aligned} \text{minimiser}_{\substack{x \in \mathbb{R}^n \\ \mathbf{y} \in \mathbb{R}^m}} \quad & \mathbb{E}_\omega [\mathbf{c}^T x + \mathbf{q}^T \mathbf{y}] \\ \text{s.c.} \quad & \mathbf{A}x = \mathbf{b} \\ & \mathbf{T}x + \mathbf{N}\mathbf{y} = \mathbf{h} \\ & x \geq \mathbf{0} \\ & \mathbf{y} \geq \mathbf{0}, \end{aligned} \quad (4.1.6)$$

dans laquelle il est implicite que $\mathbf{q}$, $\mathbf{T}$, $\mathbf{N}$ et $\mathbf{h}$ dépendent de ω .

Exemple 4.1.1 (Problème de distribution) On achète une quantité $x \geq 0$ d'un bien pour répondre à une demande future incertaine. On considère que le coût d'achat à $t = 0$ par unité vaut c_u , que le coût d'achat à $t = 1$ (en cas de stock insuffisant) est de $c_p > c_u$ et que le coût de stockage à $t = 1$ (en cas de surplus) est de c_s . Si $d > 0$ est la valeur incertaine de la demande à $t = 1$, le coût final de l'achat de quantité x est

$$F(x, d) = c_u x + c_p \max(d - x, 0) + c_s \max(x - d, 0).$$

La formulation de ce problème en programme stochastique (au sens de la définition 4.1.1) donne

$$\begin{aligned} & \text{minimiser}_{x \in \mathbb{R}} \quad \mathbb{E}_d [F(x, d)] \\ & \text{s.c.} \quad x \geq 0, \end{aligned}$$

tandis que celle avec recours s'écrit

$$\begin{aligned} & \text{minimiser}_{x \in \mathbb{R}} \quad c_u x + \mathbb{E}_d [Q(x, d)] \\ & \text{s.c.} \quad x \geq 0, \end{aligned}$$

avec $Q(x, d) = c_p \max(d - x, 0) + c_s \max(x - d, 0)$. Cette quantité peut s'interpréter comme la solution d'un programme linéaire

$$Q(x, d) = \begin{aligned} & \min_{y(d), z(d)} \quad c_p y(d) + c_s z(d) \\ & \text{s. c.} \quad y(d) \geq d - x \\ & \quad \quad z(d) \geq x - d \\ & \quad \quad y(d) \geq 0 \\ & \quad \quad z(d) \geq 0. \end{aligned}$$

La reformulation avec dépendance implicite en d , à la manière de (4.1.6), s'écrit :

$$\begin{aligned} & \text{minimiser}_{\substack{x \in \mathbb{R} \\ y \in \mathbb{R} \\ z \in \mathbb{R}}} \quad \mathbb{E}_d [c_u x + c_p y + c_s z] \\ & \text{s. c.} \quad x + y \geq d \\ & \quad \quad -x + z \geq -d \\ & \quad \quad x, y, z \geq 0. \end{aligned}$$

4.2 Approche par scénarios

4.2.1 Principe

L'approche par scénarios est l'une des techniques les plus fréquentes en programmation stochastique. Elle consiste à n'envisager qu'un nombre fini de valeurs possibles pour la variable aléatoire ω , appelés **scénarios**. On définit ainsi des scénarios $\{\omega_1, \dots, \omega_K\}$ ainsi que leurs probabilités d'occurrence $\{p_1, \dots, p_K\}$. Par conséquent, ces probabilités vérifient $\sum_{k=1}^K p_k = 1$.

L'intérêt de cette approche réside dans le fait qu'il est maintenant possible de réécrire un programme stochastique de manière déterministe, en introduisant une variable pour chaque scénario possible. Dans le cadre d'un programme en deux temps, cela signifie que l'on définit une variable y_k associée au scénario ω_k .

Proposition 4.2.1 *Sous l'hypothèse de K scénarios sur ω , le programme linéaire stochastique en deux temps (4.1.6) s'écrit de manière déterministe comme*

$$\begin{aligned} \underset{\substack{\mathbf{x} \in \mathbb{R}^n \\ \mathbf{y}_1, \dots, \mathbf{y}_K \in \mathbb{R}^m}}{\text{minimiser}} \quad & \mathbf{c}^T \mathbf{x} + \sum_{k=1}^K p_k \mathbf{q}_k^T \mathbf{y}_k \\ \text{s.c.} \quad & \mathbf{A} \mathbf{x} = \mathbf{b} \\ & \mathbf{T}_k \mathbf{x} + \mathbf{N}_k \mathbf{y}_k = \mathbf{h}_k \quad k = 1, \dots, K \\ & \mathbf{x} \geq \mathbf{0} \\ & \mathbf{y}_k \geq \mathbf{0} \quad k = 1, \dots, K. \end{aligned} \quad (4.2.1)$$

La structure très particulière du problème (4.2.1) va faciliter sa résolution. En effet, à $\mathbf{x}$ fixé, les $\mathbf{y}_k$ peuvent être calculés indépendamment en résolvant k programmes linéaires. C'est le principe de la méthode décrite dans la section suivante.

4.2.2 Décomposition de Benders pour programmes linéaires stochastiques

Dans ce qui suit, on réécrit le problème (4.2.1) comme

$$\begin{aligned} \underset{\mathbf{x} \in \mathbb{R}^n}{\text{minimiser}} \quad & \mathbf{c}^T \mathbf{x} + \sum_{k=1}^K P_k(\mathbf{x}) \\ \text{s.c.} \quad & \mathbf{A} \mathbf{x} = \mathbf{b} \\ & \mathbf{x} \geq \mathbf{0}, \end{aligned} \quad (4.2.2)$$

où chaque valeur $P_k(\mathbf{x})$ est la valeur optimale d'un programme linéaire déterminé par $\mathbf{x}$:

$$P_k(\mathbf{x}) = \underset{\mathbf{y}_k}{\text{min}} \quad p_k \mathbf{q}_k^T \mathbf{y}_k \quad \text{s. c.} \quad \begin{aligned} & \mathbf{N}_k \mathbf{y}_k = \mathbf{h}_k - \mathbf{T}_k \mathbf{x} \\ & \mathbf{y}_k \geq \mathbf{0}. \end{aligned} \quad (4.2.3)$$

L'algorithme dit de décomposition de Benders procède itérativement en partant de $\mathbf{x}^0 \in \mathbb{R}^n$, typiquement choisi comme

$$\mathbf{x}^0 \in \underset{\mathbf{x} \in \mathbb{R}^n}{\text{argmin}} \{ \mathbf{c}^T \mathbf{x} \mid \mathbf{A} \mathbf{x} = \mathbf{b}, \mathbf{x} \geq \mathbf{0} \}.$$

Pour chaque itération i , l'algorithme calcule les valeurs $\{P_k(\mathbf{x}^i)\}$ en résolvant les programmes linéaires associés, pour lesquels on suppose que leurs duals sont toujours réalisables. On se trouve alors dans l'une des situations suivantes.

a) Soit le problème (4.2.3) possède une solution, et dans ce cas il existe $\mathbf{u}_k^i$ tel que

$$\forall \mathbf{x}, P_k(\mathbf{x}) \geq [\mathbf{u}_k^i]^T (\mathbf{T}_k \mathbf{x}^i - \mathbf{T}_k \mathbf{x}) + P_k(\mathbf{x}^i).$$

b) Soit le problème est irréalisable, et il existe alors $\mathbf{u}_k^i$ tel que

$$[\mathbf{u}_k^i]^T (\mathbf{h}_k - \mathbf{T}_k \mathbf{x}) \leq 0.$$

Dans les deux cas, il s'agit de contraintes linéaires (appelées coupes) que l'on peut rajouter au problème utilisé pour calculer $\mathbf{x}^i$ de sorte à obtenir un meilleur $\mathbf{x}^{i+1}$. Ce vecteur sera en effet calculé

en résolvant

$$\begin{aligned}
 & \text{minimiser}_{\substack{\mathbf{x} \in \mathbb{R}^n \\ \mathbf{y}_1, \dots, \mathbf{y}_K \in \mathbb{R}^m}} \quad \mathbf{c}^T \mathbf{x} + \sum_{k=1}^K P_k(\mathbf{x}) \\
 & \text{s.c.} \quad \mathbf{Ax} = \mathbf{b} \\
 & \quad P_k(\mathbf{x}) \geq [\mathbf{u}_k^i]^T (\mathbf{T}_k \mathbf{x}^i - \mathbf{T}_k \mathbf{x}) + P_k(\mathbf{x}^i) \\
 & \quad \quad \forall (i, k) \text{ tel que le problème associé à } P_k(\mathbf{x}^i) \text{ ait une solution,} \\
 & \quad [\mathbf{u}_k^i]^T (\mathbf{h}_k - \mathbf{T}_k \mathbf{x}) \leq 0 \\
 & \quad \quad \forall (i, k) \text{ tel que le problème associé à } P_k(\mathbf{x}^i) \text{ soit irréalisable,} \\
 & \quad \mathbf{x} \geq \mathbf{0} \\
 & \quad \mathbf{y}_k \geq \mathbf{0} \quad k = 1, \dots, K.
 \end{aligned} \tag{4.2.4}$$

L'avantage de cette procédure est qu'elle ne requiert que de résoudre des programmes linéaires (au nombre de contraintes croissant). On peut montrer que ce processus converge vers la solution en un nombre fini d'itérations.

4.3 Mesures de risque

4.3.1 Motivation et définition générique

Dans le chapitre 3, et plus particulièrement dans la section 3.3, nous avons étudié une fonction de coût mêlant la variance d'un vecteur de rendement ainsi que sa moyenne. Ce programme peut s'écrire comme un programme stochastique, où la composante stochastique est le rendement. On obtient ainsi

$$\begin{aligned}
 & \text{minimiser}_{\mathbf{x} \in \mathbb{R}^n} \quad f(\mathbf{r}^T \mathbf{x}) \\
 & \text{s.c.} \quad \mathbf{e}^T \mathbf{x} = 1
 \end{aligned} \tag{4.3.1}$$

avec

$$f(t) = \frac{\gamma}{2} \sigma_{\mathbf{r}}^2(t) - \mathbb{E}_{\mathbf{r}}[t], \quad \sigma_{v_r}^2(t) = \mathbb{E}_{\mathbf{r}}[t^2] - \mathbb{E}_{\mathbf{r}}[t]^2.$$

A travers le terme de variance, f incorpore une information liée au risque du portefeuille. Plus généralement, les fonctions de ce type sont appelées des **mesures de risque**.

Remarque 4.3.1 L'utilisation, déjà vue au chapitre (2), du rendement moyen comme mesure conduit à un problème dont le résultat sera **neutre au risque** (risk neutral), par opposition aux autres mesures de risque, qui elles seront attentives au risque (risk averse).

4.3.2 VaR et CVaR

L'utilisation de la variance présente l'inconvénient de diminuer l'investissement dans des actifs qui peuvent prendre des valeurs extrêmes (et donc avoir une forte variance), même si cela n'arrive qu'avec une faible probabilité. Il est possible de définir d'autres mesures de risque qui soient moins sensibles à cela. C'est notamment le but de la valeur à risque (*value-at-risk*), définie ci-dessous et popularisée par J.P. Morgan.

Définition 4.3.1 Soit z une valeur aléatoire et $\alpha \in (0, 1)$ un niveau de sûreté. La **valeur à risque α de z** est la valeur δ telle que

$$\mathbb{P}(z \geq \delta) = 1 - \alpha. \tag{4.3.2}$$

On note alors $\text{VaR}_{\alpha}[z] = \delta$.

Exemple 4.3.1 i) Si $z \sim \mathcal{N}(\mu, \sigma^2)$, alors $\text{VaR}_{0.99}[z] = \mu + 2.33\sigma$.

ii) Si $z \in \{z_1, \dots, z_K\}$ avec $\mathbb{P}(z = z_k) = p_k$, alors $\text{VaR}_\alpha[z] = z_J$, où J est le plus petit entier tel que $\sum_{k=J}^K p_k \geq 1 - \alpha$.

On peut ainsi formuler des problèmes d'optimisation où l'on cherche à minimiser la valeur à risque, ce qui donne généralement des solutions plus robustes aux comportements extrêmes qu'en utilisant la variance. Dans le cas d'une approche par scénarios, on retombe à nouveau sur un programme linéaire.

Valeur à risque conditionnelle L'un des inconvénients de la valeur à risque est qu'il ne s'agit pas d'une mesure de risque cohérente, car $\text{VaR}_\alpha[z + y] \not\leq \text{VaR}_\alpha[z] + \text{VaR}_\alpha[y]$ en général. Cela signifie que diversifier un portefeuille n'apporte pas de robustesse au sens de la valeur à risque. La variante cohérente de la valeur à risque la plus répandue est définie ci-après.

Définition 4.3.2 Soit z une valeur aléatoire et $\alpha \in (0, 1)$ un niveau de sûreté. La **valeur à risque conditionnelle** α de z , notée $\text{CVaR}_\alpha[z]$, est définie par

$$\text{CVaR}_\alpha[z] = \mathbb{E}_z[z | z \geq \text{VaR}_\alpha[z]]. \quad (4.3.3)$$

L'interprétation de la valeur à risque conditionnelle est qu'elle exprime, pour une valeur à risque donnée, de combien on excèdera cette valeur en moyenne. Il s'agit d'une mesure de risque cohérente.

Exemple 4.3.2 i) Si $z \sim \mathcal{N}(\mu, \sigma^2)$, alors $\text{VaR}_{0.99}[z] = \mu + 2.67\sigma$.

ii) Si $z \in \{z_1, \dots, z_K\}$ avec $\mathbb{P}(z = z_k) = p_k$, alors

$$\text{CVaR}_\alpha[z] = \frac{1}{1 - \alpha} \sum_{k=J}^K p_k,$$

où J est le plus petit entier tel que $\sum_{k=J+1}^K p_k \in [1 - \alpha - p_J, 1 - \alpha)$.

Le calcul de la valeur à risque conditionnelle n'est en général pas possible de manière explicite, car il correspond à la résolution d'un programme linéaire stochastique en deux temps, comme le montre le résultat ci-dessous.

Théorème 4.3.1 Si z une valeur aléatoire et $\alpha \in (0, 1)$, alors

$$\begin{aligned} \text{CVaR}_\alpha[z] &= \min_{\delta} \left[\delta + \frac{1}{1 - \alpha} \mathbb{E}_z[\max\{z - \delta, 0\}] \right] \\ &= \min_{\delta, u} \quad \delta + \frac{1}{1 - \alpha} \mathbb{E}_z[u] \\ &\quad \text{s.c.} \quad u \geq z - \delta \\ &\quad \quad u \geq 0. \end{aligned} \quad (4.3.4)$$

En règle générale, l'utilisation de la valeur à risque conditionnelle conduira donc au moins à une modélisation en programme stochastique avec recours.

4.4 Conclusion du chapitre 4

La prise en compte explicite du caractère aléatoire d'une composante de l'environnement (typiquement, le rendement d'un actif) conduit à des modélisations et donc des problèmes d'optimisation plus complexes. Parmi ceux-ci, les programmes stochastiques en deux temps font partie des classes les mieux comprises sur le plan théorique ainsi que sur le plan algorithmique. Les approches par scénarios permettent également de décrire ces problèmes de manière déterministe. Dans le cas où ces programmes sont aussi linéaires, des approches telles que la décomposition de Benders peuvent exploiter la structure inhérente à la présence de scénarios.

L'un des atouts majeurs de la programmation stochastique est de pouvoir quantifier le risque de certains investissements, de manière plus flexible que ce qui peut être fait avec la variance. Cela est rendu possible grâce à la notion de mesure de risque, dont la valeur à risque et la valeur à risque conditionnelle forment deux importantes instances.

4.5 Exercices

Exercice 5 : Déviation en moyenne absolue

Comme on l'a vu plus haut, l'utilisation de la variance en tant que mesure de risque présente quelques inconvénients, notamment celui d'être sensible aux comportements aberrants, et donc de conduire à des solutions conservatrices. On s'intéresse ici à une autre métrique définie via la **déviation en moyenne absolue**.

Soit $\mathbf{r} \in \mathbb{R}^n$ un vecteur de retours sur investissement aléatoire de moyenne $\boldsymbol{\mu} = \mathbb{E}[\mathbf{r}]$. La déviation en moyenne absolue de $\mathbf{r}$ est la fonction $d : \mathbb{R}^n \rightarrow \mathbb{R}$ définie par

$$\forall \mathbf{x} \in \mathbb{R}^n, \quad d(\mathbf{x}) = \mathbb{E}[|(\mathbf{r} - \boldsymbol{\mu})^T \mathbf{x}|]. \quad (4.5.1)$$

- a) Supposons que $\mathbf{r}$ suive une loi discrète de valeurs $\{\mathbf{r}_1, \dots, \mathbf{r}_K\}$ avec $\mathbb{P}(\mathbf{r} = \mathbf{r}_k) = p_k$. Écrire alors le problème

$$\begin{aligned} & \text{minimiser}_{\mathbf{x} \in \mathbb{R}^n} \quad d(\mathbf{x}) \\ & \text{s.c.} \quad \mathbf{e}^T \mathbf{x} = 1, \end{aligned} \quad (4.5.2)$$

en fonction de $\{\mathbf{r}_k\}_{k=1, \dots, K}$ et $\{p_k\}_{k=1, \dots, K}$.

- b) En introduisant la décomposition

$$\begin{cases} (\mathbf{r}_k - \boldsymbol{\mu})^T \mathbf{x} &= w_k^+ - w_k^- \\ |(\mathbf{r}_k - \boldsymbol{\mu})^T \mathbf{x}| &= w_k^+ + w_k^-, \end{cases}$$

où w_k^+ et w_k^- sont des quantités positives ou nulles (représentant la partie positive et négative de $(\mathbf{r}_k - \boldsymbol{\mu})^T \mathbf{x}$), formuler le problème (4.5.2) comme un programme linéaire.

- c) Donner un avantage pratique de résoudre un programme linéaire tel que celui formulé dans cet exercice par rapport à la résolution de programmes quadratiques impliquant la variance.

Solution de l'exercice 5

a) Par définition de l'espérance, on a dans le cas de cette question

$$d(\mathbf{x}) = \mathbb{E} [|(\mathbf{r} - \boldsymbol{\mu})^T \mathbf{x}|] = \sum_{i=1}^K p_k |(\mathbf{r}_k - \boldsymbol{\mu})^T \mathbf{x}|.$$

Le problème s'écrit donc

$$\begin{aligned} & \text{minimiser}_{\mathbf{x} \in \mathbb{R}^n} \quad \sum_{i=1}^K p_k |(\mathbf{r}_k - \boldsymbol{\mu})^T \mathbf{x}| \\ & \text{s.c.} \quad \mathbf{e}^T \mathbf{x} = 1. \end{aligned}$$

b) En utilisant la décomposition, on peut reformuler le problème en remplaçant l'objectif par

$$\sum_{i=1}^K p_k |(\mathbf{r}_k - \boldsymbol{\mu})^T \mathbf{x}| = \sum_{i=1}^K p_k (w_k^+ - w_k^-)$$

et en rajoutant les contraintes

$$(\mathbf{r}_k - \boldsymbol{\mu})^T \mathbf{x} = w_k^+ - w_k^-.$$

On obtient donc le problème

$$\begin{aligned} & \text{minimiser}_{\substack{\mathbf{x} \in \mathbb{R}^n \\ \mathbf{w}^+ \in \mathbb{R}^K \\ \mathbf{w}^- \in \mathbb{R}^K}} \quad \sum_{i=1}^K p_k (w_k^+ - w_k^-) \\ & \text{s.c.} \quad (\mathbf{r}_k - \boldsymbol{\mu})^T \mathbf{x} - w_k^+ + w_k^- = 0 \quad k = 1, \dots, K, \\ & \quad \mathbf{e}^T \mathbf{x} = 1. \end{aligned}$$

c) Un programme linéaire peut se résoudre avec un plus grand nombre de variables qu'un programme quadratique, car il existe des algorithmes très efficaces en grande dimension.

Exercice 6 : Utilité

On considère que l'on dispose d'un capital de $w_0 > 0$ à investir pour construire un portefeuille. En tant que personne, on dispose d'une utilité qui nous est propre représentée par une fonction $U : \mathbb{R} \rightarrow \mathbb{R}$, que l'on va chercher à maximiser via la construction du portefeuille. On note $w = w_0(1 + \mathbf{r}^T \mathbf{x})$ le capital obtenu après investissement, dont on suppose qu'il dépend d'un vecteur de rendement aléatoire.

Le problème de construction de portefeuille d'utilité maximale s'écrit comme suit :

$$\begin{aligned} & \text{maximiser}_{\mathbf{x} \in \mathbb{R}^n} \quad \mathbb{E}_r [U(w)] \\ & \text{s. c.} \quad \begin{aligned} w &= w_0(1 + \mathbf{r}^T \mathbf{x}) \\ \mathbf{e}^T \mathbf{x} &= 1. \end{aligned} \end{aligned} \quad (4.5.3)$$

On se concentre dans la suite sur le cas $U : w \mapsto U(w)$.

- Si $\mathbf{r} = \boldsymbol{\mu}$ est déterministe, reformuler le problème (4.5.3) en un programme linéaire. Utiliser pour cela le fait que maximiser $\log(f(\mathbf{x}))$ revient à minimiser $f(\mathbf{x})$.
- Généraliser le résultat précédent dans le cas où $\mathbf{r} \in \{\mathbf{r}_1, \dots, \mathbf{r}_K\}$ avec $\mathbb{P}(\mathbf{r} = \mathbf{r}_k) = p_k$.

c) Si $\mathbf{r}$ suit une loi normale de moyenne $\boldsymbol{\mu}$ et de matrice de covariance $\mathbf{V}$, on a alors

$$\mathbb{E}_r [U(w)] = -\frac{c}{2}(\mathbf{x} - \boldsymbol{\mu})^T \mathbf{V}(\mathbf{x} - \boldsymbol{\mu})$$

pour $c > 0$. Montrer que dans ce cas, le problème (4.5.3) se ramène à un problème quadratique. Comparer le problème obtenu avec celui de gestion de portefeuille de la forme proposée par Markowitz.

Solution de l'exercice 6

a) Dans le cas d'un vecteur $\mathbf{r}$ déterministe, on a

$$\mathbb{E}_r [U(w)] = U(w_0(1 + \boldsymbol{\mu}^T \mathbf{x})) = \log(w_0(1 + \boldsymbol{\mu}^T \mathbf{x})).$$

Par conséquent, le problème (4.5.3) devient

$$\begin{array}{ll} \text{maximiser}_{\mathbf{x} \in \mathbb{R}^n} & \log(w_0(1 + \boldsymbol{\mu}^T \mathbf{x})) \\ \text{s. c.} & \mathbf{e}^T \mathbf{x} = 1. \end{array}$$

En utilisant le fait que maximiser un logarithme revient à minimiser la fonction, une reformulation équivalente de ce problème est

$$\begin{array}{ll} \text{minimiser}_{\mathbf{x} \in \mathbb{R}^n} & w_0 \boldsymbol{\mu}^T \mathbf{x} \\ \text{s. c.} & \mathbf{e}^T \mathbf{x} = 1, \end{array}$$

qui est bien un programme linéaire.

b) Lorsque $\mathbf{r}$ suit une distribution discrète, la fonction objectif du problème (4.5.3) s'écrit

$$\begin{aligned} \mathbb{E}_r [U(w)] &= \mathbb{E}_r [U(w_0(1 + \mathbf{r}^T \mathbf{x}))] \\ &= \sum_{k=1}^K p_k U(w_0(1 + \mathbf{r}_k^T \mathbf{x})) \\ &= \log \left(\prod_{k=1}^K w_0^{p_k} (1 + \mathbf{r}_k^T \mathbf{x})^{p_k} \right). \end{aligned}$$

Par conséquent, en utilisant le fait que maximiser un logarithme revient à minimiser la fonction argument du logarithme, on arrive à

$$\begin{array}{ll} \text{minimiser}_{\mathbf{x} \in \mathbb{R}^n} & \prod_{k=1}^K w_0^{p_k} (1 + \mathbf{r}_k^T \mathbf{x})^{p_k} \\ \text{s. c.} & \mathbf{e}^T \mathbf{x} = 1, \end{array}$$

On notera que ce problème n'est un programme linéaire que lorsque $K = 1$ et $p_1 = 1$.¹

¹Sinon, il s'agit d'un programme dit géométrique, pour lequel des techniques telles que celles vues en cours peuvent être développées.

c) En introduisant la première contrainte du problème (4.5.3) dans l'objectif, on aboutit au problème

$$\begin{array}{ll} \text{maximiser}_{x \in \mathbb{R}^n} & -\frac{c}{2}(x - \mu)^T V(x - \mu) \\ \text{s. c.} & e^T x = 1, \end{array}$$

qui est équivalent à

$$\begin{array}{ll} \text{minimiser}_{x \in \mathbb{R}^n} & \frac{c}{2}(x - \mu)^T V(x - \mu) \\ \text{s. c.} & e^T x = 1. \end{array}$$

Comme multiplier par une constante positive ne change pas les solutions d'un problème d'optimisation, ce dernier problème est aussi équivalent à

$$\begin{array}{ll} \text{minimiser}_{x \in \mathbb{R}^n} & \frac{1}{2}(x - \mu)^T V(x - \mu) \\ \text{s. c.} & e^T x = 1 \end{array}$$

En développant la fonction objectif, on obtient

$$\begin{aligned} \frac{1}{2}(x - \mu)^T V(x - \mu) &= \frac{1}{2}(x^T V x - 2\mu^T V x + \mu^T V \mu) \\ &= \frac{1}{2}x^T V x - \mu^T V x + \frac{1}{2}\mu^T V \mu. \end{aligned}$$

Le dernier terme étant constant, il n'importe pas dans le processus de minimisation. Au final, on obtient donc le problème

$$\begin{array}{ll} \text{minimiser}_{x \in \mathbb{R}^n} & \frac{1}{2}x^T V x - \mu^T V x \\ \text{s. c.} & e^T x = 1. \end{array}$$

Ce problème ressemble au problème d'allocation de Markowitz (avec $\gamma = 1$). Cependant, au lieu d'avoir un terme en $-\mu^T x$ qui tend à maximiser le rendement, on a ici un terme en $-\mu^T V x$, qui tend à maximiser une version normalisée du rendement moyen.

Bibliographie

- [1] J. R. Birge and F. Louveaux. *Introduction to Stochastic Programming*. Springer Series in Operations Research and Financial Engineering. Springer, New York, second edition, 2011.
- [2] S. Boyd and L. Vandenberghe. *Convex optimization*. Cambridge University Press, Cambridge, United Kingdom, 2004.
- [3] G. Cornuéjols, J. Peña, and R. Tütüncü. *Optimization Methods in Finance*. Cambridge University Press, Cambridge, United Kingdom, second edition, 2018.

Annexe A

Bases mathématiques

A.1 Éléments d'algèbre linéaire

Les fondements mathématiques de l'optimisation se trouvent dans l'analyse réelle, et en particulier dans le calcul différentiel. Les structures d'algèbre linéaire dans $\mathbb{R}^n$ jouent cependant un rôle prépondérant en optimisation (et en sciences des données). Cette section regroupe les résultats de base qui seront utilisés dans le cours.

Pour approfondir ces notions, on pourra consulter les liens suivants :

- <https://www.ceremade.dauphine.fr/~carlier/polyalgebre.pdf> (en français);
- <http://vmls-book.stanford.edu/vmls.pdf> (Chapitres 1 à 3, en anglais).

A.1.1 Algèbre linéaire vectorielle

On considèrera toujours l'espace des vecteurs $\mathbb{R}^n$ muni de sa structure d'espace vectoriel normé de dimension d . On définit donc les opérations suivantes :

- Pour tous $\mathbf{x}, \mathbf{y} \in \mathbb{R}^n$, la somme des vecteurs $\mathbf{x}$ et $\mathbf{y}$ est notée $\mathbf{x} + \mathbf{y} = [x_i + y_i]_{1 \leq i \leq d}$;
- Pour tout $\lambda \in \mathbb{R}$, on définit $\lambda \mathbf{x} \stackrel{n}{=} \lambda \cdot \mathbf{x} = [\lambda x_i]_{1 \leq i \leq d}$. Dans ce contexte, un tel réel λ sera appelé un *scalaire*.

A l'aide de ces opérations, nous pouvons donc construire des **combinaisons linéaires** de vecteurs de $\mathbb{R}^n$ dont le résultat sera un vecteur de $\mathbb{R}^n$, à savoir des vecteurs de la forme $\sum_{i=1}^p \lambda_i \mathbf{x}_i$, où $\mathbf{x}_i \in \mathbb{R}^n$ et $\lambda_i \in \mathbb{R}$ pour tout $i = 1, \dots, p$.

Pour ce qui est de l'espace des matrices $\mathbb{R}^{n \times d}$, on peut également le munir d'une structure d'espace vectoriel normé de dimension nd . On définit donc les opérations suivantes :

- Pour tous $\mathbf{A}, \mathbf{B} \in \mathbb{R}^{n \times d}$, la somme des matrices $\mathbf{A}$ et $\mathbf{B}$ est notée $\mathbf{A} + \mathbf{B} = [\mathbf{A}_{ij} + \mathbf{B}_{ij}]_{\substack{1 \leq i \leq d \\ 1 \leq j \leq d}}$;
- Pour tout scalaire $\lambda \in \mathbb{R}$, on définit $\lambda \mathbf{A} \stackrel{n}{=} \lambda \cdot \mathbf{A} = [\lambda \mathbf{A}_{ij}]_{\substack{1 \leq i \leq n \\ 1 \leq j \leq d}}$.

Définition A.1.1 Un ensemble $\mathcal{S} \subseteq \mathbb{R}^n$ vérifiant les conditions

1. $\mathbf{0}_d \in \mathcal{S}$;

2. $\forall (\mathbf{x}, \mathbf{y}) \in \mathcal{S}, \mathbf{x} + \mathbf{y} \in \mathcal{S};$
3. $\forall \mathbf{x} \in \mathcal{S}, \forall \lambda \in \mathbb{R}, \lambda \mathbf{x} \in \mathcal{S}.$

s'appelle un sous-espace vectoriel de $\mathbb{R}^n$.

Définition A.1.2 Soient $\mathbf{x}_1, \dots, \mathbf{x}_p$ p vecteurs de $\mathbb{R}^n$. Le **sous-espace engendré** par les vecteurs $\mathbf{x}_1, \dots, \mathbf{x}_p$, noté $\text{vect}(\mathbf{x}_1, \dots, \mathbf{x}_p)$, est le sous-espace vectoriel

$$\text{vect}(\mathbf{x}_1, \dots, \mathbf{x}_p) := \left\{ \mathbf{x} = \sum_{i=1}^p \alpha_i \mathbf{x}_i \mid \alpha_i \in \mathbb{R} \forall i \right\}.$$

On rappelle ensuite les différentes propriétés notables des familles de vecteurs.

Définition A.1.3 • Une famille libre $\{\mathbf{x}_i\}_{i=1}^k$ de vecteurs de $\mathbb{R}^n$ est telle que les vecteurs sont linéairement indépendants : pour tous scalaires $\lambda_1, \dots, \lambda_k$ tels que $\sum_{i=1}^k \lambda_i \mathbf{x}_i = \mathbf{0}$, alors $\lambda_1 = \dots = \lambda_k = 0$. On a nécessairement $k \leq n$.

- Une famille liée est une famille de vecteurs qui n'est pas libre.
- Une famille génératrice de vecteurs de $\mathbb{R}^n$ est un ensemble de vecteurs $\{\mathbf{x}_i\}$ tel que le sous-espace engendré par ces vecteurs soit égal à $\mathbb{R}^n$.
- Une base de $\mathbb{R}^n$ est une famille de vecteurs $\{\mathbf{x}_i\}_{i=1}^n$ qui est à la fois libre et génératrice. Tout vecteur de $\mathbb{R}^n$ s'écrit alors de manière unique comme combinaison linéaire des $\mathbf{x}_i$. Toute base de $\mathbb{R}^n$ comporte exactement n vecteurs.

Comme la taille d'une base de $\mathbb{R}^n$ est n , on dit que cet espace vectoriel est de dimension n . En conséquence, la dimension d'un sous-espace vectoriel est au plus n .

Exemple A.1.1 Tout vecteur $\mathbf{x}$ de $\mathbb{R}^n$ s'écrit $\mathbf{x} = \sum_{i=1}^n x_i \mathbf{e}_i$, où $\mathbf{e}_i = [0 \dots 0 \ 1 \ 0 \dots 0]^T$ est le i -ème vecteur de la base canonique (le coefficient 1 se trouvant en i -ème position).

Norme et produit scalaire L'utilisation d'une norme, et du produit scalaire associé, permet de mesurer les distances entre vecteurs, ce qui sera particulièrement utile pour montrer la convergence d'une suite de points générés par un algorithme vers une solution d'un problème d'optimisation.

Définition A.1.4 La **norme euclidienne** $\|\cdot\|$ sur $\mathbb{R}^n$ est définie pour tout vecteur $\mathbf{x} \in \mathbb{R}^n$ par

$$\|\mathbf{x}\| := \sqrt{\sum_{i=1}^n x_i^2}.$$

Remarque A.1.1 Il s'agit bien d'une norme, car elle vérifie les quatres axiomes d'une norme :

1. $\forall \mathbf{x}, \mathbf{y} \in \mathbb{R}^n, \|\mathbf{x} + \mathbf{y}\| \leq \|\mathbf{x}\| + \|\mathbf{y}\|;$
2. $\|\mathbf{x}\| = 0 \Leftrightarrow \mathbf{x} = \mathbf{0}_{\mathbb{R}^n};$
3. $\forall \mathbf{x}, \|\mathbf{x}\| \geq 0;$

$$4. \forall \mathbf{x} \in \mathbb{R}^n, \forall \lambda \in \mathbb{R}, \|\lambda \mathbf{x}\| = |\lambda| \|\mathbf{x}\|.$$

On dira qu'un vecteur $\mathbf{x} \in \mathbb{R}^n$ est unitaire si $\|\mathbf{x}\| = 1$.

Définition A.1.5 Pour tous vecteurs $\mathbf{x}, \mathbf{y} \in \mathbb{R}^n$, le **produit scalaire** dérivé de la norme euclidienne est une fonction de $\mathbf{x}$ et $\mathbf{y}$, notée $\mathbf{x}^T \mathbf{y}$, et définie par

$$\mathbf{x}^T \mathbf{y} := \sum_{i=1}^n x_i y_i.$$

Deux vecteurs $\mathbf{x}$ et $\mathbf{y}$ tels que $\mathbf{x}^T \mathbf{y} = 0$ seront dits orthogonaux.

On notera en particulier que $\mathbf{y}^T \mathbf{x} = \mathbf{x}^T \mathbf{y}$. Ce produit scalaire définit donc un "produit" entre un vecteur ligne et un vecteur colonne.

Proposition A.1.1 Soient $\mathbf{x}$ et $\mathbf{y}$ deux vecteurs de $\mathbb{R}^n$. Alors, on a les propriétés suivantes :

- i) $\|\mathbf{x} + \mathbf{y}\|^2 = \|\mathbf{x}\|^2 + 2\mathbf{x}^T \mathbf{y} + \|\mathbf{y}\|^2$;
- ii) $\|\mathbf{x} - \mathbf{y}\|^2 = \|\mathbf{x}\|^2 - 2\mathbf{x}^T \mathbf{y} + \|\mathbf{y}\|^2$;
- iii) $\|\mathbf{x}\|^2 + \|\mathbf{y}\|^2 = \frac{1}{4} (\|\mathbf{x} + \mathbf{y}\|^2 + \|\mathbf{x} - \mathbf{y}\|^2)$;
- iv) **Inégalité de Cauchy-Schwarz** :

$$\forall \mathbf{x}, \mathbf{y} \in \mathbb{R}^n, \quad \mathbf{x}^T \mathbf{y} \leq \|\mathbf{x}\| \|\mathbf{y}\|.$$

Remarque A.1.2 La dernière inégalité est un résultat très important, qui s'utilise aussi bien en algèbre linéaire qu'en analyse fonctionnelle. Comme on le verra, elle apparaît également dans la preuve des inégalités de Taylor, qui sont fondamentales en optimisation.

A.1.2 Algèbre linéaire matricielle

Sur les espaces matriciels, on définit également un produit matriciel entre matrices de dimensions compatibles. Pour toutes matrices $\mathbf{A} \in \mathbb{R}^{m \times n}$ et $\mathbf{B} \in \mathbb{R}^{n \times p}$, la matrice produit $\mathbf{AB}$ est définie comme la matrice $\mathbf{C} \in \mathbb{R}^{m \times p}$ telle que

$$\forall i = 1, \dots, m, \forall j = 1, \dots, p, \quad C_{ij} = \sum_{k=1}^n A_{ik} B_{kj}.$$

Par analogie, le produit d'une matrice $\mathbf{A} \in \mathbb{R}^{m \times n}$ avec un vecteur $\mathbf{x} \in \mathbb{R}^n$ sera le vecteur $\mathbf{y} \in \mathbb{R}^m$ donné par

$$\forall i = 1, \dots, m, \quad y_i = \sum_{j=1}^n A_{ij} x_j.$$

Remarque A.1.3 On notera que la notation du produit scalaire sur $\mathbb{R}^n$ correspond au produit de deux matrices de taille $1 \times n$ et $n \times 1$, dont le résultat est une matrice 1×1 , c'est-à-dire un scalaire.

Lorsque l'on travaille avec des matrices, on s'intéresse généralement aux sous-espaces définis ci-dessous.

Définition A.1.6 (Sous-espaces matriciels) Soit une matrice $A \in \mathbb{R}^{m \times n}$, on définit les deux sous-espaces suivants :

- Le noyau de A est le sous-espace vectoriel

$$\ker(A) := \{x \in \mathbb{R}^n \mid Ax = 0_m\}$$

- L'image de A est le sous-espace vectoriel

$$\text{Im}(A) := \{y \in \mathbb{R}^m \mid \exists x \in \mathbb{R}^n, y = Ax\}$$

La dimension de ce sous-espace vectoriel s'appelle le **rang** de A . On la note $\text{rang}(A)$ et on a $\text{rang}(A) \leq \min\{m, n\}$.

Théorème A.1.1 (Théorème du rang) Pour toute matrice $A \in \mathbb{R}^{m \times n}$, on a

$$\dim(\ker(A)) + \text{rang}(A) = n.$$

Définition A.1.7 (Normes matricielles) On définit sur $\mathbb{R}^{m \times n}$ la norme d'opérateur $\|\cdot\|$ et la norme de Frobenius $\|\cdot\|_F$ par

$$\forall A \in \mathbb{R}^{m \times n}, \begin{cases} \|A\| &:= \max_{\substack{x \in \mathbb{R}^n \\ x \neq 0_n}} \frac{\|Ax\|}{\|x\|} = \max_{\|x\|=1} \|Ax\| \\ \|A\|_F &:= \sqrt{\sum_{1 \leq i \leq m, 1 \leq j \leq n} A_{ij}^2}. \end{cases}$$

Définition A.1.8 (Matrice symétrique) Une matrice carrée $A \in \mathbb{R}^{n \times n}$ est dite **symétrique** si elle vérifie $A^T = A$.

L'ensemble des matrices symétriques de taille $n \times n$ sera noté S^n .

Définition A.1.9 (Matrice inversible) Une matrice carrée $A \in \mathbb{R}^{n \times n}$ est dite **inversible** s'il existe $B \in \mathbb{R}^{n \times n}$ telle que $BA = AB = I_n$ (où l'on rappelle que I_n désigne la matrice identité de $\mathbb{R}^{n \times n}$).

Si elle existe, une telle matrice B est unique : elle est appelée **l'inverse de A** et on la note A^{-1} .

Définition A.1.10 (Matrice (semi-)définie positive) Une matrice carrée $A \in \mathbb{R}^{n \times n}$ symétrique est dite **semi-définie positive** si

$$\forall x \in \mathbb{R}^n, \quad x^T A x \geq 0,$$

ce que l'on notera $A \succeq 0$.

Une telle matrice est dite **définie positive** lorsque $x^T A x > 0$ pour tout vecteur x non nul, ce que l'on notera $A \succ 0$.

Définition A.1.11 (Matrice orthogonale) Une matrice carrée $P \in \mathbb{R}^{n \times n}$ est dite **orthogonale** si $P^T = P^{-1}$.

Par extension, on dira que $Q \in \mathbb{R}^{m \times n}$ avec $m \leq n$ est orthogonale si $QQ^T = I_m$ (les colonnes de Q sont donc orthonormées dans $\mathbb{R}^m$).

On notera que lorsque $Q \in \mathbb{R}^{n \times n}$ est une matrice orthogonale, alors sa transposée Q^T est également orthogonale (ce résultat n'est pas vrai pour une matrice rectangulaire). On utilisera fréquemment la propriété des matrices orthogonales énoncée ci-dessous.

Lemme A.1.1 Soit une matrice $A \in \mathbb{R}^{m \times n}$ et $U \in \mathbb{R}^{m \times m}$, $V \in \mathbb{R}^{n \times n}$ des matrices orthogonales, respectivement de $\mathbb{R}^{m \times m}$ et $\mathbb{R}^{n \times n}$. On a

$$\|A\| = \|UA\| = \|AV\| \quad \text{et} \quad \|A\|_F = \|UA\|_F = \|AV\|_F,$$

c'est-à-dire que la multiplication par une matrice orthogonale ne modifie pas la norme d'opérateur.

Par corollaire immédiat du lemme précédent, on note qu'une matrice $Q \in \mathbb{R}^{m \times n}$ orthogonale avec $m \leq n$ vérifie nécessairement $\|Q\| = 1$ et $\|Q\|_F = \sqrt{m}$.

Définition A.1.12 (Valeur propre) Soit une matrice $A \in \mathbb{R}^{n \times n}$. On dit que $\lambda \in \mathbb{R}$ est une **valeur propre de A** si

$$\exists v \in \mathbb{R}^n, v \neq 0_n, \quad Av = \lambda v.$$

Le vecteur v est appelé un **vecteur propre associé à la valeur propre λ** . L'ensemble des valeurs propres de A s'appelle le **spectre de A** .

Le sous-espace engendré par les vecteurs propres associés à la même valeur propre d'une matrice s'appelle un sous-espace propre. Sa dimension correspond à l'ordre de multiplicité de la valeur propre relativement à la matrice.

Proposition A.1.2 Pour toute matrice $A \in \mathbb{R}^{n \times n}$, on a les propriétés suivantes :

- La matrice A possède n valeurs propres complexes mais pas nécessairement réelles.
- Si la matrice A est semi-définie positive (respectivement définie positive), alors ses valeurs propres sont réelles positives (respectivement strictement positives).
- Le noyau de A est engendré par les vecteurs propres associés à la valeur propre 0.

Théorème A.1.2 (Théorème spectral) Toute matrice carrée $A \in \mathbb{R}^{n \times n}$ symétrique admet une décomposition dite **spectrale** de la forme :

$$A = P\Lambda P^{-1},$$

où $P \in \mathbb{R}^{n \times n}$ est une matrice orthogonale, dont les colonnes $p_1, \dots, p_n$ forment une base orthonormée de vecteurs propres, et $\Lambda \in \mathbb{R}^{n \times n}$ est une matrice diagonale qui contient les n valeurs propres de A $\lambda_1, \dots, \lambda_n$ sur la diagonale.

Il est à noter que la décomposition spectrale n'est pas unique. En revanche, l'ensemble des valeurs propres est unique, que l'on prenne en compte les ordres de multiplicité ou non.

Géométriquement parlant, on voit ainsi que, pour tout vecteur $x \in \mathbb{R}^n$ décomposé dans la base des p_i que l'on multiplie par A , les composantes de ce vecteur associées aux plus grandes valeurs propres¹ seront augmentées, tandis que celles associées aux valeurs propres de petite magnitude seront réduites (voire annihilées dans le cas d'une valeur propre nulle).

¹On parle ici de plus grandes valeurs propres en valeur absolue, ou magnitude.

Lien avec la décomposition en valeurs singulières Soit une matrice rectangulaire $\mathbf{A} \in \mathbb{R}^{m \times n}$: dans le cas général, les dimensions de la matrice diffèrent, et on ne peut donc pas parler de valeurs propres de la matrice $\mathbf{A}$. On peut en revanche considérer les deux matrices

$$\mathbf{A}^T \mathbf{A} \in \mathbb{R}^{n \times n} \quad \text{et} \quad \mathbf{A} \mathbf{A}^T \in \mathbb{R}^{m \times m}.$$

Ces matrices sont symétriques réelles, et par conséquent diagonalisables. On peut se baser sur cette propriété pour obtenir une décomposition en valeurs singulières de $\mathbf{A}$ (ou *SVD* d'après l'acronyme anglais).