

Optimisation Convexe

Notes de Cours

Master M1 — 2025–2026

Yaë Ulrich Gaba

*« La théorie sans la pratique est inutile,
la pratique sans la théorie est aveugle. »*

— Leonhard Euler

Table des matières

Préface	1
1 Ensembles Convexes	7
1.1 Introduction et motivation	7
1.2 Définitions fondamentales	7
1.3 Exemples d'ensembles convexes	8
1.4 Cônes convexes	8
1.5 Opérations préservant la convexité	9
1.6 Théorèmes de séparation	9
1.7 Hyperplan d'appui	10
1.8 Fonction d'appui et fonction indicatrice	10
1.9 Projection sur un convexe fermé	11
1.10 Polyèdres et représentations	11
1.11 Interprétation géométrique	12
1.12 Implémentation Python	12
1.13 Exercices	13
2 Fonctions Convexes	15
2.1 Introduction	15
2.2 Définitions et premières propriétés	15
2.3 Épigraphe et caractérisation	16
2.4 Ensembles de sous-niveau	16
2.5 Caractérisations différentielles	16
2.5.1 Condition du premier ordre	16
2.5.2 Condition du second ordre	17
2.6 Inégalité de Jensen	17
2.7 Opérations préservant la convexité des fonctions	18
2.8 Convexité et lissité	18
2.9 Exemples importants	18
2.10 Continuité des fonctions convexes	19
2.11 Fonctions convexes conjuguées — aperçu	19
2.12 Implémentation Python	19
2.13 Exercices	20
3 Sous-différentiels et Sous-gradients	21
3.1 Introduction	21
3.2 Définition du sous-gradient et du sous-différentiel	21
3.3 Propriétés du sous-différentiel	22
3.4 Exemples de sous-différentiels	22

3.5	Règles de calcul sous-différentiel	23
3.6	Monotonie du sous-différentiel	23
3.7	Sous-gradient et direction de descente	23
3.8	Méthode du sous-gradient	24
3.9	Interprétation géométrique	25
3.10	Implémentation Python	25
3.11	Exercices	26
4	Dualité de Fenchel–Moreau–Rockafellar	27
4.1	Introduction	27
4.2	Conjuguée de Fenchel	27
4.3	Exemples de conjuguées	28
4.4	Théorème de biconjugaison de Fenchel–Moreau	28
4.5	Lien entre conjuguée et sous-différentiel	29
4.6	Dualité de Fenchel–Rockafellar	29
4.7	Cas particulier : $A = I$	29
4.8	Règles de calcul des conjuguées	30
4.9	Application : formulation duale du LASSO	30
4.10	Implémentation Python	30
4.11	Exercices	31
5	Conditions d’Optimalité — KKT	33
5.1	Introduction	33
5.2	Problème d’optimisation convexe sous contraintes	33
5.3	Conditions KKT	34
5.4	Qualification des contraintes	34
5.5	Exemples d’application des conditions KKT	35
5.6	Interprétation géométrique	36
5.7	Autres qualifications de contraintes	36
5.8	Sensibilité et interprétation économique	36
5.9	Implémentation Python	36
5.10	Exercices	37
6	Dualité de Lagrange	39
6.1	Introduction	39
6.2	Le lagrangien	39
6.3	Fonction duale de Lagrange	40
6.4	Problème dual	40
6.5	Dualité forte	40
6.6	Exemples de calcul dual	41
6.7	Interprétation min-max	41
6.8	Point-selle du lagrangien	41
6.9	Interprétation économique	42
6.10	Implémentation Python	42
6.11	Exercices	43

7	Descente de Gradient et Variantes	45
7.1	Introduction	45
7.2	Descente de gradient à pas fixe	45
7.3	Analyse de convergence — fonctions L -lisses	46
7.4	Convergence pour fonctions fortement convexes	46
7.5	Recherche de pas (line search)	47
7.6	Gradient accéléré de Nesterov	47
7.7	Gradient conjugué	48
7.8	Gradient stochastique (SGD)	48
7.9	Variantes modernes du SGD	48
7.10	Interprétation géométrique	49
7.11	Implémentation Python	49
7.12	Exercices	50
8	Méthodes Proximales	51
8.1	Introduction	51
8.2	Opérateur proximal	51
8.3	Méthode de gradient proximal	52
8.4	Accélération de Nesterov : FISTA	52
8.5	ADMM	53
8.6	Séparation de Douglas-Rachford	54
8.7	Applications	54
8.8	Exercices	54
9	Méthodes de Points Intérieurs	55
9.1	Introduction	55
9.2	Méthode de barrière	55
9.2.1	Formulation	55
9.2.2	Chemin central	56
9.3	Méthodes à pas court et à pas long	57
9.4	Méthode primale-duale	57
9.5	Application à la programmation linéaire	57
9.6	Exercices	58
10	Applications	59
10.1	Introduction	59
10.2	Compressed sensing	59
10.3	LASSO et régularisation ℓ_1	60
10.4	Optimisation de portefeuille	60
10.5	Transport optimal	61
10.6	Applications en apprentissage automatique	61
10.6.1	SVM dual	61
10.6.2	Régression logistique	61
10.7	Implémentation Python	61
10.8	Exercices	62

Préface

Objectifs du cours

L'optimisation convexe occupe une place centrale dans les mathématiques appliquées modernes. Elle fournit un cadre théorique rigoureux et des algorithmes efficaces pour résoudre une vaste classe de problèmes apparaissant en apprentissage automatique, traitement du signal, finance, recherche opérationnelle et bien d'autres domaines.

Ce cours de niveau Master s'adresse aux étudiants possédant de solides bases en analyse réelle, algèbre linéaire et topologie. Il couvre les fondements théoriques de la convexité, les conditions d'optimalité, la dualité, ainsi que les méthodes algorithmiques les plus importantes utilisées en pratique.

L'approche adoptée est à la fois rigoureuse sur le plan mathématique et tournée vers les applications. Chaque concept théorique est accompagné d'interprétations géométriques, d'exemples concrets et d'implémentations en Python.

Organisation du cours

Le cours est organisé en dix chapitres, structurés selon une progression logique :

1. **Ensembles convexes** — Définitions fondamentales, propriétés topologiques, opérations préservant la convexité, cônes, polyèdres, théorèmes de séparation.
2. **Fonctions convexes** — Caractérisations différentielles, inégalités fondamentales (Jensen, gradient), continuité et différentiabilité.
3. **Sous-différentiels et sous-gradients** — Analyse convexe non lisse, calcul sous-différentiel, règles de Moreau–Rockafellar.
4. **Dualité de Fenchel–Moreau–Rockafellar** — Conjugaison convexe, théorème de biconjugaison, dualité de Fenchel.
5. **Conditions d'optimalité — KKT** — Conditions nécessaires et suffisantes, qualifications des contraintes.
6. **Dualité lagrangienne** — Problème dual, saut de dualité, dualité forte, théorème de Slater, interprétation économique.
7. **Descente de gradient et variantes** — Gradient à pas fixe et adaptatif, gradient accéléré de Nesterov, gradient stochastique.
8. **Méthodes proximales et ADMM** — Opérateur proximal, algorithme ISTA/FISTA, ADMM, décomposition.

9. **Méthodes de points intérieurs** — Fonctions barrières, méthode du chemin central, complexité polynomiale.
10. **Applications** — LASSO, SVM, régression logistique, optimisation de portefeuille.

Prérequis

Les prérequis essentiels pour suivre ce cours sont :

- **Analyse réelle** — Suites, séries, continuité, différentiabilité dans $\mathbb{R}^n$, théorèmes de convergence.
- **Algèbre linéaire** — Espaces vectoriels, matrices, valeurs propres, décompositions matricielles (SVD, Cholesky), formes quadratiques.
- **Topologie** — Ouverts, fermés, compacité, connexité dans $\mathbb{R}^n$, notions d'intérieur relatif.
- **Programmation** — Connaissance de base de Python (NumPy, matplotlib). La bibliothèque `cvxpy` sera utilisée pour les implémentations.

Notations

Nous utilisons les notations suivantes tout au long du cours :

Notation	Description
$\mathbb{R}^n$	Espace euclidien de dimension n
$\langle x, y \rangle$	Produit scalaire $\sum_{i=1}^n x_i y_i$
$\ x\ $	Norme euclidienne $\sqrt{\langle x, x \rangle}$
$\ x\ _1$	Norme ℓ_1 : $\sum_i x_i $
$\ x\ _\infty$	Norme ℓ_∞ : $\max_i x_i $
$B(x, r)$	Boule ouverte de centre x et rayon r
$\overline{C}$	Adhérence de l'ensemble C
$\text{int}(C)$	Intérieur de l'ensemble C
$\text{ri}(C)$	Intérieur relatif de C
$\text{aff}(C)$	Enveloppe affine de C
$\text{conv}(S)$	Enveloppe convexe de S
$\text{cone}(S)$	Cône engendré par S
$\text{dom}(f)$	Domaine effectif de f
$\text{epi}(f)$	Épigraphe de f
f^*	Conjuguée de Fenchel de f
$\partial f(x)$	Sous-différentiel de f en x
$\nabla f(x)$	Gradient de f en x
$\nabla^2 f(x)$	Hessienne de f en x
$\arg \min_x f(x)$	Ensemble des minimiseurs de f
$A \succeq 0$	A semi-définie positive
$A \succ 0$	A définie positive
$\mathbb{S}_+^n$	Cône des matrices SDP de taille n
ι_C	Fonction indicatrice convexe de C
σ_C	Fonction d'appui de C
prox_f	Opérateur proximal de f

Conventions

- Sauf mention contraire, tous les espaces considérés sont de dimension finie sur $\mathbb{R}$.
- Les fonctions convexes considérées sont à valeurs dans $\mathbb{R} \cup \{+\infty\}$ (fonctions convexes propres).
- Le symbole $\square$ ou $\square$ marque la fin d'une démonstration.
- Les exercices sont classés par difficulté croissante : $(\star)$ basique, $(\star\star)$ intermédiaire, $(\star\star\star)$ avancé.

- Les blocs `algorithme` présentent le pseudo-code des méthodes étudiées.
- Les blocs `codeblock` contiennent du code Python exécutable.

Fil conducteur

Le fil conducteur de ce cours est le problème d'optimisation convexe général :

$$\min_{x \in \mathbb{R}^n} f(x) \quad \text{sous} \quad g_i(x) \leq 0, \quad i = 1, \dots, m, \quad h_j(x) = 0, \quad j = 1, \dots, p,$$

où f et les g_i sont convexes et les h_j sont affines. Nous étudions successivement :

- la *géométrie* de l'ensemble réalisable (chapitres 1–2),
- l'*analyse* des fonctions objectif (chapitres 2–4),
- les *conditions d'optimalité* et la *dualité* (chapitres 5–6),
- les *algorithmes* de résolution (chapitres 7–9),
- les *applications* concrètes (chapitre 10).

Approche pédagogique

Chaque chapitre suit une structure similaire :

1. **Motivation** — Pourquoi ce sujet est important.
2. **Théorie** — Définitions, théorèmes avec démonstrations complètes.
3. **Interprétation géométrique** — Illustrations TikZ et intuitions.
4. **Exemples** — Exemples détaillés et contre-exemples.
5. **Implémentation** — Code Python/cvxpy.
6. **Exercices** — Problèmes de difficulté variée.

Références principales

1. S. Boyd et L. Vandenberghe, *Convex Optimization*, Cambridge University Press, 2004.
2. R.T. Rockafellar, *Convex Analysis*, Princeton University Press, 1970.
3. J.-B. Hiriart-Urruty et C. Lemaréchal, *Fundamentals of Convex Analysis*, Springer, 2001.
4. Y. Nesterov, *Introductory Lectures on Convex Optimization*, Springer, 2004.
5. A. Beck, *First-Order Methods in Optimization*, SIAM, 2017.

6. N. Parikh et S. Boyd, *Proximal Algorithms*, Foundations and Trends in Optimization, 2014.
7. H.H. Bauschke et P.L. Combettes, *Convex Analysis and Monotone Operator Theory in Hilbert Spaces*, Springer, 2017.
8. D.P. Bertsekas, *Convex Optimization Algorithms*, Athena Scientific, 2015.

L'auteur
Mars 2026

Chapitre 1

Ensembles Convexes

Imaginez que vous tendez un élastique entre deux punaises plantées dans un morceau de carton. L'élastique suit une ligne droite. Maintenant, imaginez que la région accessible du carton est contrainte : si, quelle que soit la position des deux punaises *dans* la région, l'élastique reste entièrement dans la région, alors cette région est *convexe*. C'est une idée aussi vieille que la géométrie elle-même — Euclide connaissait les polygones convexes — mais ses conséquences en optimisation sont profondes et n'ont été pleinement exploitées qu'au XX^e siècle, par des pionniers comme Minkowski, Fenchel, Rockafellar et Boyd.

1.1 Introduction et motivation

Les ensembles convexes constituent la brique fondamentale de l'optimisation convexe. Un problème d'optimisation est dit convexe lorsque l'on minimise une fonction convexe sur un ensemble convexe. La géométrie de ces ensembles détermine la structure des solutions et guide la conception des algorithmes.

Intuition

Un ensemble est convexe si, pour tout couple de points de l'ensemble, le segment les reliant est entièrement contenu dans l'ensemble. Cette propriété simple a des conséquences profondes : tout minimum local est global, et les techniques de séparation par hyperplans deviennent possibles.

1.2 Définitions fondamentales

Définition 1.1 (Ensemble convexe). Un ensemble $C \subseteq \mathbb{R}^n$ est *convexe* si, pour tout $x, y \in C$ et tout $\theta \in [0, 1]$:

$$\theta x + (1 - \theta)y \in C.$$

Définition 1.2 (Combinaison convexe). Un point x est une *combinaison convexe* des points $x_1, \dots, x_k$ s'il existe $\theta_1, \dots, \theta_k \geq 0$ avec $\sum_{i=1}^k \theta_i = 1$ tels que :

$$x = \sum_{i=1}^k \theta_i x_i.$$

Définition 1.3 (Enveloppe convexe). L'*enveloppe convexe* d'un ensemble $S \subseteq \mathbb{R}^n$, notée $\text{conv}(S)$, est l'ensemble de toutes les combinaisons convexes de points de S :

$$\text{conv}(S) = \left\{ \sum_{i=1}^k \theta_i x_i \mid x_i \in S, \theta_i \geq 0, \sum_{i=1}^k \theta_i = 1, k \in \mathbb{N}^* \right\}.$$

Théorème 1.4 (Carathéodory). Soit $S \subseteq \mathbb{R}^n$. Tout point de $\text{conv}(S)$ peut s'écrire comme combinaison convexe d'au plus $n + 1$ points de S .

Démonstration. Soit $x \in \text{conv}(S)$. Alors $x = \sum_{i=1}^k \theta_i x_i$ avec $\theta_i > 0$, $\sum \theta_i = 1$, $x_i \in S$. Si $k > n + 1$, les vecteurs $x_2 - x_1, \dots, x_k - x_1$ sont au nombre de $k - 1 > n$, donc linéairement dépendants. Il existe $\mu_2, \dots, \mu_k$ non tous nuls tels que $\sum_{i=2}^k \mu_i (x_i - x_1) = 0$. Posons $\mu_1 = -\sum_{i=2}^k \mu_i$. Alors $\sum_{i=1}^k \mu_i x_i = 0$ et $\sum_{i=1}^k \mu_i = 0$.

Pour $t \in \mathbb{R}$, $x = \sum_{i=1}^k (\theta_i - t\mu_i) x_i$. Choisissons $t^* = \min_{i: \mu_i > 0} \theta_i / \mu_i$. Alors $\theta_i - t^* \mu_i \geq 0$ pour tout i , et au moins un coefficient s'annule, réduisant le nombre de termes. On itère jusqu'à $k \leq n + 1$. $\square$

1.3 Exemples d'ensembles convexes

Exemple 1.5 (Ensembles convexes classiques). Les ensembles suivants sont convexes :

1. **Hyperplan** : $\{x \in \mathbb{R}^n \mid \langle a, x \rangle = b\}$ avec $a \neq 0$.
2. **Demi-espace** : $\{x \in \mathbb{R}^n \mid \langle a, x \rangle \leq b\}$.
3. **Boule euclidienne** : $B(x_0, r) = \{x \mid \|x - x_0\| \leq r\}$.
4. **Ellipsoïde** : $\{x \mid (x - x_0)^T P^{-1} (x - x_0) \leq 1\}$, $P \succ 0$.
5. **Polyèdre** : $\{x \mid Ax \leq b, Cx = d\}$.
6. **Simplexe** : $\Delta_n = \{x \in \mathbb{R}_+^n \mid \sum x_i = 1\}$.
7. **Cône SDP** : $\mathbb{S}_+^n = \{X \in \mathbb{S}^n \mid X \succeq 0\}$.

Vérification pour la boule. Soient $x, y \in B(x_0, r)$ et $\theta \in [0, 1]$. Par l'inégalité triangulaire :

$$\|\theta x + (1 - \theta)y - x_0\| = \|\theta(x - x_0) + (1 - \theta)(y - x_0)\| \leq \theta \|x - x_0\| + (1 - \theta) \|y - x_0\| \leq \theta r + (1 - \theta)r = r.$$

$\square$

1.4 Cônes convexes

Définition 1.6 (Cône). Un ensemble $K \subseteq \mathbb{R}^n$ est un *cône* si pour tout $x \in K$ et tout $\lambda \geq 0$, on a $\lambda x \in K$.

Définition 1.7 (Cône convexe). Un cône K est *convexe* si K est un ensemble convexe, ce qui équivaut à : pour tout $x, y \in K$ et $\lambda, \mu \geq 0$:

$$\lambda x + \mu y \in K.$$

Définition 1.8 (Cône dual). Le *cône dual* de K est :

$$K^* = \{y \in \mathbb{R}^n \mid \langle y, x \rangle \geq 0, \forall x \in K\}.$$

Proposition 1.9 (Propriétés du cône dual). 1. K^* est un cône convexe fermé.

2. Si $K_1 \subseteq K_2$, alors $K_2^* \subseteq K_1^*$.

3. Si K est un cône convexe fermé, alors $K^{**} = K$.

Exemple 1.10 (Cônes duaux classiques). • Le cône dual de $\mathbb{R}_+^n$ est $\mathbb{R}_+^n$ (auto-dual).

- Le cône dual du cône de Lorentz $\mathcal{L}^n = \{(x, t) \in \mathbb{R}^{n-1} \times \mathbb{R} \mid \|x\| \leq t\}$ est $\mathcal{L}^n$ lui-même.
- Le cône dual de $\mathbb{S}_+^n$ est $\mathbb{S}_+^n$.

1.5 Opérations préservant la convexité

Théorème 1.11 (Opérations convexes). *Les opérations suivantes préservent la convexité :*

1. **Intersection** : si C_α est convexe pour tout $\alpha \in I$, alors $\bigcap_{\alpha \in I} C_\alpha$ est convexe.
2. **Image affine** : si C est convexe et $f(x) = Ax + b$, alors $f(C) = \{Ax + b \mid x \in C\}$ est convexe.
3. **Image inverse affine** : si C est convexe, alors $f^{-1}(C) = \{x \mid Ax + b \in C\}$ est convexe.
4. **Somme de Minkowski** : si C_1, C_2 sont convexes, alors $C_1 + C_2 = \{x + y \mid x \in C_1, y \in C_2\}$ est convexe.
5. **Projection** : si $C \subseteq \mathbb{R}^n \times \mathbb{R}^m$ est convexe, alors $\{x \in \mathbb{R}^n \mid \exists y, (x, y) \in C\}$ est convexe.
6. **Fonction perspective** : si C est convexe et $P(x, t) = x/t$ pour $t > 0$, alors $P(C)$ est convexe.

Preuve de (1). Soient $x, y \in \bigcap_{\alpha} C_\alpha$ et $\theta \in [0, 1]$. Pour tout α , $x, y \in C_\alpha$, donc $\theta x + (1 - \theta)y \in C_\alpha$ par convexité de C_α . Ainsi $\theta x + (1 - \theta)y \in \bigcap_{\alpha} C_\alpha$. $\square$

Remarque 1.12. L'union de deux ensembles convexes n'est généralement pas convexe. Par exemple, $\{0\} \cup \{1\}$ n'est pas convexe dans $\mathbb{R}$. L'intersection arbitraire préserve la convexité, mais pas l'union.

1.6 Théorèmes de séparation

Théorème 1.13 (Séparation par hyperplan). Soient $C, D \subseteq \mathbb{R}^n$ deux ensembles convexes non vides disjoints. Il existe $a \in \mathbb{R}^n$, $a \neq 0$, et $b \in \mathbb{R}$ tels que :

$$\langle a, x \rangle \leq b \quad \forall x \in C \quad \text{et} \quad \langle a, x \rangle \geq b \quad \forall x \in D.$$

Théorème 1.14 (Séparation stricte). *Soient C un ensemble convexe fermé non vide et $x_0 \notin C$. Il existe $a \in \mathbb{R}^n$, $a \neq 0$, et $b \in \mathbb{R}$ tels que :*

$$\langle a, x_0 \rangle > b > \langle a, x \rangle \quad \forall x \in C.$$

Démonstration. Soit $\bar{x} = \text{proj}_C(x_0)$ la projection de x_0 sur C (qui existe et est unique car C est convexe fermé). Posons $a = x_0 - \bar{x}$ et $b = \langle a, \bar{x} \rangle + \frac{1}{2} \|a\|^2$.

Pour tout $x \in C$, la caractérisation de la projection donne :

$$\langle x_0 - \bar{x}, x - \bar{x} \rangle \leq 0,$$

c'est-à-dire $\langle a, x \rangle \leq \langle a, \bar{x} \rangle$. De plus :

$$\langle a, x_0 \rangle = \langle a, \bar{x} \rangle + \|a\|^2 > \langle a, \bar{x} \rangle + \frac{1}{2} \|a\|^2 = b > \langle a, \bar{x} \rangle \geq \langle a, x \rangle.$$

□

Intuition

Géométriquement, le théorème de séparation affirme qu'on peut toujours placer un hyperplan entre deux ensembles convexes disjoints. C'est l'analogue en dimension supérieure du fait qu'on peut séparer deux intervalles disjoints sur la droite réelle par un point.

1.7 Hyperplan d'appui

Définition 1.15 (Hyperplan d'appui). Soit $C \subseteq \mathbb{R}^n$ un ensemble convexe. Un *hyperplan d'appui* de C en un point frontière $x_0 \in \partial C$ est un hyperplan $H = \{x \mid \langle a, x \rangle = b\}$ tel que :

$$\langle a, x_0 \rangle = b \quad \text{et} \quad \langle a, x \rangle \leq b \quad \forall x \in C.$$

Théorème 1.16 (Existence d'hyperplans d'appui). *Soit C un ensemble convexe non vide avec intérieur non vide, et $x_0 \in \partial C$. Alors il existe un hyperplan d'appui de C en x_0 .*

1.8 Fonction d'appui et fonction indicatrice

Définition 1.17 (Fonction d'appui). La *fonction d'appui* d'un ensemble C est :

$$\sigma_C(y) = \sup_{x \in C} \langle y, x \rangle.$$

Définition 1.18 (Fonction indicatrice convexe). La *fonction indicatrice* (au sens de l'analyse convexe) de C est :

$$\iota_C(x) = \begin{cases} 0 & \text{si } x \in C, \\ +\infty & \text{sinon.} \end{cases}$$

Proposition 1.19. Si C est convexe fermé non vide, alors ι_C est convexe, propre et semi-continue inférieurement. De plus, $\sigma_C = \iota_C^*$ (la conjuguée de Fenchel de ι_C).

1.9 Projection sur un convexe fermé

Théorème 1.20 (Projection convexe). *Soit $C \subseteq \mathbb{R}^n$ un ensemble convexe fermé non vide. Pour tout $y \in \mathbb{R}^n$, il existe un unique $\bar{x} \in C$ tel que :*

$$\bar{x} = \arg \min_{x \in C} \|x - y\|^2.$$

Ce point, noté $\text{proj}_C(y)$, est caractérisé par :

$$\bar{x} \in C \quad \text{et} \quad \langle y - \bar{x}, x - \bar{x} \rangle \leq 0, \quad \forall x \in C.$$

Démonstration. Existence. Soit $d = \inf_{x \in C} \|x - y\|$. Prenons $(x_k) \subset C$ avec $\|x_k - y\| \rightarrow d$. Par l'identité du parallélogramme :

$$\|x_k - x_l\|^2 = 2\|x_k - y\|^2 + 2\|x_l - y\|^2 - 4\left\|\frac{x_k + x_l}{2} - y\right\|^2.$$

Comme $\frac{x_k + x_l}{2} \in C$ (convexité), $\left\|\frac{x_k + x_l}{2} - y\right\| \geq d$, donc :

$$\|x_k - x_l\|^2 \leq 2\|x_k - y\|^2 + 2\|x_l - y\|^2 - 4d^2 \rightarrow 0.$$

(x_k) est de Cauchy, converge vers $\bar{x} \in C$ (car C est fermé).

Unicité. Si $\bar{x}_1, \bar{x}_2$ sont deux projections, le même argument montre $\|\bar{x}_1 - \bar{x}_2\| = 0$.

Caractérisation. $\bar{x}$ minimise $\varphi(x) = \|x - y\|^2$ sur C si et seulement si pour tout $x \in C$ et $t \in (0, 1]$: $\varphi(\bar{x} + t(x - \bar{x})) \geq \varphi(\bar{x})$, ce qui donne $\langle y - \bar{x}, x - \bar{x} \rangle \leq 0$. $\square$

Proposition 1.21 (Non-expansivité de la projection). La projection proj_C est non expansive :

$$\|\text{proj}_C(x) - \text{proj}_C(y)\| \leq \|x - y\|, \quad \forall x, y \in \mathbb{R}^n.$$

En particulier, proj_C est continue (lipschitzienne de constante 1).

1.10 Polyèdres et représentations

Définition 1.22 (Polyèdre). Un *polyèdre* est l'intersection d'un nombre fini de demi-espaces :

$$P = \{x \in \mathbb{R}^n \mid Ax \leq b\}, \quad A \in \mathbb{R}^{m \times n}, \quad b \in \mathbb{R}^m.$$

Un polyèdre borné est appelé *polytope*.

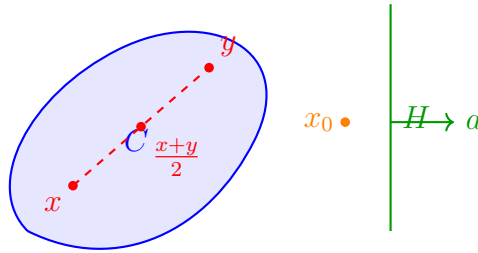
Théorème 1.23 (Minkowski–Weyl). *Un ensemble $P \subseteq \mathbb{R}^n$ est un polytope si et seulement s'il est l'enveloppe convexe d'un nombre fini de points :*

$$P = \text{conv}(\{v_1, \dots, v_k\}).$$

Plus généralement, tout polyèdre admet une représentation de la forme :

$$P = \text{conv}(\{v_1, \dots, v_k\}) + \text{cone}(\{d_1, \dots, d_l\}).$$

1.11 Interprétation géométrique



1.12 Implémentation Python

Projection sur le simplexe et vérification de convexité

```
import numpy as np
import matplotlib.pyplot as plt

def project_simplex(v):
    """Projection sur le simplexe standard."""
    n = len(v)
    u = np.sort(v)[::-1]
    cumsum = np.cumsum(u)
    rho = np.max(np.where(u + (1 - cumsum) / (np.arange(n) + 1) > 0))
    theta = (cumsum[rho] - 1) / (rho + 1)
    return np.maximum(v - theta, 0)

def is_convex_hull_member(point, vertices):
    """Verifie si point est dans conv(vertices) via LP."""
    import scipy.optimize as opt
    k = len(vertices)
    V = np.array(vertices).T # n x k
    n = V.shape[0]
    # min 0 s.t. V @ theta = point, sum(theta) = 1, theta >= 0
    c = np.zeros(k)
    A_eq = np.vstack([V, np.ones((1, k))])
    b_eq = np.append(point, 1.0)
    bounds = [(0, None)] * k
    res = opt.linprog(c, A_eq=A_eq, b_eq=b_eq, bounds=bounds)
    return res.success

# Demonstration: projection sur le simplexe
np.random.seed(42)
v = np.random.randn(3)
p = project_simplex(v)
print(f"Point original: {v}")
print(f"Projection: {p}")
print(f"Somme = {p.sum():.6f}, min = {p.min():.6f}")
```

1.13 Exercices

Exercice 1.1 (*). Montrer que l'intersection de deux demi-espaces est convexe.

Exercice 1.2 (*). Montrer que l'image d'un ensemble convexe par une application linéaire est convexe.

Exercice 1.3 (**). Soit $C = \{x \in \mathbb{R}^n \mid \|x\|_1 \leq 1\}$. Montrer que C est un polytope et déterminer ses sommets.

Exercice 1.4 (**). Soit C un ensemble convexe fermé et $x_0 \notin C$. Démontrer qu'il existe un unique point $\bar{x} \in C$ le plus proche de x_0 . Montrer que $\langle x_0 - \bar{x}, x - \bar{x} \rangle \leq 0$ pour tout $x \in C$.

Exercice 1.5 (**). Calculer le cône dual du cône de Lorentz $\mathcal{L}^n = \{(x, t) \in \mathbb{R}^{n-1} \times \mathbb{R} \mid \|x\| \leq t\}$.

Exercice 1.6 (***). Démontrer le théorème de Carathéodory en dimension 2 géométriquement, puis généraliser la preuve en dimension n .

Exercice 1.7 (***). Soit $f : \mathbb{R}^n \rightarrow \mathbb{R}$ continue. Montrer que l'ensemble de sous-niveau $\{x \mid f(x) \leq \alpha\}$ est fermé. Donner un exemple montrant qu'il n'est pas nécessairement convexe même si f est continue.

Chapitre 2

Fonctions Convexes

Si les ensembles convexes sont le décor, les fonctions convexes en sont les protagonistes. Une fonction est convexe si elle « courbe vers le haut » : la sécante entre deux points de son graphe est toujours au-dessus du graphe. Cette propriété, d'apparence modeste, a des conséquences spectaculaires : tout minimum local est global, les conditions d'optimalité du premier ordre sont suffisantes, et les algorithmes d'optimisation convergent avec des garanties théoriques solides. Rockafellar a dit un jour que les fonctions convexes sont aux inégalités ce que les fonctions affines sont aux égalités — une analogie profonde que ce chapitre va explorer.

2.1 Introduction

Les fonctions convexes sont au cœur de l'optimisation convexe. Leur propriété fondamentale — tout minimum local est global — rend les problèmes d'optimisation convexe « bien posés » et théoriquement tractables. Ce chapitre présente les définitions, caractérisations et propriétés essentielles.

2.2 Définitions et premières propriétés

Définition 2.1 (Fonction convexe). Une fonction $f : \mathbb{R}^n \rightarrow \mathbb{R} \cup \{+\infty\}$ est *convexe* si son domaine effectif $\text{dom}(f) = \{x \mid f(x) < +\infty\}$ est convexe et si pour tout $x, y \in \text{dom}(f)$ et tout $\theta \in [0, 1]$:

$$f(\theta x + (1 - \theta)y) \leq \theta f(x) + (1 - \theta)f(y).$$

Définition 2.2 (Convexité stricte). f est *strictement convexe* si l'inégalité ci-dessus est stricte pour $x \neq y$ et $\theta \in (0, 1)$.

Définition 2.3 (Forte convexité). f est *fortement convexe* de paramètre $m > 0$ (ou *m-fortement convexe*) si $f(x) - \frac{m}{2} \|x\|^2$ est convexe, c'est-à-dire pour tout x, y et $\theta \in [0, 1]$:

$$f(\theta x + (1 - \theta)y) \leq \theta f(x) + (1 - \theta)f(y) - \frac{m}{2} \theta(1 - \theta) \|x - y\|^2.$$

Hiérarchie de convexité

fortement convexe $\Rightarrow$ strictement convexe $\Rightarrow$ convexe.

Les implications inverses sont fausses en général.

2.3 Épigraphe et caractérisation

Définition 2.4 (Épigraphe). L'épigraphe d'une fonction $f : \mathbb{R}^n \rightarrow \mathbb{R} \cup \{+\infty\}$ est :

$$\text{epi}(f) = \{(x, t) \in \mathbb{R}^n \times \mathbb{R} \mid f(x) \leq t\}.$$

Théorème 2.5 (Caractérisation épigraphique). f est convexe si et seulement si $\text{epi}(f)$ est un ensemble convexe de $\mathbb{R}^{n+1}$.

Démonstration. ($\Rightarrow$) Soient $(x_1, t_1), (x_2, t_2) \in \text{epi}(f)$ et $\theta \in [0, 1]$. Alors :

$$f(\theta x_1 + (1 - \theta)x_2) \leq \theta f(x_1) + (1 - \theta)f(x_2) \leq \theta t_1 + (1 - \theta)t_2,$$

donc $(\theta x_1 + (1 - \theta)x_2, \theta t_1 + (1 - \theta)t_2) \in \text{epi}(f)$.

($\Leftarrow$) Soient $x_1, x_2 \in \text{dom}(f)$. Alors $(x_1, f(x_1)), (x_2, f(x_2)) \in \text{epi}(f)$. Par convexité de l'épigraphe :

$$(\theta x_1 + (1 - \theta)x_2, \theta f(x_1) + (1 - \theta)f(x_2)) \in \text{epi}(f),$$

ce qui signifie $f(\theta x_1 + (1 - \theta)x_2) \leq \theta f(x_1) + (1 - \theta)f(x_2)$. □

2.4 Ensembles de sous-niveau

Définition 2.6 (Ensemble de sous-niveau). L'ensemble de sous-niveau α de f est :

$$S_\alpha = \{x \in \text{dom}(f) \mid f(x) \leq \alpha\}.$$

Proposition 2.7. Si f est convexe, alors tous ses ensembles de sous-niveau sont convexes.

Réciproque fausse

La réciproque est fausse : $f(x) = e^x$ a des sous-niveaux convexes (intervalles) mais $g(x) = \sqrt{|x|}$ a aussi des sous-niveaux convexes sans être convexe. Une fonction dont les sous-niveaux sont convexes est dite *quasi-convexe*.

2.5 Caractérisations différentielles

2.5.1 Condition du premier ordre

Théorème 2.8 (Condition du premier ordre). Soit $f : \mathbb{R}^n \rightarrow \mathbb{R}$ différentiable sur un ouvert convexe Ω . Alors f est convexe sur Ω si et seulement si :

$$f(y) \geq f(x) + \langle \nabla f(x), y - x \rangle, \quad \forall x, y \in \Omega.$$

Démonstration. ($\Rightarrow$) Par convexité, pour $t \in (0, 1]$:

$$f(x + t(y - x)) \leq f(x) + t(f(y) - f(x)).$$

Donc $\frac{f(x+t(y-x))-f(x)}{t} \leq f(y) - f(x)$. En passant à la limite $t \rightarrow 0^+$: $\langle \nabla f(x), y - x \rangle \leq f(y) - f(x)$.

($\Leftarrow$) Soient $x, y \in \Omega$ et $z = \theta x + (1 - \theta)y$. Alors :

$$\begin{aligned} f(x) &\geq f(z) + \langle \nabla f(z), x - z \rangle, \\ f(y) &\geq f(z) + \langle \nabla f(z), y - z \rangle. \end{aligned}$$

En multipliant par θ et $1 - \theta$ respectivement et en sommant :

$$\theta f(x) + (1 - \theta)f(y) \geq f(z) + \langle \nabla f(z), \theta x + (1 - \theta)y - z \rangle = f(z).$$

□

Intuition

La condition du premier ordre signifie que le graphe de f est toujours au-dessus de ses tangentes. Géométriquement, la surface $z = f(x)$ est « bombée vers le haut ».

2.5.2 Condition du second ordre

Théorème 2.9 (Condition du second ordre). *Soit $f : \mathbb{R}^n \rightarrow \mathbb{R}$ deux fois différentiable sur un ouvert convexe Ω . Alors :*

1. f est convexe sur $\Omega \Leftrightarrow \nabla^2 f(x) \succeq 0$ pour tout $x \in \Omega$.
2. f est fortement convexe de paramètre $m \Leftrightarrow \nabla^2 f(x) \succeq mI$ pour tout $x \in \Omega$.
3. Si $\nabla^2 f(x) \succ 0$ pour tout x , alors f est strictement convexe (la réciproque est fausse).

Preuve de (1), $\Rightarrow$. Par la condition du premier ordre, pour tout $h \in \mathbb{R}^n$ et $t > 0$ assez petit :

$$f(x + th) \geq f(x) + t \langle \nabla f(x), h \rangle.$$

Par développement de Taylor :

$$f(x + th) = f(x) + t \langle \nabla f(x), h \rangle + \frac{t^2}{2} h^T \nabla^2 f(x) h + o(t^2).$$

Donc $h^T \nabla^2 f(x) h + o(t^2)/t^2 \geq 0$. En faisant $t \rightarrow 0$: $h^T \nabla^2 f(x) h \geq 0$. □

2.6 Inégalité de Jensen

Théorème 2.10 (Inégalité de Jensen). *Soit f convexe et $x_1, \dots, x_k \in \text{dom}(f)$ avec $\theta_1, \dots, \theta_k \geq 0$, $\sum \theta_i = 1$. Alors :*

$$f\left(\sum_{i=1}^k \theta_i x_i\right) \leq \sum_{i=1}^k \theta_i f(x_i).$$

Corollaire 2.11 (Jensen probabiliste). *Soit X une variable aléatoire à valeurs dans $\text{dom}(f)$ avec $\mathbb{E}[|X|] < \infty$. Si f est convexe :*

$$f(\mathbb{E}[X]) \leq \mathbb{E}[f(X)].$$

2.7 Opérations préservant la convexité des fonctions

Théorème 2.12 (Opérations sur les fonctions convexes). 1. **Combinaison positive** :

si $f_1, \dots, f_k$ sont convexes et $\alpha_1, \dots, \alpha_k \geq 0$, alors $\sum \alpha_i f_i$ est convexe.

2. **Maximum ponctuel** : si $f_1, \dots, f_k$ sont convexes, alors $g(x) = \max_i f_i(x)$ est convexe.

3. **Supremum** : si $(f_\alpha)_{\alpha \in I}$ est une famille de fonctions convexes, alors $g(x) = \sup_\alpha f_\alpha(x)$ est convexe.

4. **Composition affine** : si f est convexe, alors $g(x) = f(Ax + b)$ est convexe.

5. **Infimale convolution** : si f, g sont convexes, alors $(f \square g)(x) = \inf_y \{f(y) + g(x - y)\}$ est convexe.

2.8 Convexité et lissité

Définition 2.13 (L -lissité). Une fonction convexe différentiable f est L -lisse si son gradient est L -lipschitzien :

$$\|\nabla f(x) - \nabla f(y)\| \leq L \|x - y\|, \quad \forall x, y.$$

Théorème 2.14 (Caractérisations de la L -lissité). Pour f convexe et différentiable, les conditions suivantes sont équivalentes :

1. ∇f est L -lipschitzien.
2. $f(y) \leq f(x) + \langle \nabla f(x), y - x \rangle + \frac{L}{2} \|y - x\|^2$ pour tout x, y .
3. $f(y) \geq f(x) + \langle \nabla f(x), y - x \rangle + \frac{1}{2L} \|\nabla f(y) - \nabla f(x)\|^2$ pour tout x, y .
4. $\langle \nabla f(x) - \nabla f(y), x - y \rangle \geq \frac{1}{L} \|\nabla f(x) - \nabla f(y)\|^2$ pour tout x, y .

Si f est deux fois différentiable, ces conditions équivalent à $\nabla^2 f(x) \preceq LI$.

Encadrement fondamental

Pour f convexe, m -fortement convexe et L -lisse :

$$f(x) + \langle \nabla f(x), y - x \rangle + \frac{m}{2} \|y - x\|^2 \leq f(y) \leq f(x) + \langle \nabla f(x), y - x \rangle + \frac{L}{2} \|y - x\|^2.$$

Le nombre de conditionnement $\kappa = L/m$ mesure la difficulté du problème.

2.9 Exemples importants

Exemple 2.15 (Fonctions convexes classiques). 1. **Fonctions affines** : $f(x) = \langle a, x \rangle + b$ (convexe et concave).

2. **Normes** : $\|\cdot\|_p$ pour $p \geq 1$ (convexe).

3. **Quadratique** : $f(x) = \frac{1}{2}x^T Qx + q^T x + c$ avec $Q \succeq 0$. Fortement convexe ssi $Q \succ 0$, avec $m = \lambda_{\min}(Q)$, $L = \lambda_{\max}(Q)$.
4. **Exponentielle** : $f(x) = e^{ax}$ (convexe sur $\mathbb{R}$).
5. **Log-sum-exp** : $f(x) = \log(\sum_i e^{x_i})$ (convexe sur $\mathbb{R}^n$).
6. **Entropie négative** : $f(x) = \sum_i x_i \log x_i$ sur $\mathbb{R}_{++}^n$.
7. **Perte logistique** : $f(x) = \log(1 + e^{-x})$ (convexe, lisse).

2.10 Continuité des fonctions convexes

Théorème 2.16 (Continuité). *Toute fonction convexe $f : \mathbb{R}^n \rightarrow \mathbb{R}$ (finie partout) est continue. Plus précisément, toute fonction convexe propre est continue sur l'intérieur de son domaine effectif.*

Théorème 2.17 (Différentiabilité presque partout). *Toute fonction convexe $f : \mathbb{R}^n \rightarrow \mathbb{R}$ est différentiable presque partout (au sens de Lebesgue). En dimension 1, f admet des dérivées à gauche et à droite en tout point.*

2.11 Fonctions convexes conjuguées — aperçu

Définition 2.18 (Conjuguée de Fenchel). La *conjuguée de Fenchel* (ou transformée de Legendre–Fenchel) de $f : \mathbb{R}^n \rightarrow \mathbb{R} \cup \{+\infty\}$ est :

$$f^*(y) = \sup_{x \in \mathbb{R}^n} \{\langle y, x \rangle - f(x)\}.$$

Remarque 2.19. f^* est toujours convexe et semi-continue inférieurement (comme supremum de fonctions affines en y), même si f ne l'est pas. La conjugaison convexe sera étudiée en détail au chapitre 4.

2.12 Implémentation Python

Vérification de convexité et visualisation

```
import numpy as np
import matplotlib.pyplot as plt

def check_convexity_numerical(f, x_range, n_tests=1000):
    """Vérification numérique de convexité en 1D."""
    for _ in range(n_tests):
        x = np.random.uniform(*x_range)
        y = np.random.uniform(*x_range)
        t = np.random.uniform(0, 1)
        lhs = f(t * x + (1 - t) * y)
        rhs = t * f(x) + (1 - t) * f(y)
        if lhs > rhs + 1e-10:
            return False
```

```

    return True

# Exemples
functions = {
    'x^2': lambda x: x**2,
    'exp(x)': lambda x: np.exp(x),
    '|x|': lambda x: np.abs(x),
    'log(1+exp(x))': lambda x: np.log(1 + np.exp(x)),
}

fig, axes = plt.subplots(2, 2, figsize=(10, 8))
for ax, (name, f) in zip(axes.flat, functions.items()):
    x = np.linspace(-3, 3, 200)
    ax.plot(x, f(x), 'b-', linewidth=2)
    # Tangente en x=1
    x0 = 1.0
    h = 1e-5
    grad = (f(x0 + h) - f(x0 - h)) / (2 * h)
    tangent = f(x0) + grad * (x - x0)
    ax.plot(x, tangent, 'r--', alpha=0.7)
    is_cvx = check_convexity_numerical(f, (-3, 3))
    ax.set_title(f'{name} (convexe: {is_cvx})')
    ax.grid(True, alpha=0.3)
plt.tight_layout()
plt.savefig('convex_functions.pdf')

```

2.13 Exercices

Exercice 2.1 (*). Montrer que $f(x) = \max(x_1, \dots, x_n)$ est convexe sur $\mathbb{R}^n$.

Exercice 2.2 (*). Montrer que si f est convexe et g est affine, alors $f \circ g$ est convexe.

Exercice 2.3 (**). Montrer qu'une fonction $f : \mathbb{R}^n \rightarrow \mathbb{R}$ différentiable est m -fortement convexe si et seulement si $\langle \nabla f(x) - \nabla f(y), x - y \rangle \geq m \|x - y\|^2$ pour tout x, y .

Exercice 2.4 (**). Montrer que $f(x) = \log(\sum_{i=1}^n e^{x_i})$ est convexe et calculer sa constante de lissité L par rapport à la norme ℓ_2 .

Exercice 2.5 (**). Soit f convexe, propre et semi-continue inférieurement. Montrer que f est minorée par une fonction affine.

Exercice 2.6 (***). Montrer l'équivalence des quatre caractérisations de la L -lissité du théorème 2.14.

Exercice 2.7 (***). Soit f convexe sur $\mathbb{R}^n$ et x^* un minimiseur. Montrer que si f est m -fortement convexe :

$$f(x) - f(x^*) \geq \frac{m}{2} \|x - x^*\|^2, \quad \forall x \in \mathbb{R}^n.$$

Chapitre 3

Sous-différentiels et Sous-gradients

3.1 Introduction

Voici un problème concret : vous voulez minimiser $\|x\|_1$ sous contraintes linéaires — un problème central en apprentissage parcimonieux. Mais la norme ℓ_1 n'est pas différentiable à l'origine : elle possède un « coin ». Le gradient n'existe pas, et pourtant c'est précisément là que se trouve le minimum. Comment écrire une condition d'optimalité sans gradient ?

La réponse, développée par Jean-Jacques Moreau et R. Tyrrell Rockafellar dans les années 1960, est le *sous-différentiel* : au lieu d'un unique plan tangent, on considère *tous* les hyperplans qui restent en-dessous du graphe de la fonction. Aux points de différentiabilité, il n'y en a qu'un (le gradient classique). Aux coins, il y en a une infinité, formant un ensemble convexe compact. Cette généralisation ouvre la voie à une théorie de l'optimisation non lisse d'une remarquable cohérence.

Intuition

Pour une fonction différentiable, le gradient définit l'unique hyperplan tangent qui minore la fonction. Pour une fonction convexe non lisse, il peut exister *plusieurs* hyperplans minorants en un point — leur ensemble forme le sous-différentiel.

3.2 Définition du sous-gradient et du sous-différentiel

Définition 3.1 (Sous-gradient). Soit $f : \mathbb{R}^n \rightarrow \mathbb{R} \cup \{+\infty\}$ une fonction convexe. Un vecteur $g \in \mathbb{R}^n$ est un *sous-gradient* de f en $x \in \text{dom}(f)$ si :

$$f(y) \geq f(x) + \langle g, y - x \rangle, \quad \forall y \in \mathbb{R}^n.$$

Définition 3.2 (Sous-différentiel). Le *sous-différentiel* de f en x , noté $\partial f(x)$, est l'ensemble de tous les sous-gradients de f en x :

$$\partial f(x) = \{g \in \mathbb{R}^n \mid f(y) \geq f(x) + \langle g, y - x \rangle, \forall y \in \mathbb{R}^n\}.$$

Si $x \notin \text{dom}(f)$, on pose $\partial f(x) = \emptyset$.

Théorème 3.3 (Existence des sous-gradients). Soit $f : \mathbb{R}^n \rightarrow \mathbb{R} \cup \{+\infty\}$ convexe propre. Alors $\partial f(x) \neq \emptyset$ pour tout $x \in \text{ri}(\text{dom}(f))$. En particulier, si $\text{dom}(f)$ est ouvert, $\partial f(x) \neq \emptyset$ pour tout $x \in \text{dom}(f)$.

Démonstration. Soit $x \in \text{ri}(\text{dom}(f))$. Le point $(x, f(x))$ est sur le bord de $\text{epi}(f)$, qui est convexe. Par le théorème de l'hyperplan d'appui, il existe $(g, -\alpha) \neq 0$ tel que :

$$\langle g, y - x \rangle - \alpha(t - f(x)) \leq 0, \quad \forall (y, t) \in \text{epi}(f).$$

En prenant $y = x$ et $t \rightarrow +\infty$, on obtient $\alpha \geq 0$. Si $\alpha = 0$, alors $\langle g, y - x \rangle \leq 0$ pour tout $y \in \text{dom}(f)$, ce qui contredit $x \in \text{ri}(\text{dom}(f))$ (sauf si $g = 0$). Donc $\alpha > 0$, et g/α est un sous-gradient. $\square$

3.3 Propriétés du sous-différentiel

Théorème 3.4 (Propriétés fondamentales). *Soit $f : \mathbb{R}^n \rightarrow \mathbb{R} \cup \{+\infty\}$ convexe propre.*

1. $\partial f(x)$ est un ensemble convexe fermé (éventuellement vide).
2. Si f est différentiable en x , alors $\partial f(x) = \{\nabla f(x)\}$.
3. $\partial f(x)$ est borné si $x \in \text{int}(\text{dom}(f))$.
4. x^* minimise f si et seulement si $0 \in \partial f(x^*)$.

Preuve de (4). ($\Rightarrow$) Si x^* minimise f , alors $f(y) \geq f(x^*)$ pour tout y , c'est-à-dire $f(y) \geq f(x^*) + \langle 0, y - x^* \rangle$, donc $0 \in \partial f(x^*)$.

($\Leftarrow$) Si $0 \in \partial f(x^*)$, alors $f(y) \geq f(x^*) + \langle 0, y - x^* \rangle = f(x^*)$ pour tout y , donc x^* est un minimiseur global. $\square$

Condition d'optimalité sous-différentielle

$$x^* \in \arg \min_{x \in \mathbb{R}^n} f(x) \iff 0 \in \partial f(x^*).$$

C'est la généralisation de $\nabla f(x^*) = 0$ au cas non lisse.

3.4 Exemples de sous-différentiels

Exemple 3.5 (Valeur absolue). Pour $f(x) = |x|$ sur $\mathbb{R}$:

$$\partial f(x) = \begin{cases} \{-1\} & \text{si } x < 0, \\ [-1, 1] & \text{si } x = 0, \\ \{1\} & \text{si } x > 0. \end{cases}$$

Exemple 3.6 (Norme ℓ_1). Pour $f(x) = \|x\|_1 = \sum_i |x_i|$ sur $\mathbb{R}^n$:

$$\partial f(x) = \{g \in \mathbb{R}^n \mid g_i = \text{sign}(x_i) \text{ si } x_i \neq 0, g_i \in [-1, 1] \text{ si } x_i = 0\}.$$

Exemple 3.7 (Maximum de fonctions lisses). Pour $f(x) = \max(f_1(x), \dots, f_k(x))$ avec f_i différentiables et convexes :

$$\partial f(x) = \text{conv}\{\nabla f_i(x) \mid i \in I(x)\},$$

où $I(x) = \{i \mid f_i(x) = f(x)\}$ est l'ensemble des indices actifs.

Exemple 3.8 (Fonction indicatrice). Pour $f = \iota_C$ (indicatrice d'un convexe fermé C) :

$$\partial \iota_C(x) = N_C(x) = \{g \in \mathbb{R}^n \mid \langle g, y - x \rangle \leq 0, \forall y \in C\},$$

qui est le cône normal à C en x .

3.5 Règles de calcul sous-différentiel

Théorème 3.9 (Règle de somme de Moreau–Rockafellar). *Soient $f_1, f_2 : \mathbb{R}^n \rightarrow \mathbb{R} \cup \{+\infty\}$ convexes propres. Si $\text{ri}(\text{dom}(f_1)) \cap \text{ri}(\text{dom}(f_2)) \neq \emptyset$, alors :*

$$\partial(f_1 + f_2)(x) = \partial f_1(x) + \partial f_2(x).$$

Théorème 3.10 (Règle de la chaîne). *Soit f convexe et $g(x) = f(Ax + b)$ avec $A \in \mathbb{R}^{m \times n}$, sous la condition de qualification $Ax + b \in \text{ri}(\text{dom}(f))$. Alors :*

$$\partial g(x) = A^T \partial f(Ax + b).$$

Théorème 3.11 (Sous-différentiel du maximum). *Soient $f_1, \dots, f_k$ convexes. Pour $f(x) = \max_i f_i(x)$:*

$$\partial f(x) = \text{conv} \left(\bigcup_{i \in I(x)} \partial f_i(x) \right),$$

où $I(x) = \{i \mid f_i(x) = f(x)\}$.

Proposition 3.12 (Mise à l'échelle). Pour $\alpha > 0$: $\partial(\alpha f)(x) = \alpha \partial f(x)$.

3.6 Monotonie du sous-différentiel

Théorème 3.13 (Monotonie). *Soit f convexe. L'opérateur ∂f est monotone : pour tout $x, y \in \text{dom}(\partial f)$, $g_x \in \partial f(x)$, $g_y \in \partial f(y)$:*

$$\langle g_x - g_y, x - y \rangle \geq 0.$$

Si f est m -fortement convexe, ∂f est fortement monotone :

$$\langle g_x - g_y, x - y \rangle \geq m \|x - y\|^2.$$

Démonstration. Par définition des sous-gradients :

$$\begin{aligned} f(y) &\geq f(x) + \langle g_x, y - x \rangle, \\ f(x) &\geq f(y) + \langle g_y, x - y \rangle. \end{aligned}$$

En sommant : $0 \geq \langle g_x, y - x \rangle + \langle g_y, x - y \rangle = -\langle g_x - g_y, x - y \rangle$. □

3.7 Sous-gradient et direction de descente

Proposition 3.14. Soit f convexe et $x \notin \arg \min f$. Alors pour tout $g \in \partial f(x)$ avec $g \neq 0$, la direction $d = -g$ est une direction de descente :

$$f(x - tg) < f(x)$$

pour $t > 0$ assez petit (si f est continue en x).

Sous-gradient vs gradient

Contrairement au gradient, la direction $-g$ avec $g \in \partial f(x)$ n'est *pas* nécessairement la direction de plus forte descente. De plus, la méthode de sous-gradient ne garantit pas une décroissance monotone de f .

3.8 Méthode du sous-gradient

Méthode du sous-gradient

1. Initialiser $x_0 \in \text{dom}(f)$, $f_{\text{best}} = +\infty$.
2. Pour $k = 0, 1, 2, \dots$:
 - (a) Choisir $g_k \in \partial f(x_k)$.
 - (b) Mettre à jour : $x_{k+1} = x_k - \alpha_k g_k$.
 - (c) $f_{\text{best}} = \min(f_{\text{best}}, f(x_k))$.

Choix du pas :

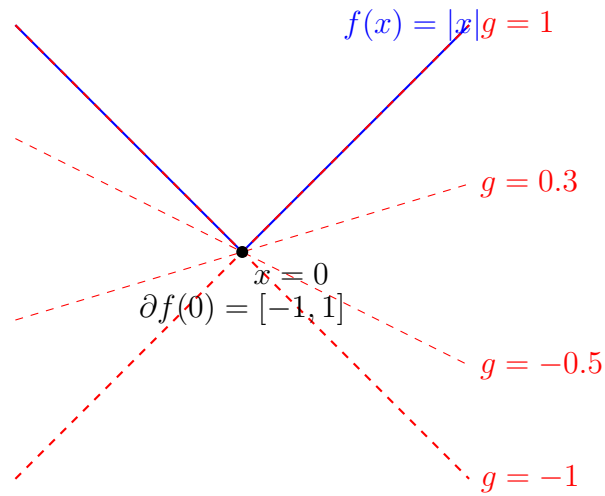
- Pas constant : $\alpha_k = \alpha$.
- Pas décroissant : $\alpha_k \rightarrow 0$ avec $\sum \alpha_k = +\infty$.
- Polyak : $\alpha_k = \frac{f(x_k) - f^*}{\|g_k\|^2}$ (si f^* connu).

Théorème 3.15 (Convergence du sous-gradient). Avec le pas $\alpha_k = c/\sqrt{k+1}$ et $\|g_k\| \leq G$ pour tout k :

$$f_{\text{best}}^{(k)} - f^* \leq \frac{\|x_0 - x^*\|^2 + c^2 G^2 \sum_{i=0}^k 1}{2c\sqrt{k+1}} = O\left(\frac{1}{\sqrt{k}}\right).$$

Le taux optimal est $O(1/\sqrt{k})$, beaucoup plus lent que la descente de gradient pour les fonctions lisses.

3.9 Interprétation géométrique



3.10 Implémentation Python

Méthode du sous-gradient pour la norme ℓ_1

```
import numpy as np
import matplotlib.pyplot as plt

def subgradient_l1(A, b, x0, n_iter=500, c=1.0):
    """Minimise ||Ax - b||_1 par sous-gradient."""
    x = x0.copy()
    n = len(x)
    f_best = np.inf
    x_best = x.copy()
    history = []

    for k in range(n_iter):
        r = A @ x - b
        f_val = np.sum(np.abs(r))
        if f_val < f_best:
            f_best = f_val
            x_best = x.copy()
        history.append(f_best)

        # Sous-gradient de ||Ax - b||_1
        g = A.T @ np.sign(r)
        # Pas de Polyak-like
        alpha = c / np.sqrt(k + 1)
        x = x - alpha * g

    return x_best, history

# Test
np.random.seed(42)
```

```
m, n = 50, 20
A = np.random.randn(m, n)
x_true = np.zeros(n)
x_true[:5] = np.random.randn(5)
b = A @ x_true + 0.1 * np.random.randn(m)

x0 = np.zeros(n)
x_sol, hist = subgradient_l1(A, b, x0, n_iter=1000)

plt.semilogy(hist)
plt.xlabel('Iteration')
plt.ylabel('Best f value')
plt.title('Subgradient method for L1 minimization')
plt.grid(True, alpha=0.3)
plt.savefig('subgradient_l1.pdf')
```

3.11 Exercices

Exercice 3.1 (*). Calculer le sous-différentiel de $f(x) = \max(0, x)$ en tout point $x \in \mathbb{R}$.

Exercice 3.2 (*). Calculer $\partial f(x)$ pour $f(x) = \|x\|_2$ sur $\mathbb{R}^n$.

Exercice 3.3 (**). Soit $f(x) = \max(x_1, x_2)$ sur $\mathbb{R}^2$. Calculer $\partial f(x)$ pour (x_1, x_2) avec $x_1 > x_2$, $x_1 < x_2$, et $x_1 = x_2$.

Exercice 3.4 (**). Montrer la règle de somme : si f et g sont convexes avec $\text{dom}(f) \cap \text{int}(\text{dom}(g)) \neq \emptyset$, alors $\partial(f + g)(x) = \partial f(x) + \partial g(x)$.

Exercice 3.5 (**). Montrer que le cône normal $N_C(x) = \partial \iota_C(x)$ est un cône convexe fermé.

Exercice 3.6 (***). Montrer que f est convexe si et seulement si l'opérateur ∂f est monotone.

Exercice 3.7 (***). Implémenter la méthode du sous-gradient avec pas de Polyak pour minimiser $f(x) = \|Ax - b\|_1 + \lambda \|x\|_\infty$ et comparer avec CVXPY.

Chapitre 4

Dualité de Fenchel–Moreau–Rockafellar

La transformée de Legendre, connue des physiciens depuis le XVIII^e siècle pour passer du lagrangien au hamiltonien, a été généralisée par Werner Fenchel dans les années 1940 en un outil d’une puissance remarquable : la *conjugaison convexe*. L’idée est de représenter une fonction convexe non pas par son graphe, mais par l’ensemble de ses hyperplans d’appui — une sorte de dualité géométrique où chaque fonction convexe a un alter ego. Le théorème de biconjugaison de Fenchel–Moreau affirme que pour une fonction convexe semi-continue inférieurement, appliquer deux fois cette transformation redonne la fonction de départ. Cette involution est la clé de la dualité en optimisation convexe.

4.1 Introduction

La conjugaison convexe, ou transformée de Legendre–Fenchel, est un outil fondamental de l’analyse convexe. Elle établit une correspondance entre fonctions convexes et permet de construire des problèmes duaux puissants. Le théorème de biconjugaison de Fenchel–Moreau caractérise les fonctions convexes semi-continues inférieurement.

4.2 Conjuguée de Fenchel

Définition 4.1 (Conjuguée de Fenchel). Soit $f : \mathbb{R}^n \rightarrow \mathbb{R} \cup \{+\infty\}$. La *conjuguée de Fenchel* (ou *transformée de Legendre–Fenchel*) de f est :

$$f^*(y) = \sup_{x \in \mathbb{R}^n} \{ \langle y, x \rangle - f(x) \}.$$

Proposition 4.2 (Propriétés immédiates). 1. f^* est toujours convexe et semi-continue inférieurement (comme supremum de fonctions affines en y).

2. **Inégalité de Fenchel–Young** : $f(x) + f^*(y) \geq \langle x, y \rangle$ pour tout x, y .

3. Si f est propre, alors f^* ne prend jamais la valeur $-\infty$.

Preuve de l’inégalité de Fenchel–Young. Par définition de f^* :

$$f^*(y) = \sup_z \{ \langle y, z \rangle - f(z) \} \geq \langle y, x \rangle - f(x),$$

d’où $f(x) + f^*(y) \geq \langle x, y \rangle$. □

4.5 Lien entre conjuguée et sous-différentiel

Théorème 4.5 (Équivalence conjuguée–sous-différentiel). *Soit f convexe propre et s.c.i. Les assertions suivantes sont équivalentes :*

1. $y \in \partial f(x)$,
2. $x \in \partial f^*(y)$,
3. $f(x) + f^*(y) = \langle x, y \rangle$.

Démonstration. (1) $\Rightarrow$ (3) : Si $y \in \partial f(x)$, alors pour tout z : $f(z) \geq f(x) + \langle y, z - x \rangle$, soit $\langle y, z \rangle - f(z) \leq \langle y, x \rangle - f(x)$. Donc $f^*(y) = \langle y, x \rangle - f(x)$.

(3) $\Rightarrow$ (1) : $f^*(y) = \langle y, x \rangle - f(x)$ signifie que le supremum dans la définition de f^* est atteint en x , donc $y \in \partial f(x)$.

(1) $\Leftrightarrow$ (2) : Par $f^{**} = f$ et symétrie. □

4.6 Dualité de Fenchel–Rockafellar

Théorème 4.6 (Dualité de Fenchel–Rockafellar). *Soient $f, g : \mathbb{R}^n \rightarrow \mathbb{R} \cup \{+\infty\}$ convexes propres et $A \in \mathbb{R}^{m \times n}$. Considérons le problème primal :*

$$(P) \quad p^* = \inf_{x \in \mathbb{R}^n} \{f(x) + g(Ax)\}$$

et le problème dual :

$$(D) \quad d^* = \sup_{y \in \mathbb{R}^m} \{-f^*(-A^T y) - g^*(y)\}.$$

Alors :

1. (Dualité faible) $d^* \leq p^*$.
2. (Dualité forte) Si $\text{ri}(\text{dom}(f)) \cap A^{-1}(\text{ri}(\text{dom}(g))) \neq \emptyset$ (condition de qualification), alors $d^* = p^*$ et le supremum dual est atteint.

Preuve de la dualité faible. Pour tout x et y , par l'inégalité de Fenchel–Young :

$$\begin{aligned} f(x) &\geq \langle -A^T y, x \rangle - f^*(-A^T y) = -\langle y, Ax \rangle - f^*(-A^T y), \\ g(Ax) &\geq \langle y, Ax \rangle - g^*(y). \end{aligned}$$

En sommant : $f(x) + g(Ax) \geq -f^*(-A^T y) - g^*(y)$. En prenant l'infimum en x et le supremum en y : $p^* \geq d^*$. □

4.7 Cas particulier : $A = I$

Corollaire 4.7 (Dualité de Fenchel). *Pour f, g convexes propres avec $\text{ri}(\text{dom}(f)) \cap \text{ri}(\text{dom}(g)) \neq \emptyset$:*

$$\inf_x \{f(x) + g(x)\} = \sup_y \{-f^*(-y) - g^*(y)\} = -\inf_y \{f^*(-y) + g^*(y)\}.$$

4.8 Règles de calcul des conjuguées

Proposition 4.8 (Règles de calcul). 1. **Séparabilité** : si $f(x) = \sum_i f_i(x_i)$, alors $f^*(y) = \sum_i f_i^*(y_i)$.

2. **Mise à l'échelle** : $(\alpha f)^*(y) = \alpha f^*(y/\alpha)$ pour $\alpha > 0$.

3. **Translation** : si $g(x) = f(x - a)$, alors $g^*(y) = f^*(y) + \langle y, a \rangle$.

4. **Ajout de linéaire** : si $g(x) = f(x) + \langle b, x \rangle + c$, alors $g^*(y) = f^*(y - b) - c$.

5. **Composition linéaire** : si $g(x) = f(Ax)$ avec A inversible, alors $g^*(y) = f^*(A^{-T}y)$.

6. **Infimale convolution** : $(f \square g)^* = f^* + g^*$.

4.9 Application : formulation duale du LASSO

Le problème LASSO primal est :

$$\min_x \frac{1}{2} \|Ax - b\|^2 + \lambda \|x\|_1.$$

En écrivant $f(x) = \lambda \|x\|_1$ et $g(z) = \frac{1}{2} \|z - b\|^2$ avec $z = Ax$:

$$\begin{aligned} f^*(y) &= \iota_{\{y \mid \|y\|_\infty \leq \lambda\}}(y), \\ g^*(w) &= \frac{1}{2} \|w\|^2 + \langle w, b \rangle. \end{aligned}$$

Le dual de Fenchel–Rockafellar donne :

$$\max_{\nu} \left\{ -\frac{1}{2} \|\nu\|^2 + \langle \nu, b \rangle - \frac{1}{2} \|b\|^2 \right\} \quad \text{sous} \quad \|A^T \nu\|_\infty \leq \lambda.$$

Soit, après simplification :

$$\max_{\nu} \left\{ -\frac{1}{2} \|b - \nu\|^2 \right\} \quad \text{sous} \quad \|A^T \nu\|_\infty \leq \lambda.$$

4.10 Implémentation Python

Calcul de conjuguées et vérification de Fenchel–Young

```
import numpy as np
import cvxpy as cp

def fenchel_conjugate_numerical(f, y, n, bounds=(-10, 10)):
    """Calcul numérique de  $f^*(y) = \sup_x \{\langle y, x \rangle - f(x)\}$ ."""
    x = cp.Variable(n)
    objective = cp.Maximize(y @ x - f(x))
    prob = cp.Problem(objective)
    prob.solve()
```

```

    return prob.value

# Exemple:  $f(x) = \|x\|_2^2 / 2$ 
n = 5
y_test = np.random.randn(n)

f = lambda x: 0.5 * cp.sum_squares(x)
f_star_val = fenchel_conjugate_numerical(f, y_test, n)
f_star_exact = 0.5 * np.sum(y_test**2)

print(f"f*(y) numerique: {f_star_val:.6f}")
print(f"f*(y) analytique: {f_star_exact:.6f}")

# Verification Fenchel-Young:  $f(x) + f^*(y) \geq \langle x, y \rangle$ 
x_test = np.random.randn(n)
f_x = 0.5 * np.sum(x_test**2)
lhs = f_x + f_star_exact
rhs = np.dot(x_test, y_test)
print(f"f(x) + f*(y) = {lhs:.4f} >= <x,y> = {rhs:.4f}: {lhs >= rhs - 1e-10}")

# Dualite de Fenchel pour le LASSO
m, nn = 50, 20
A = np.random.randn(m, nn)
b = np.random.randn(m)
lam = 1.0

# Primal
x = cp.Variable(nn)
primal = cp.Problem(cp.Minimize(
    0.5 * cp.sum_squares(A @ x - b) + lam * cp.norm1(x)))
primal.solve()
print(f"\nLASSO primal: {primal.value:.6f}")

# Dual
nu = cp.Variable(m)
dual = cp.Problem(cp.Maximize(
    -0.5 * cp.sum_squares(b - nu)),
    [cp.norm_inf(A.T @ nu) <= lam])
dual.solve()
print(f"LASSO dual: {dual.value + 0.5*np.sum(b**2):.6f}")
    
```

4.11 Exercices

Exercice 4.1 (*). Calculer la conjuguée de $f(x) = \frac{1}{p} \|x\|_p^p$ pour $p > 1$. Vérifier que $f^*(y) = \frac{1}{q} \|y\|_q^q$ avec $1/p + 1/q = 1$.

Exercice 4.2 (*). Montrer que si $f(x) = \iota_C(x)$, alors $f^* = \sigma_C$.

Exercice 4.3 (**). Montrer que $f^{**} \leq f$ toujours, et que $f^{**} = f$ si et seulement si f est

Chapitre 5

Conditions d'Optimalité — KKT

En optimisation sans contraintes, la condition d'optimalité est simple : le gradient s'annule. Mais dès qu'on ajoute des contraintes — inégalités, égalités — cette condition naïve ne suffit plus. Les conditions de Karush-Kuhn-Tucker (KKT), découvertes indépendamment par William Karush en 1939 (dans un mémoire de master resté longtemps ignoré) et par Harold Kuhn et Albert Tucker en 1951, généralisent la condition du gradient nul en introduisant les *multiplicateurs de Lagrange* pour les contraintes. En optimisation convexe, sous des conditions de qualification très générales, les conditions KKT sont non seulement nécessaires mais aussi *suffisantes* — elles caractérisent complètement les solutions optimales.

5.1 Introduction

Les conditions de Karush-Kuhn-Tucker (KKT) généralisent la condition classique $\nabla f(x^*) = 0$ aux problèmes d'optimisation avec contraintes. En optimisation convexe, sous des conditions de qualification très générales, les conditions KKT sont à la fois nécessaires et suffisantes pour l'optimalité.

5.2 Problème d'optimisation convexe sous contraintes

Considérons le problème général :

$$\begin{aligned} \min_{x \in \mathbb{R}^n} \quad & f(x) \\ \text{sous} \quad & g_i(x) \leq 0, \quad i = 1, \dots, m, \\ & h_j(x) = 0, \quad j = 1, \dots, p, \end{aligned} \tag{P}$$

où f et les g_i sont convexes et les h_j sont affines : $h_j(x) = a_j^T x - b_j$.

Définition 5.1 (Ensemble réalisable). L'*ensemble réalisable* est :

$$\mathcal{F} = \{x \in \mathbb{R}^n \mid g_i(x) \leq 0, \ i = 1, \dots, m, \ h_j(x) = 0, \ j = 1, \dots, p\}.$$

Définition 5.2 (Contrainte active). La contrainte $g_i(x) \leq 0$ est *active* en x si $g_i(x) = 0$. L'ensemble des indices actifs est $\mathcal{A}(x) = \{i \mid g_i(x) = 0\}$.

5.3 Conditions KKT

Définition 5.3 (Conditions KKT). Un point $x^* \in \mathcal{F}$ satisfait les *conditions KKT* s'il existe des multiplicateurs $\lambda^* \in \mathbb{R}^m$ et $\nu^* \in \mathbb{R}^p$ tels que :

1. **Stationnarité** : $0 \in \partial f(x^*) + \sum_{i=1}^m \lambda_i^* \partial g_i(x^*) + \sum_{j=1}^p \nu_j^* a_j$.
2. **Réalisabilité primale** : $g_i(x^*) \leq 0$ et $h_j(x^*) = 0$.
3. **Réalisabilité duale** : $\lambda_i^* \geq 0$ pour tout i .
4. **Complémentarité** : $\lambda_i^* g_i(x^*) = 0$ pour tout i .

Si f et les g_i sont différentiables, la condition de stationnarité s'écrit :

$$\nabla f(x^*) + \sum_{i=1}^m \lambda_i^* \nabla g_i(x^*) + \sum_{j=1}^p \nu_j^* a_j = 0.$$

Conditions KKT (cas différentiable)

$$\nabla f(x^*) + \sum_{i=1}^m \lambda_i^* \nabla g_i(x^*) + A^T \nu^* = 0, \quad \lambda^* \geq 0, \quad \lambda^* \circ g(x^*) = 0, \quad g(x^*) \leq 0, \quad Ax^* = b.$$

Intuition

La condition de complémentarité $\lambda_i^* g_i(x^*) = 0$ signifie que soit la contrainte est inactive ($g_i(x^*) < 0$ et $\lambda_i^* = 0$), soit elle est active ($g_i(x^*) = 0$ et $\lambda_i^* \geq 0$ potentiellement non nul). Le multiplicateur λ_i^* mesure la « force » avec laquelle la contrainte pousse la solution.

5.4 Qualification des contraintes

Définition 5.4 (Condition de Slater). La *condition de Slater* est satisfaite s'il existe un point strictement réalisable $\hat{x}$ tel que :

$$g_i(\hat{x}) < 0, \quad i = 1, \dots, m, \quad h_j(\hat{x}) = 0, \quad j = 1, \dots, p.$$

(Pour les contraintes affines g_i , il suffit que $g_i(\hat{x}) \leq 0$.)

Théorème 5.5 (KKT — conditions nécessaires et suffisantes). *Considérons le problème (P) avec f, g_i convexes et h_j affines.*

1. **Suffisance (toujours)** : si (x^*, λ^*, ν^*) satisfait les conditions KKT, alors x^* est une solution optimale de (P).
2. **Nécessité (sous Slater)** : si la condition de Slater est satisfaite et x^* est optimal, alors il existe (λ^*, ν^*) tel que (x^*, λ^*, ν^*) satisfait les conditions KKT.

Preuve de la suffisance. Supposons que (x^*, λ^*, ν^*) satisfait KKT. Pour tout $x \in \mathcal{F}$:

$$\begin{aligned}
 f(x) &\geq f(x^*) + \langle g_f, x - x^* \rangle \quad \text{où } g_f \in \partial f(x^*) \\
 &= f(x^*) - \sum_i \lambda_i^* \langle g_{g_i}, x - x^* \rangle - \sum_j \nu_j^* a_j^T (x - x^*) \quad (\text{stationnarité}) \\
 &\geq f(x^*) - \sum_i \lambda_i^* (g_i(x) - g_i(x^*)) - \sum_j \nu_j^* (h_j(x) - h_j(x^*)) \quad (\text{sous-gradient de } g_i) \\
 &\geq f(x^*) - \sum_i \lambda_i^* g_i(x) + \sum_i \lambda_i^* g_i(x^*) \quad (\text{car } h_j(x) = h_j(x^*) = 0) \\
 &\geq f(x^*) \quad (\text{car } \lambda_i^* \geq 0, g_i(x) \leq 0, \lambda_i^* g_i(x^*) = 0).
 \end{aligned}$$

□

5.5 Exemples d'application des conditions KKT

Exemple 5.6 (Programmation quadratique). Considérons :

$$\min_x \frac{1}{2} x^T Q x + c^T x \quad \text{sous} \quad Ax \leq b.$$

Les conditions KKT sont :

$$Qx^* + c + A^T \lambda^* = 0, \quad \lambda^* \geq 0, \quad Ax^* \leq b, \quad \lambda^* \circ (Ax^* - b) = 0.$$

Exemple 5.7 (Projection sur un convexe). La projection $\bar{x} = \text{proj}_C(y)$ résout :

$$\min_x \frac{1}{2} \|x - y\|^2 \quad \text{sous} \quad x \in C.$$

Si $C = \{x \mid g_i(x) \leq 0\}$, les conditions KKT donnent :

$$\bar{x} - y + \sum_i \lambda_i^* \nabla g_i(\bar{x}) = 0, \quad \lambda_i^* \geq 0, \quad \lambda_i^* g_i(\bar{x}) = 0.$$

Exemple 5.8 (SVM à marge dure). Le SVM primal :

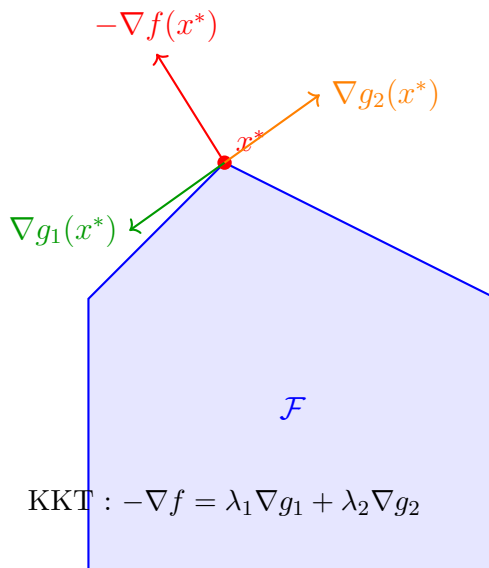
$$\min_{w,b} \frac{1}{2} \|w\|^2 \quad \text{sous} \quad y_i(\langle w, x_i \rangle + b) \geq 1, \quad i = 1, \dots, N.$$

Les conditions KKT :

$$w^* = \sum_i \alpha_i^* y_i x_i, \quad \sum_i \alpha_i^* y_i = 0, \quad \alpha_i^* \geq 0, \quad \alpha_i^* (y_i(\langle w^*, x_i \rangle + b^*) - 1) = 0.$$

La complémentarité montre que seuls les points avec $y_i(\langle w^*, x_i \rangle + b^*) = 1$ (vecteurs de support) ont $\alpha_i^* > 0$.

5.6 Interprétation géométrique



5.7 Autres qualifications de contraintes

Définition 5.9 (Indépendance linéaire (LICQ)). La *LICQ* est satisfaite en x^* si les gradients des contraintes actives $\{\nabla g_i(x^*)\}_{i \in \mathcal{A}(x^*)} \cup \{a_j\}_{j=1}^p$ sont linéairement indépendants.

Proposition 5.10. Sous LICQ, les multiplicateurs KKT sont uniques.

Définition 5.11 (Qualification de Mangasarian–Fromovitz (MFCQ)). La *MFCQ* est satisfaite en x^* si les $\{a_j\}$ sont linéairement indépendants et il existe $d \in \mathbb{R}^n$ tel que :

$$\langle \nabla g_i(x^*), d \rangle < 0, \quad \forall i \in \mathcal{A}(x^*), \quad a_j^T d = 0, \quad \forall j.$$

5.8 Sensibilité et interprétation économique

Théorème 5.12 (Sensibilité). Soit $p^*(u, v)$ la valeur optimale du problème perturbé :

$$\min f(x) \quad \text{sous} \quad g_i(x) \leq u_i, \quad h_j(x) = v_j.$$

Sous des conditions de régularité, les multiplicateurs KKT satisfont :

$$\lambda_i^* = - \frac{\partial p^*}{\partial u_i} \Big|_{(0,0)}, \quad \nu_j^* = - \frac{\partial p^*}{\partial v_j} \Big|_{(0,0)}.$$

Remarque 5.13. λ_i^* représente le « prix marginal » de la i -ème contrainte : relâcher la contrainte d'une unité améliore la valeur optimale de λ_i^* .

5.9 Implémentation Python

Vérification des conditions KKT avec CVXPY

```

import numpy as np
import cvxpy as cp

# Programmation quadratique
n = 5
np.random.seed(42)
Q = np.random.randn(n, n)
Q = Q.T @ Q + 0.1 * np.eye(n) # PD
c = np.random.randn(n)
A = np.random.randn(3, n)
b = np.ones(3)

x = cp.Variable(n)
constraints = [A @ x <= b]
prob = cp.Problem(cp.Minimize(0.5 * cp.quad_form(x, Q) + c @ x),
                  constraints)
prob.solve()

x_star = x.value
lambda_star = constraints[0].dual_value

print("Solution x*:", x_star)
print("Multipliateurs lambda*:", lambda_star)

# Verification stationnarite
grad_L = Q @ x_star + c + A.T @ lambda_star
print(f"||grad L|| = {np.linalg.norm(grad_L):.2e}")

# Verification complementarite
slack = b - A @ x_star
comp = lambda_star * slack
print(f"Complementarite: max|lambda*g| = {np.max(np.abs(comp)):.2e}")

# Verification faisabilite
print(f"Faisabilite primale: max(Ax-b) = {np.max(A @ x_star - b):.2e}")
print(f"Faisabilite duale: min(lambda) = {np.min(lambda_star):.2e}")

```

5.10 Exercices

Exercice 5.1 (★). Résoudre $\min x_1^2 + x_2^2$ sous $x_1 + x_2 \geq 1$ en utilisant les conditions KKT.

Exercice 5.2 (★). Vérifier que la condition de Slater est satisfaite pour $\min \|x\|^2$ sous $\|x\|_\infty \leq 1$.

Exercice 5.3 (★★). Dériver les conditions KKT du problème SVM à marge souple :

$$\min_{w,b,\xi} \frac{1}{2} \|w\|^2 + C \sum_i \xi_i \quad \text{sous} \quad y_i(\langle w, x_i \rangle + b) \geq 1 - \xi_i, \quad \xi_i \geq 0.$$

Exercice 5.4 (**). Montrer que pour un problème de programmation linéaire, les conditions KKT sont équivalentes aux conditions de complémentarité primale-duale.

Exercice 5.5 (***). Démontrer le théorème de sensibilité 5.12 pour un problème de programmation linéaire.

Exercice 5.6 (***). Montrer que sous la condition LICQ, les multiplicateurs KKT sont uniques. Donner un exemple où LICQ n'est pas satisfaite et les multiplicateurs ne sont pas uniques.

Chapitre 6

Dualité de Lagrange

La dualité lagrangienne est l'une des idées les plus unificatrices de l'optimisation. Son principe : relaxer les contraintes en les pénalisant dans l'objectif via des *multiplicateurs de Lagrange*, puis maximiser sur ces multiplicateurs. Le problème dual ainsi obtenu est toujours convexe (même si le primal ne l'est pas), fournit une borne inférieure garantie, et coïncide avec le primal sous des conditions de *dualité forte* — lesquelles sont automatiquement satisfaites en optimisation convexe dès qu'une condition de qualification de Slater est vérifiée. Cette théorie, héritière des travaux de Lagrange (1797), Kuhn-Tucker (1951) et Rockafellar (1970), est le pont entre la formulation primal et les conditions KKT.

6.1 Introduction

La dualité lagrangienne est un principe fondamental qui associe à tout problème d'optimisation (primal) un problème dual. Le dual fournit des bornes inférieures sur la valeur optimale primale et, sous des conditions de convexité et de qualification, la dualité forte permet de résoudre le primal via le dual.

6.2 Le lagrangien

Définition 6.1 (Lagrangien). Le *lagrangien* du problème (P) est la fonction $L : \mathbb{R}^n \times \mathbb{R}_+^m \times \mathbb{R}^p \rightarrow \mathbb{R}$ définie par :

$$L(x, \lambda, \nu) = f(x) + \sum_{i=1}^m \lambda_i g_i(x) + \sum_{j=1}^p \nu_j h_j(x).$$

Les variables $\lambda \in \mathbb{R}_+^m$ et $\nu \in \mathbb{R}^p$ sont les *variables duales* ou *multiplicateurs de Lagrange*.

Proposition 6.2 (Le lagrangien comme relaxation). Pour tout $\lambda \geq 0$ et ν , et tout x réalisable :

$$L(x, \lambda, \nu) \leq f(x).$$

En effet, $\lambda_i g_i(x) \leq 0$ et $\nu_j h_j(x) = 0$.

6.3 Fonction duale de Lagrange

Définition 6.3 (Fonction duale). La *fonction duale de Lagrange* est :

$$q(\lambda, \nu) = \inf_{x \in \mathbb{R}^n} L(x, \lambda, \nu) = \inf_x \left\{ f(x) + \sum_i \lambda_i g_i(x) + \sum_j \nu_j h_j(x) \right\}.$$

Théorème 6.4 (Propriétés de la fonction duale). 1. q est concave (comme infimum de fonctions affines en (λ, ν)).

2. **Dualité faible** : $q(\lambda, \nu) \leq p^*$ pour tout $\lambda \geq 0$ et ν , où p^* est la valeur optimale primale.

3. Le domaine $\text{dom}(q) = \{(\lambda, \nu) \mid q(\lambda, \nu) > -\infty\}$ est convexe.

Preuve de la dualité faible. Pour tout x réalisable et tout (λ, ν) avec $\lambda \geq 0$:

$$q(\lambda, \nu) = \inf_z L(z, \lambda, \nu) \leq L(x, \lambda, \nu) = f(x) + \sum_i \underbrace{\lambda_i g_i(x)}_{\leq 0} + \sum_j \underbrace{\nu_j h_j(x)}_{=0} \leq f(x).$$

En prenant l'infimum sur les x réalisables : $q(\lambda, \nu) \leq p^*$. □

6.4 Problème dual

Définition 6.5 (Problème dual de Lagrange). Le *problème dual* est :

$$(D) \quad d^* = \sup_{\lambda \geq 0, \nu} q(\lambda, \nu) = \sup_{\lambda \geq 0, \nu} \inf_x L(x, \lambda, \nu).$$

Définition 6.6 (Saut de dualité). Le *saut de dualité* (*duality gap*) est $p^* - d^* \geq 0$. Si $p^* = d^*$, on dit que la *dualité forte* a lieu.

Relation primal–dual

$$d^* = \sup_{\lambda \geq 0, \nu} \inf_x L(x, \lambda, \nu) \leq \inf_x \sup_{\lambda \geq 0, \nu} L(x, \lambda, \nu) = p^*.$$

L'inégalité est l'inégalité max-min.

6.5 Dualité forte

Théorème 6.7 (Dualité forte — Slater). Si le problème primal est convexe (f, g_i convexes, h_j affines), que p^* est fini, et que la condition de Slater est satisfaite, alors :

1. La dualité forte a lieu : $p^* = d^*$.
2. Le supremum dual est atteint : il existe (λ^*, ν^*) optimal.

Idée de la preuve. On considère l'ensemble :

$$\mathcal{G} = \{(u, v, t) \in \mathbb{R}^m \times \mathbb{R}^p \times \mathbb{R} \mid \exists x : g_i(x) \leq u_i, h_j(x) = v_j, f(x) \leq t\}.$$

$\mathcal{G}$ est convexe (convexité de f, g_i et affinité de h_j). Le point $(0, 0, p^*)$ est sur le bord de $\mathcal{G}$. Par le théorème de l'hyperplan d'appui, il existe un hyperplan séparant, ce qui donne les multiplicateurs de Lagrange optimaux. □

6.6 Exemples de calcul dual

Exemple 6.8 (Programmation linéaire). Primal : $\min c^T x$ sous $Ax \leq b, x \geq 0$.

$$L(x, \lambda, \mu) = c^T x + \lambda^T (Ax - b) - \mu^T x = (c + A^T \lambda - \mu)^T x - \lambda^T b.$$

$$q(\lambda, \mu) = \inf_x L = \begin{cases} -\lambda^T b & \text{si } c + A^T \lambda - \mu = 0, \\ -\infty & \text{sinon.} \end{cases}$$

Dual : $\max -b^T \lambda$ sous $A^T \lambda \leq c, \lambda \geq 0$, soit $\max b^T y$ sous $A^T y \leq c, y \leq 0$ (en posant $y = -\lambda$).

Après reformulation standard : le dual de la PL est une PL.

Exemple 6.9 (Norme minimale). Primal : $\min \|x\|^2$ sous $Ax = b$.

$$L(x, \nu) = \|x\|^2 + \nu^T (Ax - b).$$

$$\inf_x L = \inf_x \{\|x\|^2 + \nu^T Ax\} - \nu^T b.$$

La condition d'optimalité $2x + A^T \nu = 0$ donne $x = -\frac{1}{2} A^T \nu$.

$$q(\nu) = \frac{1}{4} \nu^T A A^T \nu - \frac{1}{2} \nu^T A A^T \nu - \nu^T b = -\frac{1}{4} \nu^T A A^T \nu - \nu^T b.$$

$$\text{Dual : } \max -\frac{1}{4} \nu^T A A^T \nu - \nu^T b.$$

Exemple 6.10 (Entropie maximale). $\min \sum_i x_i \log x_i$ sous $Ax = b, \mathbf{1}^T x = 1, x \geq 0$.

Dual : $\max -\log \left(\sum_i \exp(-a_i^T \nu - \mu) \right)$ sous les contraintes duales, où a_i est la i -ème colonne de A .

6.7 Interprétation min-max

Théorème 6.11 (Théorème du minimax de Von Neumann–Sion). *Si $\mathcal{X}$ est compact convexe, $\mathcal{Y}$ est convexe, et $\Phi(x, y)$ est convexe-concave (convexe en x , concave en y), alors :*

$$\min_{x \in \mathcal{X}} \sup_{y \in \mathcal{Y}} \Phi(x, y) = \sup_{y \in \mathcal{Y}} \min_{x \in \mathcal{X}} \Phi(x, y).$$

Remarque 6.12. Le lagrangien $L(x, \lambda, \nu)$ est convexe en x (somme de fonctions convexes) et affine (donc concave) en (λ, ν) . Sous les conditions de Slater, le minimax s'applique, donnant la dualité forte.

6.8 Point-selle du lagrangien

Définition 6.13 (Point-selle). (x^*, λ^*, ν^*) est un *point-selle* du lagrangien si :

$$L(x^*, \lambda, \nu) \leq L(x^*, \lambda^*, \nu^*) \leq L(x, \lambda^*, \nu^*), \quad \forall x, \forall \lambda \geq 0, \nu.$$

Théorème 6.14 (Équivalence point-selle et KKT). *Sous les conditions de Slater, (x^*, λ^*, ν^*) est un point-selle du lagrangien si et seulement si x^* est primal optimal, (λ^*, ν^*) est dual optimal, et la dualité forte a lieu.*

6.9 Interprétation économique

Intuition

Considérons le problème de maximisation de profit d'une entreprise sous contraintes de ressources :

$$\max_x \text{profit}(x) \quad \text{sous} \quad \text{ressource}_i(x) \leq b_i.$$

Le multiplicateur λ_i^* représente le *prix implicite* de la ressource i : c'est le gain marginal obtenu en augmentant la disponibilité de la ressource i d'une unité. Un marché concurrentiel fixerait naturellement ce prix.

6.10 Implémentation Python

Dualité lagrangienne et saut de dualité

```
import numpy as np
import cvxpy as cp

# Exemple: min ||x||^2 s.t. Ax = b
np.random.seed(42)
m, n = 3, 10
A = np.random.randn(m, n)
b = np.random.randn(m)

# Primal
x = cp.Variable(n)
primal = cp.Problem(cp.Minimize(cp.sum_squares(x)), [A @ x == b])
primal.solve()
p_star = primal.value
print(f"Valeur primale p* = {p_star:.6f}")
print(f"Multiplicateurs: {primal.constraints[0].dual_value}")

# Dual analytique: max -1/4 nu^T A A^T nu - nu^T b
nu = cp.Variable(m)
M = A @ A.T
dual = cp.Problem(cp.Maximize(-0.25 * cp.quad_form(nu, M) - b @ nu))
dual.solve()
d_star = dual.value
print(f"Valeur duale d* = {d_star:.6f}")
print(f"Saut de dualite = {p_star - d_star:.2e}")

# Programmation lineaire: verification dualite forte
c = np.random.randn(n)
A_ineq = np.random.randn(m, n)
b_ineq = np.abs(np.random.randn(m)) + 1

x_lp = cp.Variable(n)
primal_lp = cp.Problem(cp.Minimize(c @ x_lp),
```

```

[A_ineq @ x_lp <= b_ineq, x_lp >= 0])
primal_lp.solve()

y_lp = cp.Variable(m)
dual_lp = cp.Problem(cp.Maximize(b_ineq @ y_lp),
                      [A_ineq.T @ y_lp + c >= 0, y_lp <= 0])
dual_lp.solve()

print(f"\nPL primal: {primal_lp.value:.6f}")
print(f"PL dual:    {dual_lp.value:.6f}")
print(f"Gap:        {abs(primal_lp.value - dual_lp.value):.2e}")

```

6.11 Exercices

Exercice 6.1 (★). Écrire le dual de Lagrange de $\min c^T x$ sous $Ax = b, x \geq 0$.

Exercice 6.2 (★). Calculer la fonction duale pour $\min \frac{1}{2}x^T Qx + c^T x$ sous $Ax = b$ avec $Q \succ 0$.

Exercice 6.3 (★★). Montrer que la dualité forte a lieu pour la programmation linéaire sans nécessiter la condition de Slater.

Exercice 6.4 (★★). Dériver le dual du problème SVM à marge souple et montrer qu'il s'agit d'un QP borné.

Exercice 6.5 (★★). Montrer que (x^*, λ^*, ν^*) est un point-selle du lagrangien si et seulement si les conditions KKT sont satisfaites.

Exercice 6.6 (★★★). Démontrer le théorème de Slater en utilisant le théorème de l'hyperplan d'appui appliqué à l'ensemble $\mathcal{G}$ défini dans la preuve.

Exercice 6.7 (★★★). Soit $p^*(u)$ la valeur optimale du problème perturbé $\min f(x)$ sous $g_i(x) \leq u_i$. Montrer que p^* est convexe en u . En déduire l'interprétation des multiplicateurs comme sous-gradients de p^* .

Chapitre 7

Descente de Gradient et Variantes

7.1 Introduction

Augustin-Louis Cauchy, en 1847, propose une idée d'une simplicité désarmante : pour minimiser une fonction, il suffit de suivre la direction opposée au gradient. Comme une bille qui roule sur une surface, on descend la pente la plus raide, pas après pas, jusqu'à atteindre le fond de la vallée. Cette méthode, la *descente de gradient*, est restée longtemps un outil théorique. Mais l'explosion de l'apprentissage automatique au XXI^e siècle l'a propulsée au rang d'algorithme le plus exécuté au monde : chaque fois qu'un réseau de neurones apprend, c'est une variante de la descente de gradient qui ajuste ses milliards de paramètres. Les questions fondamentales — quel pas choisir ? comment accélérer la convergence ? que faire quand le gradient est trop coûteux à calculer exactement ? — ont conduit aux méthodes de Nesterov, d'Adam, et au gradient stochastique, qui forment aujourd'hui la boîte à outils incontournable de l'optimisation moderne.

7.2 Descente de gradient à pas fixe

Descente de gradient

1. Choisir $x_0 \in \mathbb{R}^n$, pas $\alpha > 0$.
2. Pour $k = 0, 1, 2, \dots$:
$$x_{k+1} = x_k - \alpha \nabla f(x_k).$$
3. Arrêter quand $\|\nabla f(x_k)\| < \varepsilon$.

Intuition

À chaque itération, on se déplace dans la direction opposée au gradient (direction de plus forte descente locale) avec un pas α . Le gradient pointe dans la direction de plus forte montée, donc $-\nabla f(x_k)$ est la direction de plus forte descente.

7.3 Analyse de convergence — fonctions L -lisses

Théorème 7.1 (Convergence pour fonctions convexes L -lisses). *Soit f convexe et L -lisse. Avec le pas $\alpha = 1/L$:*

$$f(x_k) - f^* \leq \frac{L \|x_0 - x^*\|^2}{2k}.$$

Le taux de convergence est $O(1/k)$.

Démonstration. Par la L -lissité, pour $\alpha = 1/L$:

$$f(x_{k+1}) \leq f(x_k) + \langle \nabla f(x_k), x_{k+1} - x_k \rangle + \frac{L}{2} \|x_{k+1} - x_k\|^2.$$

Avec $x_{k+1} - x_k = -\frac{1}{L} \nabla f(x_k)$:

$$f(x_{k+1}) \leq f(x_k) - \frac{1}{2L} \|\nabla f(x_k)\|^2.$$

Par la condition du premier ordre : $f(x_k) - f^* \leq \langle \nabla f(x_k), x_k - x^* \rangle \leq \|\nabla f(x_k)\| \|x_k - x^*\|$, donc :

$$\|\nabla f(x_k)\|^2 \geq \frac{(f(x_k) - f^*)^2}{\|x_k - x^*\|^2}.$$

De plus, $\|x_{k+1} - x^*\|^2 = \|x_k - x^*\|^2 - \frac{2}{L} \langle \nabla f(x_k), x_k - x^* \rangle + \frac{1}{L^2} \|\nabla f(x_k)\|^2$.

En posant $\delta_k = f(x_k) - f^*$ et $R_k = \|x_k - x^*\|^2$:

$$\delta_{k+1} \leq \delta_k - \frac{1}{2L} \|\nabla f(x_k)\|^2, \quad R_{k+1} \leq R_k - \frac{2}{L} \delta_k + \frac{1}{L^2} \|\nabla f(x_k)\|^2.$$

En sommant la décroissance suffisante : $\sum_{i=0}^{k-1} \frac{1}{2L} \|\nabla f(x_i)\|^2 \leq \delta_0$, et par un argument télescopique sur R_k :

$$\sum_{i=0}^{k-1} \delta_i \leq \frac{L}{2} R_0.$$

Comme δ_k est décroissante : $k\delta_k \leq \sum_{i=0}^{k-1} \delta_i \leq \frac{L}{2} R_0$. □

7.4 Convergence pour fonctions fortement convexes

Théorème 7.2 (Convergence linéaire). *Soit f m -fortement convexe et L -lisse. Avec $\alpha = 1/L$:*

$$f(x_k) - f^* \leq \left(1 - \frac{m}{L}\right)^k (f(x_0) - f^*).$$

Le taux est linéaire (convergence géométrique) avec facteur $1 - 1/\kappa$ où $\kappa = L/m$ est le nombre de conditionnement.

Démonstration. Par la L -lissité et le pas $\alpha = 1/L$: $f(x_{k+1}) \leq f(x_k) - \frac{1}{2L} \|\nabla f(x_k)\|^2$.

Par la forte convexité : $\|\nabla f(x_k)\|^2 \geq 2m(f(x_k) - f^*)$ (inégalité de Polyak–Łojasiewicz).
Donc :

$$f(x_{k+1}) - f^* \leq (f(x_k) - f^*) - \frac{m}{L} (f(x_k) - f^*) = \left(1 - \frac{m}{L}\right) (f(x_k) - f^*).$$

□

Résumé des taux de convergence

Hypothèse	Taux	Complexité pour ε
Convexe, L -lisse	$O(1/k)$	$O(L/\varepsilon)$
m -fort. convexe, L -lisse	$O((1 - m/L)^k)$	$O(\kappa \log(1/\varepsilon))$
Convexe, non lisse	$O(1/\sqrt{k})$	$O(1/\varepsilon^2)$

7.5 Recherche de pas (line search)

Définition 7.3 (Condition d'Armijo (backtracking)). Choisir α_k comme le plus grand $\alpha = \beta^j \hat{\alpha}$ ($j = 0, 1, \dots$) satisfaisant :

$$f(x_k - \alpha \nabla f(x_k)) \leq f(x_k) - c\alpha \|\nabla f(x_k)\|^2,$$

avec $c \in (0, 1/2)$ et $\beta \in (0, 1)$.

Proposition 7.4. La descente de gradient avec backtracking d'Armijo conserve les mêmes taux de convergence $O(1/k)$ et $O((1 - m/L)^k)$ (avec des constantes différentes).

7.6 Gradient accéléré de Nesterov

Gradient accéléré de Nesterov (NAG)

1. Initialiser $x_0 = y_0 \in \mathbb{R}^n$.
2. Pour $k = 0, 1, 2, \dots$:

$$\begin{aligned} x_{k+1} &= y_k - \frac{1}{L} \nabla f(y_k), \\ y_{k+1} &= x_{k+1} + \frac{k}{k+3} (x_{k+1} - x_k). \end{aligned}$$

Théorème 7.5 (Convergence de Nesterov). Soit f convexe et L -lisse. L'algorithme de Nesterov satisfait :

$$f(x_k) - f^* \leq \frac{2L \|x_0 - x^*\|^2}{(k+1)^2}.$$

Le taux $O(1/k^2)$ est optimal parmi les méthodes du premier ordre (borne inférieure de Nemirovski–Yudin).

Optimalité du taux $O(1/k^2)$

Nemirovski et Yudin (1983) ont montré qu'aucune méthode du premier ordre (n'utilisant que les valeurs de f et ∇f) ne peut converger plus vite que $O(1/k^2)$ pour la classe des fonctions convexes L -lisses. L'accélération de Nesterov atteint cette borne : elle est *optimale*.

7.7 Gradient conjugué

Gradient conjugué non linéaire (Fletcher–Reeves)

1. $d_0 = -\nabla f(x_0)$.
2. Pour $k = 0, 1, 2, \dots$:
 - (a) $\alpha_k = \arg \min_{\alpha > 0} f(x_k + \alpha d_k)$.
 - (b) $x_{k+1} = x_k + \alpha_k d_k$.
 - (c) $\beta_{k+1} = \frac{\|\nabla f(x_{k+1})\|^2}{\|\nabla f(x_k)\|^2}$.
 - (d) $d_{k+1} = -\nabla f(x_{k+1}) + \beta_{k+1} d_k$.

Remarque 7.6. Pour une fonction quadratique $f(x) = \frac{1}{2}x^T A x - b^T x$ avec $A \succ 0$, le gradient conjugué converge en au plus n itérations (convergence finie).

7.8 Gradient stochastique (SGD)

Descente de gradient stochastique

Pour $f(x) = \frac{1}{N} \sum_{i=1}^N f_i(x)$:

1. Pour $k = 0, 1, 2, \dots$:
 - (a) Tirer i_k uniformément dans $\{1, \dots, N\}$.
 - (b) $x_{k+1} = x_k - \alpha_k \nabla f_{i_k}(x_k)$.

Théorème 7.7 (Convergence du SGD). *Pour f convexe, L -lisse, avec $\mathbb{E}[\|\nabla f_i(x) - \nabla f(x)\|^2] \leq \sigma^2$ et $\alpha_k = c/\sqrt{k}$:*

$$\mathbb{E}[f(\bar{x}_k)] - f^* = O\left(\frac{1}{\sqrt{k}}\right),$$

où $\bar{x}_k = \frac{1}{k} \sum_{i=1}^k x_i$. Pour f m -fortement convexe avec $\alpha_k = 1/(mk)$:

$$\mathbb{E}[f(\bar{x}_k)] - f^* = O\left(\frac{1}{k}\right).$$

7.9 Variantes modernes du SGD

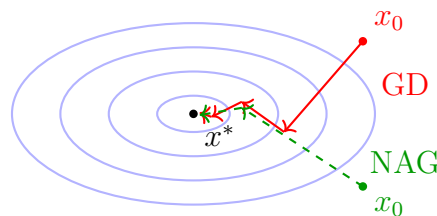
Adam (Kingma & Ba, 2015)

1. Paramètres : $\beta_1 = 0.9$, $\beta_2 = 0.999$, $\varepsilon = 10^{-8}$, α .
2. $m_0 = 0$, $v_0 = 0$.

3. Pour $k = 1, 2, \dots$:

$$\begin{aligned} g_k &= \nabla f_{i_k}(x_k), \\ m_k &= \beta_1 m_{k-1} + (1 - \beta_1) g_k, \\ v_k &= \beta_2 v_{k-1} + (1 - \beta_2) g_k^2, \\ \hat{m}_k &= m_k / (1 - \beta_1^k), \\ \hat{v}_k &= v_k / (1 - \beta_2^k), \\ x_{k+1} &= x_k - \alpha \hat{m}_k / (\sqrt{\hat{v}_k} + \varepsilon). \end{aligned}$$

7.10 Interprétation géométrique



7.11 Implémentation Python

Comparaison GD, Nesterov et SGD

```
import numpy as np
import matplotlib.pyplot as plt

def gradient_descent(f, grad_f, x0, L, n_iter):
    x = x0.copy()
    history = [f(x)]
    for k in range(n_iter):
        x = x - (1/L) * grad_f(x)
        history.append(f(x))
    return x, history

def nesterov_accelerated(f, grad_f, x0, L, n_iter):
    x = x0.copy()
    y = x0.copy()
    history = [f(x)]
    for k in range(n_iter):
        x_new = y - (1/L) * grad_f(y)
        y = x_new + (k / (k + 3)) * (x_new - x)
        x = x_new
        history.append(f(x))
    return x, history

# Quadratique mal conditionnee
```

```

np.random.seed(42)
n = 50
eigvals = np.logspace(-2, 2, n) # kappa = 10000
Q_diag = eigvals
b = np.random.randn(n)

f = lambda x: 0.5 * np.sum(Q_diag * x**2) - np.sum(b * x)
grad_f = lambda x: Q_diag * x - b
L = np.max(Q_diag)
f_star = -0.5 * np.sum(b**2 / Q_diag)

x0 = np.zeros(n)
_, hist_gd = gradient_descent(f, grad_f, x0, L, 500)
_, hist_nag = nesterov_accelerated(f, grad_f, x0, L, 500)

plt.semilogy(np.array(hist_gd) - f_star, label='GD')
plt.semilogy(np.array(hist_nag) - f_star, label='Nesterov')
plt.xlabel('Iteration k')
plt.ylabel('$f(x_k) - f^*$')
plt.legend()
plt.title('GD vs Nesterov (kappa=10000)')
plt.grid(True, alpha=0.3)
plt.savefig('gd_vs_nesterov.pdf')

```

7.12 Exercices

Exercice 7.1 (*). Montrer que pour $f(x) = \frac{1}{2} \|Ax - b\|^2$, le gradient est $\nabla f(x) = A^T(Ax - b)$ et la constante de lissité est $L = \|A^T A\|$.

Exercice 7.2 (**). Démontrer l'inégalité de Polyak–Łojasiewicz : si f est m -fortement convexe, alors $\|\nabla f(x)\|^2 \geq 2m(f(x) - f^*)$.

Exercice 7.3 (**). Implémenter le backtracking d'Armijo et comparer avec le pas fixe $1/L$ sur un problème de régression logistique.

Exercice 7.4 (**). Montrer que le gradient conjugué converge en au plus n itérations pour une quadratique en dimension n .

Exercice 7.5 (***). Démontrer la borne de convergence $O(1/k^2)$ du gradient accéléré de Nesterov en utilisant une fonction de Lyapunov.

Exercice 7.6 (***). Montrer la borne inférieure de Nemirovski–Yudin : il existe une fonction convexe L -lisse telle que toute méthode du premier ordre satisfait $f(x_k) - f^* \geq \frac{cL\|x_0 - x^*\|^2}{k^2}$ pour $k \leq (n - 1)/2$.

Chapitre 8

Méthodes Proximales

8.1 Introduction

Dans les années 1960, Jean-Jacques Moreau, mathématicien à Montpellier, introduit l'*opérateur proximal* : pour une fonction convexe f , le proximal de x est le point qui minimise f plus un terme de pénalité quadratique autour de x . L'idée semble anodine, mais elle révèle une puissance considérable : là où la descente de gradient exige la différentiabilité, l'opérateur proximal fonctionne pour toute fonction convexe semi-continue inférieurement, même non lisse. Cette observation a ouvert la voie aux *méthodes proximales*, qui dominent aujourd'hui l'optimisation appliquée : ISTA et FISTA pour le LASSO en statistique, ADMM pour l'optimisation distribuée, l'algorithme de Douglas–Rachford pour la faisabilité. Chacun exploite la décomposition d'un problème complexe en sous-problèmes proximaux simples.

Intuition

L'opérateur proximal est une généralisation de la projection sur un ensemble convexe. Là où la descente de gradient fait un pas puis projette, l'opérateur proximal combine les deux en une seule étape : il cherche un point proche de l'argument qui minimise aussi la fonction. C'est un compromis entre rester près du point courant et diminuer la fonction objectif.

8.2 Opérateur proximal

Définition 8.1 (Opérateur proximal). Soit $g : \mathbb{R}^n \rightarrow \mathbb{R} \cup \{+\infty\}$ convexe, s.c.i., propre. L'**opérateur proximal** de g est :

$$\text{prox}_g(v) = \arg \min_{x \in \mathbb{R}^n} \left\{ g(x) + \frac{1}{2} \|x - v\|^2 \right\}.$$

Proposition 8.2 (Existence et unicité). Pour toute fonction g convexe, s.c.i., propre, $\text{prox}_g(v)$ existe et est unique pour tout $v \in \mathbb{R}^n$.

Théorème 8.3 (Caractérisation). $x^* = \text{prox}_g(v)$ si et seulement si :

$$v - x^* \in \partial g(x^*),$$

où ∂g désigne le sous-différentiel de g .

Démonstration. La condition d'optimalité de $\min_x g(x) + \frac{1}{2} \|x - v\|^2$ s'écrit : $0 \in \partial g(x^*) + (x^* - v)$, soit $v - x^* \in \partial g(x^*)$. $\square$

Opérateurs proximaux classiques

- $g = 0$: $\text{prox}_g(v) = v$.
- $g = \iota_C$ (indicatrice d'un convexe C) : $\text{prox}_g(v) = \Pi_C(v)$ (projection).
- $g = \lambda \|\cdot\|_1$ (LASSO) : $[\text{prox}_g(v)]_i = \text{sign}(v_i) \max(|v_i| - \lambda, 0)$ (seuillage doux).
- $g = \frac{\lambda}{2} \|\cdot\|^2$: $\text{prox}_g(v) = \frac{v}{1+\lambda}$.
- $g = \lambda \|\cdot\|_2$ (group LASSO) : $\text{prox}_g(v) = v \max\left(1 - \frac{\lambda}{\|v\|}, 0\right)$.

Définition 8.4 (Résolvante de Moreau). L'identité de **Moreau** relie l'opérateur proximal de g à celui de sa conjuguée g^* :

$$\text{prox}_g(v) + \text{prox}_{g^*}(v) = v.$$

8.3 Méthode de gradient proximal

Définition 8.5 (Problème composite). On considère le problème :

$$\min_{x \in \mathbb{R}^n} F(x) = f(x) + g(x),$$

où f est convexe différentiable (∇f est L -lipschitzienne) et g est convexe, s.c.i., propre (potentiellement non différentiable).

Gradient proximal (ISTA)

1. Choisir x_0 , pas $\gamma \in (0, 1/L]$.

2. Pour $k = 0, 1, 2, \dots$:

$$x_{k+1} = \text{prox}_{\gamma g}(x_k - \gamma \nabla f(x_k)).$$

Théorème 8.6 (Convergence d'ISTA). Avec $\gamma = 1/L$, l'algorithme ISTA vérifie :

$$F(x_k) - F(x^*) \leq \frac{L \|x_0 - x^*\|^2}{2k}.$$

La convergence est en $\mathcal{O}(1/k)$.

8.4 Accélération de Nesterov : FISTA

FISTA (Fast ISTA)

1. Choisir $x_0 = y_0$, $t_0 = 1$, $\gamma = 1/L$.
2. Pour $k = 0, 1, 2, \dots$:

$$\begin{aligned} x_{k+1} &= \text{prox}_{\gamma g}(y_k - \gamma \nabla f(y_k)), \\ t_{k+1} &= \frac{1 + \sqrt{1 + 4t_k^2}}{2}, \\ y_{k+1} &= x_{k+1} + \frac{t_k - 1}{t_{k+1}}(x_{k+1} - x_k). \end{aligned}$$

Théorème 8.7 (Convergence de FISTA). *L'algorithme FISTA vérifie :*

$$F(x_k) - F(x^*) \leq \frac{2L \|x_0 - x^*\|^2}{(k+1)^2}.$$

La convergence est en $\mathcal{O}(1/k^2)$, ce qui est **optimal** parmi les méthodes de premier ordre.

FISTA n'est pas monotone

Contrairement à ISTA, la suite $F(x_k)$ générée par FISTA n'est pas nécessairement décroissante. L'extrapolation peut temporairement augmenter la valeur de l'objectif.

8.5 ADMM

Définition 8.8 (ADMM). L'**Alternating Direction Method of Multipliers** résout :

$$\min_{x,z} f(x) + g(z) \quad \text{s.c.} \quad Ax + Bz = c.$$

Les itérations sont :

$$\begin{aligned} x_{k+1} &= \arg \min_x \left\{ f(x) + \frac{\rho}{2} \|Ax + Bz_k - c + u_k\|^2 \right\}, \\ z_{k+1} &= \arg \min_z \left\{ g(z) + \frac{\rho}{2} \|Ax_{k+1} + Bz - c + u_k\|^2 \right\}, \\ u_{k+1} &= u_k + Ax_{k+1} + Bz_{k+1} - c, \end{aligned}$$

où u est la variable duale échelonnée et $\rho > 0$.

Théorème 8.9 (Convergence d'ADMM). *Sous des hypothèses de convexité de f et g (propres, fermées, convexes) et la faisabilité du problème, ADMM converge :*

- **Résidu primal** : $Ax_k + Bz_k - c \rightarrow 0$.
- **Résidu dual** : $\rho A^\top B(z_k - z_{k-1}) \rightarrow 0$.
- **Objectif** : $f(x_k) + g(z_k) \rightarrow p^*$.

Remarque 8.10. ADMM est particulièrement efficace lorsque les sous-problèmes en x et z admettent des solutions explicites (par exemple, moindres carrés + seuillage doux pour le LASSO).

8.6 Séparation de Douglas-Rachford

Définition 8.11 (Douglas-Rachford). Pour résoudre $\min_x f(x) + g(x)$, l'algorithme de Douglas-Rachford itère :

$$\begin{aligned}\hat{x}_k &= \text{prox}_f(z_k), \\ z_{k+1} &= z_k + \text{prox}_g(2\hat{x}_k - z_k) - \hat{x}_k.\end{aligned}$$

Proposition 8.12 (Lien avec ADMM). ADMM appliqué à $\min f(x) + g(z)$ s.c. $x = z$ est équivalent à Douglas-Rachford appliqué à la formulation duale.

8.7 Applications

Exemple 8.13 (LASSO – regression parcimonieuse). Le problème LASSO :

$$\min_x \frac{1}{2} \|Ax - b\|^2 + \lambda \|x\|_1$$

est un problème composite avec $f(x) = \frac{1}{2} \|Ax - b\|^2$ (lisse, $L = \|A^\top A\|$) et $g(x) = \lambda \|x\|_1$. ISTA donne :

$$x_{k+1} = S_{\lambda/L} \left(x_k - \frac{1}{L} A^\top (Ax_k - b) \right),$$

où S_τ est le seuillage doux composante par composante.

Exemple 8.14 (Débruitage par variation totale). Pour le débruitage d'un signal $y \in \mathbb{R}^n$:

$$\min_x \frac{1}{2} \|x - y\|^2 + \lambda \|Dx\|_1,$$

où D est la matrice de différences finies. ADMM décompose ce problème en un sous-problème quadratique et un seuillage doux.

8.8 Exercices

Exercice 8.1 ($\star$ – Opérateurs proximaux). Calculer l'opérateur proximal de $g(x) = \lambda \|x\|_1$ et de $g(x) = \iota_{[0, +\infty)^n}(x)$ (indicatrice de l'orthant positif).

Exercice 8.2 ($\star\star$ – ISTA pour le LASSO). Implémenter ISTA pour le LASSO avec $A \in \mathbb{R}^{50 \times 200}$, $x^* \in \mathbb{R}^{200}$ creux (10 composantes non nulles), et $b = Ax^* + \epsilon$. Tracer la convergence et comparer avec FISTA.

Exercice 8.3 ($\star\star$ – Identité de Moreau). Démontrer l'identité de Moreau : $\text{prox}_g(v) + \text{prox}_{g^*}(v) = v$. En déduire $\text{prox}_{\lambda \|\cdot\|_\infty}$.

Exercice 8.4 ($\star\star\star$ – ADMM pour le LASSO). Dériver les itérations ADMM pour le problème LASSO. Montrer que le sous-problème en x est un système linéaire et celui en z un seuillage doux. Implémenter et comparer avec FISTA.

Exercice 8.5 ($\star\star\star$ – Convergence accélérée). Montrer que la suite t_k de FISTA vérifie $t_k \geq (k+1)/2$ et en déduire le taux de convergence $\mathcal{O}(1/k^2)$.

Chapitre 9

Méthodes de Points Intérieurs

9.1 Introduction

En 1984, Narendra Karmarkar, ingénieur chez AT&T Bell Labs, publie un algorithme qui secoue le monde de l'optimisation : une méthode de programmation linéaire en temps polynomial qui, contrairement au simplexe de Dantzig, ne parcourt pas les sommets du polyèdre mais traverse son *intérieur*. La communauté est électrisée — et divisée. Certains voient une révolution, d'autres un coup médiatique (le simplexe, bien que de complexité exponentielle dans le pire cas, est rapide en pratique). Mais l'histoire tranchera : les méthodes de points intérieurs, raffinées par Nesterov et Nemirovski dans les années 1990, s'avèrent décisives pour l'optimisation semi-définie, l'optimisation conique, et les problèmes de grande taille. L'idée est élégante : remplacer les contraintes d'inégalité par une barrière logarithmique, puis suivre le *chemin central* vers l'optimum.

Intuition

Imaginez que vous êtes dans une pièce (le polyèdre des contraintes) et que vous cherchez le point le plus bas (l'optimum). Le simplexe longe les murs et les coins. Les méthodes de points intérieurs, elles, marchent au milieu de la pièce en se rapprochant progressivement de la solution optimale, qui se trouve souvent près d'un mur. Une barrière invisible vous empêche de toucher les murs, mais cette barrière s'affaiblit au fil des itérations.

9.2 Méthode de barrière

9.2.1 Formulation

Définition 9.1 (Problème avec inégalités). On considère le problème convexe :

$$\min_{x \in \mathbb{R}^n} f_0(x) \quad \text{s.c.} \quad f_i(x) \leq 0, \quad i = 1, \dots, m, \quad Ax = b,$$

où $f_0, f_1, \dots, f_m$ sont convexes et deux fois différentiables.

Définition 9.2 (Fonction barrière logarithmique). La **fonction barrière logarithmique** est :

$$\phi(x) = - \sum_{i=1}^m \ln(-f_i(x)),$$

définie sur l'intérieur strict du domaine réalisable $\{x : f_i(x) < 0, \forall i\}$.

Définition 9.3 (Problème barrière). Le **problème barrière** pour un paramètre $t > 0$ est :

$$\min_x t f_0(x) + \phi(x) \quad \text{s.c.} \quad Ax = b.$$

On note $x^*(t)$ la solution, appelée **point central** pour le paramètre t .

9.2.2 Chemin central

Définition 9.4 (Chemin central). Le **chemin central** est l'ensemble des points $\{x^*(t) : t > 0\}$. Quand $t \rightarrow +\infty$, $x^*(t) \rightarrow x^*$, la solution du problème original.

Théorème 9.5 (Borne de sous-optimalité). Pour tout $t > 0$, le point central $x^*(t)$ vérifie :

$$f_0(x^*(t)) - p^* \leq \frac{m}{t},$$

où p^* est la valeur optimale et m le nombre de contraintes d'inégalité.

Démonstration. Les conditions KKT du problème barrière montrent que $x^*(t)$ est réalisable pour le problème original, avec des multiplicateurs $\lambda_i^*(t) = -1/(t f_i(x^*(t)))$. Par la dualité faible :

$$f_0(x^*(t)) - p^* \leq \sum_{i=1}^m \lambda_i^*(t) \cdot (-f_i(x^*(t))) = \sum_{i=1}^m \frac{1}{t} = \frac{m}{t}.$$

□

Méthode de barrière

1. Choisir x_0 strictement réalisable, $t_0 > 0$, facteur $\mu > 1$, tolérance $\varepsilon > 0$.
2. Pour $k = 0, 1, 2, \dots$:
 - (a) **Centrage** : résoudre (par Newton) $\min_x t_k f_0(x) + \phi(x)$ s.c. $Ax = b$, en partant de x_k , pour obtenir $x_{k+1} = x^*(t_k)$.
 - (b) **Test d'arrêt** : si $m/t_k < \varepsilon$, **stop**.
 - (c) **Mise à jour** : $t_{k+1} = \mu \cdot t_k$.

Complexité de la méthode de barrière

- Nombre d'itérations extérieures (pas de barrière) : $\lceil \frac{\ln(m/(t_0 \varepsilon))}{\ln \mu} \rceil$.
- Choix typique : $\mu = 10$ à 20 , ce qui donne peu d'itérations extérieures.
- Chaque étape de centrage nécessite quelques pas de Newton (6–40 en pratique).

9.3 Méthodes à pas court et à pas long

Définition 9.6 (Méthode à pas court). La méthode à **pas court** choisit μ proche de 1 (par exemple $\mu = 1 + 1/\sqrt{m}$) et effectue un seul pas de Newton par étape de centrage. Le nombre d'itérations est $\mathcal{O}(\sqrt{m} \ln(m/\varepsilon))$.

Définition 9.7 (Méthode à pas long). La méthode à **pas long** choisit μ grand et effectue plusieurs pas de Newton pour recentrer. Le nombre total de pas de Newton est aussi $\mathcal{O}(\sqrt{m} \ln(m/\varepsilon))$ mais avec une constante plus grande.

Théorème 9.8 (Complexité polynomiale). *Les méthodes de points intérieurs résolvent un programme linéaire à n variables et m contraintes en $\mathcal{O}(\sqrt{m} \ln(1/\varepsilon))$ pas de Newton, chacun coûtant $\mathcal{O}(n^3)$ (ou moins avec l'exploitation de la structure creuse). La complexité totale est donc polynomiale.*

9.4 Méthode primale-duale

Définition 9.9 (Système KKT modifié). La méthode **primale-duale** résout simultanément les conditions KKT perturbées :

$$\begin{cases} \nabla f_0(x) + \sum_i \lambda_i \nabla f_i(x) + A^\top \nu = 0, \\ -\lambda_i f_i(x) = 1/t, \quad i = 1, \dots, m, \\ Ax = b. \end{cases}$$

On applique la méthode de Newton à ce système.

Remarque 9.10. La méthode primale-duale est plus efficace en pratique que la méthode de barrière pure car elle ne nécessite pas de résoudre exactement le problème de centrage à chaque étape. Les solveurs modernes (MOSEK, CVXOPT) utilisent cette approche.

Théorème 9.11 (Convergence primale-duale). *La méthode primale-duale atteint une précision ε en $\mathcal{O}(\sqrt{m} \ln(1/\varepsilon))$ itérations de Newton, sous des hypothèses standards de régularité.*

9.5 Application à la programmation linéaire

Exemple 9.12 (PL sous forme standard). Pour le programme linéaire $\min c^\top x$ s.c. $Ax = b, x \geq 0$, la fonction barrière est :

$$\phi(x) = -\sum_{j=1}^n \ln(x_j).$$

Le système de Newton pour le centrage s'écrit :

$$\begin{pmatrix} X^{-2} & A^\top \\ A & 0 \end{pmatrix} \begin{pmatrix} \Delta x \\ \Delta \nu \end{pmatrix} = - \begin{pmatrix} tc - X^{-1}e \\ 0 \end{pmatrix},$$

où $X = \text{diag}(x)$ et $e = (1, \dots, 1)^\top$.

9.6 Exercices

Exercice 9.1 (★ – Barrière logarithmique). Pour le PL $\min -x_1 - x_2$ s.c. $x_1 + x_2 \leq 1$, $x_1, x_2 \geq 0$, écrire la fonction barrière et calculer $x^*(t)$ pour $t = 1, 10, 100$. Vérifier que $x^*(t) \rightarrow x^*$.

Exercice 9.2 (★★ – Chemin central). Tracer le chemin central pour le problème $\min x_1$ s.c. $x_1^2 + x_2^2 \leq 1$ pour $t \in \{0.1, 1, 10, 100\}$. Montrer qu'il converge vers $(-1, 0)$.

Exercice 9.3 (★★ – Pas de Newton). Dériver le système de Newton pour la méthode de barrière appliquée au problème quadratique contraint : $\min \frac{1}{2}x^\top Qx + c^\top x$ s.c. $Ax \leq b$.

Exercice 9.4 (★★★ – Borne de sous-optimalité). Démontrer en détail que $f_0(x^*(t)) - p^* \leq m/t$. En déduire le nombre d'itérations extérieures nécessaires pour atteindre une précision ε avec facteur μ .

Exercice 9.5 (★★★ – Primale-duale). Implémenter la méthode primale-duale de points intérieurs pour le PL sous forme standard. Tester sur un problème à 50 variables et 20 contraintes. Comparer le nombre d'itérations avec la méthode de barrière classique.

Chapitre 10

Applications

10.1 Introduction

L'optimisation convexe n'est pas un exercice théorique : c'est un langage universel pour formuler et résoudre des problèmes concrets. Le LASSO en statistique, la décomposition en composantes principales robuste, l'allocation de portefeuille de Markowitz, le compressed sensing en traitement du signal, l'entraînement des machines à vecteurs de support — tous se formulent comme des problèmes d'optimisation convexe et se résolvent par les algorithmes des chapitres précédents. Stephen Boyd, dans son célèbre manuel de 2004, résume l'idée : « reconnaître qu'un problème est convexe, c'est le résoudre. » Ce chapitre met cette philosophie en pratique.

Intuition

L'optimisation convexe n'est pas une théorie abstraite : c'est un langage universel pour formuler et résoudre efficacement des problèmes dans de nombreux domaines. La clé est de *reconnaître* la structure convexe cachée dans un problème appliqué, puis d'appliquer l'algorithme adapté.

10.2 Compressed sensing

Définition 10.1 (Compressed sensing). Le **compressed sensing** (acquisition comprimée) vise à reconstruire un signal $x^* \in \mathbb{R}^n$ creux (*sparse*) à partir de $m \ll n$ mesures linéaires $b = Ax^* + \epsilon$, où $A \in \mathbb{R}^{m \times n}$.

Théorème 10.2 (Reconstruction par ℓ_1). Si x^* est s -creux et A satisfait la **propriété d'isométrie restreinte** (RIP) d'ordre $2s$ avec constante $\delta_{2s} < \sqrt{2} - 1$, alors la solution de :

$$\min_x \|x\|_1 \quad \text{s.c.} \quad \|Ax - b\|_2 \leq \epsilon$$

vérifie $\|\hat{x} - x^*\|_2 \leq C\epsilon$ pour une constante C .

Remarque 10.3. En pratique, on résout souvent la formulation LASSO équivalente :

$$\min_x \frac{1}{2} \|Ax - b\|_2^2 + \lambda \|x\|_1,$$

qui est un problème composite résoluble par ISTA/FISTA (chapitre 8).

Exemple 10.4 (Reconstruction d'un signal creux). Soit $x^* \in \mathbb{R}^{500}$ avec 20 composantes non nulles, $A \in \mathbb{R}^{100 \times 500}$ gaussienne, et $b = Ax^*$. Le LASSO avec $\lambda = 0.1$ reconstruit exactement x^* , alors que les moindres carrés (sous-déterminés) échouent.

10.3 LASSO et régularisation ℓ_1

Définition 10.5 (LASSO). Le **LASSO** (Least Absolute Shrinkage and Selection Operator) :

$$\min_{\beta \in \mathbb{R}^p} \frac{1}{2n} \|y - X\beta\|^2 + \lambda \|\beta\|_1,$$

effectue simultanément l'estimation et la sélection de variables en régression linéaire.

Proposition 10.6 (Propriétés du LASSO). 1. La pénalisation ℓ_1 produit des solutions **creuses** (beaucoup de coefficients exactement nuls).

2. Le paramètre λ contrôle le degré de parcimonie : $\lambda \rightarrow +\infty$ donne $\hat{\beta} = 0$; $\lambda \rightarrow 0$ donne les moindres carrés.

3. Le chemin de régularisation $\lambda \mapsto \hat{\beta}(\lambda)$ est linéaire par morceaux.

Théorème 10.7 (Conditions d'optimalité du LASSO). $\hat{\beta}$ est solution du LASSO si et seulement si :

$$\frac{1}{n} X^\top (X\hat{\beta} - y) + \lambda v = 0,$$

où $v \in \partial \|\hat{\beta}\|_1$, c'est-à-dire $v_j = \text{sign}(\hat{\beta}_j)$ si $\hat{\beta}_j \neq 0$ et $|v_j| \leq 1$ sinon.

10.4 Optimisation de portefeuille

Définition 10.8 (Modèle de Markowitz). Le **modèle de Markowitz** cherche le portefeuille de variance minimale pour un rendement espéré donné :

$$\min_w w^\top \Sigma w \quad \text{s.c.} \quad \mu^\top w \geq r_{\min}, \quad \mathbf{1}^\top w = 1, \quad w \geq 0,$$

où $w \in \mathbb{R}^n$ est le vecteur de poids, Σ la matrice de covariance, μ le vecteur de rendements espérés.

Remarque 10.9. C'est un **programme quadratique** convexe (QP) car Σ est semi-définie positive. Il se résout par les méthodes de points intérieurs ou par le simplexe pour QP.

Proposition 10.10 (Frontière efficiente). L'ensemble des portefeuilles optimaux (en faisant varier $r_{\min}$) forme la **frontière efficiente** dans le plan (écart-type, rendement). C'est une courbe convexe.

Exemple 10.11 (Portefeuille à 3 actifs). Avec $\mu = (0.12, 0.10, 0.07)^\top$ et

$$\Sigma = \begin{pmatrix} 0.04 & 0.006 & 0.002 \\ 0.006 & 0.025 & 0.004 \\ 0.002 & 0.004 & 0.01 \end{pmatrix},$$

le portefeuille de variance minimale (sans contrainte de rendement) est $w^* \approx (0.14, 0.28, 0.58)^\top$ avec $\sigma^* \approx 8.1\%$.

10.5 Transport optimal

Définition 10.12 (Problème de Kantorovich). Étant données deux mesures de probabilité discrètes $\mu = \sum_{i=1}^n a_i \delta_{x_i}$ et $\nu = \sum_{j=1}^m b_j \delta_{y_j}$, le **transport optimal** de Kantorovich cherche :

$$\min_{\Pi \geq 0} \sum_{i,j} C_{ij} \Pi_{ij} \quad \text{s.c.} \quad \Pi \mathbf{1} = a, \quad \Pi^\top \mathbf{1} = b,$$

où $C_{ij} = c(x_i, y_j)$ est le coût de transport.

Remarque 10.13. C'est un programme linéaire. La distance de Wasserstein associée définit une métrique sur l'espace des mesures de probabilité.

Définition 10.14 (Régularisation entropique). La régularisation de **Sinkhorn** ajoute un terme entropique :

$$\min_{\Pi \geq 0} \sum_{i,j} C_{ij} \Pi_{ij} + \varepsilon \sum_{i,j} \Pi_{ij} \ln \Pi_{ij} \quad \text{s.c.} \quad \Pi \mathbf{1} = a, \quad \Pi^\top \mathbf{1} = b.$$

L'algorithme de Sinkhorn résout ce problème par des multiplications matricielles alternées. Pour des mesures discrètes à n points, la complexité par itération est $\mathcal{O}(n^2)$, et le nombre d'itérations pour atteindre une précision ε est $\mathcal{O}(1/\varepsilon)$ (Altschuler et al., 2017).

10.6 Applications en apprentissage automatique

10.6.1 SVM dual

Définition 10.15 (SVM – formulation duale). Le **Support Vector Machine** (SVM) linéaire se formule en dual :

$$\max_{\alpha} \sum_{i=1}^n \alpha_i - \frac{1}{2} \sum_{i,j} \alpha_i \alpha_j y_i y_j x_i^\top x_j \quad \text{s.c.} \quad 0 \leq \alpha_i \leq C, \quad \sum_i \alpha_i y_i = 0,$$

où $(x_i, y_i) \in \mathbb{R}^d \times \{-1, +1\}$ sont les données et $C > 0$. C'est un QP convexe.

Proposition 10.16 (Vecteurs de support). Les **vecteurs de support** sont les points x_i pour lesquels $\alpha_i^* > 0$. L'hyperplan séparateur est $w^* = \sum_i \alpha_i^* y_i x_i$.

10.6.2 Régression logistique

Définition 10.17 (Régression logistique régularisée). La régression logistique avec régularisation ℓ_2 :

$$\min_{\beta} \sum_{i=1}^n \ln(1 + e^{-y_i x_i^\top \beta}) + \frac{\lambda}{2} \|\beta\|^2.$$

La fonction objectif est fortement convexe ($\lambda > 0$) et lisse. Elle se résout efficacement par Newton, L-BFGS, ou gradient proximal (avec régularisation ℓ_1 pour la parcimonie).

Théorème 10.18 (Convexité de la log-vraisemblance). La fonction de perte logistique $\ell(\beta) = \ln(1 + e^{-y x^\top \beta})$ est convexe en β pour tout (x, y) . Avec la régularisation ℓ_2 , le problème est λ -fortement convexe, garantissant l'unicité de la solution.

10.7 Implémentation Python

LASSO et portefeuille avec CVXPY

```

import numpy as np
import cvxpy as cp

# --- LASSO ---
n, p = 100, 200
X = np.random.randn(n, p)
beta_true = np.zeros(p)
beta_true[:10] = np.random.randn(10)
y = X @ beta_true + 0.1 * np.random.randn(n)

beta = cp.Variable(p)
lam = 0.1
prob = cp.Problem(cp.Minimize(
    0.5 * cp.sum_squares(X @ beta - y)
    + lam * cp.norm1(beta)))
prob.solve()
print("LASSO: nnz =", np.sum(np.abs(beta.value) > 1e-4))

# --- Portefeuille de Markowitz ---
mu = np.array([0.12, 0.10, 0.07])
Sigma = np.array([[0.04, 0.006, 0.002],
                  [0.006, 0.025, 0.004],
                  [0.002, 0.004, 0.01]])

w = cp.Variable(3)
ret = mu @ w
risk = cp.quad_form(w, Sigma)
prob2 = cp.Problem(cp.Minimize(risk),
    [ret >= 0.09, cp.sum(w) == 1, w >= 0])
prob2.solve()
print("Portefeuille:", np.round(w.value, 3))

```

10.8 Exercices

Exercice 10.1 (★ – LASSO). Pour $X \in \mathbb{R}^{20 \times 50}$ et $y = X\beta^* + \epsilon$ avec β^* 5-creux, résoudre le LASSO pour $\lambda \in \{0.01, 0.1, 1\}$. Tracer le nombre de composantes non nulles en fonction de λ .

Exercice 10.2 (★★ – Portefeuille). Résoudre le problème de Markowitz pour 5 actifs avec données historiques. Tracer la frontière efficiente pour $r_{\min}$ variant de 5% à 15%.

Exercice 10.3 (★★ – SVM). Formuler et résoudre le SVM dual pour un jeu de données linéairement séparable à 2 dimensions. Identifier les vecteurs de support.

Exercice 10.4 (★★★ – Transport optimal). Implémenter l'algorithme de Sinkhorn pour deux distributions discrètes sur $\mathbb{R}^2$ (à 50 points chacune). Tracer le plan de transport optimal et étudier l'influence de ε .

Exercice 10.5 (★★★ – Compressed sensing). Générer $A \in \mathbb{R}^{m \times n}$ gaussienne, x^* s -creux, et $b = Ax^*$. Étudier expérimentalement la transition de phase : pour $n = 200$, faire varier m et s et déterminer quand la reconstruction ℓ_1 réussit.

Exercice 10.6 ($\star\star$ – Propriétés de la constante RIP). Soit $A \in \mathbb{R}^{m \times n}$ et δ_s la constante d'isométrie restreinte (RIP) d'ordre s , définie comme le plus petit $\delta \geq 0$ tel que :

$$(1 - \delta) \|x\|_2^2 \leq \|Ax\|_2^2 \leq (1 + \delta) \|x\|_2^2$$

pour tout vecteur s -creux x .

1. Montrer que $\delta_s \leq \delta_{s'}$ pour $s \leq s'$ (monotonie).
2. Montrer que si $\delta_{2s} < 1$, alors A est injective sur l'ensemble des vecteurs s -creux.
3. Soit A une matrice gaussienne $m \times n$ avec entrées i.i.d. $\mathcal{N}(0, 1/m)$. En utilisant une borne de concentration sur la norme d'un sous-ensemble de colonnes, montrer que $\delta_s < \delta$ avec haute probabilité dès que $m \geq C \delta^{-2} s \ln(n/s)$.

Exercice 10.7 ($\star\star\star$ – Chemin de régularisation du LASSO). Considérer le problème LASSO : $\hat{\beta}(\lambda) = \arg \min_{\beta} \frac{1}{2n} \|y - X\beta\|^2 + \lambda \|\beta\|_1$.

1. Montrer que pour $\lambda \geq \lambda_{\max} = \frac{1}{n} \|X^\top y\|_\infty$, la solution est $\hat{\beta}(\lambda) = 0$.
2. En utilisant les conditions KKT, montrer que l'ensemble actif $\mathcal{A}(\lambda) = \{j : \hat{\beta}_j(\lambda) \neq 0\}$ est constant par morceaux en λ et que les changements ont lieu en un nombre fini de valeurs $\lambda_1 > \lambda_2 > \dots > 0$.
3. Sur chaque intervalle $(\lambda_{k+1}, \lambda_k)$, montrer que $\hat{\beta}_{\mathcal{A}}(\lambda)$ est une fonction affine de λ :

$$\hat{\beta}_{\mathcal{A}}(\lambda) = (X_{\mathcal{A}}^\top X_{\mathcal{A}})^{-1} (X_{\mathcal{A}}^\top y - n\lambda s_{\mathcal{A}})$$

où $s_{\mathcal{A}} = \text{sign}(\hat{\beta}_{\mathcal{A}})$. En déduire que le chemin $\lambda \mapsto \hat{\beta}(\lambda)$ est linéaire par morceaux.

Bibliographie

- [1] Boyd, S. and Vandenberghe, L., *Convex Optimization*, Cambridge, 2004.
- [2] Rockafellar, R.T., *Convex Analysis*, Princeton, 1970.
- [3] Bauschke, H.H. and Combettes, P.L., *Convex Analysis and Monotone Operator Theory in Hilbert Spaces*, 2nd ed., Springer, 2017.
- [4] Nesterov, Y., *Introductory Lectures on Convex Optimization*, Springer, 2004.