



N° Réf :.....

Centre Universitaire
Abd Elhafid Boussouf Mila

Institut des Sciences et Technologie

Département de Mathématiques et Informatique

Mémoire préparé en vue de l'obtention du diplôme de Master

En : Mathématiques

Spécialité : Mathématiques Appliquées

Méthodes d'optimisation non linéaire sans contraintes

Préparé par :

Lalaoui Bahidja
Zerafi Safa

Les membres de jury :

Bouballouta Khadidja (M.C.B)
Fadel Wahida (M.A.A)
Bouzekria Fahima (M.A.A)

C.U.Abd Elhafid Boussouf	Président
C.U.Abd Elhafid Boussouf	Rapporteur
C.U.Abd Elhafid Boussouf	Examineur

Année Universitaire : 2020/2021

Remerciements

*Nous remercions d'abord et avant tout le bon **Dieu** qui nous avons donné la force et la patience pour mener notre travail à terme.*

*Nous tenons d'abord à remercier infiniment notre encadreur **Madame Fadel Wahida** maitre de centre universitaire Abdelhafid Boussouf-Mila d'avoir accepté de diriger notre travail, ses conseils bienveillant, qui nous a encourager à fournir plus d'efforts pour être à la hauteur de leur attente.*

Nos plus sincères remerciements vont aux membres du jury pour l'honneur qu'ils acceptant de juger ce travail.

Encore et avec une grande fierté et honneur que nous tenons à présenter nos remerciement a tout la famille administrative du département mathématique, à tous ces enseignants qui nous ont permis d'acquérir des connaissances.

Dédicaces

J'ai le grand plaisir de dédier ce modeste travail :

*A la lumière de mes jours, la source de mes efforts, ma vie et mon âme, la prunelle de mes yeux, a la femme qui peut être fier et trouver ici le résultat de longues années de sacrifices et de privations pour m'aider à avancer dans la vie ma chère maman "**Rehifa**".*

*A femme qui m'a encouragé durant toutes mes études, et qui sans elle, ma réussite n'aura pas en lieu. Qu'elle trouve ici mon amour et mon affection, mon chère mère "**Salih**a".*

*A l'homme, mon précieux offre du dieu, qui doit ma vie, ma réussite et tout mon respect mon chère père "**Elhani**".*

*A mes frères "**Djamel**", "**Oussama**" et mon âme sœur "**Marwa**" qui m'avez toujours soutenu et encouragé durant ces années d'étude.*

*A ma chère binôme "**Bahidja Lalaoui**" qui a été pour moi la sœur que la providence ma enviée et qui a partagé avec moi le travail, d'avoir eu la patience et le courage pour achever ce travail, et à tout sa famille.*

A ma famille, mes proches et à ceux qui me donnent de l'amour et de la vivacité.

A mes amis qui m'ont toujours encouragé et à qui je souhaité plus de succès.

****Safa****

Dédicaces

Je dédie ce modeste travail qui le fruit de mon étude pendant cinq ans, premièrement aux personnes les plus chère à mon cœur :

*Ma très chère mère "**Zineb**" le symbole de tendresse et de la gentillesse.*

*Celui qui nous a protégé et nous a appris comment vivre, mon très chère père "**Salah**".*

*Mes frères "**Hamza et Abd rezak**", mes sœurs "**Sabrina et Meriem**", ses enfants, et ma familles.*

Je profite cette occasion pour dédier ce mémoire à tout ma promo, et tous mes amies.

En fin, à tout qui me donnent de l'énergie et la confiance pour réaliser toujours les meilleurs résultats.

****Bahidja****

Résumer

Soit $f : \mathbb{R}^n \rightarrow \mathbb{R}$. On cherche dans ce mémoire à résoudre le problème de minimisation sans contrainte suivant :

$$\left\{ \min_{x \in \mathbb{R}^n} f(x) \right\}.$$

Pour cela, en premier lieu, il est essentiel de présenter les conditions d'optimalité (nécessaires et suffisantes), qui sont utiles par la suite, puis nous avons introduit quelques méthodes de résolution qu'elles permettaient de déterminer un minimum pour le problème étudié :

Méthode du gradient, gradient conjugué, méthode de Newton, méthode de relaxation.

A la fin, nous nous sommes focalisés sur une étude comparative entre les méthodes étudiées.

Mot clés : optimisation sans contraintes, méthode du gradient, méthode de Newton, gradient conjugué.

Abstract

Our aim in this memory is to solve non-constraints optimization problem :

$$\left\{ \min_{x \in \mathbb{R}^n} f(x) \right\}.$$

For that, in the first place, we have given main and important definitions.

Secondly, it is necessarily necessary to present the necessary and sufficient condition of optimality, which are the main ones to solve the studied problem.

Then, we introduced some resolution methods that allowed to determine a minimum for the studied problem which are : Gradient method, Conjugate gradient, Newton's method and relaxation method.

Key words : non-constraints optimization problem, Gradient method, Conjugate gradient method, Newton's method.

ملخص

لتكن f دالة معرفة من $\mathbb{R}^n$ في $\mathbb{R}$ بدون قيود.
نتناول في هذه المذكرة مشكل المثالية التالي :

$$\left\{ \min_{x \in \mathbb{R}^n} f(x) \right\} .$$

حيث وفي البداية من الضروري عرض الشروط اللازمة و الكافية للمثالية بالنسبة للمشكل المطروح والتي تستعمل لاحقا.
بعدها قدمنا بعض الطرق التي تسمح بايجاد الحل الادنى على غرار : طريقة التدرج، طريقة التدرج المرفق، طريقة نيوتن.
وفي الاخير قدمنا دراسة مقارنة بين الطرق المدروسة.
الكلمات المفتاحية : طريقة التدرج، التدرج المرفق، طريقة نيون

Table des matières

1	<i>Généralités et outils de base</i>	7
1.1	<i>Définitions de base</i>	7
1.2	<i>Différentiabilité</i>	8
1.2.1	<i>Différentiabilité : du premier ordre</i>	8
1.2.2	<i>Différentiabilité : du second ordre</i>	13
1.3	<i>Développement de Taylor</i>	14
1.4	<i>Convexité</i>	16
1.4.1	<i>Ensemble convexe</i>	16
1.4.2	<i>Combinaison convexe</i>	17
2	<i>Minimisation sans contraintes</i>	19
2.1	<i>Résultats d'existence et d'unicité</i>	20
2.2	<i>Conditions d'optimalité</i>	23
2.2.1	<i>Conditions nécessaires d'optimalité du premier ordre</i>	23
2.2.2	<i>Conditions nécessaires d'optimalité du seconde ordre</i>	24
2.2.3	<i>Conditions suffisantes d'optimalité</i>	25

3	Méthodes de résolutions	28
3.1	Définition (algorithmes)	28
3.1.1	convergence d'un algorithme	29
3.1.2	Taux de convergence d'un algorithme	29
3.2	Méthode de gradient	30
3.3	Méthode du gradient conjugué	33
3.3.1	Le principe général d'une méthode à direction conjuguées	34
3.3.2	Méthode de gradient conjugué dans le cas qua- dratique	36
3.3.3	Méthode du gradient conjugué dans le cas non quadratique	42
3.4	Méthode de Newton	44
3.4.1	Principe de la Méthode de Newton	44
3.5	Méthode de quasi Newton ou quasi-Newtonniennes	46
3.5.1	Formules de mise à jour de l'approximation du hessien :	47
3.5.2	Méthode de correction de rang un	47
3.5.3	Méthode de Davidon Fletcher Powell (DFP)	48
3.5.4	Méthode de Broyden, Fletcher, Goldfarb et Shanno (BFGS)	50
3.5.5	Les méthodes de classe Broyden	53
3.6	Méthode de relaxation	54
3.7	Étude comparative des méthodes d'optimisation sans contraintes	54
	Bibliographie	57

Introduction

Les mathématiques pures ont pour objectif le développement des connaissances mathématiques pour elles-mêmes sans aucun intérêt a priori pour les applications, sans aucune motivation d'autre sciences. L'objet de la recherche mathématique peut ainsi être une meilleure compréhension d'une série d'exemples particuliers abstraits, lesquels s'appuie et se développé la réflexion mathématique, la généralisation d'un aspect d'une discipline ou la mise en évidence de liens entre diverses disciplines des mathématiques.

Les mathématiques appliquées sont une branche des mathématiques qui s'intéressent à l'application du savoir mathématique aux autres domaines, l'analyse numérique, les mathématique de l'ingénierie, l'optimisation linéaire et non linéaire, la théorie de l'information, ..., ainsi qu'une bonne partie de l'informatique sont autant de domaines d'application des mathématiques.

L'optimisation est une branche des mathématiques cherchant à modéliser, à analyser et à résoudre analytiquement ou numériquement les problèmes qui consistent à minimiser ou maximiser une fonction sur un ensemble.

Les problèmes que rencontrent les gestionnaires des plus hauts niveaux jusqu'aux petites entreprises, se présentent sous forme de données, de contraintes, dont on doit tenir compte et d'un objectif ou plusieurs objectifs à atteindre.

On doit commencer par interpréter tous les paramètres et les transformer sous des formes qu'on peut gérer, en cherchant des approches mathématiques pour les résoudre.

Les techniques d'optimisation sont des procédés systématique permettant d'approcher, voire de trouver la technique de fabrication la plus rapide, le cout le plus bas ou l'impôt le plus juste. La technique d'optimisation la plus simple est sans conteste celle de l'essai-erreur. Il suffit de faire varier quelques paramètres et de conserver le candidat dès que celui-ci, tout en respectant les contraintes du problème, s'avère meilleur que le modèle le plus performant actuellement disponible.

Avec un peu d'habitude, d'intuition ou de chance, une amélioration peut être obtenue. Mais s'agit-il vraiment d'un optimum ?

Partie intégrante des mathématiques appliquées, l'optimisation se veut de résoudre des problèmes scientifiques et industriels. L'optimisation est une méthode (ensemble de méthode) qui permet d'obtenir le meilleur résultat approché d'un problème de recherche de minimum d'effort à fournir pour une machine (par exemple) ou le besoin d'obtenir un bénéfice maximal dans une gestion de production ou dans la gestion d'un portefeuille. L'optimiser c'est donner les meilleures conditions de fonctionnement, de rendement" (le petit Robert).

L'optimisation joue un rôle important en recherche opérationnelle (domaine à la frontière entre l'informatique, les mathématiques et l'économie), dans les mathématiques appliquées (fondamentales pour l'industrie et l'ingénierie), en analyse et en analyse numérique, en statistique pour l'estimation du maximum de vraisemblance d'une distribution, pour la recherche de stratégies dans le cadre de la théorie des jeux, ou encore en théorie du contrôle et de la commande.

D'un point de vue mathématique, l'optimisation consiste à rechercher le minimum ou maximum d'une fonction avec ou sans contraintes.

L'optimisation possède ses racines au 18 ième siècle dans les travaux de :

** Taylor, Newton, la grange, qui ont élaboré les bases des développements limités.*

** Cauchy ([1847]) fut le premier à mettre en œuvre une méthode d'optimisation, méthode du pas de descente, pour la résolution de problèmes sans contraintes.*

Il faut attendre le milieu du vingtième siècle, avec l'émergence des calculateurs et surtout la fin de la seconde guerre mondiale pour voir apparaître des avancées spectaculaires en termes de techniques d'optimisation. Annoter, ces avancées ont été essentiellement obtenues en Grande Bretagne.

Dans ce travail, nous allons intéresser à l'étude de problèmes d'optimisation non linéaires sans contraintes pour ce faire, nous avons opté pour le plan du travail

suivant :

** Le premier chapitre commence par traiter des notions mathématiques fondamentales, on définit et on introduit les outils fonctionnels de base nécessaire pour l'optimisation sans contraintes.*

** Dans le deuxième chapitre nous nous intéressons à l'étude de quelque résultats théoriques (existence et unicité des solutions) aussi les conditions d'optimalité (nécessaires et suffisantes).*

** Et dans le dernier chapitre, nous consacrerons de quelques méthodes (algorithmes) d'optimisation les plus utilisées dans la pratique.*

Généralités et outils de base

1.1 Définitions de base

Définition 1.1.1. Soit $A \in \mathbb{R}^n$ une matrice symétrique

1. On dit que A est semi définie positive si $x^t A x \geq 0 \quad \forall x \in \mathbb{R}^n$.
2. On dit que A est définie positive si $x^t A x > 0 \quad \forall x \in \mathbb{R}^n / \{0\}$.

Définition 1.1.2. Étant donnée $x \in \mathbb{R}^n$, on appelle voisinage de x , noté $V(x)$ tout sous ensemble ouvert contenant x . De façon équivalente, $V(x)$ est un voisinage de x s'il contient une boule de centre x .

Définition 1.1.3. Un ensemble S est dit ouvert s'il coïncide avec son intérieure, c'est -à-dire si $S = \text{int}(S)$.

1.2 Différentiabilité

1.2.1 Différentiabilité : du premier ordre

Dérivée partielle

Définition 1.2.1. [3] soit $f : \mathbb{R}^n \rightarrow \mathbb{R}$ une fonction continue. La fonction notée $\nabla_i f(x) : \mathbb{R}^n \rightarrow \mathbb{R}$, également notée $\frac{\partial f(x)}{\partial x_i}$ est appelée i^{me} dérivée partielle de f et est définie par

$$\lim_{\alpha \rightarrow 0} \frac{f(x_1, \dots, x_i + \alpha, \dots, x_n) - f(x_1, \dots, x_i, \dots, x_n)}{\alpha}$$

Cette limite peut ne pas exister.

Exemple 1.2.1. Soit $f : \mathbb{R}^3 \rightarrow \mathbb{R}^2$ définie par :

$$f(x) = \begin{pmatrix} x_1^3 + x_2^4 + x_3^2 \\ 5x_1^2 + 2x_2^3 + x_3 \end{pmatrix}.$$

$$\frac{\partial f_1}{\partial x_1} = 3x_1^2, \quad \frac{\partial f_1}{\partial x_2} = 4x_2^3, \quad \frac{\partial f_1}{\partial x_3} = 2x_3$$

$$\frac{\partial f_2}{\partial x_1} = 10x_1, \quad \frac{\partial f_2}{\partial x_2} = 6x_2^2, \quad \frac{\partial f_2}{\partial x_3} = 1$$

Gradient

Si les dérivées partielles $\frac{\partial f}{\partial x_i}$ existent pour tout i , le gradient de f est défini de la façon suivante.

Définition 1.2.2. [24] [23] Soit $f : \mathbb{R}^n \rightarrow \mathbb{R}$ une fonction continue. La fonction notée $\nabla f(x) : \mathbb{R}^n \rightarrow \mathbb{R}^n$ est appelée le gradient de f et est définie par

$$\nabla f(x) = \left(\frac{\partial f(x)}{\partial x_i} \right)_{i=1, \dots, n} = \begin{pmatrix} \frac{\partial f(x)}{\partial x_1} \\ \vdots \\ \frac{\partial f(x)}{\partial x_n} \end{pmatrix}$$

Elle ne pas existe pour certains $x \in \mathbb{R}^n$. On note dans $\mathbb{R}$,

$$\nabla f(x) = f'(x).$$

Le gradient jouera un rôle essentiel dans le développement et l'analyse des algorithmes d'optimisation.

Exemple 1.2.2. Soit $f(x_1, x_2, x_3) = 5x_1^2 - 3(x_2 - x_3^2)^2 + 25(x_1 + x_3^2)$. Le gradient de f est donnée par :

$$\nabla f(x_1, x_2, x_3) = \begin{pmatrix} 10x_1 + 25 \\ -6x_2 - 6x_3^2 \\ x_3(6x_2 + 50) \end{pmatrix}$$

Soit $f(x) = 15x^2 + \exp(3x + 1)$, on calcule le gradient de f on obtient :

$$\nabla f(x) = f'(x) = 30x + 3\exp(3x + 1)$$

Remarque 1.2.1. Nous rappellerons aussi la formule :

$$\frac{\partial f}{\partial h}(x) = \langle \nabla f(x), h \rangle, \forall x \in \Omega \forall h \in \mathbb{R}^n$$

Proposition 1.2.1. pour tout $x \in \mathbb{R}^n$ le gradient $\nabla f(x)$ est l'unique vecteur tel que pour tout $h \in \mathbb{R}^n$ on a :

$$\nabla f(x)(h) = \langle f(x), h \rangle$$

où pour deux vecteurs $u = (u_1, \dots, u_n)$ et $v = (v_1, \dots, v_n)$ de $\mathbb{R}^n$ on a noté $\langle u, v \rangle$ le produit scalaire usuel $\sum_{j=1}^n u_j v_j$.

Proposition 1.2.2. (Gradient de la composée) Supposons qu'on a deux ouverts $\Omega \subset \mathbb{R}^n$ et $U \subset \mathbb{R}$ et deux fonctions $f : \Omega \rightarrow \mathbb{R}$ et $g : U \rightarrow \mathbb{R}$ avec on plus $f(\Omega)$ (on peut alors définir $\text{gof} : \Omega \rightarrow \mathbb{R}$). Supposons que f, g sont de classe C^1 alors gof est aussi de classe C^1 avec en plus

$$\nabla(\text{gof})(x) = g'(f(x))\nabla f(x) \forall x \in \Omega.$$

Exemple 1.2.3. $f(x_1, x_2) = 2x_1x_2^2$, $g(x) = 3x + 2$

$$\begin{aligned}\nabla(gof)(x) &= 6(x_1x_2^2) \begin{pmatrix} 2x_2^2 \\ 4x_1x_2 \end{pmatrix} \\ &= \begin{pmatrix} 12x_1x_2^4 \\ 24x_1^2x_2^3 \end{pmatrix}.\end{aligned}$$

Dérivée directionnelle

Définition 1.2.3. [10] On appelle dérivée directionnelle de f dans la direction d au point x , notée $\delta f(x, d)$, la limite (éventuellement $+\infty$ et $-\infty$) du rapport :

$$\frac{f(x + \alpha d) - f(x)}{\alpha}$$

lorsque α tend vers 0. Autrement dit :

$$\delta f(x, d) = \lim_{\alpha \rightarrow 0} \frac{f(x + \alpha d) - f(x)}{\alpha} = \nabla^T f(x) d$$

Remarque 1.2.2. Si $\|d\| = 1$: la dérivée directionnelle est le taux d'accroissement de f dans la direction d au point x .

Remarque 1.2.3. Pour tout $x \in \Omega$ et $d \in \mathbb{R}^n$ on note

$$\frac{\partial f}{\partial d}(x) = \lim_{\alpha \rightarrow 0} \frac{1}{\alpha} [f(x + \alpha d) - f(x)] = g'(0)$$

(c'est la dérivée directionnelle de f en x de direction d) où on a noté :

$$g(\alpha) = f(x + \alpha d).$$

Exemple 1.2.4. soit $f(x_1, x_2, x_3) = 5x_1x_3 - x_1^3x_2 + 2x_1x_2x_3$ et soit

$$d = \begin{pmatrix} d_1 \\ d_2 \\ d_3 \end{pmatrix}$$

La dérivée directionnelle de f dans la direction d est

$$(d_1 \ d_2 \ d_3) \nabla f(x_1, x_2, x_3) = d_1(5x_3 - 3x_1^2x_2 + 2x_2x_3) + d_2(-x_1^3 + 2x_1x_3) + d_3(5x_1 + 2x_1x_2)$$

où $\nabla f(x_1, x_2, x_3)$ est donnée par

$$\nabla f(x_1, x_2, x_3) = \begin{pmatrix} 5x_3 - 3x_1^2x_2 + 2x_2x_3 \\ -x_1^3 + 2x_1x_3 \\ 5x_1 + 2x_1x_2 \end{pmatrix}$$

Fonction différentiable

Définition 1.2.4. [3] Soit $f : \mathbb{R}^n \rightarrow \mathbb{R}$ une fonction continue. Si, pour tout $d \in \mathbb{R}^n$, la dérivée directionnelle de f dans la direction d existe, alors la fonction f est dite différentiable.

Remarque 1.2.4. Cette notion est parfois appelée $\langle \langle \text{Gateaux-différentiabilité} \rangle \rangle$ en ce sens que d'autres type de différentiabilité peuvent être définis (comme la différentiabilité au sens Fréchet). La dérivée directionnelle donne des informations sur la pente de la fonction dans la direction d , tout comme la dérivée donne des informations sur la pente des fonctions à une variable. Notamment, la fonction est croissante dans la direction d si la dérivée directionnelle est strictement positive et décroissante si elle est strictement négative. Dans ce dernier cas, nous dirons qu'il s'agit d'une direction de descente.

Définition 1.2.5. (différentiable au sens de Gateaux) Soit f une application de $\mathbb{R}^n$ à valeur dans $\mathbb{R}$, v un élément de $\mathbb{R}^n$ et x un point de $\mathbb{R}^n$. On dit que f admet une dérivée en x suivant le vecteur v si la limite suivante existe :

$$\frac{\partial f(x)}{\partial v} = \lim_{\alpha \rightarrow 0} \frac{f(x + \alpha v) - f(x)}{\alpha}$$

On dit que f est **Gateaux-différentiable** en x si les dérivées directionnelles $\frac{\partial f(x)}{\partial v}$ sont définies pour tout v et si de plus l'application $v \mapsto \frac{\partial f(x)}{\partial v}$ est linéaire.

Direction de descente

Définition 1.2.6. [2] Soient $f : \mathbb{R}^n \rightarrow \mathbb{R}$ et $x \in \mathbb{R}^n$. Le vecteur $d \in \mathbb{R}^n$ est une direction de descente pour f à partir du point x si $\alpha \mapsto f(x + \alpha d)$ est décroissante

en $x = 0$, c'est à dire s'il existe $\eta > 0$ tel que :

$$f(x + \alpha d) < f(x), \forall x \in]0, \eta].$$

Définition 1.2.7. Soit d un vecteur non nul de $\mathbb{R}^n$. On dit que d est une direction de descente de f en $x \in \mathbb{R}^n$ si $d^T \nabla f(x) < 0$.

Le terminologie $\langle \langle \text{direction de descente} \rangle \rangle$ est justifiée par la théorème suivante.

Théorème 1.2.1. Soit $f : \mathbb{R}^n \rightarrow \mathbb{R}$ une fonction différentiable. Soient $x \in \mathbb{R}^n$ tel que $\nabla f(x) \neq 0$ et $d \in \mathbb{R}^n$. Si d est une direction de descente, alors il existe $\eta > 0$ tel que

$$f(x + \alpha d) < f(x),$$

$\forall 0 < \alpha \leq \eta$ De plus, pour tout $\beta < 1$, il existe $\eta > 0$ tel que

$$f(x + \alpha d) < f(x) + \alpha \beta \nabla f(x)^T d$$

pour tout $0 < \alpha \leq \eta$.

Théorème 1.2.2. (plus forte pente) Soit $f : \mathbb{R}^n \rightarrow \mathbb{R}$ une fonction différentiable. Soient $x \in \mathbb{R}^n$ et $d^* = \nabla f(x)$. Alors pour tout $d \in \mathbb{R}^n$ tel que $\|d\| = \|\nabla f(x)\|$, on a

$$d^T \nabla f(x) \leq d^{*T} \nabla f(x) = \nabla f(x)^T \nabla f(x).$$

Exemple 1.2.5. soit : $f(x) = \frac{1}{2}x_1^2 + x_2^2$ et soit $x = (1.1)^T$. Nous considrons trois directions

$$d_1 = \nabla f(x) = \begin{pmatrix} 1 \\ 4 \end{pmatrix}, d_2 = \begin{pmatrix} 1 \\ 1 \end{pmatrix} \text{ et } d_3 = \begin{pmatrix} -1 \\ 3 \end{pmatrix}.$$

La dérivée directionnelle de f dans chacune de ces directions vaut :

$$d_1^T \nabla f(x) = 17,$$

$$d_2^T \nabla f(x) = 5,$$

$$d_3^T \nabla f(x) = 11.$$

Théorème 1.2.3. (plus forte descente) Soit $f : \mathbb{R}^n \rightarrow \mathbb{R}$ une fonction différentiable. Soient $x \in \mathbb{R}^n$ et $d^* = -\nabla f(x)$. Alors pour tout $d \in \mathbb{R}^n$ tel que $\|d\| = \|\nabla f(x)\|$, on a

$$-\nabla f(x)^T \nabla f(x) = d^{*T} \nabla f(x) \leq d^T \nabla f(x)$$

et la direction opposée au gradient est celle où la fonction a la plus forte descente.

Exemple 1.2.6. Soit : $f(x) = \frac{1}{2}x_1^2 + 2x_2^2$ et soit $x = (1,1)^T$. Nous considérons trois directions

$$d_1 = -\nabla f(x) = \begin{pmatrix} -1 \\ -4 \end{pmatrix}, \quad d_2 = \begin{pmatrix} -1 \\ -1 \end{pmatrix} \quad \text{et} \quad d_3 = \begin{pmatrix} 1 \\ -3 \end{pmatrix}.$$

La dérivée directionnelle de f dans chacune de ces directions vaut :

$$d_1^T \nabla f(x) = -17,$$

$$d_2^T \nabla f(x) = -5,$$

$$d_3^T \nabla f(x) = -11.$$

1.2.2 Différentiabilité : du second ordre

Nous pouvons effectuer la même analyse de différentiabilité faite sur la fonction dans la section (1.2) pour chacune des fonctions $\nabla_i f(x)$ de la dérivée partielle. La $j^{\text{ème}}$ dérivée partielle de $\nabla_i f(x)$ est la dérivée seconde de f par rapport aux variables i et j , car

$$\frac{\partial \nabla_i f(x)}{\partial x_j} = \frac{\partial \left(\frac{\partial f(x)}{\partial x_i} \right)}{\partial x_j} = \frac{\partial^2 f(x)}{\partial x_i \partial x_j}$$

Il est courant d'organiser ces dérivées secondes dans une matrice $n \times n$ dont l'élément de la ligne i et de colonne j soit $\frac{\partial^2 f(x)}{\partial x_i \partial x_j}$. Cette matrice est appelée matrice hessienne ou hessien.

Matrice Hessienne

Définition 1.2.8. [24] [23] Soit $f : \mathbb{R}^n \rightarrow \mathbb{R}$ une fonction deux fois différentiable. La fonction notée $\nabla^2 f(x) : \mathbb{R}^n \rightarrow \mathbb{R}^{n \times n}$ est appelée matrice hessienne ou hessien

de f est définie par

$$\nabla^2 f(x) = \begin{pmatrix} \frac{\partial^2 f(x)}{\partial x_1^2} & \frac{\partial^2 f(x)}{\partial x_1 \partial x_2} & \cdots & \frac{\partial^2 f(x)}{\partial x_1 \partial x_n} \\ \frac{\partial^2 f(x)}{\partial x_2 \partial x_1} & \frac{\partial^2 f(x)}{\partial x_2^2} & \cdots & \frac{\partial^2 f(x)}{\partial x_2 \partial x_n} \\ \vdots & & \ddots & \vdots \\ \frac{\partial^2 f(x)}{\partial x_n \partial x_1} & \frac{\partial^2 f(x)}{\partial x_n \partial x_2} & \cdots & \frac{\partial^2 f(x)}{\partial x_n^2} \end{pmatrix}$$

On note dans $\mathbb{R}$,

$$\nabla^2 f(x) = f''(x).$$

Remarque 1.2.5. Si f deux fois différentiable en x alors $\nabla^2 f(x)$ est une matrice symétrique, c'est-à-dire

$$\frac{\partial}{\partial x_i} \left(\frac{\partial f}{\partial x_j} \right) = \frac{\partial}{\partial x_j} \left(\frac{\partial f}{\partial x_i} \right)$$

.

Exemple 1.2.7. Soit $f : f(x_1, x_2, x_3) = 2x_1^2 + 15(x_1 - x_2^2)^2 + 20(x_2 + x_3^2)^2$ la matrice hessien de f est :

$$\nabla^2 f(x_1, x_2, x_3) = \begin{pmatrix} 34 & (-60)x_2 & 0 \\ (-60)x_2 & (-60)x_1 + 180x_2^2 + 40 & 80x_3 \\ 0 & 80x_3 & 80x_2 + 240x_3^2 \end{pmatrix}$$

Soit : $f(x) = 5x^3 + \exp(2x^2)$ la hessienne de f est :

$$\nabla^2 f(x) = f''(x) = 30x + 4\exp(2x^2) + 16x^2 \exp(2x^2)$$

.

1.3 Développement de Taylor

La formule de Taylor est un outil important en convexité. Nous la rappelons dans le cas général. Soit $\Omega \subset \mathbb{R}^n$ ouvert, $f : \Omega \rightarrow \mathbb{R}$. $\alpha \in \Omega$ et $h \in \mathbb{R}$ tels que $[\alpha, \alpha + h]$. Alors

1. Si $f \in \mathbb{C}^1(\Omega)$ alors

i) Formulation de Taylor à l'ordre 1 avec reste intégral

$$f(\alpha + h) = f(\alpha) + \int_0^1 \langle \nabla f(\alpha + th), h \rangle dt.$$

ii) Formulation de Taylor-Maclaurin à l'ordre 1

$$f(\alpha + h) = f(\alpha) + \langle \nabla f(\alpha), h \rangle + o(\|h\|).$$

iii) Formulation de Taylor-Young à l'ordre 1

$$f(\alpha + h) = f(\alpha) + \langle \nabla f(\alpha), h \rangle + o(\|h\|).$$

2. Si $f \in \mathbb{C}^2(\Omega)$ alors

i) Formulation de Taylor à l'ordre 2 avec reste intégral

$$f(\alpha + h) = f(\alpha) + \langle \nabla f(\alpha), h \rangle + \int_0^1 (1-t) \langle \nabla^2 f(\alpha + th)h, h \rangle dt.$$

ii) Formulation de Taylor-Maclaurin à l'ordre 2

$$f(\alpha + h) = f(\alpha) + \langle \nabla f(\alpha), h \rangle + \frac{1}{2} \langle \nabla^2 f(\alpha)h, h \rangle + o(\|h\|^2).$$

avec $0 < \theta < 1$

iii) Formulation de Taylor-Young à l'ordre 2

$$f(\alpha + h) = f(\alpha) + \langle \nabla f(\alpha), h \rangle + \frac{1}{2} \langle \nabla^2 f(\alpha)h, h \rangle + o(\|h\|^2).$$

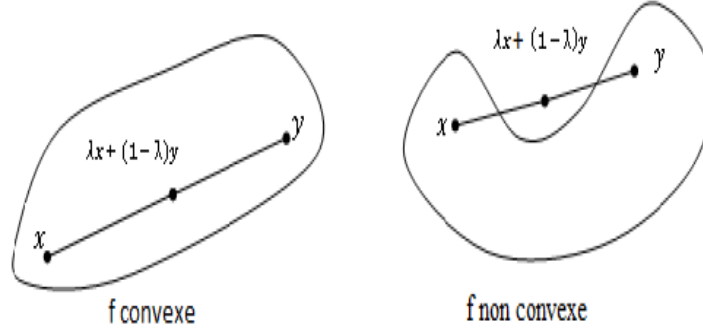
Remarque 1.3.1. La notation $o(\|h\|^k)$ pour $k \in \mathbb{N}^*$ signifie une expression qui tend vers 0 plus vite que $\|h\|^k$ (c'est à dire, si on la divise par $\|h\|^k$, le résultat tend vers 0 quand $\|h\|$ tend vers 0)

Théorème 1.3.1. [7] Soit $f : \Omega \rightarrow \mathbb{R}$ de classe $C^{m+1}(\Omega)$. Si le segment $[\alpha, \alpha + h]$ est contenu dans Ω , on a :

$$f(\alpha + h) = f(\alpha) + f'(\alpha) \cdot h + \dots + \frac{1}{n!} f^n(\alpha) (h)^n + \int_0^1 \frac{(1-t)^n}{n!} f^{(n+1)}(\alpha + th) (h)^{(n+1)} dt.$$

1.4 Convexité

1.4.1 Ensemble convexe



Définition 1.4.1. Un ensemble $C \subset \mathbb{R}^n$ est dit convexe si pour tout couple $(x, y) \in C^2$ et $\forall \lambda \in [0, 1]$ on a :

$$\lambda x + (1 - \lambda)y \in C$$

.

Propriétés 1.4.1. 1. Soit $X_1, X_2, \dots, X_m \subset \mathbb{R}^n$, m ensembles convexes alors on a : $Y = \bigcap_{i=1}^m X_i$ est un ensemble convexe.

2. La réunion d'ensembles convexes n'est pas forcément convexe.

3. Soit $X_1, X_2, \dots, X_m \subset \mathbb{R}^n$, m ensembles convexes,

$$\forall (\lambda_1, \lambda_2, \dots, \lambda_m) \in \mathbb{R}_+^m,$$

$$Y = \sum_{i=1}^m \lambda_i X_i = \left\{ \sum_{i=1}^m \lambda_i x_i, x_i \in X_i, i = 1, \dots, m \right\}$$

1.4.2 Combinaison convexe

On appelle combinaison convexe de n points $x_1, x_2, \dots, x_n$, tout point obtenu par la formule [6] :

$$y = \sum_{i=1}^n \lambda_i x_i, \text{ avec } \sum_{i=1}^n \lambda_i = 1.$$

Théorème 1.4.1. *un ensemble $C \subset \mathbb{R}^n$ est convexe si et seulement si toute combinaison convexe des points de C appartient à C .*

Enveloppe convexe

L'enveloppe convexe d'un sous ensemble $C \subset X$ quelconque est le plus petit convexe contenant C , elle est notée $\text{conv}(X)$ [19].

Fonction convexe

Définition 1.4.2. [3] [15] Une fonction $f : \mathbb{R}^n \rightarrow \mathbb{R}$ est dite convexe si pour tout $x, y \in \mathbb{R}^n$ et pour tout $\lambda \in [0, 1]$, on a

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

La définition (1.4.2) est illustrée dans la figure (1.1). Le segment de droite reliant les points $(x, f(x))$ et $(y, f(y))$ se trouve totalement de droite du graphe de f . Le point $z = \lambda x + (1 - \lambda)y$ est situé quelque part entre x et y . Le point de coordonnées $(z, \lambda f(x) + (1 - \lambda)f(y))$ se trouve sur le segment de droite entre les points $(x, f(x))$ et $(y, f(y))$. Pour que la fonction soit convexe, il faut que ce point se trouve toujours (c'est-à-dire pour tout x, y et $0 \leq \lambda \leq 1$) au-dessus du graphe de la fonction.

La figure (1.2) illustre une fonction non convexe, ou l'on voit qu'il existe un x et un y tels que la reliant $(x, f(x))$ et $(y, f(y))$ n'est pas située totalement au-dessus du graphe de la fonction.

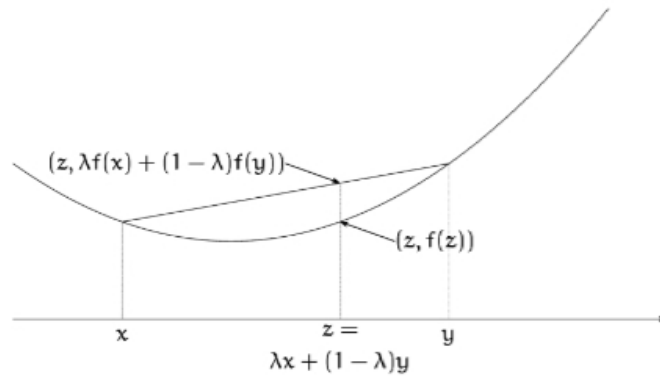


FIGURE 1.1 – Illustration de la définition 1.4.2

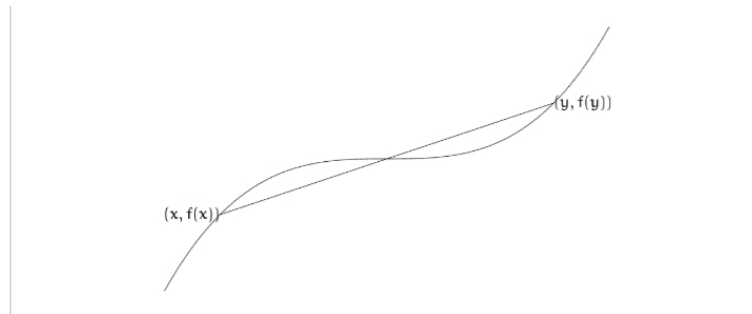
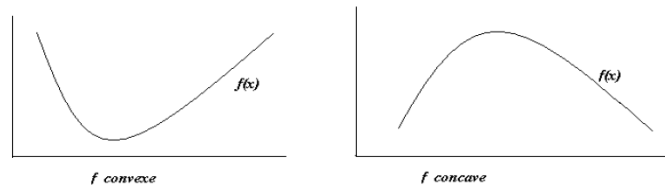


FIGURE 1.2 – Illustration d'un contre-exemple à la définition 1.4.2

Fonctions concave

Définition 1.4.3. [3] Une fonction $f : \mathbb{R}^n \rightarrow \mathbb{R}$ est dite concave si $-f$ est une fonction convexe, c'est-à-dire si pour tout $x, y \in \mathbb{R}^n, x \neq y$, et pour tout $\lambda \in [0, 1]$, on a

$$f(\lambda x + (1 - \lambda)y) \geq \lambda f(x) + (1 - \lambda)f(y).$$



Fonction strictement convexe

Définition 1.4.4. [3] Une fonction $f : \mathbb{R}^n \rightarrow \mathbb{R}$ est dit strictement convexe si pour tout $x, y \in \mathbb{R}^n$, $x \neq y$, et pour tout $\lambda \in]0, 1[$, on a

$$f(\lambda x + (1 - \lambda)y) < \lambda f(x) + (1 - \lambda)f(y).$$

Proposition 1.4.1. des fonctions convexes :

1. Soit $f : X \rightarrow \mathbb{R}^n$ avec X un ensemble convexe. Si f est une fonction convexe, alors : $f(\sum_{i=1}^n \lambda_i x_i) \leq \sum_{i=1}^n \lambda_i f(x_i)$ avec $\lambda_i \geq 0$, $\sum_{i=1}^n \lambda_i = 1$, $x_i \in X$, $i = 1, \dots, n$.
2. Soit $f_1, f_2, \dots, f_n$ n fonction convexe définies sur un ensemble convexe X , alors : $f(x) = \sum_{i=1}^n \lambda_i f_i(x)$ est une fonction convexe sur X , avec $\lambda_i \geq 0$, $i = 1, \dots, n$. En d'autres termes, une combinaison linéaire à coefficients positifs de fonctions convexes est une fonction convexe.

Chapitre 2

Minimisation sans contraintes

soit $f : \mathbb{R}^n \rightarrow \mathbb{R}$. On appelle problème de minimisation sans contraintes le problème suivant

$$(P) \left\{ \min_{x \in \mathbb{R}^n} f(x) \right\}. \quad (2.1)$$

L'étude de ces problèmes est importante pour des raisons diverses. Beaucoup de problèmes d'optimisation avec contraintes sont transformés en des suites de problèmes d'optimisation sans contraintes (multiplicateur de Lagrange, méthodes des pénalités, . . .). L'étude des problèmes d'optimisation sans contraintes trouve aussi des applications dans la résolution des systèmes non linéaires. Une grande classe d'algorithmes que nous allons considérer pour le problème d'optimisation sans contraintes ont la forme générale suivante

$$\begin{cases} x_0 \text{ point initial} \\ x_{k+1} = x_k + \alpha_k d_k \end{cases}$$

le vecteur d_k s'appelle la direction de descente, α_k le pas de la méthode à la k -ième itération. En pratique, on s'arrange presque toujours pour avoir l'inégalité suivante

$$f(x_{k+1}) \leq f(x_k),$$

qui assure la décroissance suffisante de la fonction objectif f . de tels algorithmes sont souvent appelés méthodes de descente. Essentiellement la différence entre

ces algorithmes réside dans le choix de la direction de descente d_k , cette direction étant choisie nous sommes plus ou moins ramenés à un problème unidimensionnel pour la détermination de α_k . Pour s'approcher de la solution optimale du problème (P) (dans le cas générale, c'est un point en lequel ont lieu peut être avec une certaine précision les conditions nécessaires d'optimalité de f), on se déplace naturellement à partir du point x_k dans la direction de la décroissance de la fonction f . L'optimisation sans contraintes a les propriétés suivantes :

- toutes les méthodes nécessitent un point de départ x_0 .
- les méthodes déterministes convergent vers le minimum local le plus proche.
- plus vous saurez sur la fonction (gradient, hessien) plus la minimisation sera efficace.

considérons le problème d'optimisation sans contraintes (P).

Définition 2.0.5.

1) $\hat{x} \in \mathbb{R}^n$ est un minimum local de (P) si et seulement s'il existe un voisinage $V_\varepsilon(\hat{x})$ tel que :

$$f(\hat{x}) \leq f(x) : \forall x \in V_\varepsilon(\hat{x})$$

2) $\hat{x} \in \mathbb{R}^n$ est un minimum local strict de (P) si et seulement s'il existe un voisinage $V_\varepsilon(\hat{x})$ tel que :

$$f(\hat{x}) < f(x) : \forall x \in V_\varepsilon(\hat{x}), x \neq \hat{x}$$

3) $\hat{x} \in \mathbb{R}^n$ est un minimum global de (P) si et seulement si

$$f(\hat{x}) \leq f(x) : \forall x \in \mathbb{R}^n$$

2.1 Résultats d'existence et d'unicité

Avant d'étudier les propriétés de la solution (ou des solutions) de (P) il faut s'assurer de leur existence. Nous donnerons ensuite des résultats d'unicité.

Définition 2.1.1. On dit que $f : \mathbb{R}^n \rightarrow \mathbb{R}$ est **coercive** si

$$\lim_{\|x\| \rightarrow +\infty} f(x) = +\infty.$$

Ici $\|\cdot\|$ désigne une norme quelconque de $\mathbb{R}^n$. On notera $\|\cdot\|_p$ ($p \in \mathbb{N}$) la norme l_p de $\mathbb{R}^n$:

$$\forall x = (x_1, \dots, x_n) \in \mathbb{R}^n \quad \|x\|_p = \left[\sum_{i=1}^n |x_i|^p \right]^{\frac{1}{p}}.$$

La norme infinie de $\mathbb{R}^n$ est

$$\forall x = (x_1, \dots, x_n) \in \mathbb{R}^n \quad \|x\|_\infty = \max_{1 \leq i \leq n} |x_i|.$$

Théorème 2.1.1. (Existence) Soit $f : \mathbb{R}^n \rightarrow \mathbb{R} \cup \{+\infty\}$ propre, continue et coercive. Alors (P) admet au moins une solution.

Preuve 2.1.1. [21]

* Soit $d = \inf(f)$; $d < +\infty$ car f est **propre**. Soit $(x_p)_{p \in \mathbb{N}} \in \mathbb{R}^n$ une suite minimisante, c'est-à-dire telle que

$$\lim_{p \rightarrow +\infty} f(x_p) = d.$$

Montrons que (x_p) est bornée.

Si ce n'était pas le cas on pourrait extraire de cette suite une sous-suite (encore notée (x_p)) telle $\lim_{p \rightarrow +\infty} f(x_p) = +\infty$. Par **coercivité** de f on aurait

$$\lim_{p \rightarrow +\infty} f(x_p) = +\infty \text{ ce qui contredit le fait que } \lim_{p \rightarrow +\infty} f(x_p) = d < +\infty.$$

Comme (x_p) est bornée, on peut alors en extraire une sous-suite (encore notée (x_p)) qui converge vers $\hat{x} \in \mathbb{R}^n$ par **continuité** de f , on a alors

$$d = \lim_{p \rightarrow +\infty} f(x_p) = f(\hat{x}).$$

En particulier $d > -\infty$ est $\hat{x}$ est une solution du problème (P).

Théorème 2.1.2. (Unicité) Soit $f : \mathbb{R}^n \rightarrow \mathbb{R} \cup \{+\infty\}$ strictement convexe. Alors le problème (P) admet au plus une solution.

Preuve 2.1.2. [21]

* Supposons que f admette au moins un minimum m et soit $x_1 \neq x_2$ (dans $\mathbb{R}^n$) réalisant ce minimum : $f(x_1) = f(x_2) = m$. Par stricte convexité de la fonction f on a alors :

$$f\left(\frac{x_1 + x_2}{2}\right) < \frac{1}{2}(f(x_1) + f(x_2)) = m;$$

ceci contredit le fait que m est le minimum. Donc $x_1 = x_2$.

Donnons pour terminer un critère pour qu'une fonction soit strictement convexe et coercive.

Théorème 2.1.3. Soit f une fonction $\mathcal{C}^1$ de $\mathbb{R}^n$ dans $\mathbb{R}$. On suppose qu'il existe $\alpha > 0$ tel que

$$\forall (x, y) \in \mathbb{R}^n \times \mathbb{R}^n \quad \langle \nabla f(x) - \nabla f(y), x - y \rangle \geq \alpha \|x - y\|^2. \quad (2.2)$$

Alors f est strictement convexe et coercive, en particulier le problème (P) admet solution unique .

Preuve 2.1.3. [21]

* La condition (2.2) implique que ∇f est monotone et que f est convexe. De plus on a la stricte convexité de f . Enfin f est coercive : en effet, appliquons la formulation de Taylor avec reste intégral

$$f(y) = f(x) + \int_0^1 \frac{d}{dt} f(x + t(y - x)) dt = f(x) + \int_0^1 \langle \nabla f(x + t(y - x)), y - x \rangle dt.$$

Donc

$$f(y) = f(x) + \langle \nabla f(x), y - x \rangle + \int_0^1 \langle \nabla f(x + t(y - x)) - \nabla f(x), y - x \rangle dt. \quad (2.3)$$

D'après (2.3) on obtient

$$f(y) \geq f(x) + \langle \nabla f(x), y - x \rangle + \int_0^1 t\alpha \|x - y\|^2 dt.$$

Finalement

$$f(y) \geq f(x) - \|\nabla f(x)\| \|y - x\| + \frac{\alpha}{2} \|x - y\|^2.$$

Fixons $x = 0$ par exemple ; il est alors clair que f est coercive.

Par conséquent, f admet un minimum unique $\hat{x}$ sur $\mathbb{R}^n$ caractérisé par

$\nabla f(\hat{x}) = 0$. La condition (2.2) nous amène à la définition suivante :

Définition 2.1.2. (Fonction élliptique) On dit que $f : \mathbb{R}^n \rightarrow \mathbb{R}$ est *élliptique* si la condition (2.2) est vérifiée, c'est-à-dire $\exists \alpha > 0$ tel que $\forall (x, y) \in \mathbb{R}^n \times \mathbb{R}^n : \langle \nabla f(x) - \nabla f(y), x - y \rangle \geq \alpha \|x - y\|^2$. α est la constante d'ellipticité.

Proposition 2.1.1. [21] Une fonction $f : \mathbb{R}^n \rightarrow \mathbb{R}$ deux fois différentiable sur $\mathbb{R}^n$ est élliptique si et seulement si

$$\forall (x, y) \in \mathbb{R}^n \times \mathbb{R}^n : \langle \nabla^2 f(x)y, y \rangle \geq \alpha \|y\|^2.$$

Pour la preuve on utilise de nouveau la formule de Taylor appliquée à la fonction $\varphi : t \rightarrow \varphi(t) = f(x + ty)$.

Il faut maintenant donner des conditions pour pouvoir calculer la (ou les) solutions.

On va rechercher à montrer que cette solution est solution de certaines équations, de sorte qu'il sera plus facile de la calculer.

2.2 Conditions d'optimalité

2.2.1 Conditions nécessaires d'optimalité du premier ordre

Théorème 2.2.1. [4] soit $f : \mathbb{R}^n \rightarrow \mathbb{R}$ tel que f est différentiable au point $\hat{x} \in \mathbb{R}^n$. soit $d \in \mathbb{R}^n$ tel que $\nabla f(\hat{x})^T \cdot d < 0$. Alors il existe $\delta > 0$ tel que $f(\hat{x} + \alpha d) < f(\hat{x})$ pour tout $\alpha \in]0, \delta[$. La direction d s'appelle dans ce cas direction de descente.

Preuve 2.2.1. [21] comme f est différentiable en $\hat{x}$ alors :

$$f(\hat{x} + \alpha d) = f(\hat{x}) + \alpha \nabla f(\hat{x})^T \cdot d + \alpha \|d\| \lambda(\hat{x}, \alpha d)$$

ou $\lambda(\hat{x}, \alpha d) \rightarrow 0$ pour $\alpha \rightarrow 0$. Ceci implique :

$$\frac{f(\hat{x} + \alpha d) - f(\hat{x})}{\alpha} = \nabla f(\hat{x})^T \cdot d + \|d\| \lambda(\hat{x}, \alpha d) \quad \alpha \neq 0.$$

et comme $\nabla f(\hat{x})^T \cdot d < 0$ et $\lambda(\hat{x}, \alpha d) \rightarrow 0$ pour $\alpha \rightarrow 0$, il existe $\delta > 0$ tel que :

$$\nabla f(\hat{x})^T \cdot d + \|d\| \lambda(\hat{x}, \alpha d) < 0 \quad \text{pour tout } \alpha \in]0, \delta[.$$

et par conséquent on obtient :

$$f(\hat{x} + \alpha d) < f(\hat{x}) \quad \text{pour tout } \alpha \in]0, \delta[.$$

Corolaire 2.2.1. *soit $f : \mathbb{R}^n \rightarrow \mathbb{R}$ différentiable en $\hat{x}$, si $\hat{x}$ est un minimum local alors $\nabla f(\hat{x}) = 0$.*

Preuve 2.2.2. [10] *On suppose que $\nabla f(\hat{x}) \neq 0$. Si on pose $d = -\nabla f(\hat{x})$, on obtient :*

$$\nabla f(\hat{x})^T \cdot d = -\|\nabla f(\hat{x})\|^2 < 0.$$

et par le théorème précédent, il existe $\delta > 0$ tel que :

$$f(\hat{x} + \alpha d) < f(\hat{x}) \quad \text{pour tout } \alpha \in]0, \delta[.$$

mais ceci est contradictoire avec le fait que $\hat{x}$ est un minimum local d'où $\nabla f(\hat{x}) = 0$.

2.2.2 Conditions nécessaires d'optimalité du seconde ordre

Théorème 2.2.2. [8] *soit $f : \mathbb{R}^n \rightarrow \mathbb{R}$ deux fois différentiable, $\hat{x}$ est un minimum local. Alors $\nabla f(\hat{x}) = 0$ et $\nabla^2 f(\hat{x})$ est semi définie positive.*

Preuve 2.2.3. [21] *le corollaire ci-dessus montre la première proposition, pour la deuxième proposition on a :*

$$f(\hat{x} + \alpha d) = f(\hat{x}) + \alpha \nabla f(\hat{x})^T \cdot d + \frac{1}{2} \alpha^2 \nabla^2 f(\hat{x}) d + \alpha^2 \|d\| \lambda(\hat{x}, \alpha d), \quad \alpha \neq 0.$$

comme $\hat{x}$ est un minimum local alors $f(\hat{x} + \alpha d) \geq f(\hat{x})$ pour α suffisamment petit, d'où :

$$\frac{1}{2} d^T \nabla^2 f(\hat{x}) d + \|d\|^2 \lambda(\hat{x}, \alpha d) \geq 0 \quad \text{pour } \alpha \text{ petit.}$$

En passant à la limite quand $\alpha \rightarrow 0$, on obtient que $d^T \nabla^2 f(\hat{x}) d \geq 0$, d'où $\nabla^2 f(\hat{x})$ est semi définie positive.

2.2.3 Conditions suffisantes d'optimalité

Les conditions données précédemment sont nécessaires (si f n'est pas convexe), c'est-à-dire qu'elles doivent être satisfaites pour tout minimum local, cependant, tout point vérifiant ces conditions n'est pas nécessairement un minimum local. Le théorème suivant établit une condition suffisante pour qu'un point soit un minimum local, si f est deux fois différentiable.

Théorème 2.2.3. [8] Soit $f : \mathbb{R}^n \rightarrow \mathbb{R}$ deux fois différentiable en $\hat{x}$.

Si $\nabla f(\hat{x}) = 0$ et $\nabla^2 f(\hat{x})$ est définie positive, alors $\hat{x}$ est minimum local strict.

Preuve 2.2.4. [20] f est deux fois différentiable au point $\hat{x}$. Alors $\forall x \in \mathbb{R}^n$, on obtient :

$$f(x) = f(\hat{x}) + \nabla f(\hat{x})^T (x - \hat{x}) + \frac{1}{2} (x - \hat{x})^T \nabla^2 f(\hat{x}) (x - \hat{x}) + \|x - \hat{x}\|^2 \lambda(\hat{x}; x - \hat{x}) \quad (2.4)$$

où $\lambda(\hat{x}; x - \hat{x}) \rightarrow 0$ pour $x \rightarrow \hat{x}$. Supposons que $\hat{x}$ n'est pas un minimum local strict. Donc il existe une suite $\{x_k\}$ convergente vers $\hat{x}$ telle que

$$f(x_k) \leq f(\hat{x}), x_k \neq \hat{x}, \forall k.$$

Posons $d_k = \frac{(x_k - \hat{x})}{\|x_k - \hat{x}\|}$. Donc $\|d_k\| = 1$ et on obtient à partir de (2.4) :

$$\frac{f(x_k) - f(\hat{x})}{\|x_k - \hat{x}\|^2} = \frac{1}{2} d_k^T \nabla^2 f(\hat{x}) d_k + \lambda(\hat{x}; x_k - \hat{x}) \leq 0, \forall k \quad (*)$$

et comme $\|d_k\| = 1, \forall k$ alors $\exists \{d_k\}_{k \in N_1 \subset \mathbb{N}}$ telle que $d_k \rightarrow d$ pour $k \rightarrow \infty$ et $k \in N_1$. On a bien sur $\|d\| = 1$. Considérons donc $\{d_k\}_{k \in N_1}$ et le fait que $\lambda(\hat{x}; x - \hat{x}) \rightarrow 0$ pour $k \rightarrow \infty$ et $k \in N_1$. Alors (*) donne : $d^T \nabla^2 f(\hat{x}) d \leq 0$, ce qui contredit le fait que $\nabla^2 f(\hat{x})$ est définie positive car $\|d\| = 1$ (donc $d \neq 0$). Donc $\hat{x}$ est un minimum local stricte.

Exemple 2.2.1. Soit $f(x, y) = mx^2 + y^2$ définie sur $\mathbb{R}^2$ dans $\mathbb{R}$, $m \neq 0$ pour déterminer les extrêmes de f , il faut d'abord déterminer les points stationnaires de f et la matrice Hessienne va déterminer la nature du point stationnaire.

*) Calcule de $\nabla f(x, y) = \left(\frac{\partial f(x, y)}{\partial x}, \frac{\partial f(x, y)}{\partial y} \right) = (2mx, 2y)$

$$(2mx, 2y) = (0, 0) \Rightarrow \begin{cases} 2mx = 0 \\ 2y = 0 \end{cases} \Leftrightarrow \begin{cases} x = 0, \\ y = 0 \end{cases}$$

$S(f) = (0, 0)$. Ce point stationnaire est le seul extrémum éventuel.

a) Nature du point stationnaire $(0, 0)$

$$\nabla^2 f(x, y) = \begin{pmatrix} 2m & 0 \\ 0 & 2 \end{pmatrix}$$

$$\nabla^2 f(0, 0) = \begin{pmatrix} 2m & 0 \\ 0 & 2 \end{pmatrix}$$

* Si $m > 0$, alors $\nabla^2 f(x_0)$ est défini positive donc par la condition suffisante $x_0 = (0, 0)$ est minimum local strict.

* Si $m < 0$, alors les valeurs propre n'ont pas le même signe donc par la condition nécessaire, $x_0 = (0, 0)$ ne peut pas être un minimum local.

Cas convexe

Théorème 2.2.4. Soit $f : \mathbb{R}^n \rightarrow \mathbb{R}$ tel que f est convexe et différentiable. Alors $\hat{x}$ est un minimum global de f si et seulement si

$$\nabla f(\hat{x}) = 0.$$

Remarque 2.2.1. Les théorèmes (2.2.4) et (2.2.3) demeurent vrais sinon remplace $\mathbb{R}^n$ par un ouvert S de $\mathbb{R}^n$

Remarque 2.2.2. *Dans le cas où f est convexe, alors tout minimum local est aussi globale. De plus si f est strictement convexe, alors tout minimum local devient non seulement global mais aussi unique.*

Méthodes de résolutions

Dans ce chapitre nous allons introduire une classe importante d'algorithmes de résolution des problèmes d'optimisation sans contraintes. Bien entendu, nous ne pouvons pas être exhaustifs; nous présentons les méthodes "de base" les plus classique. et avec cela, la plupart de ces algorithmes exploitent les conditions d'optimalité dont on a vu qu'elles permettaient (au mieux) de déterminer des minima locaux. La question de la détermination de minima globaux est difficile et dépasse le cadre que nous nous sommes fixés. Néanmoins, nous décrirons dans la section suivante, un algorithme probabiliste permettant de "déterminer" x un minimum global.

Remarquons aussi que nous avons fait l'hypothèse de différentiabilité. Nous n'en parlerons par ici. Nous commencerons par quelques définition.

3.1 Définition (algorithmes)

Un algorithme est défini par une application A de $\mathbb{R}^n$ dans $\mathbb{R}^n$ permettant la génération d'une suite d'éléments de $\mathbb{R}^n$ par la formule :

$$\begin{cases} x_0 \in \mathbb{R}^n \text{ donné, } k = 0 \text{ étape d'initialisation,} \\ x_{k+1} = A(x_k), k = k + 1 \text{ itération } k \end{cases}$$

Ecrire un algorithme n'est ni plus ni moins que se donner une suite $(x_k)_{k \in \mathbb{N}}$ de $\mathbb{R}^n$; étudier la convergence de l'algorithme, c'est étudier la convergence de la suite $(x_k)_{k \in \mathbb{N}}$ [5].

3.1.1 convergence d'un algorithme

Définition 3.1.1. On dit que l'algorithme A **converge** si la suite $(x_k)_{k \in \mathbb{N}}$ engendrée par l'algorithme converge vers une limite $\hat{x}$.

3.1.2 Taux de convergence d'un algorithme

Définition 3.1.2. Soit $(x_k)_{k \in \mathbb{N}}$ une suite de limite $\hat{x}$ définie par la donnée d'un algorithme convergent A . On dit que la convergence de A est :

- **linéaire** : si l'erreur $e_k = \|x_k - \hat{x}\|$ décroît linéairement :

$$\exists C \in [0, 1[, \exists k_0, \forall k \geq k_0 \quad e_{k+1} \leq C e_k.$$

- **super-linéaire** : si l'erreur e_k décroît de la manière suivante :

$$e_{k+1} \leq \alpha_k e_k.$$

où α_k est une suite positive convergente vers 0. Si α_k est une suite géométrique, la convergence de l'algorithme est dite **géométrique**.

- **d'ordre p** : si l'erreur e_k décroît de la manière suivante :

$$\exists C \geq 0, \exists k_0, \forall k \geq k_0 \quad e_{k+1} \leq C [e_k]^p.$$

Si $p=2$, la convergence de l'algorithme est dit **quadratique**.

- **locale** : si elle n'a lieu que pour des points de départ x_0 dans un voisinage de $\hat{x}$. Dans le cas contraire la convergence est globale.

Remarque 3.1.1. La "classification" précédente des vitesses (taux) de convergence renvoie à la notion de comparaison des fonctions au voisinage de $+\infty$. En effet, si on suppose que l'erreur e_k ne s'annule pas, une convergence linéaire revient à dire que $\frac{e_{k+1}}{e_k} = o(1)$, alors qu'une convergence super-linéaire est équivalente à $\frac{e_{k+1}}{e_k} = o(e_k^{p-2})$. On a bien entendu intérêt à ce que la vitesse de convergence d'un algorithme soit la plus élevée possible (afin d'obtenir la solution avec un minimum d'itérations pour une précision donnée).

3.2 Méthode de gradient

La méthode (ou algorithme) du gradient fait partie de la classe plus grande des méthodes appelées méthodes de descente.

On veut minimiser une fonction f . Pour cela on se donne un point de départ arbitraire x_0 . Pour construire l'itéré suivante x_1 il faut penser qu'on veut rapprocher du minimum de f , on veut donc que $f(x_1) < f(x_0)$. On cherche alors x_1 sous la forme $x_1 = x_0 + \rho_0 d_0$ où $d_0 \in \mathbb{R}^n$ et $\rho_0 > 0$. En pratique donc, on cherche d_0 et ρ_0 pour que $f(x_0 + \rho_0 d_0) < f(x_0)$. Quand d_0 existe on dit que c'est une **direction de descente** et ρ_0 est le **pas de descente**. La direction et le pas de descente peuvent être fixés ou changer à chaque itération. Le schéma général d'une méthode de descente est suivant :

$$\begin{cases} x_0 \text{ donné dans } \mathbb{R}^n \\ x_{k+1} = x_k + \alpha_k d_k, \text{ avec } d_k \in \mathbb{R}^n \text{ et } \alpha_k > 0 \end{cases}$$

où α_k et d_k choisis de telle sorte que $f(x_k + \alpha_k d_k) \leq f(x_k)$.

Une idée pour trouver la direction de descente est de faire un développement de Taylor de f au premier ordre au voisinage de x_{k+1} est donné par :

$$\begin{aligned} f(x_{k+1}) &= f(x_k + \alpha_k d_k) \\ &= f(x_k) + \alpha_k \nabla f(x_k) d_k + o(\alpha_k d_k). \end{aligned}$$

Or, puisque l'on désire avoir : $f(x_{k+1}) \leq f(x_k)$, une solution évidente consiste à prendre :

$$d_k = -\nabla f(x_k).$$

La méthode obtenue ainsi obtenue s'appelle méthode de gradient. Donc le schéma général de la méthode est :

$$\begin{cases} x_0 \text{ donné dans } \mathbb{R}^n \\ x_{k+1} = x_k - \alpha_k \nabla f(x_k), \text{ avec } \alpha_k > 0. \end{cases}$$

Algorithme de Gradient :

1- Initialisation

$k=0$: choix de x_0 et de $\alpha_0 > 0$.

2- Itération k

le pas α_k est choisi constant ou variable

$$x_{k+1} = x_k - \alpha_k \nabla f(x_k).$$

3- critère d'arrêt

si $\|x_{k+1} - x_k\| < \varepsilon$, stop.

sinon, on pose $k = k + 1$ et on retourne à 2.

Théorème 3.2.1. soit f une fonction de classe C^1 de $\mathbb{R}^n$ dans $\mathbb{R}$, $\hat{x}$ est un minimum de f , suppose que :

1) f est α -élliptique, c'est à dire :

$$\exists \alpha > 0, \forall (x, y) \in \mathbb{R}^n \times \mathbb{R}^n : \langle \nabla f(x) - \nabla f(y), x - y \rangle \geq \alpha \|x - y\|^2.$$

2) L'application ∇f est Lipschitzienne, c'est à dire :

$$\exists L > 0, \forall (x, y) \in \mathbb{R}^n \times \mathbb{R}^n : \|\nabla f(x) - \nabla f(y)\| \leq L \|x - y\|.$$

s'il existe deux réels a et b tels que $0 < a < \alpha_k < b < \frac{2\alpha}{L^2}$, pour tout $k \geq 0$, alors la méthode du gradient converge pour tout choix de x_0 de façon géométrique c'est à dire :

$$\exists \beta \in]0, 1[, \|x_k - \hat{x}\| \leq \beta^k \|x_0 - \hat{x}\|.$$

Méthode du gradient à pas fixe

La méthode du gradient à pas fixe est définie par :

$$\begin{cases} x_0 \text{ donné dans } \mathbb{R}^n \\ x_{k+1} = x_k - \alpha \nabla f(x_k). \end{cases}$$

C'est une méthode de descente utilisant un pas fixe $\alpha_k = \alpha$ (indépendant de k). La convergence de la méthode est assurée sous les hypothèse du théorème précédent, en particulier pour un choix de pas α vérifiant :

$$0 < \alpha < \frac{2\alpha}{L^2}.$$

C'est la méthode de gradient la plus simple à utiliser.

Exemple 3.2.1. soit la fonction f définie de $\mathbb{R}$ dans $\mathbb{R}$ par :

$$f(x) = 4x^2 + 6x - 2.$$

L'algorithme de descente est :

$$\begin{cases} x_0 \text{ donné dans } \mathbb{R} \\ x_{k+1} = x_k + \alpha d_k. \end{cases}$$

f est différentiable donc : $d_k = -\nabla f(x) = -f'(x) = -8x - 6$. D'ou

$$\begin{cases} x_0 \text{ donné dans } \mathbb{R} \\ x_{k+1} = x_k - \alpha f'(x_k). \end{cases}$$

calculons le pas α on a :

$$\begin{aligned} |\nabla f(x) - \nabla f(y)| &= |(8x + 6) - (8y - 6)| \\ &= |8x + 6 - 8y - 6| \\ &= |8(x - y)| \\ &\leq 8|x - y|, \end{aligned}$$

donc ∇f est Lipschitzienne de constante $L = 8$. Et

$$\langle \nabla f(x) - \nabla f(y), x - y \rangle = \langle 8(x - y), x - y \rangle$$

$$\begin{aligned} &= 8(x - y)^2 \\ &= 8|x - y|^2 \end{aligned}$$

donc f est α -élliptique avec $\alpha = 8$. Comme $\alpha \in]0, \frac{2\alpha}{L^2}[$ donc $\alpha \in]0, \frac{1}{4}[$.

si on pose $x_0 = -0.5$ $\alpha = 0.1$

Itération 1

$$\begin{aligned} x_1 &= x_0 - \alpha f'(x_0) \\ &= -0.5 - 0.1(8(-0.5) + 6) \\ &= -0.560.2 \\ &= -0.7 \end{aligned}$$

d'ou $x_1 = -0.7$

critère d'arrêt

$$|x_1 - x_0| = |-0.7 + 0.5| = 0.2 < \varepsilon.$$

Stop

Méthode du gradient à pas optimal

Ici on choisit à chaque étape α_k de façon que :

$$f(x_k - \alpha_k \nabla f(x_k)) = \inf_{\alpha \in \mathbb{R}} f(x_k - \alpha \nabla f(x_k)).$$

Théorème 3.2.2. Si f est α -élliptique sur $\mathbb{R}^n$, si ∇f est uniformément lipschitzien de constante de Lipschitz L , l'algorithme de gradient à pas optimal est bien défini et converge vers la solution optimale.

Remarque 3.2.1. les directions de descente sont orthogonales, c'est à dire :

$$\langle \nabla f(x_k), \nabla f(x_{k+1}) \rangle = 0.$$

3.3 Méthode du gradient conjugué

Les méthodes du gradient conjugué sont utilisées pour résoudre les problèmes d'optimisation non linéaires sans contraintes spécialement les problèmes de grandes tailles. On l'utilise aussi pour résoudre les grands systèmes linéaires. Elles reposent

sur le concept des directions conjuguées parce que les gradients successifs sont orthogonaux entre eux et aux directions précédentes. L'idée initiale était de trouver une suite de directions de descente permettant de résoudre le problème

$$\min \{f(x) : x \in \mathbb{R}^n\} \quad (P)$$

Où f est régulière (continument différentiable). Dans ce chapitre on va décrire toutes ces méthodes, mais avant d'accéder à ces dernières, on va d'abord donner le principe général d'une méthode à directions conjuguées.

3.3.1 Le principe général d'une méthode à direction conjuguées

Définition 3.3.1. (conjugaison) : Soit A une matrice symétrique $n \times n$, définie positive. On dit que deux vecteurs x et y de $\mathbb{R}^n$ sont A -conjugués (ou conjugués par rapport à A) s'ils vérifient

$$x^T A y = 0 \quad (3.1)$$

Description de la méthode

Soit $d_0, d_1, \dots, d_n$ une famille des vecteurs A -conjugué. On appelle alors méthode de direction conjuguée toute méthode itérative appliquée à une fonction quadratique strictement convexe de n variables :

$q(x) = \frac{1}{2}x^T A x + b^T x + c$; avec $x \in \mathbb{R}^n$ et $A \in \mathbb{M}_{n \times n}$ est symétrique et définie positive, $b \in \mathbb{R}^n$ et $c \in \mathbb{R}$, conduisant à l'optimum en n étapes au plus. Cette méthode est de la forme suivante :

x_0 point initial

$$x_{k+1} = x_k + \alpha_k d_k \quad (3.2)$$

où α_k est optimal et $d_0, d_1, \dots, d_n$ possédant la propriété d'être mutuellement conjuguées par rapport à la fonction quadratique. Si l'on note $g_k = \nabla q(x_k)$, la

méthode se construit comme suit :

Calcul de α_k : comme α_k minimise q dans la direction d_k , on a, $\forall k$:

$$q'(\alpha_k) = d_k^T \nabla q(x_{k+1}) = 0$$

$$d_k^T \nabla q(x_{k+1}) = d_k^T (Ax_{k+1} + b) = 0$$

soit :

$$d_k^T A(x_k + \alpha_k d_k) + d_k^T b = 0$$

d'où l'on tire :

$$\alpha_k = \frac{-d_k^T (Ax_k + b)}{d_k^T A d_k} \quad (3.3)$$

Comment construire les directions A -conjuguées ? des directions A -conjuguées $d_0, \dots, d_k$ peuvent être générées à partir d'un ensemble de vecteurs linéairement indépendants $\xi_0, \dots, \xi_k$ en utilisant la procédure dite de Gram-Schmidt, de telle sorte que pour tout i entre 0 et k , le sous-espace généré par $d_0, \dots, d_i$ soit égale au sous-espace généré par $\xi_0, \dots, \xi_i$. Alors d_{i+1} est construite comme suit :

$$d_{i+1} = \xi_{i+1} + \sum_{m=0}^i \varphi(i+1) m^{d_m}$$

Nous pouvons noter que si d_{i+1} est construite d'une telle manière, elle est effectivement linéairement indépendante avec $d_0, \dots, d_i$.

En effet, le sous-espace généré par les directions $d_0, \dots, d_i$ est le même que le sous-espace généré par les directions $\xi_0, \dots, \xi_i$, et ξ_{i+1} est linéairement indépendant de $\xi_0, \dots, \xi_i$. ξ_{i+1} ne fait donc pas partie du sous-espace généré par les combinaisons linéaires de la forme $\sum_{m=0}^i \varphi(i+1) m^{d_m}$, de sorte que d_{i+1} n'en fait pas partie non plus et est donc linéairement indépendante des $d_0, \dots, d_i$. Les coefficients $\varphi(i+1)m$, eux sont choisis de manière à assurer la A -conjugaison des $d_0, \dots, d_{i+1}$.

3.3.2 Méthode de gradient conjugué dans le cas quadratique

La méthode du gradient conjugué quadratique est obtenue en appliquant la procédure de Gram-Schmidt aux gradients $\nabla q(x_0), \dots, \nabla q(x_{n-1})$, c'est-à-dire en posant $\xi_0 = -\nabla q(x_0), \dots, \xi_{n-1} = -\nabla q(x_{n-1})$.

En outre, nous avons :

$$\begin{aligned}\nabla q(x) &= Ax + b \\ \text{et } \nabla^2 q(x) &= A.\end{aligned}$$

Notons que la méthode se termine si $\nabla q(x_k) = 0$. La particularité intéressante de la méthode du gradient conjugué est que le membre de droite de l'équation donnant la valeur de d_{k+1} dans la procédure de Gram-Schmidt peut être grandement simplifié. Notons que la méthode du gradient conjugué est inspirée de celle du gradient (plus profonde pente) [21].

Algorithme de la méthode de gradient conjugué pour les fonctions quadratiques

On suppose ici que la fonction à minimiser est quadratique sous la forme : $q(x) = \frac{1}{2}x^T Ax + b^T x + c$. Si l'on note $g_k = \nabla q(x_k)$, l'algorithme prend la forme suivante .

Cet algorithme consiste à générer une suite d'itérés x_k sous la forme :

$$x_{k+1} = x_k + \alpha_k d_k \tag{3.4}$$

L'idée de la méthode est :

1-construire itérativement des directions $d_0, \dots, d_k$ mutuellement conjuguées : A chaque étape k la direction d_k est obtenue comme combinaison linéaire du gradient en x_k et de la direction précédente d_{k-1} c'est-à-dire

$$d_{k+1} = -\nabla q(x_{k+1}) + \beta_{k+1} d_k \tag{3.5}$$

Les coefficients

β_{k+1} étant choisis de telle manière que d_k soit conjuguée avec toutes les directions précédentes. autrement dit :

$$\begin{aligned} d_{k+1}^T Ad_k = 0 &\Rightarrow (-\nabla q(x_{k+1}) + \beta_{k+1} d_k)^T Ad_k = 0 \\ &\Rightarrow -\nabla^T q(x_{k+1}) Ad_k + \beta_{k+1} d_k^T Ad_k = 0 \\ &\Rightarrow \beta_{k+1} = \frac{\nabla^T q(x_{k+1}) Ad_k}{d_k^T Ad_k} = \frac{q_{k+1}^T Ad_k}{d_k^T Ad_k} \end{aligned} \quad (3.6)$$

2-déterminer le pas α_k : En particulier, une façon de choisir α_k consiste à résoudre le problème d'optimisation unidimensionnelle suivant :

$$\alpha_k = \min q(x_k + \alpha d_k), \alpha > 0 \quad (3.7)$$

on en déduit :

$$\alpha_k = \frac{-d_k^T g_k}{d_k^T Ad_k} = -\frac{1}{Ad_k} g_k \frac{d_k^T}{d_k^T} = \frac{-d_k^T g_k}{d_k^T Ad_k} \quad (3.8)$$

Le pas α_k obtenu ainsi s'appelle le pas optimal.

Algorithme 3.3.1. (Algorithme du gradient conjugué "quadratique")

Etape 0 : (initialisation) Soit x_0 le point de départ, $g_0 = \nabla q(x_0) = Ax_0 + b$, poser $d_0 = -g_0$

poser $k = 0$ et aller à l'étape 1.

Etape 1 :

si $g_k = 0$: *STOP* ($x^* = x_k$). "Test d'arrêt" si non aller à étape 2.

Etape 2 :

Prendre $x_{k+1} = x_k + \alpha_k d_k$ avec :

$$\begin{aligned} \alpha_k &= \frac{-d_k^T g_k}{d_k^T Ad_k} \\ d_{k+1} &= -g_{k+1} + \beta_{k+1} d_k \\ \beta_{k+1} &= \frac{g_{k+1}^T Ad_k}{d_k^T Ad_k} \end{aligned}$$

Poser $k = k + 1$ et aller à l'étape 1.

La validité de l'algorithme du gradient conjugué linéaire

On va maintenant montrer que l'algorithme ci-dessus définit bien une méthode de directions conjuguées [16].

Théorème 3.3.1. *A une itération k quelconque de l'algorithme où l'optimum de $q(x)$ n'est pas encore atteint (c'est-à-dire $g_i \neq 0, i = 0, \dots, k$) on a :*

a)

$$\alpha_k = \frac{g_k^T g_k}{d_k^T A d_k} \neq 0 \quad (3.9)$$

b)

$$\beta_{k+1} = \frac{g_k^T [g_{k+1} - g_k]}{g_k^T g_k} \quad (3.10)$$

$$= \frac{g_{k+1}^T g_{k+1}}{g_k^T g_k} \quad (3.11)$$

c) Les directions $d_0, d_1, \dots, d_{k+1}$ engendrées par l'algorithme sont mutuellement conjuguées [22].

Preuve 3.3.1. [16] On raisonne par récurrence sur k en supposant que $d_0, d_1, \dots, d_k$ sont mutuellement conjuguées.

a) Montrons d'abord l'équivalence de (3.9) et de (3.10) On a :

$$d_k = -g_k + \beta_k d_{k-1}$$

Donc (3.9) s'écrit :

$$\begin{aligned} \alpha_k &= \frac{-d_k^T g_k}{d_k^T A d_k} \\ &= \frac{-[-g_k + \beta_k d_{k-1}]^T g_k}{d_k^T A d_k} \\ &= \frac{g_k^T g_k}{d_k^T A d_k} - \beta_k \frac{d_{k-1}^T g_k}{d_k^T A d_k} \end{aligned}$$

Comme $(d_0, d_1, \dots, d_{k-1})$ sont mutuellement conjuguées, x_k est l'optimum de $q(x)$ sur la variété $\mathcal{V}^k$ passant par x_0 et engendrée par $(d_0, d_1, \dots, d_{k-1})$.

Donc $d_{k-1}^T g_k = 0$ d'où l'on déduit (3.10).

b) Pour démontrer (3.11) remarquons que :

$$\begin{aligned} g_{k+1} - g_k &= A(x_{k+1} - x_k) = \alpha_k Ad_k \\ \Rightarrow Ad_k &= \frac{1}{\alpha_k} [g_{k+1} - g_k] \end{aligned}$$

On a alors :

$$g_{k+1}^T Ad_k = \frac{1}{\alpha_k} g_{k+1}^T [g_{k+1} - g_k]$$

et en utilisant (3.10)

$$\alpha_k = \frac{g_k^T g_k}{d_k^T Ad_k}$$

il vient

$$\begin{aligned} g_{k+1}^T Ad_k &= \frac{d_k^T Ad_k}{g_k^T g_k} g_{k+1}^T [g_{k+1} - g_k] \\ \Rightarrow \frac{g_{k+1}^T Ad_k}{d_k^T Ad_k} &= \frac{g_{k+1}^T [g_{k+1} - g_k]}{g_k^T g_k} \end{aligned}$$

$$\beta_{k+1} = \frac{g_{k+1}^T Ad_k}{d_k^T Ad_k} = \frac{g_{k+1}^T [g_{k+1} - g_k]}{g_k^T g_k}$$

ce qui démontre (3.11).

Alors

$$g_{k+1}^T g_k = 0$$

car

$$g_k = d_k - \beta_k d_{k-1}.$$

Appartient au sous-espace engendré par $(d_0, d_1, \dots, d_k)$ et que g_{k+1} est orthogonal à ce sous-espace.

c) Montrons enfin que d_{k+1} est conjuguée par rapport à $(d_1, d_2, \dots, d_k)$.

On a bien $d_{k+1}^T Ad_k$ car, en utilisant $d_{k+1} = -g_{k+1} + \beta_k d_k$ on aura :

$$\begin{aligned} d_{k+1}^T Ad_k &= (-g_{k+1} + \beta_{k+1} d_k)^T Ad_k \\ &= -g_{k+1}^T Ad_k + \beta_{k+1} d_k^T Ad_k \\ &= -g_{k+1}^T Ad_k + \frac{g_{k+1}^T Ad_k}{d_k^T Ad_k} d_k^T Ad_k \\ &= 0 \end{aligned}$$

Vérifions maintenant que :

$$d_{k+1}^T Ad_i = 0 \text{ pour } i = 0, 1, \dots, k-1.$$

On a :

$$d_{k+1}^T Ad_i = -g_{k+1}^T Ad_i + \beta_{k+1} d_k^T Ad_i.$$

Le seconde terme est nul l'hypothèse de récurrence $((d_0, d_1, \dots, d_k)$ sont mutuellement conjuguées).

Montrons qu'il en est de même du premier terme. Puisque $x_{i+1} = x_i + \alpha_i d_i$ et que $\alpha_i \neq 0$ on a :

$$\begin{aligned} Ad_i &= \frac{1}{\alpha_i} (Ax_{i+1} - Ax_i) \\ &= \frac{1}{\alpha_i} (g_{i+1} - g_i) \end{aligned}$$

En écrivant :

$$\begin{aligned} g_{i+1} &= d_{i+1} - \beta_i d_i \\ g_i &= d_i - \beta_{i-1} d_{i-1} \end{aligned}$$

on voit que Ad_i est combinaison linéaire de d_{i+1} , d_i et de d_{i-1} seulement.

Mais puisque $(d_0, d_1, \dots, d_k)$ sont mutuellement conjuguées, on sait que le point x_{k+1} est l'optimum de $q(x)$ sur la variété V^{k+1} , engendrée par $(d_0, d_1, \dots, d_k)$.

Donc g_{k+1} est orthogonal au sous-espace engendré par $(d_0, d^1, \dots, d_k)$ et comme Ad_i appartient à ce sous-espace pour $i = 0, 1, \dots, k-1$, on en déduit $g_{k+1}^T Ad_i = 0$ ce qui achève la démonstration.

Théorème 3.3.2. Dans ce cas d_k est une direction de descente puisque

$$\begin{aligned} d_k^T \nabla q(x_k) &= (-\nabla q(x_k) + \beta_{k+1} d_{k-1}^T \nabla q(x_k)) \\ &= -\nabla q(x_k)^T \nabla q(x_k) + \beta_{k+1} d_{k+1}^T \nabla q(x_k) \\ &= -\|\nabla q(x_k)\|^2 (\text{car } d_{k-1}^T \nabla q(x_k) = 0) \\ d_k^T \nabla q(x_k) &= -\|\nabla q(x_k)\|^2 < 0 \end{aligned}$$

Les avantages de la méthode du gradient conjugué quadratique

1. la consommation mémoire de l'algorithme est minime : on doit stocker les quatre vecteurs x_k, g_k, d_k, Ad_{k+1}
(bien sur x_{k+1} prend la place de x_k au niveau de son calcul avec des remarques analogues pour $g_{k+1}, d_{k+1}, Ad_{k+1}$) et les scalaires α_k, β_{k+1} .
2. L'algorithme du gradient conjugué linéaire est surtout utile pour résoudre des grands systèmes creux, en effet il suffit de savoir appliquer la matrice A à un vecteur.
3. La convergence peut être assez rapide : si A admet seulement r ($r < n$) valeurs propres distincts la convergence a lieu en au plus r itérations.

Différentes formules de β_{k+1} dans le cas quadratique

Formule de Hestenes-Stiefel Cette méthode été découverte en 1952 par Hestenes et Stiefels

$$\beta_{k+1}^{HS} = \frac{g_{k+1}^T [g_{k+1} - g_k]}{d_k^T [g_{k+1} - g_k]} \quad (3.12)$$

Formule de Polak-Ribière-Polyak Cette méthode été découverte en 1969 par Polak, Ribière et Poyak

$$\beta_{k+1}^{PRP} = \frac{g_{k+1}^T [g_{k+1} - g_k]}{\|g_k\|^2} \quad (3.13)$$

Formule de Fletcher-Reeves Cette méthode été découverte en 1964 par Fletcher et Reeves

$$\beta_{k+1}^{FR} = \frac{\|g_{k+1}\|^2}{\|g_k\|^2} \quad (3.14)$$

Formule de la descente conjuguée Cette méthode été découverte en 1987 par Fletcher

$$\beta_{k+1}^{CD} = \frac{\|g_{k+1}\|^2}{-d_k^T g_k} \quad (3.15)$$

formule de Dai-Yuan

$$\beta_{k+1}^{DY} = \frac{\|g_{k+1}\|^2}{d_k^T g_k} \quad (3.16)$$

[17] [14] [9] [13]

Remarque 3.3.1. Dans le cas quadratique on a vu que :

$$\beta_{k+1}^{HS} = \beta_{k+1}^{PRP} = \beta_{k+1}^{FR} = \beta_{k+1}^{DY}.$$

Dans le cas non quadratique, ces quantités ont en général des valeurs différentes.

3.3.3 Méthode du gradient conjugué dans le cas non quadratique

On s'intéresse dans cette section à la minimisation d'une fonction $f : \mathbb{R}^n \rightarrow \mathbb{R}$ non nécessairement quadratique :

$$\min f(x); \quad x \in \mathbb{R}^n \quad (3.17)$$

Les méthodes du gradient conjugué génèrent des suites $\{x_k\}_{k=0,1,2,\dots}$ de la forme suivante :

$$x_{k+1} = x_k + \alpha_k d_k \quad (3.18)$$

Le pas $\alpha_k \in \mathbb{R}$ étant déterminé par une recherche linéaire. La direction d_k est définie par la formule de récurrence suivante ($\beta_k \in \mathbb{R}$)

$$d_k = \begin{cases} -g_k & \text{si } k = 1 \\ -g_k + \beta_k d_{k-1} & \text{si } k \geq 2 \end{cases} \quad (3.19)$$

Ces méthodes sont des extensions de la méthode du gradient conjugué linéaire du cas quadratique, si

β_k prend l'une des valeurs

$$\beta_k^{FR} = \frac{\|g_k\|^2}{\|g_{k-1}\|^2}$$

$$\beta_k^{CD} = \frac{\|g_k\|^2}{-d_{k-1}^T g_{k-1}}$$

où $y_{k-1} = g_k - g_{k-1}$.

Algorithme 3.3.2. *Algorithme de La méthode du gradient conjugué pour les fonctions quelconques*

Etape 0 : (*initialisation*) Soit x_0 le point de départ, $g_0 = \nabla f(x_0)$, poser $d_0 = -g_0$

Poser $k = 0$ et aller à l'étape 1.

Etape 1 : Si $g_k = 0$:

STOP ($x^* = x_k$). "Test d'arrêt"

Si non aller à l'étape 2.

Etape 2 : Définir $x_{k+1} = x_k + \alpha_k d_k$ avec :

α_k : calculer par la recherche linéaire

$$d_{k+1} = -g_{k+1} + \beta_{k+1} d_k$$

où

β_{k+1} : défini selon la méthode.

Poser $k = k + 1$ et aller à l'étape 1.

Remarque 3.3.2. Dans le cas quadratique avec recherche linéaire exacte, on a vu que $\beta_k^{PRP} = \beta_k^{FR} = \beta_k^{CD} = \beta_k^{DY}$.

3.4 Méthode de Newton

La méthode de Newton n'est pas une méthode d'optimisation à proprement parler. C'est en réalité une méthode utilisée pour résoudre des équations non linéaires de la forme $f(x) = 0$ où f est une fonction de $\mathbb{R}^n$ dans $\mathbb{R}$. Nous allons d'abord la décrire puis montrer comment on peut l'appliquer à la recherche de minimum.

3.4.1 Principe de la Méthode de Newton

Considérons le problème d'optimisation sans contraintes (P)

$$(P) : \min_{x \in \mathbb{R}^n} f(x)$$

où

$$f : \mathbb{R}^n \rightarrow \mathbb{R}.$$

Le principe de la méthode de Newton consiste à minimiser successivement les approximations du second ordre de f , plus précisément si

$$f(x) = f(x_k) + \nabla f(x_k)^T(x - x_k) + \frac{1}{2}(x - x_k)^T \nabla^2 f(x_k)(x - x_k) + o(\|x - x_k\|^2),$$

posons

$$q(x) = f(x_k) + \nabla f(x_k)^T(x - x_k) + \frac{1}{2}(x - x_k)^T \nabla^2 f(x_k)(x - x_k).$$

Soit x_{k+1} l'optimum q , alors il vérifie $\nabla q(x+1) = 0$, soit en remplaçant :

$$\nabla f(x_k) + \nabla^2 f(x_k)(x_{k+1} - x_k) = 0,$$

donc

$$x_{k+1} = x_k - [\nabla^2 f(x_k)]^{-1} \nabla f(x_k)$$

Exemple 3.4.1. Soit $f : \mathbb{R}^2 \rightarrow \mathbb{R}$ une fonction définie par :

$$f(x_1, x_2) = 2x_1^2 + 2x_2^2 - 2x_1 - 4x_2$$

Soit $x^0 = \begin{pmatrix} x_1^0 \\ x_2^0 \end{pmatrix} = \begin{pmatrix} 4 \\ 6 \end{pmatrix}$, on trouve les itérations 1 et 2 par la méthode de Newton d'optimisation

$$\begin{cases} x^0 \text{ donné} \\ x^{k+1} = x^k - (\nabla^2 f(x^k))^{-1} \nabla f(x^k) \end{cases}$$

On a

$$\nabla f(x_1^0, x_2^0) = \begin{pmatrix} 4x_1^0 - 2 \\ 4x_2^0 - 4 \end{pmatrix} = \begin{pmatrix} 14 \\ 20 \end{pmatrix}, \nabla^2 f(x_1^0, x_2^0) = \begin{pmatrix} 4 & 0 \\ 0 & 4 \end{pmatrix}.$$

On pose $d^k = -(\nabla^2 f(x_1^k, x_2^k))^{-1} \nabla f(x_1^k, x_2^k)$ alors :

$$\begin{cases} x^0 \text{ donné} \\ x^{k+1} = x^k + d^k \end{cases}$$

Donc

$$(\nabla^2 f(x_1^0, x_2^0))d^0 = -\nabla f(x_1^0, x_2^0)$$

alors

$$\begin{pmatrix} 4 & 0 \\ 0 & 4 \end{pmatrix} \begin{pmatrix} d_1^0 \\ d_2^0 \end{pmatrix} = \begin{pmatrix} -14 \\ -20 \end{pmatrix} \Rightarrow \begin{cases} 4d_1^0 = -14 \\ 4d_2^0 = -20 \end{cases}$$

D'où $(d_1^0 = \frac{-7}{2}, d_2^0 = -5) \Rightarrow d^0 = \begin{pmatrix} \frac{-7}{2} \\ -5 \end{pmatrix}$, ce qui donne :

$$x^1 = x^0 + d^0 = \begin{pmatrix} 4 \\ 6 \end{pmatrix} + \begin{pmatrix} \frac{-7}{2} \\ -5 \end{pmatrix} = \begin{pmatrix} \frac{1}{2} \\ 1 \end{pmatrix}.$$

$x^2 = x^1 + d^1 = \begin{pmatrix} \frac{1}{2} \\ 1 \end{pmatrix} + \begin{pmatrix} 0 \\ 0 \end{pmatrix} = \begin{pmatrix} \frac{1}{2} \\ 1 \end{pmatrix}$. Ainsi le point $\begin{pmatrix} \frac{1}{2} \\ 1 \end{pmatrix}$ est la solution approché de la fonction f .

3.5 Méthode de quasi Newton ou quasi-Newtonniennes

Une méthode de quasi Newton est une méthode de type :

$$\begin{cases} x_{k+1} = x_k + \lambda_k d_k, \\ d_k = -B_k g_k. \end{cases} \quad (3.20)$$

Où B_k est une matrice destinée à approcher l'inverse du Hessien de f (respectivement le Hessien de f) en x_k . Le problème posé est : quelle stratégie à adopter pour faire cette approximation ? On peut par exemple poser $B_0 = I$, mais comment ensuite mettre à jour l'approximation B_k au cours des itérations ? L'idée est la suivante : on sait que au point x_k , le gradient et le hessien de f vérifient la relation

$$g_{k+1} = g_k + \nabla^2 f(x_k)(x_{k+1} - x_k) + \epsilon(x_{k+1} - x_k).$$

Si on suppose que l'approximation quadratique est bonne, on peut alors négliger le reste et considérer que l'on a

$$g_{k+1} \simeq g_k + \nabla^2 f(x_k)(x_{k+1} - x_k);$$

cela conduit à la notion de relation de quasi-Newton

Définition 3.5.1. On dit que les matrice B_{k+1} et $\nabla^2 f(x_{k+1})$ vérifient une relation de quasi Newton si on a

$$\nabla^2 f(x_{k+1})(x_{k+1} - x_k) = \nabla f(x_{k+1}) - \nabla f(x_k),$$

ou

$$x_{k+1} - x_k = B_{k+1}[\nabla f(x_{k+1}) - \nabla f(x_k)].$$

Il reste un problème à résoudre : comment mettre à jour B_k tout en assurant $B_k > 0$? C'est ce que nous allons voir maintenant

3.5.1 Formules de mise à jour de l'approximation du hessien :

Le principe de la mise à jour consiste, à une itération donnée de l'algorithme

$$\begin{cases} x_{k+1} = x_k + \lambda_k d_k \\ d_k = -B_k g_k, \end{cases}$$

à appliquer une formule du type

$$B_{k+1} = B_k + \nabla_k, \quad (3.21)$$

avec ∇_k symétrique, assurant la relation de quasi-Newton

$$x_{k+1} - x_k = B_{k+1}[g_{k+1} - g_k],$$

ainsi que B_{k+1} , sous l'hypothèse que B_k . La formule (3.21) permet d'utiliser les nouvelles informations obtenues lors de l'étape k de l'algorithme, c'est à dire essentiellement le gradient $g_{k+1} = \nabla f(x_{k+1})$ au point x_{k+1} , obtenu par recherche linéaire (exacte ou approchée) dans la direction d_k . Il existe différentes formules du type (3.20). Suivant que ∇_k est de rang un ou deux, on parlera de correction de rang un ou de rang deux.

3.5.2 Méthode de correction de rang un

Étant donné que $[\nabla^2 f(x_k)]^{-1}$ est symétrique, la formule de mise à jour de l'approximation du Hessien B_k est la suivante :

$$B_{k+1} = B_k + a_k u_k u_k^T, \quad u_k \in \mathbb{R}$$

donc la condition de quasi -Newton s'écrit comme suit

$$s_k = (B_k + a_k u_k u_k^T) y_k,$$

ou encore

$$s_k - B_k y_k = a_k u_k u_k^T y_k.$$

D'où l'on déduit que u_k est proportionnel à $s_k - B_k y_k$, avec un facteur qui peut être pris en compte dans a_k . Un choix évident pour vérifier cette dernière équation est de prendre $u_k = s_k - B_k y_k$ et a_k tel que $a_k(u_k^T y_k) = 1$; on obtient :

$$B_{k+1} = B_k + \frac{(s_k - B_k y_k)(s_k - B_k y_k)^T}{(s_k - B_k y_k)^T y_k} \quad (3.22)$$

Théorème 3.5.1. Si f est quadratique, de matrice Hessienne $\nabla^2 f(x)$ définie positive et si $s_1, s_2, \dots, s_n$ sont des vecteurs indépendants, alors la méthode de correction de rang un converge au plus dans $(n+1)$ itérations et $(B_{n+1})^{-1} = \nabla^2 f(x)$.

Avantages : Cette méthode présente l'avantage, que le point x_{k+1} n'a pas besoin d'être choisi comme le minimum exact, c'est à dire qu'on n'a pas besoin d'effectuer des recherches linéaire exactes. **Inconvénients :** Même si la fonction est quadratique, et même si son Hessien est défini positif, il se peut que la matrice B_k ne soit pas définie positive. Le dénominateur $(s_k - B_k y_k)^T y_k$ peut devenir nul ou très petit, ce qui rend le procédé instable.

3.5.3 Méthode de Davidon Fletcher Powell (DFP)

Cette méthode a été proposée par Davidon en 1959 et développé plus tard en 1963 par Fletcher. La formule de mise à jour de DFP est une formule de correction de rang deux. De façon plus précise construisons B_{k+1} en fonction de B_k de la forme :

$$B_{k+1} = B_k + A_k + \Delta_k, \quad (3.23)$$

avec Δ_k et A_k deux rang un tel que

$$A_k = a_k u_k u_k^T, \quad \Delta_k = b_k v_k v_k^T,$$

a_k, b_k sont des constantes, u_k, v_k sont deux vecteurs de $\mathbb{R}^n$. B_{k+1} doit satisfaire la condition quasi-Newton c'est à dire

$$x_{k+1} - x_k = B_{k+1}[g_{k+1} - g_k].$$

Si on pose par suite

$$s_k = x_{k+1} - x_k, \quad y_k = g_{k+1} - g_k,$$

donc

$$s_k = B_{k+1}y_k = (B_k + a_k u_k u_k^T + b_k v_k v_k^T)y_k, \quad (3.24)$$

par suite

$$a_k u_k u_k^T y_k + b_k v_k v_k^T y_k = s_k - B_k y_k.$$

Un choix évident pour satisfaire cette équation est de prendre

$$u_k = s_k, \quad v_k = B_k y_k, \quad a_k u_k^T y_k = 1, \quad b_k v_k^T y_k = -1,$$

d'où

$$B_{k+1} = B_k + \frac{s_k s_k^T}{s_k^T y_k} = \frac{B_k y_k y_k^T B_k}{y_k^T B_k y_k} \quad (3.25)$$

Remarque 3.5.1. Le résultat suivant montre que sous certaines conditions, la formule (3.25) conserve la définie positivité des matrices B_k .

Théorème 3.5.2. On considère la méthode définie par

$$\begin{cases} x_{k+1} = x_k + \lambda d_k \\ d_k = -B_k g_k, \end{cases}$$

où λ_k optimal, $B_0 > 0$ est donnée ainsi que x_0 , alors les matrices B_k sont définies positives.

Théorème 3.5.3. Appliqué à une forme quadratique f , l'algorithme DFP décrit par la relation

$$B_{k+1} = B_k + \frac{s_k s_k^T}{s_k^T y_k} - \frac{B_k y_k y_k^T B_k}{y_k^T B_k y_k},$$

engendre des directions conjuguées $d_1, d_2, \dots, d_k$ vérifiant

$$\begin{aligned} d_i^T \nabla^2 f(x) d_j &= 0, & 1 \leq i \leq j \leq k, \\ B_{k+1} \nabla^2 f(x) d_i &= d_i, & 1 \leq i \leq k. \end{aligned} \quad (3.26)$$

La méthode DFP se comporte donc dans le cas quadratique, comme une méthode de directions conjuguées. On peut aussi remarquer qu'on a pour $k = n - 1$ la relation

$$B_{k+1} = B_k + \frac{s_k s_k^T}{s_k^T y_k} - \frac{B_k y_k y_k^T B_k}{y_k^T B_k y_k},$$

engendre des directions conjuguées $d_1, d_2, \dots, d_k$ vérifiant

$$\begin{aligned} d_i^k \nabla^2 f(x) d_j &= 0, & 1 \leq i \leq j \leq k, \\ B_{k+1} \nabla^2 f(x) d_i &= d_i, & 1 \leq i \leq k. \end{aligned} \tag{3.27}$$

La méthode DFP se comporte donc dans le cas quadratique, comme une méthode de directions conjuguées. On peut aussi remarquer qu'on a pour $k = n - 1$ la relation

$$B_{n+1} \nabla^2 f(x) d_i = d_i, \quad i = 1, \dots, n - 1.$$

Et comme les d_i sont linéairement indépendants (car mutuellement conjugués) on en déduit que $B_{n+1} = \nabla^2 f(x)^{-1}$.

Avantages :

1. Pour des fonctions quadratiques (avec une recherche linéaire exacte) :
 - . L'algorithme converge dans au plus n étapes avec $B_{n+1} = \nabla^2 f(x)^{-1}$.
 - . Elles engendrent des directions conjuguées.
2. Pour les fonctions quelconques :
 - . La matrice B_k reste définie positive, ce qui est nécessaire pour que la direction soit une direction de descente.

Inconvénients :

- . La méthode DFP est sensible à la précision de la recherche linéaire.

3.5.4 Méthode de Broyden, Fletcher, Goldfarb et Shanno (BFGS)

La formule de mise à jour de Broyden, Fletcher, Goldfarb et Shanno (BFGS) est une formule de correction de rang deux, qui s'obtient à partir de la formule

DFP en intervertissant les rôles de s_k et y_k . La formule obtenue permet de mettre à jour une approximation B_k de Hessien lui même et non de son inverse comme dans le cas de la méthode DFP. On exigera que posée dans les mêmes propriétés, à savoir B_{k+1} reste définie positive si B_k l'est et bien sur l'équation d'approximation de quasi-Newton doit être vérifiée, c'est à dire

$$B_{k+1}s_k = y_k.$$

On obtient donc

$$B_{k+1} = B_k + \frac{y_k y_k^T}{y_k^T s_k} - \frac{B_k s_k s_k^T B_k}{s_k^T B_k s_k} \quad (3.28)$$

1. Notons que la direction d_k est obtenue par une résolution d'un système linéaire. En particulier la mise à jour de B_k est faite directement sur le facteur de Cholesky C_k où $B_k = C_k C_k^T$ ce qui ramène le calcul de d_k au même coût que pour la formule de DFP.
2. La méthode BFGS possède les mêmes propriétés que la méthode DFP dans le cas quadratique. Les directions engendrées sont conjuguées. Cette méthode est reconnue comme étant beaucoup moins sensible que la méthode DFP aux imprécisions dans la recherche linéaire, du point de vue de vitesse de convergence. Elle est donc tout à fait adaptée quand la recherche linéaire est faite de façon économique, avec par exemple la règle de Goldstein ou la règle de wolfe et Powell.
3. La relation (3.28) permet de construire une approximation de la matrice Hessienne elle même (et non pas son inverse).

En effet : Posons

$$C_k = \frac{y_k y_k^T}{y_k^T s_k} - \frac{B_k s_k s_k^T B_k}{s_k^T B_k s_k} \quad (3.29)$$

Nous avons

$$\nabla^2 f(x_{k+1}) = [B_{k+1}]^{-1} = [B_k + C_k]^{-1}.$$

Par application de la formule de Sherman-Morrison-Woodbury suivante

$$(A + ab^T)^{-1} = A^{-1} - \frac{A^{-1}ab^T A^{-1}}{1 + b^T A^{-1}a}. \quad (3.30)$$

Où A est une matrice inversible, et b est un vecteur de $\mathbb{R}^n$, et en supposant que $b^T A^{-1} a \neq -1$, alors on a

$$\nabla^2 f(x_{k+1}) = [B_{k+1}]^{-1} = [B_k + C_k]^{-1} = \left[B_k + \frac{y_k y_k^T}{y_k^T s_k} - \frac{B_k s_k s_k^T B_k}{s_k^T B_k s_k} \right]^{-1}$$

Posons

$$A = B_k + \frac{y_k y_k^T}{y_k^T s_k}, \quad a = -\frac{B_k s_k}{s_k^T B_k s_k}, \quad b^T = s_k^T B_k,$$

donc

$$\nabla^2 f(x_{k+1}) = \left[B_k + \frac{y_k y_k^T}{y_k^T s_k} \right]^{-1} + \frac{\left[B_k + \frac{y_k y_k^T}{y_k^T s_k} \right]^{-1} \frac{B_k s_k}{s_k^T B_k s_k} s_k^T B_k \left[B_k + \frac{y_k y_k^T}{y_k^T s_k} \right]^{-1}}{1 - s_k^T B_k \left[B_k + \frac{y_k y_k^T}{y_k^T s_k} \right]^{-1} \frac{B_k s_k}{s_k^T B_k s_k}}, \quad (3.31)$$

on doit calculer $\left[B_k + \frac{y_k y_k^T}{y_k^T s_k} \right]^{-1}$ pour cela applique la formule de Sherman-Morrison-Woodbury une deuxième fois, on pose :

$$A = B_k, \quad a = \frac{y_k}{y_k^T s_k}, \quad b^T = y_k^T$$

$$\begin{aligned} \left[B_k + \frac{y_k y_k^T}{y_k^T s_k} \right]^{-1} &= [B_k]^{-1} - \frac{[B_k]^{-1} \frac{y_k}{y_k^T s_k} y_k^T [B_k]^{-1}}{1 + y_k^T [B_k]^{-1} \frac{y_k}{y_k^T s_k}} \\ &= [B_k]^{-1} - \frac{[B_k]^{-1} y_k y_k^T [B_k]^{-1}}{y_k^T s_k + y_k^T [B_k]^{-1} y_k}. \end{aligned}$$

Remplaçons cette dernière dans la formule (3.5.4) et d'après un calcul on obtient

$$\nabla^2 f(x_{k+1}) = [B_{k+1}]^{-1} \quad (3.32)$$

$$\begin{aligned} &= [B_k]^{-1} + \frac{y_k y_k^T}{y_k^T s_k} - \frac{[B_k]^{-1} \frac{y_k}{y_k^T s_k} y_k^T [B_k]^{-1}}{1 + y_k^T [B_k]^{-1} \frac{y_k}{y_k^T s_k}} \\ &= \nabla^2 f(x_k) + \frac{y_k y_k^T}{y_k^T s_k} - \frac{\nabla^2 f(x_k) s_k s_k^T \nabla^2 f(x_k)}{s_k^T \nabla^2 f(x_k) s_k}. \end{aligned}$$

Algorithme de Broyden, Fletcher, Goldfarb et Shanno

1. Choisir x_0 et $\nabla^2 f(x_0)$ définie positive quelconque (par exemple $\nabla^2 f(x_0) = I$)

2. A l'itération k , calculer la direction de déplacement

$$d_k = -\nabla^2 f(x_k)^{-1} \nabla f(x_k),$$

déterminer le pas optimal λ_k et poser

$$x_{k+1} = x_k + \lambda_k d_k$$

3. Poser $s_k = \lambda_k d_k$ et $y_k = \nabla f(x_{k+1}) - \nabla f(x_k)$ puis calculer

$$\nabla^2 f(x_{k+1}) = \nabla^2 f(x_k) + \frac{y_k y_k^T}{y_k^T s_k} - \frac{\nabla^2 f(x_k) s_k s_k^T \nabla^2 f(x_k)}{s_k^T \nabla^2 f(x_k) s_k} \quad (3.33)$$

4. Faire $k + 1 \rightarrow k$. Retourner en 2 sauf si le critère d'arrêt est vérifié. Comme critère d'arrêt on retiendra par exemple $\|g_{k+1}\| < \epsilon$.

3.5.5 Les méthodes de classe Broyden

Une méthode de classe Broyden est une méthode de quasi-Newton où l'approximation de l'inverse du Hessien prend la formule suivant :

$$B_{k+1} = B_k - \frac{B_k s_k s_k^T B_k}{s_k^T B_k s_k} + \frac{y_k y_k^T}{y_k^T s_k} + \phi (s_k^T B_k s_k) v_k v_k^T, \quad (3.34)$$

tel que

$$\phi \in [0, 1], v_k^T = \left[\frac{y_k}{y_k^T s_k} - \frac{B_k s_k}{s_k^T B_k s_k} \right]$$

. Si $\phi = 0$ on obtient la méthode BFGS, car la formule (3.32) devient

$$B_{k+1} = B_k - \frac{B_k s_k s_k^T B_k}{s_k^T B_k s_k} + \frac{y_k y_k^T}{y_k^T s_k},$$

c'est exactement la formule de l'approximation de la méthode BFGS.

. Si $\phi = 1$ on obtient la méthode DFP.

3.6 Méthode de relaxation

La dernière méthode que nous présentons permet de ramener un problème de minimisation dans $\mathbb{R}^n$ à la résolution successive de n problèmes de minimisation dans $\mathbb{R}$ (à chaque itération). On cherche à minimiser $J : \mathbb{R}^n \rightarrow \mathbb{R}$; posons $X = (x_1, \dots, x_n)$. Le principe de la méthode est le suivant : étant donné un itéré X^k de coordonnées $(x_1^k, \dots, x_n^k)$, on fixe toutes les composantes sauf la première et on minimise sur la première :

$$\min f(x, x_2^k, x_3^k, \dots, x_n^k), x \in \mathbb{R}.$$

On obtient ainsi la première coordonnée de l'itéré suivant X^{k+1} que l'on note x_1^{k+1} ; on peut, pour effectuer cette minimisation dans $\mathbb{R}$, utiliser par exemple la méthode de Newton dans $\mathbb{R}$. On recommence ensuite en fixant la première coordonnée à X_1^{k+1} et les $n - 2$ dernières comme précédemment. On minimise sur la deuxième coordonnée et ainsi de suite. L'algorithme obtenu est le suivant :

Méthode de relaxation successive

1- Initialisation $k = 0$: choix de $X^0 \in \mathbb{R}^n$.

2- Itération k

pour i variant de 1 à n , on calcule la solution x_i^{k+1} de

$$\min J(x_1^{k+1}, x_2^{k+1}, \dots, x_{i-1}^{k+1}, x, x_{i+1}^k, x_n^k), \in \mathbb{R}.$$

3- Critère d'arrêt Si $\|x_{k+1} - x_k\| < \varepsilon$, STOP

Sinon, on pose $k = k + 1$ et on retourne à 2.

3.7 Étude comparative des méthodes d'optimisation sans contraintes

Dans cette partie, nous présentons une étude comparative entre quelques méthodes de l'optimisation sans contraintes, en étudiant le problème quadra-

tique suivant :

$$\text{minimiser } f(x) = x_1^2 + x_2^2 - 2x_2 - x_1x_2$$

La méthode de gradient :

Pour résoudre ce problème on prend comme point de départ $x_0 = (2, 3)$ d'arrêt $\|\nabla f(x_k)\| < 0.01$. Le résultat est donné dans le tableau suivant :

k	x_k	$f(x_k)$	$d_k = -\nabla f(x_k)$	ρ_k	$\ \nabla f(x_k)\ $	x_{k+1}
1	(5.00,7.00)	5.00	(2.00,3.00)	0.92	3.60	(1.16,1.24)
2	(1.16,1.24)	-1.03	(1.08,-0.68)	0.56	1.27	(0.56,1.62)
3	(0.56,1.62)	-1.20	(-0.5,0.68)	0.3	0.84	(0.71,1.42)
4	(0.71,1.42)	-0.24	(0.00,0.13)	0.04	0.01	(-)

Exécution de la méthode du gradient pour le problème (P) depuis $x_0 = (3, 4)$

On remarque que la méthode du gradient n'a pas eu besoin de beaucoup d'itération pour trouver la solution optimale. Cela est dû au fait que la direction donnée menant au minimum est proche de la direction donnée par le gradient, permettant à la méthode de progresser rapidement. Si nous résolvons le même problème en partant du point $(3, 25)$ nous obtenons le deuxième tableau :

Exécution de la méthode du gradient pour le problème (P) depuis $x_0 = (3, 25)$

On remarque que la méthode a eu besoin de presque deux fois plus d'itération que lors de la première exécution. Il s'ensuit qu'il y a une forte dépendance du comportement de la méthode du gradient à son point de départ. En effet, une conséquence du raisonnement précédent est que, pour un problème donné, il est possible qu'elle converge vite depuis un point initial et exerce une lenteur depuis un autre point. Nous exécutons maintenant la méthode du gradient

k	x_k	$f(x_k)$	$d_k = -\nabla f(x_k)$	ρ_k	$\ \nabla f(x_k)\ $	x_{k+1}
1	(3.00,25.00)	509	(-19.00,45.00)	0.32	48.85	(9.08,10.6)
2	(9.08,10.6)	77.35	(7.56,10.12)	0.9	12.63	(2.28,1.5)
3	(2.28,1.5)	1.02	(3.06,-1.28)	0.36	3.31	(1.18,1.96)
4	(1.18,1.96)	-0.99	(0.4,0.74)	0.84	1.79	(0.85,1.33)
5	(0.85,1.33)	-1.29	(0.37,-0.19)	0.35	0.4	(0.73,1.39)
6	(0.73,1.39)	-1.32	(0.07,0.05)	0.94	0.08	(0.67,1.35)
7	(0.67,1.35)	-1.33	(-0.01,0.03)	0.34	0.03	(0.67,1.34)
8	(0.67,1.34)	-1.33	(0.00,0.01)	0.5	0.01	(-)

conjugué et la méthode du Newton (La condition d'arrêt est la même que précédemment $\|\nabla f(x_k)\| < 0.01$)

La méthode du gradient conjugué :

k	x_k	$f(x_k)$	ρ_k	x_{k+1}	la direction conjuguée
1	(5.00,7.00)	5.00	0.92	(1.16,1.24)	(2.00,3.00)
2	(1.16,1.24)	-1.03	0.56	(0.56,1.62)	(1.08,-0.68)
3	(0.56,1.62)	-1.20	(-)	(-)	(0,0)

Exécution de la méthode du gradient conjugué pour le problème (P) depuis $x_0 = (3, 4)$.

On remarque que la méthode n'a pas eu besoin de beaucoup d'itération pour trouver la solution optimal.

La méthode du Newton :

Exécution de la méthode du Newton pour le problème (P) depuis $x_0 = (3, 4)$

k	x_k	$\nabla f(x_k)$	$\nabla^2 f(x_k)$	$\nabla^2 f(x_k)^{-1}$	x_{k+1}
1	(3.00,4.00)	(2.00,3.00)	[(2,-1),(-1,2)]	[(0.66,0.33),(0.33,0.66)]	(0.67,1.33)
2	(0.67,1.33)	(0.0,0.0)	(-)	(-)	(-)

On remarque que cette méthode converge d'une manière finie en une seule itération (et ceci indépendamment du point initial)

Conclusion :

En conséquence, et après avoir résolu d'autres exemples quadratiques, on conclut que chacun des algorithmes décrit précédemment peut être utilisé pour résoudre ce type de problème, cependant, la méthode de Newton et la méthode du gradient conjugué restent les mieux adaptées, de part leur convergence finie en un nombre fixe d'itérations.

Bibliographie

- [1] X. Antoine, P. Dreyfuss et Yannick Introduction à l'optimisation. Aspects Théoriques et Algorithmes, *ENSMN-ENSEM 2 A(2006-2007)*.
- [2] M. Bergounioux, Optimisation et contrôle des systèmes linéaires, *Dunod, Paris, (2001)*.
- [3] M. Bierlaire, Introduction à l'optimisation différentiable, *PPUR, 2006*.
- [4] M. s. Bazarraa and H. D. Sherali, et C . M. Shety, Nonlinear programming, Theory and Algorithms, *John Wiley and Sons, New York, Second ed, (1993)*.
- [5] J.F. Bonnans, J.C. Gilbert et C. Lemaréchal - C. Sagastizabal, Optimisation numérique Aspects théoriques et pratiques , *Collection Mathématiques et Applications, Vol. 27, Springer-verlag, 1998*.
- [6] G.Cohen, Convexité et Optimisation, *Ecole Nationale des Ponts et Chaussées et Inria. Edition (2000-2006)*.
- [7] P. G, Ciarlet, Introduction à l'analyse numérique matricielle est à l'optimisation, *Masson, (1982)*.
- [8] Y. H. Dai and Y. Yuan, Nonlinear conjugate gradient methods, *Shanghai Science and Technology Publisher, Shanghai, 2000*.
- [9] Y. H. Dai and Y. Yuan, Some properties of a new conjugate gradient method, in : *Advantes in Nonlinear programing, ed. Kluwer, Boston, (1998)*.

- [10] Y. H. Dai and Y. Yuan, Anonlinear conjugate gradient method with a strong global, *SIAMJ. Optim.* 10, (1999).
- [11] Bertsekas. Dimitri. P, Nonlinear programing , *Second printing, Athema Scientific.*
- [12] P.Faurre ,Analyse Numérique Notes d'optimisation, *Collection X, Ellipses, Paris, 1988.*
- [13] R. Fletcher, Practical Methods of Optimization vol. 1 : Unconstrained Optimization,*John Wiley and Sons, New York, 1987.*
- [14] J. C, Gilbert and J. Nocedal, Global convergance properties of conjugate gradient methods for optimization, *SIAM. J. Optimization. vol 2 No. 1(1992)*
- [15] J-B. Hiriart-Urry, Optimisation et analyse convexe, exercices corrigés, *EDP, 2006.*
- [16] D ,Kauth, Optimisation numerique methodes du gradient conjugue lineaire, *University de Fribourg, le 5 novembre(1992).*
- [17] G. Lui, J.Han and H. Yin, Global convergence of the Fletcher-Reeves algorithm with inexact line search, *Appl. Math. J-Chinese Univ-Ser. B. 10 (1995).*
- [18] M.Michel,Programmatiob Mathématique des Théories et Algorithmes,*Collection Théchnique des Télécommunications a volume 1.*
- [19] Stéphane.Mottelet,Optimisation non-linéaire,*Université de Technologie de compiégne,2003.*
- [20] J. J, Moré, B. S, Garbow and K. E, Hillstrom, Testing unconstrained optimization softwar, *ACM Transactions on Mathematical Softwar , 7 (2006).*
- [21] J. Nocedal, and S. J. Wright, Numerical Optimization, *Springer Series in Operations Research and Financial Engineering, second ed, 2006.*

- [22] *M. J. D. Powell*, Nonconvex minimization calculations and the conjugate gradient method, Numerical Analysis, Lecture Notes in Mathematics, Vol. 1066, *Springer-Verlag, Berlin*, (1984), pp. 122–141.
- [23] *L. Pujo-Menjouet*, Cours de calcul différentiel, *Université Claude Bernard, Lyon, France*.
- [24] *J. Royer*, Calcul différentiel et intégral, *Université Toulouse 3*, (2013-2014).