

ACADEMY OF ROMANIAN SCIENTISTS



Dr. Neculai ANDREI

**METODE AVANSATE
DE
GRADIENT CONJUGAT
PENTRU
OPTIMIZARE FĂRĂ RESTRICȚII**

**Editura
ACADEMIEI OAMENILOR DE ȘTIINȚĂ DIN ROMÂNIA
București, 2009**

Copyright © Editura Academiei Oamenilor de Știință din România 2009

Toate drepturile asupra acestei ediții sunt rezervate Editurii.

Tipar:

Editura Academiei Oamenilor de Știință din România

Splaiul Independenței nr. 54 sector 5, 050094, București, ROMÂNIA

Tel. 00-4021/314.74.91. Fax. 00-4021/314.75.39

Web-site: www.aos.ro, e-mail: aosromania@yahoo.com

Descrierea CIP a Bibliotecii Naționale a României

ANDREI, NECULAI

Metode avansate de gradient conjugat pentru optimizare fără restricții/

dr. Neculai Andrei - București: Editura Academiei Oamenilor Știință din România, 2009

ISBN 978-606-92161-0-1

519.8

Referenți științifici și revizia lucrării: Membrii secției **Știința și Tehnologia Informației**
din cadrul **Academiei Oamenilor de Știință din România**

Șef serviciu Editură: ing. Liviu Mihai **Sima**

Redactor: prof. Andrei D. **Petrescu**

Procesare text: dr. Neculai **Andrei**

Coperta: ec. Ovidiu **Oprea**

C.Z. pentru biblioteci mari: 517.392:512

C.Z. pentru biblioteci mici: 512

© Copyright 2009

Toate drepturile prezentei ediții sunt rezervate Editurii. Nicio parte din această lucrare nu poate fi reprodusă, stocată sau transmisă sub indiferent ce formă fără acordul prealabil scris al Editurii.

ADVANCED CONJUGATE GRADIENT METHODS FOR UNCONSTRAINED OPTIMIZATION

Unconstrained optimization consists in finding values for some variables so that a given function in these variables takes the minimum or maximum possible value. When the number of variables is large enough, this operation turns out to be quite complex. This book presents the state-of-the-art of conjugate gradient algorithms for unconstrained optimization. Both the classical and the latest conjugate gradient methods for solving the unconstrained optimization problems are considered. As we know, unconstrained optimization is not merely an intellectual exercise. Its purpose is to solve practical unconstrained optimization problems on computers. Therefore, the book contains a large number of conjugate gradient algorithms advanced techniques of computational linear algebra which form the kernel of any industrial optimization algorithm.

The monograph is organized in 6 chapters and two annexes. The first chapter has an introductory character, where the main aspects of the descent methods for unconstrained optimization are presented. In chapter 2 the linear conjugate gradient methods for solving linear algebraic systems of equations are discussed. A special attention is given to the discretization of partial differential equations. With chapter 3 the book enters into the subject of nonlinear conjugate gradient methods for unconstrained optimization. The theory of classical conjugate gradient algorithms is presented both of those methods with $\|g_k\|^2$ into the denominator of β_k , and for the methods with $g_{k+1}^T y_k$ into the denominator of β_k . Chapter 4 introduces the hybrid conjugate gradient methods as projections of classical conjugate gradient methods or as convex combinations of these methods. Chapter 5 contains a large number of new conjugate gradient algorithms, some of them based on the very recent results of the author. The last chapter is dedicated to present the computational performances of conjugate gradient algorithms for solving 7 applications from MINPACK2 collection. The first annex reviews a number of 50 nonlinear conjugate gradient algorithms, and the second one articulates some open problems in these methods.

The book addresses to a large number of scientists, theorists, as well as practitioners engineers and decision makers, researchers, computer scientists or professional bodies working in academy, business, industry and government facing design optimization models and algorithms, who are interested in the theory of conjugate gradient methods for unconstrained optimization, in computational methods associated to conjugate gradient algorithms, in numerical experiments and comparisons, as well as in behavior of these algorithms for solving real applications from different areas of activity.

Dr. Neculai Andrei

Din partea autorului

Elaborată într-o manieră riguros matematică cu ilustrații computaționale de mare anvergură, monografia constituie o sinteză a celor mai recente rezultate din domeniul metodelor de gradient conjugat pentru optimizare neliniară fără restricții din care multe aparțin autorului.

Efortul depus în această monografie constă în a sugera și dovedi cititorului că metodele de gradient conjugat sunt o alternativă foarte serioasă la celelalte metode de optimizare fără restricții bazate pe conceptul de aproximație pătratică, în ceea ce privește rezolvarea problemelor de optimizare fără restricții de mari dimensiuni așa cum acestea apar în aplicațiile practice. Ideea este că această clasă de algoritmi beneficiază de rezultate teoretice de convergență globală și complexitate computațională foarte tari și în același timp pot fi cu ușurință implementați în programe de calcul care implică cerințe modeste de memorie. Importanța acestor algoritmi rezidă și în faptul că aceștia se pot include foarte ușor în alte programe de calcul dedicate rezolvării aplicațiilor complexe care implică optimizarea fără restricții.

Autorul mulțumește Editurii Academiei Oamenilor de Știință pentru profesionalismul și sollicitudinea cu care a susținut apariția acestei lucrări.

București, 20 Iunie 2009



Prefață

Optimizarea fără restricții reprezintă un capitol foarte important al programării matematice. Aceasta se referă la optimizarea unei funcții (suficient de netedă), în sensul de a calcula acele puncte în care funcția respectivă își atinge valorile extreme, de minim sau maxim. Importanța optimizării fără restricții rezultă din faptul că foarte multe fenomene reale din Natură se pot modela ca probleme de acest tip. Pe de altă parte, problemele de optimizare cu restricții, care constituie o clasă de probleme mult mai generale și mai bogată, se pot soluționa prin rezolvarea succesivă a unui șir de probleme de optimizare fără restricții. De exemplu, metodele barieră din optimizarea cu restricții au condus la o revigorare a domeniului optimizării fără restricții, în care metodele de direcții conjugate ocupă un loc central.

Programarea matematică în general și optimizarea fără restricții în particular au atins o anumită maturitate dispunând de un corp de rezultate teoretice și computaționale bine stabilit. În acest sens, pentru rezolvarea acestei clase de probleme se cunosc o multitudine de algoritmi bine justificați matematic, pentru care se cunosc proprietățile de convergență și complexitate computațională și care au fost implementați în programe de calcul deosebit de sofisticate și eficiente capabile să rezolve probleme complexe de optimizare fără restricții cu un număr mare de variabile.

*În programarea matematică s-au individualizat două mari clase de probleme: **programarea liniară** și **programarea matematică neliniară**. Programarea liniară este o problemă în care toate funcțiile (funcția de optimizat și restricțiile) sunt liniare. Chiar dacă programarea liniară și optimizarea neliniară s-au dezvoltat în paralel, rareori și ocazional având puncte de contact, anul 1984 a constituit începutul unei „revoluții” în programarea matematică, prin reconsiderarea la un nivel superior a unor tehnici de optimizare bazate pe conceptul de penalizare și care au devenit cunoscute ca metode de punct interior. Aceasta a constituit **prima unificare din programarea matematică**, în sensul că aceleași tehnici pot fi utilizate pentru rezolvarea ambelor clase de probleme. Cel mai important avantaj al metodelor de punct interior asupra metodei simplex pentru programarea liniară este eficiența cu care pot fi rezolvate probleme de mari dimensiuni și abilitatea de a utiliza direcții de căutare aproximative, care permit utilizarea unor metode de gradient conjugat bazate pe structura problemei. Un alt avantaj este că metodele de punct interior se extind foarte bine la problemele de optimizare convexă, în timp ce metoda simplex, care se bazează foarte mult pe structura liniară a problemei, nu. Această observație a condus la construcția unui cadru matematic foarte general al metodelor de punct interior pentru rezolvarea unei mari varietăți de probleme de optimizare nediferențiale convexe, ce deschide*

foarte multe perspective teoretice și computaționale. Așa au apărut o serie de noi probleme de optimizare ca programarea semidefinită, programarea pe conuri, optimizarea combinatorială neconvexă, minimizarea normelor matriceale, aproximarea Chebychev logaritmică etc. toate acestea încadrându-se în așa-numita clasă a **inegalităților matriceale liniare** și care a condus în cele din urmă la **a doua unificare dintre optimizare, teoria sistemelor și control optimal**. Astfel, s-a observat că tehnici și metode proprii unor domenii foarte diferite ca optimizarea neliniară, teoria sistemelor dinamice sau controlul optimal utilizează aceleași concepte matematice cu interpretări proprii.

Scopul lucrării. Întregul nostru efort de-a lungul acestei monografii a fost de a prezenta riguros matematic, cu suficiente detalii, într-un mod accesibil și prietenos, principalele aspecte ale metodelor de gradient conjugat pentru optimizarea fără restricții. Având în vedere faptul că scopul final al oricărei teorii, în particular al optimizării fără restricții, este de a dezvolta scheme computaționale robuste și de încredere pentru rezolvarea problemelor concrete, lucrarea își propune prezentarea principiilor care stau la baza metodelor de gradient conjugat, a algoritmilor corespunzători împreună cu caracteristicile lor de convergență și complexitate, precum și capacitățile lor de a rezolva probleme practice reale cu un număr mare de variabile.

Rigoare și formalism. În prezentarea materialului referitor la metodele avansate de gradient conjugat pentru optimizare fără restricții am adoptat ideea ca în locul unui formalism de tipul definiție – teoremă – demonstrație să utilizăm o prezentare narativă în care acest formalism este scufundat și ilustrat prin comentarii și exemple numerice relevante, care arată puterea computațională a algoritmilor. În acest sens am urmărit prezentarea riguroasă a propozițiilor, lemelor și teoremelor importante care stabilesc proprietățile de convergență și complexitate computațională a algoritmilor, precum și ilustrarea acestor rezultate prin rezolvarea unor colecții de probleme de test și a unor aplicații reale din diferite domenii industriale. Principalul efort în organizarea acestui material a fost de a se evita pe cât posibil apelul la concepte definite în secțiunile ulterioare ale lucrării, argumentarea circulară, precum și utilizarea unor noțiuni, tehnici și metode considerate ca bine cunoscute în domeniu.

Structura lucrării. Cartea este organizată în 6 capitole și două anexe. În același timp aceasta conține o bibliografie cu cele mai relevante și cele mai recente lucrări din domeniu, precum și un index de termeni și un index de autori.

♦ Primul capitol are un caracter introductiv. Acesta prezintă conceptele fundamentale cu care se operează în domeniul optimizării fără restricții, metodele de direcții descendente utilizate pentru rezolvarea acestei clase de probleme, precum și câteva rezultate de bază asupra metodelor de direcții descendente.

♦ Capitolul doi prezintă metoda gradientului conjugat pentru rezolvarea sistemelor de ecuații algebrice liniare cu matricea coeficienților simetrică și pozitiv definită. Se arată că viteza de convergență a metodei de gradient conjugat este puternic dependentă de distribuția valorilor proprii a matricei sistemului. Totodată, în acest

capitol se insistă asupra discretizării ecuațiilor diferențiale cu derivate parțiale și a sistemelor liniare obținute prin discretizare, inclusiv structura de valori proprii.

♦ Capitolul trei este dedicat metodelor clasice de gradient conjugat pentru optimizare fără restricții. Principiul pe care se bazează metodele de gradient conjugat constă în faptul că minimizarea unei funcții f pe un subspațiu dat este echivalentă cu minimizarea succesivă de-a lungul unor direcții conjugate care generează acel subspațiu. Aceste metode se reprezintă sub forma $x_{k+1} = x_k + \alpha_k d_k$, unde α_k este lungimea pasului, iar d_k este direcția de deplasare calculată sub forma $d_{k+1} = -g_{k+1} + \beta_k d_k$. Astfel, se prezintă teoria convergenței algoritmilor de gradient conjugat, separat pentru acele metode cu $\|g_k\|^2$ la numărătorul parametrului β_k și respectiv cu $g_{k+1}^T y_k$ la numărătorul lui β_k . Secțiunea 3.6 a acestui capitol este originală și prezintă o metodă de accelerare a algoritmilor de gradient conjugat [Andrei, 2008h].

♦ Capitolul patru prezintă metodele hibride și parametrizate. Se cunosc două tehnici de elaborare a algoritmilor hibrizi. Prima utilizează conceptul de proiecție a algoritmilor clasici de gradient conjugat. A doua tehnică este originală și se bazează pe conceptul de combinație convexă, care deschide o nouă direcție de cercetare în ceea ce privește elaborarea de algoritmi de optimizare fără restricții [Andrei, 2007g,h,i; 2008a,b,j,k,l,n,s; 2009a,b].

♦ În capitolul cinci se prezintă noi metode de gradient conjugat pentru optimizare fără restricții. Cele mai importante sunt: metode de gradient conjugat BFGS preconditionate fără memorie date de Shanno [1978a,b], o metodă de gradient conjugat cu descendență garantată a lui Hager și Zhang [2005], metode de gradient conjugat scalat și scalat accelerate elaborate de Andrei [2007a,b,c], metode de gradient conjugat cu descendență suficientă [Andrei, 2006b], metoda gradientului conjugat cu ecuația secantei modificată [Andrei, 2005e; 2008b], metoda gradientului conjugat cu aproximația produsului Hessian / vector [Andrei, 2008o,u] și metode de gradient conjugat cu descendență și conjugare garantate [Andrei, 2008v].

♦ Capitolul șase prezintă un număr de șase aplicații de optimizare fără restricții rezolvare cu algoritmi de gradient conjugat prezentați de-a lungul acestei lucrări, precum și comparații între ei și algoritmi quasi-Newton BFGS cu memorie limitată și Newton truncat.

Proiecte computaționale. Ideea pe care am avut-o în elaborarea acestei lucrări a fost de a scrie un text matematic, riguros și logic prezentat, care în final să se poată utiliza în dezvoltarea unor programe de calcul capabile să rezolve probleme de optimizare fără restricții, modele ale unor fenomene fizice reale. Pentru realizarea acestui demers algoritmi asociați metodelor de gradient conjugat sunt prezentați pe pași cu suficiente detalii computaționale. În același timp, lucrarea conține o multitudine de studii numerice care implică minimizarea unui tren de peste 750 de funcții de test de optimizare fără restricții cu numărul de variabile din domeniul [1000,10000].

Optimizare aplicată. Există o întreagă dispută asupra termenului „aplicat”. Aceasta arată diferențele conceptuale care există între cei care sunt interesați de teoria optimizării și practiționiști, interesați mai mult de rezolvarea unor probleme concrete. În această monografie am adoptat un punct de vedere rabinic în sensul de a mulțumi și teoreticianul și practiționistul. În acest sens lucrarea prezintă rezultatele furnizate de o multitudine de algoritmi de gradient conjugat în ceea ce privește rezolvarea unor aplicații concrete ca: torsiunea elasto-plastică a unei bare, distribuția presiunii într-un lagăr de alunecare (cuzinet), proiectarea optimă a unei bare cu rigiditate torsională maximă, rezolvarea ecuațiilor Ginzburg-Landau pentru superconductoare neomogene formate din straturi de plumb și staniu în absența unui câmp magnetic, combustia staționară a unui material solid, configurația de energie minimă a unui grup de atomi, etc.

Adresabilitate. Parcurgerea și înțelegerea acestei lucrări cere o anumită maturitate matematică în sensul abilității de a opera cu anumite concepte și obiecte matematice din algebra liniară, calculul diferențial, analiza convexă, topologie, geometrie diferențială și ecuații diferențiale. Ca atare, monografia se adresează celor interesați de teoria optimizării fără restricții și în special de metodele de gradient conjugat: matematicieni, analiști din domeniul cercetărilor operaționale și al științei managementului, planificatori, programatori, ingineri, cercetători și doctoranzi în specialitățile care includ utilizarea tehnicilor de optimizare, studenți interesați de aceste probleme.

Mulțumiri. Autorul mulțumește Fundației „Alexander von Humboldt” pentru generozitatea și solitudinea cu care a știut să-l sprijine de-a lungul timpului atât în perioada stagiului de peste trei ani petrecuți în câteva Universități din Germania, cât și după aceea. Aceasta ne-a permis o implicare directă în cercetarea științifică de cel mai înalt nivel, posibilitatea audierii și prezentării unor rezultate din domeniile modelării matematice și ale optimizării și, în general, încadrarea în acea clasă rafinată a oamenilor de știință pentru care intelectualismul ca doctrină prin care cunoașterea – parțial sau total – se obține prin rațiune pură, reprezintă un mod de existență. Scrierea acestei monografii apare deci ca o datorie de a recompensa, măcar în parte, generozitatea acestei Fundații, care în esența ei nici nu are nevoie de așa ceva. De asemenea, autorul mulțumește celor câțiva colegi și prieteni, care de-a lungul timpului cât a durat scrierea acestei monografii au știut să-l încurajeze cu solitudine și umor. În același timp, autorul este recunoscător soției sale Mihaela, care a înțeles efortul său și a știut să-l scutească de o serie de activități lungi și obositoare, precum și copiilor Teodor-Marius, Denisa și Cosmin pentru afecțiunea și înțelegerea cu care l-au înconjurat și sprijinit în efortul de a concretiza o muncă depusă pe parcursul a mai bine de 40 de ani.



CUPRINS

<i>Simboluri și notații</i>	11
1. Metode de direcții descendente pentru optimizare fără restricții	13
1.1. Algoritmul general de direcții descendente	15
1.2. Metode de direcții descendente	18
1.2.1. Metoda pasului descendent	18
1.2.2. Metoda Newton și variantele sale	19
1.2.3. Metode quasi-Newton	25
1.2.4. Metoda regiunii de încredere	29
1.3. Studiu numeric (Newton trunchiat versus LBFGS)	31
2. Metoda gradientului conjugat pentru rezolvarea sistemelor algebrice liniare	35
2.1. Metoda pasului descendent	36
2.2. Metoda direcțiilor conjugate	41
2.3. Metoda gradientului conjugat	43
2.4. Discretizarea ecuațiilor diferențiale cu derivate parțiale	55
3. Metode clasice de gradient conjugat pentru optimizare fără restricții	77
3.1. Metode pentru funcții pătratice	78
3.2. Metode clasice de gradient conjugat	85
3.3. Teoria convergenței algoritmilor de gradient conjugat	88
3.4. Metode de gradient conjugat cu $\ g_k\ ^2$ la numărătorul lui β_k .	96
3.5. Metode de gradient conjugat cu $g_{k+1}^T y_k$ la numărătorul lui β_k .	108
3.6. Accelerarea metodelor de gradient conjugat	119
4. Metode de gradient conjugat hibride și parametrizate	129
4.1. Metode de gradient conjugat hibride bazate pe conceptul de proiecție	130
4.2. Metode de gradient conjugat hibride bazate pe conceptul de combinație convexă	144
4.3. Metode de gradient conjugat parametrizate	159

5. Noi metode de gradient conjugat pentru optimizare fără restricții	161
5.1. Metode de gradient conjugat BFGS preconditionate fără memorie	161
5.2. O metodă de gradient conjugat cu descendență garantată	171
5.3. Metode de gradient conjugat scalat	178
5.4. Metode de gradient conjugat scalat accelerate	190
5.5. Metode de gradient conjugat cu descendență suficientă	197
5.6. Metoda gradientului conjugat cu ecuația secantei modificată	208
5.7. Metoda gradientului conjugat cu aproximația produsului Hessian / vector	223
5.8. Noi algoritmi de gradient conjugat accelerat ca modificări ale schemelor clasice	231
5.9. Metode de gradient conjugat cu descendență și conjugare garantate	238
5.10. Precondiționare	248
6. Aplicații	251
6.1. Torsiunea elasto-plastică a unei bare	251
6.2. Distribuția presiunii într-un lagăr (cuzinet) de alunecare	254
6.3. Proiectarea optimă a unei bare cu rigiditate torsională maximă	256
6.4. Rezolvarea problemei Ginzburg-Landau	259
6.5. Combustia staționară a unui material solid	262
6.6. Configurația de energie minimă a unui grup de atomi	264
6.7. Suprafețe minimale	266
Anexa 1	
50 de algoritmi de gradient conjugat pentru optimizare fără restricții	273
Anexa 2	
Probleme deschise în algoritmii de gradient conjugat pentru optimizare fără restricții	295
<i>Bibliografie</i>	305
<i>Index de termeni</i>	317
<i>Index de autori</i>	321



SIMBOLURI ȘI NOTAȚII

$x \in X$	x aparține mulțimii X .
$x \notin X$	x nu aparține mulțimii X .
R	Mulțimea numerelor reale (axa reală).
$ x $	Modulul numărului x .
$\lfloor j \rfloor$	Cel mai mare întreg mai mic decât sau egal cu j .
R^n	Spațiul real n – dimensional.
$ls\{d_1, \dots, d_k\}$	Subspațiul liniar generat de vectorii $d_1, \dots, d_k$.
$K_k(A, b)$	k – spațiul Krylov al matricei A corespunzător vectorului b . $K_k(A, b) = ls\{b, Ab, \dots, A^{k-1}b\}$.
$x \in R^n$	Vector (coloană) din spațiul R^n .
x^T	Transpusul vectorului x (vector linie).
$x^T y$	Produsul scalar al vectorului $x \in R^n$ cu vectorul $y \in R^n$, definit ca: $x^T y = \sum_{i=1}^n x_i y_i$.
$\cos\langle x, y \rangle$	Cosinusul unghiului dintre vectorii x și y . $\cos\langle x, y \rangle = \frac{x^T y}{\ x\ \ y\ }$.
$x \perp y$	Vectorul x este perpendicular pe vectorul y : $x^T y = 0$.
$\ x\ _2$	Norma l_2 (norma euclidiană) a vectorului x : $\ x\ _2 = \left(\sum_{i=1}^n x_i^2\right)^{1/2}$. Deseori se notează sub forma $\ x\ $.
$\ x\ _Q$	Norma vectorului x în raport cu matricea simetrică Q : $\ x\ _Q = (x^T Q x)^{1/2}$, (norma elipsoid).
$\ x - y\ _2$	Distanța dintre două puncte $x, y \in R^n$, distanța euclidiană: $\ x - y\ _2 = \sqrt{(x_1 - y_1)^2 + \dots + (x_n - y_n)^2}$.
$A = [a_{ij}]$	Matrice cu elementele a_{ij} . a_{ij} este elementul din linia i și coloana j . Dacă $i = 1, \dots, m$ și $j = 1, \dots, n$, atunci A este o matrice cu m linii și n coloane.
I_n	Matricea unitate de ordin n . Se mai notează cu I .

A^T	Transpusa matricei A .
A^{-1}	Inversa matricei A . $AA^{-1} = I$.
$R^{m \times n}$	Mulțimea tuturor matricelor $m \times n$ dimensionale cu elemente numere reale.
$\ A\ _2$	Norma spectrală a matricei A : $\ A\ _2 = (\lambda_{\max}(A^T A))^{1/2}$.
$\ A\ _F$	Norma Frobenius a matricei A : $\ A\ _F = \left(\sum_{i=1}^m \sum_{j=1}^n a_{ij} ^2 \right)^{1/2}$.
$\kappa(A)$	Numărul de condiționare a matricei A : $\kappa(A) = \sigma_{\max} / \sigma_{\min}$, unde $\sigma_{\min}$ și $\sigma_{\max}$ sunt valoarea singulară minimă, respectiv maximă a matricei A . $\kappa(A) = \ A\ \ A^{-1}\ $.
$diag(a_1, \dots, a_n)$	Matrice diagonală cu elementele de pe diagonala principală egale cu $a_1, \dots, a_n$.
$f: X \rightarrow Y$	Funcția f definită pe mulțimea X cu valori în mulțimea Y .
$dom f$	Domeniul funcției f .
$\nabla f(x)$	Gradientul funcției f în x se definește ca vectorul $\nabla f(x) = \left[\frac{\partial f(x)}{\partial x_1}, \dots, \frac{\partial f(x)}{\partial x_n} \right]^T$. În limbajul curent gradientul funcției f în punctul x_k se notează cu g_k .
$\nabla^2 f(x)$	Hessianul funcției f se definește ca $n \times n$ -matricea simetrică cu elementele: $[\nabla^2 f(x)]_{ij} = \frac{\partial^2 f(x)}{\partial x_i \partial x_j}, \quad i, j = 1, \dots, n.$
$T_k(x)$	Polinomul Cebyshev de prima speță. $T_k(x) = \cos(k \arccos(x))$.
TST	Matrice Toeplitz, simetrică, tridiagonală.
DFP	Actualizarea quasi-Newton Davidon-Fletcher-Powell.
BFGS	Actualizarea quasi-Newton Broyden-Fletcher-Goldfarb-Shanno.
LBFGS	Actualizarea BFGS cu memorie limitată.
■	Sfârșitul demonstrației.



Capitolul 1

METODE DE DIRECȚII DESCENDENTE PENTRU OPTIMIZARE FĂRĂ RESTRICȚII

Scopul acestui capitol este de a prezenta cu suficiente detalii metodele de bază ale optimizării fără restricții care utilizează informațiile date de gradientul și eventual de Hessianul funcției de minimizat. Nu vom intra în detalii. De aceea recomandăm cititorului interesat consultarea lucrărilor [Andrei, 1999a,b; 2000a,b; 2002; 2003; 2004a,b; 2009c] unde va găsi o multitudine de aspecte teoretice, computaționale și practice privind algoritmi dedicați rezolvării acestei clase de probleme.

Orice algoritm de optimizare este compunerea a două aplicații: prima determină direcția de deplasare către minim, o direcție descendentă. A doua determină lungimea pasului de deplasare de-a lungul direcției calculate. În acest capitol vom discuta ambele aceste aspecte ale unui algoritm de optimizare insistând numai asupra acelor metode care s-au dovedit eficiente mai ales pentru rezolvarea problemelor neliniare complexe cu un număr mare de variabile.

Ideea care stă la baza oricărui algoritm de optimizare, inclusiv a celor de optimizare fără restricții, este foarte simplă. Se consideră un *punct inițial* din care se demarează calculele. În acest punct se determină o direcție care este direcționată către minimul funcției de minimizat, numită *direcție de descendență*. De-a lungul acestei direcții se execută o deplasare convenabilă, numită *căutare liniară*, astfel obținându-se un nou punct din care se pot relua calculele. Toate aceste elemente care definesc o metodă descendentă de optimizare au propriile lor principii și subtilități care ilustrează bogăția și maturitatea domeniului optimizării fără restricții. Ceea ce individualizează un algoritm față de altul este modul în care se calculează direcția de deplasare. Lungimea pasului se calculează într-o manieră standard utilizând anumite condiții care asigură o reducere suficient de mică a valorilor funcției de minimizat de-a lungul direcției de descendență.

După cum am spus, algoritmi de optimizare fără restricții, proiectați pentru a rezolva problema $\min_{x \in R^n} f(x)$, unde $f : R^n \rightarrow R$, suficient de netedă, se identifică după modul în care calculează direcția de deplasare. Se cunosc o multitudine apreciabilă de algoritmi de optimizare fără restricții, care se pot organiza în următoarele clase de metode de optimizare:

1. Metode de căutare directă (nederivative)
2. Metoda pasului descendent (Cauchy)
3. Metoda Newton
 - Metoda Newton modificată
 - Metoda Newton discretă
 - Metoda Newton compusă
 - Metoda Newton trunchiată
4. Metode quasi-Newton
 - Metoda quasi-Newton cu memorie limitată
 - Metode quasi-Newton partiționate
5. Metoda regiunii de încredere
6. Metode de direcții conjugate

Principii de optimizare fără restricții

Construcția algoritmilor de optimizare fără restricții se bazează pe următoarele două principii fundamentale care stau la baza tuturor metodelor de direcții descendente de optimizare fără restricții [Andrei, 2009c].

1. Primul principiu fundamentează ideea că orice problemă de optimizare într-un punct dat din domeniul funcției de minimizat se poate foarte bine aproxima printr-o problemă de optimizare pătratică. Cu alte cuvinte, într-un punct curent x_k se calculează o aproximație pătratică a problemei Q_k care prin rezolvare furnizează un nou punct x_{k+1} din care procesul se poate imediat repeta până când un anumit criteriu de oprire a iterațiilor este îndeplinit. Această idee este foarte bună și se bazează pe faptul că modelele matematice ale fenomenelor fizice suficient de netede nu sunt puternic neliniare și deci se pot aproxima cu o acuratețe bună prin modele pătratice. Reprezentative pentru acest principiu sunt metoda Newton și variantele sale, metodele quasi-Newton și metoda regiunii de încredere.
2. Al doilea principiu, complet diferit de primul se bazează pe *conceptul de direcție conjugată*. Dată o matrice simetrică și pozitiv definită A , atunci direcțiile $d_1 \neq 0$ și $d_2 \neq 0$ sunt conjugate în raport cu matricea A , dacă $d_1^T A d_2 = 0$. Evident că acest concept de conjugare se poate imediat extinde la cazul funcțiilor generale neliniare prin considerarea în locul matricei A a Hessianului acesteia în punctul curent. Această idee a condus la clasa metodelor de direcții conjugate cu care a demarat optimizarea fără restricții de mari dimensiuni (milioane de variabile). Metodele de direcții conjugate pot împrumuta anumite idei din metodele Newton sau quasi-Newton în sensul de a încerca înglobarea informației de ordinul doi dată de Hessianul funcției de minimizat sau de o aproximare a acestuia, obținându-se astfel algoritmi foarte puternici de optimizare fără restricții.

Să considerăm problema $\min f(x)$, unde $f : R^n \rightarrow R$ este o funcție de două ori continuu diferențiabilă pe R^n cu domeniul $\text{dom } f$. Deoarece f este diferențiabilă,

atunci o condiție necesară pentru ca punctul x^* să fie soluția optimă a problemei este:

$$\nabla f(x^*) = 0.$$

Aceasta se numește *condiția necesară de optimalitate de primul ordin*. Punctele care o satisfac se numesc *puncte staționare* sau *critice*. Ca atare, rezolvarea problemei a fost redusă la determinarea unei soluții pentru sistemul $\nabla f(x^*) = 0$, care, după cum se vede, este un sistem algebric de n ecuații în n variabile. Ca atare, minimizarea funcțiilor este intim legată de rezolvarea sistemelor algebrice neliniare. Un punct staționar nu este întotdeauna un minim al funcției. De exemplu, funcția $f(x) = -x^4$ satisface condițiile necesare în punctul 0, care după cum se vede este un punct de maxim al lui f . Pentru a obține un punct de minim trebuie ca derivata de ordinul doi să fie strict pozitivă. Aceasta este *condiția suficientă de optimalitate*.

De obicei pentru rezolvarea problemei vom folosi o metodă iterativă, care generează un șir de puncte $x_0, x_1, \dots, x_k, \dots \in \text{dom } f$ cu proprietatea că $f(x^k) \rightarrow f^*$ când $k \rightarrow \infty$. Un astfel de șir cu această proprietate se numește *șir minimizant*. Astfel conceput, un algoritm care generează un șir minimizant se oprește de îndată ce un anumit criteriu de oprire a iterațiilor este îndeplinit. De exemplu, se pot utiliza următoarele criterii de oprire a iterațiilor: $|f(x_k) - f^*| \leq \varepsilon_f$, unde $\varepsilon_f > 0$ este o toleranță specificată, sau distanța dintre estimațiile variabilelor $\|x_{k+1} - x_k\|_2 \leq \varepsilon_x$, unde $\varepsilon_x > 0$, sau norma gradientului funcției de minimizat $\|\nabla f(x_k)\|_2 \leq \varepsilon_g$ sau $\|\nabla f(x_k)\|_\infty \leq \varepsilon_g$, unde $\varepsilon_g > 0$, de-a lungul iterațiilor. Evident că se pot utiliza și combinații ale acestor criterii, etc.

1.1. ALGORITMUL GENERAL DE DIRECȚII DESCENDENTE

Să considerăm deci problema

$$\min f(x) \tag{1.1.1}$$

unde $f: R^n \rightarrow R$ este o funcție continuu diferențiabilă cu domeniul $\text{dom } f$. Valoarea optimă a acestei probleme o notăm cu f^* , adică $\inf_x f(x) = f^*$. După cum am văzut, date punctul x_k și direcția descendentă de deplasare d_k , atunci schema iterativă după care se calculează punctul de optim x^* pentru (1.1.1) este

$$x_{k+1} = x_k + \alpha_k d_k, \tag{1.1.2}$$

unde α_k este *lungimea pasului de deplasare* de-a lungul direcției d_k . Menționăm că o direcție descendentă în punctul x_k satisface condiția $\nabla f(x_k)^T d_k < 0$, numită

condiția de descendență, precum și condiția de a fi mărginită față de zero, $\|d_k\| \neq 0$. Prima condiție înseamnă că aceasta trebuie să facă un unghi ascuțit cu gradientul cu semn schimbat al funcției în punctul curent x_k . Dacă $\alpha_k \neq 0$, vedem că a doua condiție face ca (1.1.2) să se încadreze în *principiul fundamental de construcție a algoritmilor de analiză numerică, care recomandă ca o estimatie a unei soluții a unei probleme să se calculeze ca o mică modificare de mare acuratețe a estimăției anterioare*.

Pentru demararea calculelor, metoda cere specificarea unui punct inițial x_0 . Punctul de plecare x_0 trebuie să se afle în $\text{dom } f$ și în plus mulțimea de nivel constant

$$S = \{x \in \text{dom } f : f(x) \leq f(x_0)\}$$

trebuie să fie închisă. Impunerea condiției de închidere a mulțimii S asigură eliminarea posibilității convergenței algoritmului într-un punct de pe frontiera lui $\text{dom } f$. Dacă $\text{dom } f = \mathbb{R}^n$, atunci condiția este satisfăcută pentru orice x_0 . Fie

$$\varphi(\alpha) = f(x_k + \alpha d_k). \quad (1.1.3)$$

Deci, acum problema este ca plecând din punctul x_k să se determine lungimea pasului α_k , în direcția d_k , astfel încât $\varphi(\alpha_k) < \varphi(0)$.

Dacă α_k se determină astfel încât funcția f este minimizată în direcția d_k , adică

$$f(x_k + \alpha_k d_k) = \min_{\alpha > 0} f(x_k + \alpha d_k) \quad (1.1.4)$$

sau

$$\varphi(\alpha_k) = \min_{\alpha > 0} \varphi(\alpha), \quad (1.1.5)$$

atunci avem o *căutare liniară exactă* sau încă o *căutare liniară optimă*, iar α_k este *lungimea optimă* a pasului de deplasare.

Pe de altă parte, dacă α_k se alege astfel încât se obține o reducere acceptabilă a funcției obiectiv, adică reducerea $f(x_k) - f(x_k + \alpha_k d_k) > 0$ este acceptabilă, atunci avem o *căutare liniară inexactă, acceptabilă sau aproximativă*.

Menționăm că deoarece în situațiile practice concrete, atât din considerente teoretice cât și datorită costului de calcul prohibitiv – deci din considerente practice, mărimea optimă a lungimii pasului, cu foarte mici excepții (cazul funcțiilor pătratice de exemplu), nu poate fi determinată, rezultă că de obicei se utilizează o căutare liniară inexactă, care implementează algoritmi ce determină o reducere acceptabilă a valorilor funcției de minimizat. Este clar că din punctul de vedere al costului computațional este mult mai bine să ne mulțumim cu o reducere acceptabilă a valorilor funcției de-a lungul direcției de căutare, decât cu reducerea optimă (maximală). În principiu, structura unui algoritm general de direcții descendente este următoarea.

Algoritmul 1.1.1. (Algoritm general de optimizare fără restricții)

Pasul 1.	<i>Inițializare.</i> Se consideră un punct inițial $x_0 \in R^n$, precum și toleranța de terminare a iterațiilor algoritmului de minimizare ε suficient de mică. Se pune $k = 0$.
Pasul 2.	Se determină <i>direcția de descendență</i> d_k .
Pasul 3.	Se calculează <i>lungimea pasului</i> α_k astfel încât, de exemplu: $f(x_k + \alpha_k d_k) = \min_{\alpha \geq 0} f(x_k + \alpha d_k),$ sau utilizând condițiile Wolfe (căutare aproximativă).
Pasul 4.	Se pune $x_{k+1} = x_k + \alpha_k d_k.$
Pasul 5.	<i>Continuarea iterațiilor.</i> Dacă $\ \nabla f(x_{k+1})\ \leq \varepsilon$, atunci stop; altfel se pune $k = k + 1$ și se continuă cu pasul 2. ♦

Câteva precizări sunt necesare.

1. Observăm că aceasta este structura generală a unui algoritm de determinare a unui optim local al problemei (1.1.1). Determinarea optimului global este o problemă mult mai complicată și încă nu avem la dispoziție algoritmi generali pentru rezolvarea acestei probleme. Totuși, dacă funcția de minimizat f este convexă, atunci se demonstrează că orice optim local al funcției f este de asemenea un optim global.
2. Referitor la inițializarea algoritmului (vezi pasul 1), în general orice punct x_0 poate fi acceptat ca punct inițial. Pentru o problemă concretă de optimizare fără restricții un astfel de punct inițial se poate preciza din anumite considerații fizice, ceea ce constituie o foarte bună metodă de inițializare. În general, pentru probleme de optimizare (fără restricții) de obicei odată cu prezentarea problemei se furnizează și punctul inițial, care într-un anumit fel sens va caracteriza performanțele algoritmului de optimizare utilizat. Cu alte cuvinte, algoritmi de optimizare fără restricții se comportă diferit la inițializări diferite. Cel mai pregnant exemplu este metoda Newton. Bine inițializată, adică inițializată lângă soluție, algoritmul este pătratic convergent la punctul de optim. Inițializată departe de soluție, metoda Newton poate fi divergentă. Deseori, problemele sunt însoțite de anumite „puncte patologice” care pun algoritmi de optimizare în dificultăți majore. Pentru cazul metodelor de optimizare fără restricții, în afara cazurilor practice când din considerații fizice se poate stabili un punct inițial, de obicei nu se utilizează o procedură de inițializare care să determine punctul x_0 . Această practică se utilizează în cazul problemelor de optimizare (liniare sau neliniare) cu restricții pentru care se pot imagina proceduri foarte sofisticate de inițializare.
3. Calculul direcției de descendență d_k , din pasul 2, este elementul care caracterizează algoritmul de rezolvare a problemei (1.1.1). Altfel spus, metodele de optimizare se clasifică după modul de generare a direcției d_k . Cerința minimală asupra direcției d_k este aceea de descendență, adică direcția trebuie să fie astfel încât o mică deplasare de-a lungul acesteia să conducă la reducerea valorilor funcției f . Condiția de

descendență este $\nabla f(x_k)^T d_k < 0$. Aceasta se poate justifica foarte simplu prin utilizarea aproximării:

$$f(x_k + \alpha d_k) - f(x_k) \cong \alpha \nabla f(x_k)^T d_k,$$

de unde vedem imediat că negativitatea membrului drept din aproximarea de mai sus garantează reducerea valorilor funcției de minimizat când ne deplasăm cu un pas suficient de mic α de-a lungul direcției d_k .

4. Determinarea lungimii pasului de deplasare α_k din pasul 3 al algoritmului reprezintă un element crucial în ceea ce privește eficiența algoritmului de optimizare. Pentru calculul lui α_k se pot utiliza două tehnici generale: căutare liniară sau metoda regiunii de încredere. În esență, metoda de căutare liniară este o problemă de minimizare uni-dimensională. Procedurile moderne se bazează pe o căutare liniară economică. Acestea sunt descrise într-o serie de lucrări ca: Andrei [1999a (partea a IV-a), 2009c], [Dennis și Schnabel, 1983], [Peressini, Sullivan și Uhl, 1988]. Metoda regiunii de încredere încearcă determinarea lungimii pasului printr-o procedură de căutare a acesteia într-un domeniu definit de parametrul regiunii de încredere.
5. Testul de continuare a iterațiilor din pasul 5 implică utilizarea unor criterii de convergență care califică un punct ca soluția optimă a problemei.

1.2. METODE DE DIRECȚII DESCENDENTE

Algoritmul general de optimizare fără restricții se concretizează prin precizarea elementelor care-l definesc, în special în ceea ce privește direcția de căutare a optimului. Am văzut mai sus că se cunosc șase metode generale de optimizare fără restricții. Metodele de căutare directă nederivative încearcă determinarea punctului de minim utilizând numai valorile funcției de minimizat. Deși aceste metode sunt active (vezi de exemplu Powell [1998, 2000, 2001, 2002]), totuși acestea sunt utilizate numai pentru situația în care numărul de variabile este redus, expresia algebrică a gradientului este greu de obținut sau evaluarea gradientului este consumatoare de timp. Detalii asupra acestor metode se pot găsi în Andrei [2009c]. În cele ce urmează vom prezenta pe scurt metodele de direcții descendente care și-au dovedit eficacitatea în ceea ce privește rezolvarea problemelor reale de optimizare fără restricții.

1.2.1. METODA PASULUI DESCENDENT

Aceasta este una dintre cele mai vechi metode de optimizare fără restricții fiind propusă de Augustin Cauchy în 1848. Importanța ei este mai mult teoretică decât practică și într-un anumit sens se află la originea tuturor metodelor de direcții descendente. La fiecare iterație direcția recomandată de metoda pasului descendent este $-\nabla f(x_k)$, adică a gradientului cu semn schimbat al funcției de minimizat.

Iterațiile se calculează sub forma $x_{k+1} = x_k + \alpha_k d_k$, unde α_k este lungimea pasului de deplasare determinată din condiția de reducere a valorilor funcției $\varphi_k(\alpha) = f(x_k - \alpha \nabla f(x_k))$. Metoda este foarte simplu de implementat, cerințele

de memorie fiind de ordinul $O(n)$, unde n este numărul de variabile ale funcției f . Dacă α_k este determinat prin căutare liniară exactă, atunci metoda pasului descendent evoluează în pași perpendiculari. Adică, dacă $\{x_k\}$ este șirul de puncte generat de metodă, atunci vectorul care unește x_k cu x_{k+1} este ortogonal celui care unește x_{k+1} cu x_{k+2} . Această comportare în pași perpendiculari face ca metoda să fie lentă și deci computațional ineficientă. Pentru cazul funcțiilor pătratice convexe, rata de convergență este liniară. În general, rata de convergență nu este mai mare decât $[(\kappa-1)/(\kappa+1)]^2$, unde κ este numărul de condiționare a Hessianului funcției de minimizat [Andrei, 2004a]. O variantă a algoritmului este pasul descendent relaxat în care $x_{k+1} = x_k + \theta_k \alpha_k d_k$, unde θ_k este o variabilă aleatoare uniform distribuită în intervalul $[0,1]$ [Andrei, 2004c; 2005a,b]. În 1988 Barzilai și Borwein au propus un nou algoritm de gradient descendent care pentru lungimea pasului utilizează o aproximație în două puncte a ecuației secantei din cadrul metodelor quasi-Newton. Utilizând o strategie de globalizare bazată pe căutarea liniară nemonotonă, Raydan [1997] a arătat convergența globală a acestui algoritm pentru funcții nepătratice. Utilizând căutarea liniară cu backtracking, metoda a fost modificată și accelerată prin modificarea lungimii pasului ca în [Andrei, 2006a].

1.2.2. METODA NEWTON ȘI VARIANTELE SALE

Toate metodele Newton se bazează pe principiul de aproximare locală a funcției de minimizat printr-un model pătratic. În punctul x_k de-a lungul direcției d_k modelul pătratic al funcției f are expresia

$$f(x_k + d) \cong f(x_k) + g_k^T d + \frac{1}{2} d^T G_k d, \quad (1.2.1)$$

unde $g_k = \nabla f(x_k)$ și $G_k = \nabla^2 f(x_k)$. Minimul membrului drept din (1.2.1) este realizat de vectorul d_k care minimizează funcția pătratică

$$m_p(d) = g_k^T d + \frac{1}{2} d^T G_k d. \quad (1.2.2)$$

Deci, direcția Newton d_k satisface sistemul algebric liniar de n ecuații cu tot atâtea necunoscute

$$G_k d = -g_k, \quad (1.2.3)$$

cunoscut ca *sistemul Newton*. Această direcție este bine definită dacă matricea G_k este nesingulară. Când punctul inițial x_0 cu care se demarează calculele este suficient de aproape de soluția x^* a problemei se demonstrează că metoda Newton este pătratic convergentă [Andrei, 2004a,b]. Aceasta înseamnă că numărul de cifre semnificative precise din soluție se dublează la fiecare iterație. Principalele avantaje ale metodei Newton sunt:

1. Inițializată într-un punct suficient de apropiat de soluția optimă și dacă G_k este nesingulară, atunci metoda converge pătratic la soluție.
2. Pentru funcții pătratice, metoda Newton furnizează o soluție exactă într-o singură iterație.
3. Pentru componentele afine ale gradientului funcției de minimizat, la fiecare iterație, metoda este exactă.

Dezavantajele metodei Newton sunt:

1. Pentru multe probleme nu este global convergentă.
2. La fiecare iterație implică evaluarea Hessianului funcției de minimizat.
3. Implică rezolvarea la fiecare iterație a unui sistem algebric liniar care poate fi prost condiționat.

Aceste caracteristici ale metodei Newton se pot utiliza pentru elaborarea de noi algoritmi de optimizare fără restricții. Astfel, orice combinație a următoarelor trei tehnici poate constitui o modificare a metodei Newton:

1. Introducerea la fiecare iterație a unei căutări liniare de-a lungul direcției Newton, adică adoptarea iterației $x_{k+1} = x_k + \alpha_k d_k$, unde α_k este lungimea pasului determinată printr-o procedură de căutare liniară. Aceasta este o modificare de globalizare a metodei Newton.
2. Considerarea la fiecare iterație a unei aproximații (simetrice și pozitiv definite) a matricei Hessian, sau a unei aproximații a inversei matricei Hessian. Acestea au condus la metodele quasi-Newton.
3. Rezolvarea inexactă la fiecare iterație a sistemului Newton, adică determinarea aproximativă direcției Newton d_k din sistemul Newton.

Aceasta a condus la metoda Newton trunchiată sau inexactă.

Menționăm faptul că metoda Newton este foarte rar utilizată în aplicațiile practice concrete. Mai importante sunt variantele acesteia pe care le vom discuta imediat.

Metoda Newton modificată

Dacă Hessianul G_k al funcției de minimizat nu este o matrice pozitiv definită, atunci vom considera o modificare a acestei matrice de forma $E(x_k) = G_k + \mu_k I$, unde parametrul μ_k este ales astfel încât $E(x_k)$ să fie pozitiv definită. Pentru a determina μ_k observăm că pentru orice $\mu > 0$, $x^T (G_k + \mu I) x = x^T G_k x + \mu \|x\|^2$. Dacă definim

$$\hat{\mu}_k = \max_{\|x\|=1} |x^T G_k x|, \quad (1.2.4)$$

atunci se poate arăta că $\hat{\mu}_k$ este egal cu valoarea absolută a celei mai mari valori proprii a Hessianului funcției de minimizat. Dacă $\mu_k > \hat{\mu}_k$ și $x \neq 0$, atunci

$$x^T (G_k + \mu_k I) x = \|x\|^2 \left[\frac{x}{\|x\|} G_k \frac{x}{\|x\|} + \mu_k \right] > \|x\|^2 (-\hat{\mu}_k + \mu_k) > 0,$$

ceea ce arată că $E(x_k)$ este pozitiv definită.

Metoda Newton discretă

După cum se vede metoda Newton (standard) așa cum a fost prezentată, la fiecare iterație implică evaluarea a n funcții (componentele gradientului) și $n(n+1)/2$ funcții care definesc Hessianul G_k . Metoda Newton discretă înlocuiește evaluarea funcțiilor Hessianului prin determinarea unei aproximări prin diferențe finite a elementelor Hessianului. Fiecare coloană i a lui G_k se poate aproxima prin vectorul

$$\frac{g(x_k + h_i e_i) - g(x_k)}{h_i},$$

unde h_i este un pas convenabil ales pentru a balansa erorile de rotunjire (proportionale cu $1/h_i$) cu erorile de trunchiere (proportionale cu h_i) [Andrei, 1999a, capitolul 15]. Într-o aritmetică exactă metoda Newton discretă converge pătratic la soluție dacă fiecare h_i tinde la zero odată cu norma gradientului. Totuși, erorile de rotunjire limitează mărimea lui h_i și deci acuratețea soluției. Când gradientul devine foarte mic este posibilă o pierdere a acurateții prin deteriorarea numerică a numărătorului expresiei de mai sus. Deci, metoda Newton discretă nu este recomandată pentru rezolvarea problemelor de mari dimensiuni, cu excepția cazurilor în care Hessianul are o structură de matrice rară cunoscută și această structură este utilizată în schema de aproximare a elementelor Hessianului.

Metoda Newton compusă

Dacă nu ținem seama de efortul de calcul implicat de evaluarea funcției f , a gradientului g_k și a Hessianului G_k , atunci numărul de operații algebrice pe iterație este de ordinul $O(n^3)$ fiind reclamat de factorizarea matricei G_k sau a unei modificări a ei $G_k + \mu_k I$. Evident, pentru valori mari ale lui n factorizarea este o problemă foarte serioasă. O tehnică de reducere a efortului de calcul pe iterație este dată de așa numita *metoda Newton simplificată* în care la fiecare iterație se rezolvă sistemul algebric liniar $\nabla^2 f(x_0) d_k = -g_k$ în privința direcției d_k . Evident, în cazul metodei Newton modificate se rezolvă sistemul $(\nabla^2 f(x_0) + \mu_0 I) d_k = -g_k$. Observăm că metoda Newton simplificată implică factorizarea inițială a lui $G_0 = \nabla^2 f(x_0)$ și apoi utilizarea acestei factorizări de-a lungul tuturor iterațiilor. În acest caz efortul de calcul este de ordinul $O(n^2)$ operații algebrice, dar prețul plătit este că algoritmul are o convergență liniară.

Pentru a depăși situațiile extreme date de metoda Newton și metoda Newton simplificată este foarte natural să imaginăm o variantă a metodei Newton care între doi pași Newton consideră m pași Newton simplificați. Astfel, prin *metoda Newton compusă de ordinul m* înțelegem un proces iterativ în care pentru $i = 0, 1, \dots, m$ se rezolvă sistemele algebrice liniare

$$\nabla^2 f(x_k)d_i = -g(x_k + d_0 + d_1 + \dots + d_{i-1}) \quad (1.2.5)$$

în privința lui d_i , și estimația soluției de la pasul $k+1$ se calculează sub forma

$$x_{k+1} = x_k + \alpha_k(d_0 + d_1 + \dots + d_m), \quad (1.2.6)$$

pentru $k = 0, 1, \dots$. Ortega și Rheinboldt [1970, capitolul 10] demonstrează că în condiții standard metoda Newton compusă este superliniar convergentă. O variantă interesantă este dată de *metoda Newton compusă de ordinul 1* în care un pas al metodei Newton este compus cu un pas al metodei Newton simplificate sub forma: pentru $k = 0, 1, \dots$ se rezolvă sistemul liniar $\nabla^2 f(x_k)d_N = -g(x_k)$ în privința lui d_N , după care se rezolvă sistemul $\nabla^2 f(x_k)d_S = -g(x_k + d_N)$ în privința lui d_S cu care se calculează următoarea estimație $x_{k+1} = x_k + \alpha_k(d_N + d_S)$. Această variantă a metodei Newton compusă este atribuită lui Traub [1964]. Ortega și Rheinboldt [1970] demonstrează convergența cubică a acestei metode.

Metoda Newton compusă este eficientă mai ales pentru cazurile în care n este mare și efortul de calcul pentru evaluarea funcției f este mic. O situație de acest fel este întâlnită în cazul metodelor de punct interior de urmărire a traiectoriei pentru rezolvarea problemelor de programare liniară sau neliniară [Andrei, 1999a,b; 2009c].

Observăm că fiecare iterație a metodei Newton compusă de ordinul m se poate privi ca o iterație majoră și este rezultatul a $(m+1)$ iterații minore (interioare). Numărul mediu de operații algebrice pe o iterație interioară este de ordinul $O((n^3 + mn^2)/(m+1))$. Vedem că pentru m mare acesta este $O(n^2)$. Mai mult, viteza medie de convergență pentru iterațiile interioare este $(m+2)^{1/(m+1)}$, care pentru m mare tinde la 1. Deci, pentru valori mari ale lui m metoda Newton compusă de ordinul m se comportă asemănător metodei Newton simplificate. Ca atare, implementările eficiente în programe de calcul ale metodei Newton compuse nu numai că nu trebuie să modifice pe m la fiecare iterație, dar trebuie să mențină pe m relativ mic [Potra și Rheinboldt, 1986], [Potra, 1987].

Metoda Newton trunchiată

Această variantă a metodei Newton a fost introdusă de Dembo, Eisenstat și Steihaug [1982] și a fost analizată de Dembo și Steihaug [1983], Deuffhard [1990] și implementată în programe foarte sofisticate de Nash [1985], Nash și Sofer [1990], Schlick și Overton [1987] și Schlick și Fogelson [1992a, b]. Ideea de bază a acestei metode este că departe de soluție nu este necesar să rezolvăm sistemul Newton cu o foarte mare acuratețe. Ca atare, iterațiile care definesc o soluție a sistemului Newton se pot trunchia de îndată ce s-a obținut o soluție aproximativă a acestui sistem. De aici provine și numele de trunchiată sau inexactă a metodei. Deci pentru $k = 0, 1, \dots$ se determină d_k astfel încât

$$\|\nabla^2 f(x_k)d_k + g(x_k)\| \leq \eta_k \|g(x_k)\|, \quad (1.2.7)$$

după care se calculează următoarea estimatie $x_{k+1} = x_k + \alpha_k d_k$.

Problemele privind analiza acestei variante a metodei Newton sunt legate de modul de definire a șirului de „forțare” $\{\eta_k\}$. Acesta lucrează ca un „regulator” care controlează reziduul $\|\nabla^2 f(x_k)d_k + g(x_k)\|$. Analiza dată de Dembo, Eisenstat și Steihaug [1982] a condus la următorul rezultat. Fie f o funcție continuu diferențiabilă și $\{\eta_k\}$ un șir de numere reale cu $0 \leq \eta_k \leq \eta < 1$. Atunci metoda Newton trunchiată, inițializată într-un punct aflat într-o vecinătate a soluției, care generează șirul $\{x_k\}$ definit de

$$\nabla^2 f(x_k)d_k = -g(x_k) + r_k, \quad (1.2.8)$$

unde

$$\frac{r_k}{\|g(x_k)\|} \leq \eta_k, \quad x_{k+1} = x_k + d_k$$

este bine definit și converge cel puțin liniar la x^* . Dacă $\{\eta_k\}$ converge la zero, atunci $\{x_k\}$ converge superliniar la x^* . Dacă $\{\eta_k\} = O(\|g(x_k)\|^p)$, pentru $0 \leq p \leq 1$, atunci $\{x_k\}$ converge la x^* cu ordinul $1+p$.

Rezultatul de mai sus garantează convergența pătratică dacă la fiecare iterație pasul Newton aproximativ d_k verifică relația

$$\frac{\|\nabla^2 f(x_k)d_k + g(x_k)\|}{\|g(x_k)\|} \leq \min\{1, \|g(x_k)\|\}. \quad (1.2.9)$$

În practică deseori cantitatea din membrul stâng din (1.2.9), numită „reziduu relativ”, este calculată în timpul rezolvării aproximative a sistemului Newton. Astfel (1.2.9) furnizează un criteriu convenabil de oprire a iterațiilor în procesul de rezolvare a sistemului Newton. Una dintre cele mai bune metode de rezolvare aproximativă a sistemului Newton este metoda de gradient conjugat preconditionată, pe care o vom prezenta cu detalii în capitolele următoare ale acestei lucrări. Cea mai bună alegerea lui η_k este necunoscută. Singura cerință este ca $\eta_k \rightarrow 0$. Practic s-au utilizat următoarele două alegeri: $\eta_k = 1/k$ sau $\eta_k = \|g(x_k)\|$. A doua alegere implică o convergență pătratică [Dembo și Steihaug, 1983]. În aplicațiile practice ambele aceste alegeri sunt implementate. Metoda Newton trunchiată cunoaște două abordări.

- A) *Abordarea Schlick-Fogelson* [1992a,b]. Rezolvarea aproximativă a sistemului Newton (1.2.3) preconditionat, în care matricea Hessian este calculată analitic sau prin diferențe finite.
- B) *Abordarea Nash* [1985]. Aproximarea matricei Hessian utilizând actualizarea BFGS și rezolvarea aproximativă a sistemului Newton corespunzător.

În ambele abordări rezolvarea aproximativă a sistemului Newton se face prin metode de direcții conjugate. O justificare a acestei selecții este dată de faptul că metoda de direcții conjugate are proprietatea că modelul pătratic asociat funcției de minimizat în punctul curent este minimizat pe subspațiul generat de aceste direcții.

În continuare să prezentăm un algoritm de tip Newton trunchiat în care la fiecare iterație externă se calculează o aproximație B_k a Hessianului printr-o schemă de actualizare de tip BFGS, de exemplu – cunoscută ca abordarea Nash.

Algoritm 1.2.1. (TN: Newton trunchiat general – Nash)

Pasul 1.	<i>Inițializare.</i> Se consideră un punct inițial x_0 . Se pune $k = 0$. Se selectează toleranța ε de calcul a minimumului funcției f , precum și parametrii ρ și σ din condițiile de căutare liniară Wolfe.
Pasul 2.	Se calculează: $f(x_k)$, $g_k = \nabla f(x_k)$, precum și o aproximație B_k a Hessianului $\nabla^2 f(x_k)$ (de exemplu BFGS).
Pasul 3.	Se execută testul de oprire a iterațiilor. Dacă $\ g_k\ _\infty \leq \varepsilon(1 + f(x_k))$ sau $\ g_k\ _\infty \leq \varepsilon$, atunci stop, altfel se continuă cu pasul 4.
Pasul 4.	<i>Iterațiile interioare.</i> Se calculează o soluție aproximativă d_k a sistemului $B_k d_k = -g_k$ dată de algoritmul 1.2.2.
Pasul 5.	<i>Determinarea lungimii pasului.</i> Se calculează α_k care verifică condițiile Wolfe $f(x_k + \alpha d_k) \leq f(x_k) + \rho \alpha \nabla f(x_k)^T d_k,$ $\nabla f(x_{k+1})^T d_k \geq \sigma \nabla f(x_k)^T d_k.$
Pasul 6.	Se calculează $x_{k+1} = x_k + \alpha_k d_k$, se pune $k = k + 1$ și se continuă cu pasul 2. ♦

În pasul 4 al algoritmului 1.2.1, care implementează iterațiile interioare, direcția d_k se calculează prin algoritmul de gradient conjugat în care B_k acționează ca o matrice de preconditionare.

Algoritm 1.2.2. (Iterația interioară: Gradient conjugat)

Pasul 1.	Se consideră: $d_k^0 = 0$, $r_0 = -g_k$, $p_0 = r_0$ și $\delta_0 = r_0^T r_0$. Se pune $i = 0$.
Pasul 2.	Se calculează $q_i = B_k p_i$. Dacă $p_i^T q_i \leq \varepsilon \delta_i$, atunci dacă $i = 0$ se pune $d_k = p_0$, sau altfel dacă $i \neq 0$, atunci se pune $d_k = d_k^i$, stop.
Pasul 3.	Se calculează:

	$a_i = \frac{r_i^T r_i}{p_i^T q_i}, \quad d_k^{i+1} = d_k^i + a_i p_i, \quad r_{i+1} = r_i - a_i q_i.$ Dacă $\frac{\ r_{i+1}\ }{\ g_k\ } \leq \eta_k$, atunci se pune $d_k = d_k^i$ și stop.
Pasul 4.	Continuarea iterațiilor de gradient conjugat. Se calculează $b_i = \frac{r_{i+1}^T r_{i+1}}{r_i^T r_i}, \quad p_{i+1} = r_{i+1} + b_i p_i, \quad \delta_{i+1} = r_{i+1}^T r_{i+1} + b_i^2 \delta_i,$ se pune $i = i + 1$ și se continuă cu pasul 2. ♦

Alegerea șirului $\{\eta_k\}$ este una dintre cele mai dificile probleme din cadrul algoritmului Newton trunchiat. Dembo și Steihaug [1983] demonstrează că dacă

$$\eta_k = \min \left\{ \frac{1}{k}, \|g_k\|^t \right\},$$

unde $0 < t \leq 1$, atunci metoda Newton trunchiată converge cu o rată de convergență de ordinul $O(1+t)$. Această alegere a șirului $\{\eta_k\}$ conduce la un algoritm adaptiv pentru rezolvarea problemei de minimizare a funcției f . Când suntem departe de soluție $\|g_k\|$ este mare și ca atare criteriul $\|r_{i+1}\| \leq \eta_k \|g_k\|$ cere un efort de calcul redus. Când $g_k \rightarrow 0$, ceea ce implică $\eta_k \rightarrow 0$, criteriul de oprire a iterațiilor interne $\|r_{i+1}\| \leq \eta_k \|g_k\|$ forțează ca direcția dată de algoritmul de gradient conjugat 1.2.2 să fie apropiată de direcția Newton atât ca modul cât și ca direcție propriu-zisă. Aceasta asigură convergența „aproape pătratică”. Observăm simplitatea algoritmului 1.2.2. Acesta nu implică decât produse scalare și un produs matrice-vector, în care apare aproximarea B_k a matricei Hessian.

1.2.3. METODE QUASI-NEWTON

Să notăm cu d_k^N direcția Newton, adică $d_k^N = -\nabla^2 f(x_k)^{-1} g_k$. Un rezultat clasic, foarte important, stabilit de Dennis și Moré [1974, 1977] arată că algoritmul $x_{k+1} = x_k + \alpha_k d_k$ este superliniar convergent dacă și numai dacă

$$\alpha_k d_k = d_k^N + o(\|d_k^N\|). \quad (1.2.10)$$

Deci, pentru a atinge o viteză mărită de convergență este necesar ca algoritmul să aproximeze asimptotic pasul Newton. Acesta este principiul lui Dennis și Moré. Problema este cum putem face aceasta fără a evalua la fiecare iterație matricea Hessian. Răspunsul a fost dat de Davidon [1959] și apoi dezvoltat și popularizat de Fletcher și Powell [1963]. Ideea este de a demara calculele cu o aproximare oarecare a matricei Hessian și la fiecare iterație să se actualizeze această matrice prin utilizarea curburii problemei măsurată de-a lungul direcției de deplasare. Făcută cu grijă această actualizare conduce la metode deosebit de eficiente și

robuste, cunoscute sub numele de *metode quasi-Newton* sau încă *metode de metrică variabilă*. Detalii asupra acestor metode se găsesc în [Andrei, 2009c].

Această schemă computațională a fost una novatoare care a revoluționat domeniul optimizării fără restricții, furnizând o alternativă foarte serioasă pentru metoda Newton. De-a lungul timpului s-au propus mai multe variante de metode quasi-Newton, dar din 1970 s-a observat că, în general, metoda BFGS (Broyden, Fletcher, Goldfarb, Shanno) este cea mai bună. Algoritmul metodei BFGS este următorul.

Algoritmul 1.2.3. (Algoritmul BFGS)

Pasul 1.	<i>Inițializare.</i> Se consideră un punct inițial $x_0 \in R^n$, o matrice simetrică și pozitiv definită B_0 , care aproximează într-un anumit sens matricea Hessian, precum și toleranța de terminare a iterațiilor algoritmului de minimizare ε suficient de mică. Se pune $k = 0$.
Pasul 2.	Se determină direcția de căutare d_k ca soluție a sistemului $B_k d_k = -g_k$.
Pasul 3.	Se determină lungimea pasului α_k și se calculează $x_{k+1} = x_k + \alpha_k d_k$.
Pasul 4.	Se testează un criteriu de oprire a iterațiilor. De exemplu dacă $\ \nabla f(x_{k+1})\ \leq \varepsilon$, atunci stop.
Pasul 5.	Se actualizează aproximația matricei Hessian B_k sub forma $B_{k+1} = B_k - \frac{B_k s_k s_k^T B_k}{s_k^T B_k s_k} + \frac{y_k y_k^T}{y_k^T s_k}, \quad (1.2.11)$ unde $s_k = x_{k+1} - x_k$ și $y_k = g_{k+1} - g_k$. Se pune $k = k+1$ și se continuă cu pasul 2. ♦

Observăm că cele două matrice de corecție care apar în (1.2.11) au fiecare rangul unu. Ca atare, teorema de deplasare a valorilor proprii [Wilkinson, 1965] ne arată că prima matrice de corecție de rang unu care se scade, descrește valorile proprii – zicem că le deplasează la stânga. A doua care se adună deplasează valorile proprii la dreapta. Deci, trebuie să existe un anumit echilibru în ceea ce privește deplasarea valorilor proprii, altfel aproximația (1.2.11) ar putea deveni fie singulară, fie arbitrar de mare ceea ce conduce la eșecul metodei. Convergența globală a metodei BFGS se obține din studiul acestor deplasări ale valorilor proprii. Utilizând operatorii urma și determinantul unei matrice, Powell [1976] a măsurat efectul celor două matrice de corecție asupra lui B_k arătând că dacă f este convexă atunci pentru orice matrice simetrică și pozitiv definită B_0 și orice punct inițial x_0 metoda BFGS realizează $\liminf_{k \rightarrow \infty} \|g_k\| = 0$. Dacă în plus șirul $\{x_k\}$ converge la o soluție a problemei în care Hessianul este o matrice pozitiv definită, atunci rata de convergență a metodei este superliniară.

Analiza lui Powell a fost extinsă de Byrd, Nocedal și Yuan [1987] la clasa Broyden a metodelor quasi-Newton în care (1.2.11) este înlocuit de o expresie ceva mai generală:

$$B_{k+1} = B_k - \frac{B_k s_k s_k^T B_k}{s_k^T B_k s_k} + \frac{y_k y_k^T}{y_k^T s_k} + \Phi (s_k^T B_k s_k) v_k v_k^T, \quad (1.2.12)$$

unde scalarul $\Phi \in [0,1]$ și

$$v_k = \frac{y_k}{y_k^T s_k} - \frac{B_k s_k}{s_k^T B_k s_k}. \quad (1.2.13)$$

Dacă în (1.2.12) considerăm $\Phi = 0$, atunci obținem metoda (sau actualizarea) BFGS, în timp ce $\Phi = 1$ furnizează metoda DFP (Davidon, Fletcher, Powell). Pentru cazul problemelor convexe, Byrd, Nocedal și Yuan [1987] au demonstrat convergența globală și superliniară a tuturor metodelor din clasa Broyden cu excepția metodei DFP. Modul de abordare eșuează când $\Phi = 1$ și lasă cazul nerezolvat. Studiile numerice intensive au arătat că metoda quasi-Newton BFGS este cea mai bună în clasa Broyden.

Metoda quasi-Newton cu memorie limitată

Metoda quasi-Newton cu memorie limitată se referă în principal la metoda BFGS cu memorie limitată. Ideea este că în loc de a memora matricele H_k care aproximează inversa matricei Hessian, se memorează un anumit număr m de perechi de vectori (s_i, y_i) care le definesc. Produsul $H_k g_k$ care definește direcția de căutare este obținut prin executarea unui număr de produse scalare care implică vectorul g_k și cele mai recente m perechi de vectori (s_i, y_i) . După calculul unei noi iterații, cea mai veche pereche din mulțimea perechilor (s_i, y_i) se elimină și una nouă este adăugată. Astfel, algoritmul întotdeauna păstrează cele mai recente m perechi (s_i, y_i) cu care se construiește aproximarea inversei Hessianului la iterația următoare.

Să detaliem procesul de actualizare cu memorie limitată. Presupunem că iterația curentă este x_k și se dispune de m perechi (s_i, y_i) , $i = k-m, \dots, k-1$. Pentru început, așa după cum a recomandat Oren și Spedicato [1976], se formează matricea „inițială” $H_k^{(0)} = \gamma_k I$, unde

$$\gamma_k = \frac{s_k^T y_k}{y_k^T y_k}. \quad (1.2.14)$$

Pentru $k+1 > m$ actualizarea H_{k+1} se calculează astfel:

$$\begin{aligned} H_{k+1} &= (V_k^T V_{k-1}^T \cdots V_{k-m+1}^T) H_k^{(0)} (V_{k-m+1} \cdots V_{k-1} V_k) \\ &\quad + \rho_{k-m+1} (V_k^T \cdots V_{k-m+2}^T) s_{k-m+1} s_{k-m+1}^T (V_{k-m+2} \cdots V_k) \\ &\quad + \rho_{k-m+2} (V_k^T \cdots V_{k-m+3}^T) s_{k-m+2} s_{k-m+2}^T (V_{k-m+3} \cdots V_k) \end{aligned}$$

$$\begin{aligned}
& + \dots \\
& + \rho_{k-1} V_k^T s_{k-1} s_{k-1}^T V_k \\
& + \rho_k s_k s_k^T,
\end{aligned} \tag{1.2.15}$$

unde

$$V_k = I - \rho_k y_k s_k^T \quad \text{și} \quad \rho_k = \frac{1}{y_k^T s_k}.$$

Formula de actualizare (1.2.15) este cunoscută ca formula L-BFGS. Nash și Nocedal [1991] arată că valorile lui m din intervalul $3 \leq m \leq 7$ dau cele mai bune rezultate. Algoritmul acestei actualizări este următorul.

Algoritmul 1.2.4. (Algoritmul general L-BFGS cu actualizarea matricei H_k)

Pasul 1.	Inițializare. Se alege $x_0 \in R^n$, parametrul m , și matricea $H_0 \in R^{n \times n}$, precum și toleranța $\varepsilon > 0$ de oprire a iterațiilor. Se pune $k = 0$.
Pasul 2.	Dacă $\ g_k\ \leq \varepsilon$, stop.
Pasul 3.	Se calculează: $d_k = -H_k g_k,$ $x_{k+1} = x_k + \alpha_k d_k,$ unde α_k este lungimea pasului de deplasare, determinată prin căutare liniară exactă sau inexactă (de exemplu Wolfe).
Pasul 4.	Fie $\hat{m} = \min\{k, m-1\}$. Dacă $y_k^T s_k \leq 0$, atunci punem $H_{k+1} = I$ (pasul descendent) și eliminăm toate perechile (s_i, y_i) , $i = k - \hat{m}, \dots, k$. Dacă $y_k^T s_k > 0$, atunci punem $ \begin{aligned} H_{k+1} = & (V_k^T V_{k-1}^T \dots V_{k-\hat{m}}^T) H_0 (V_{k-\hat{m}} \dots V_{k-1} V_k) \\ & + \rho_{k-\hat{m}} (V_k^T \dots V_{k-\hat{m}+1}^T) s_{k-\hat{m}} s_{k-\hat{m}}^T (V_{k-\hat{m}+1} \dots V_k) \\ & + \dots \\ & + \rho_{k-1} V_k^T s_{k-1} s_{k-1}^T V_k \\ & + \rho_k s_k s_k^T. \end{aligned} $
Pasul 5.	Punem $k = k+1$ și se continuă cu pasul 2. ♦

Metodele quasi-Newton cu memorie limitată se pot privi ca extensii ale metodelor gradientului conjugat în ceea ce privește accelerarea convergenței acestora. Principalul avantaj al acestor metode constă în faptul că utilizarea memoriei poate fi controlată de utilizator. În același timp, metodele quasi-Newton cu memorie limitată se pot privi și ca metode quasi-Newton în care utilizarea memoriei este restricționată. Simplitatea acestor metode se manifestă în aceea că acestea nu implică cunoașterea structurii matricei Hessian, sau a separabilității acesteia.

Metoda quasi-Newton partiționată

Griewank și Toint [1982a] au demonstrat că orice funcție f cu Hessianul matrice rară este parțial separabilă, adică se poate scrie sub forma:

$$f(x) = \sum_{i=1}^{ne} f_i(x), \quad (1.2.16)$$

unde fiecare din cele ne funcții „elementare” f_i depind numai de câteva variabile. Aceasta constituie fundamentul pentru așa numitele metode quasi-Newton partiționate pentru optimizare de mari dimensiuni. Ideea care stă la baza acestei metode este de a exploata structura parțial separabilă (1.2.16) și de a actualiza o aproximație B_k^i a Hessianului fiecărei funcții elementare f_i . Aceste matrice se pot asambla pentru a defini aproximația B_k a Hessianului funcției f . Complicația cu această abordare este că chiar dacă $\nabla^2 f(x^*)$ este pozitiv definită, câteva funcții elementare pot fi concave, astfel încât metoda BFGS nu poate fi întotdeauna aplicată. Direcțiile de căutare în cadrul metodei Newton partiționate sunt determinate ca soluții ale sistemului liniar

$$\left(\sum_{i=1}^{ne} B_k^i \right) d_k = -g_k, \quad (1.2.17)$$

în interiorul unei regiuni de încredere, utilizând o iterație de gradient conjugat trunchiată. Dacă se detectează o direcție de curbă negativă, atunci iterațiile de gradient conjugat sunt oprite (în aceasta constă trunchierea) și d_k este selectat ca această direcție de curbă negativă.

Metodele quasi-Newton partiționate lucrează bine în practică și reprezintă o contribuție importantă în domeniul optimizării neliniare. Multe probleme practice sunt formulate ca în (1.2.16), sau se pot transforma la această structură. În cazul problemelor convexe convergența globală a acestor metode este similară cu a metodei BFGS, tot ceea ce trebuie să presupunem este convexitatea tuturor funcțiilor elementare f_i . În această ipoteză, Griewank [1991] arată că metoda quasi-Newton partiționată este global convergentă chiar dacă sistemul (1.2.17) este rezolvat inexact.

1.2.4. METODA REGIUNII DE ÎNCREDERE

După cum am văzut, ideea pe care se bazează metoda Newton este de a aproxima funcția de minimizat $f(x)$ în jurul punctului curent x_k printr-un model pătratic

$$q^{(k)}(s) = f(x_k) + g_k^T s + \frac{1}{2} s^T G_k s, \quad (1.2.18)$$

și de a utiliza minimumul s_k al lui $q^{(k)}(s)$ pentru a defini următoarea aproximație a minimumului x^* sub forma

$$x_{k+1} = x_k + s_k. \quad (1.2.19)$$

Astfel proiectată metoda garantează numai convergența locală, adică dacă s este suficient de mic, atunci metoda este global convergentă. Prin introducerea căutării liniare metoda devine global convergentă. În acest context, vedem că metoda Newton este o compunere de două aplicații. Prima utilizează modelul pătratic $q^{(k)}(s)$ pentru a determina o direcție de căutare. A doua calculează o lungime convenabilă α_k a pasului de deplasare de-a lungul acestei direcții. Deși acest mod de abordare este foarte bun asigurând o bună rată de convergență atât pentru metoda Newton cât și pentru variantele ei, totuși aceasta nu utilizează suficient de bine modelul pătratic. Cu alte cuvinte, modelul pătratic nu este utilizat în determinarea lungimii pasului. Pe de altă parte, după cum am văzut, acest mod de operare eșuează dacă Hessianul nu este o matrice pozitiv definită.

Metoda regiunii de încredere nu numai că înlocuiește căutarea liniară pentru a obține convergența globală, dar depășește problemele cauzate de nepozitiv definirea matricei Hessian. În plus, aceasta produce o reducere mult mai pronunțată în valorile funcției de minimizat decât o face căutarea liniară.

În metoda regiunii de încredere pentru început se definește o *regiune* în jurul punctului curent

$$\Omega_k = \{x : \|x - x_k\| \leq \Delta_k\}, \quad (1.2.20)$$

unde Δ_k este *raza* regiunii Ω_k , în care există încrederea că modelul pătratic $q^{(k)}(s)$ constituie o bună aproximare a funcției f . În continuare se determină un pas s_k care aproximează minimumul modelului pătratic $q^{(k)}(s)$ în regiunea de încredere de mai sus. Cu alte cuvinte, se determină s_k astfel încât $x_k + s_k$ este cea mai bună aproximație a minimumului lui $q^{(k)}(s)$ pe sfera generalizată

$$\{x_k + s : \|s\| \leq \Delta_k\} \quad (1.2.21)$$

centrată în punctul x_k de rază Δ_k . Dacă pasul s_k nu este acceptabil, atunci se reduce raza regiunii de încredere și se procedează ca mai sus pentru a determina un nou minim în noua regiune de încredere. Observăm că astfel definită metoda reține rata de convergență locală rapidă a metodelor Newton sau quasi-Newton având în același timp proprietatea de convergență globală. Deoarece pasul este restricționat la regiunea de încredere, metoda se mai numește *metoda pasului restricționat*.

Cu acestea, un model al subproblemei metodei regiunii de încredere este

$$\min q^{(k)}(s) = f(x_k) + g_k^T s + \frac{1}{2} s^T B_k s$$

referitor la

$$\|s\| \leq \Delta_k,$$

(1.2.22)

unde $\Delta_k > 0$ este raza regiunii de încredere și B_k este o aproximație simetrică a matricei Hessian G_k . De obicei, în problema (1.2.22) se utilizează norma l_2 (norma euclidiană), așa încât s_k este minimumul lui $q^{(k)}(s)$ pe sfera centrată în x_k

de rază Δ_k . Dacă se utilizează alte norme, atunci se obțin alte forme ale regiunii de încredere. Dacă în (1.2.22) se pune $B_k = \nabla^2 f(x_k)$, atunci se obține chiar *metoda Newton a regiunii de încredere*.

Astfel prezentată, metoda regiunii de încredere ridică mai multe chestiuni: cum se alege Δ_k , cum se rezolvă subproblema (1.2.22), care sunt criteriile de oprire a iterațiilor, etc. detalii descrise în [Andrei, 2009c].

1.3. STUDIU NUMERIC (NEWTON TRUNCHIAT VERSUS L-BFGS)

În această secțiune vom prezenta un studiu numeric în care vom detalia profilul de performanță Dolan-Moré [2002] al algoritmului Newton trunchiat în abordarea B a lui Nash în comparație cu algoritmul quasi-Newton BFGS cu memorie limitată (LBFGS). Metoda Newton trunchiată este implementată în pachetul TN elaborat de Nash [1985]. Metoda quasi-Newton BFGS cu memorie limitată (LBFGS) este prezentată de Liu și Nocedal [1989] și implementată de Nocedal [1980]. În LBFGS condițiile Wolfe de căutare liniară sunt implementate în maniera dată de Moré și Thuente [1990, 1992, 1994]. Detalii privind pachetele TN și LBFGS, modul de utilizare a acestora, precum și experiența computațională cu aceste pachete se găsesc în lucrările [Andrei, 2007k] și respectiv [Andrei, 2008r].

Pentru aceasta considerăm un set de 75 de funcții de test, în formă generalizată sau extinsă [Andrei, 2008d]. Pentru fiecare funcție considerăm un număr de 10 experimente numerice în care numărul de variabile este crescător: $n = 1000, 2000, \dots, 10000$. Deci în total rezolvăm un tren de 750 de probleme de test de optimizare fără restricții. Toți algoritmi sunt implementați în Fortran și utilizează aceeași procedură de determinare a lungimii pasului dată de condițiile Wolfe, precum și același criteriu de oprire a iterațiilor care se referă la norma infinit a gradientului: $\|g_k\|_\infty \leq 10^{-6}$. Numărul maxim de iterații admis de fiecare algoritm este limitat la 10000. În acest studiu numeric vom compara LBFGS pentru 4 valori ale parametrului m cu TN. Din cele 750 de probleme le reținem numai acelea pentru care diferența dintre valorile optime ale funcției obiectiv, determinate de fiecare algoritm în parte, este mai mică decât 10^{-3} , adică dacă f_i^{LBFGS} și f_i^{TN} sunt valorile optime ale funcției f_i , $i = 1, \dots, 750$, determinate de algoritmul LBFGS și respectiv TN, atunci în comparațiile numerice vom reține numai acele experimente numerice pentru care $|f_i^{LBFGS} - f_i^{TN}| \leq 10^{-3}$.

Figurile 1.3.1 – 1.3.4 arată profilurile de performanță Dolan-Moré [2002] ale acestor algoritmi. În figura 1.3.1 vedem profilul de performanță al lui LBFGS ($m = 3$) versus TN. În tabelul din această figură vedem că din cele 682 de probleme care satisfac criteriul de mai sus, din punctul de vedere al timpului de calcul, algoritmul LBFGS ($m = 3$) a fost mai bun în 470 de probleme.

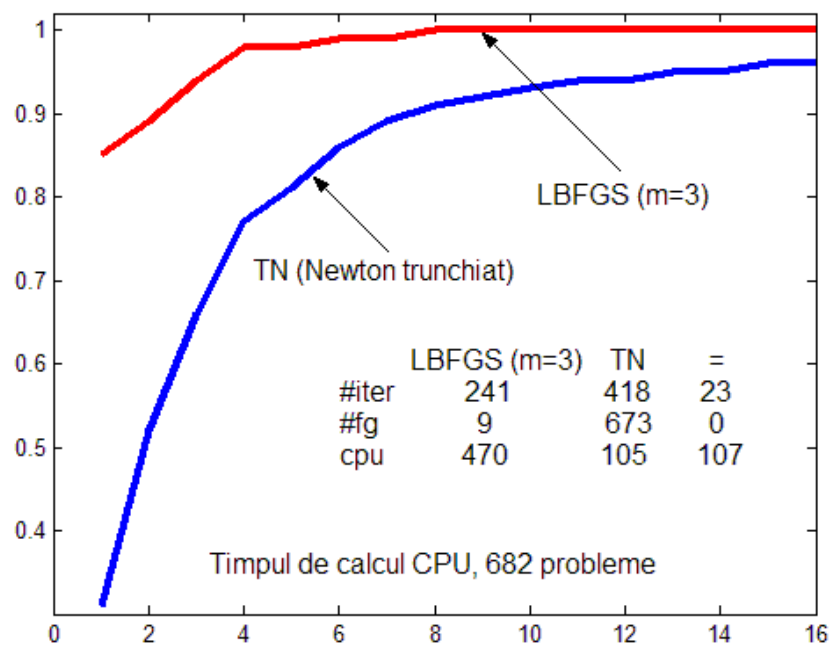


Fig. 1.3.1. Profilul de performanță L-BFGS ($m = 3$) versus TN (Newton Trunchiat).

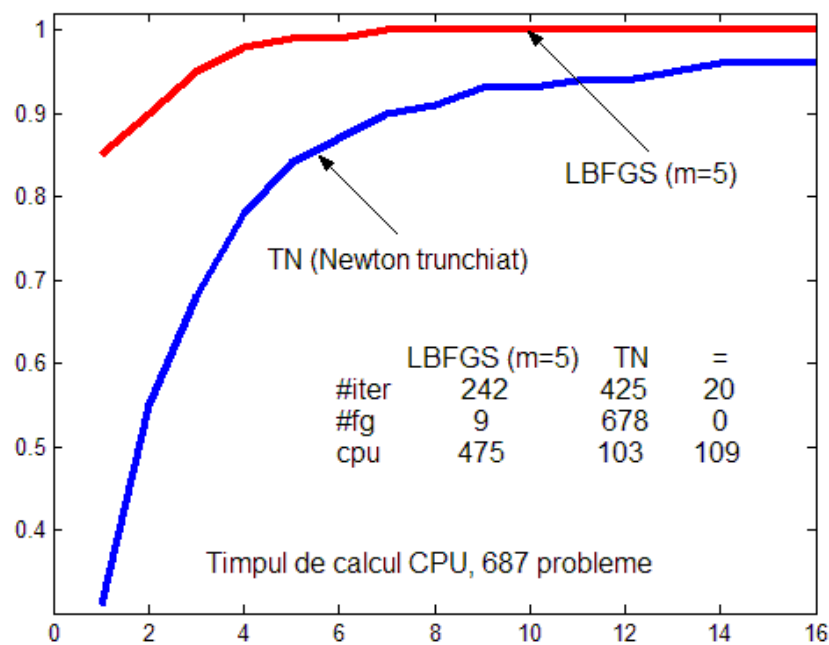


Fig. 1.3.2. Profilul de performanță L-BFGS ($m = 5$) versus TN (Newton Trunchiat).

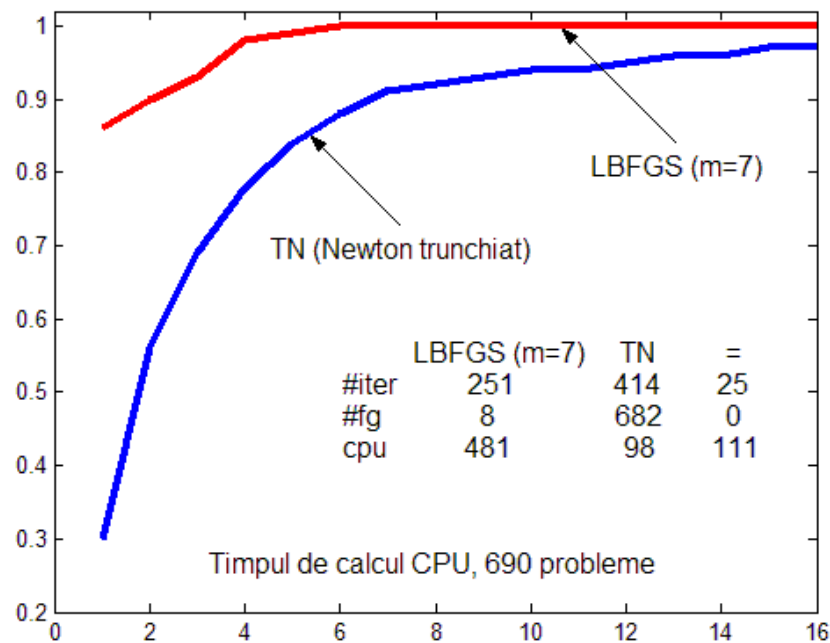


Fig. 1.3.3. Profilul de performanță L-BFGS ($m = 7$) versus TN (Newton Trunchiat).

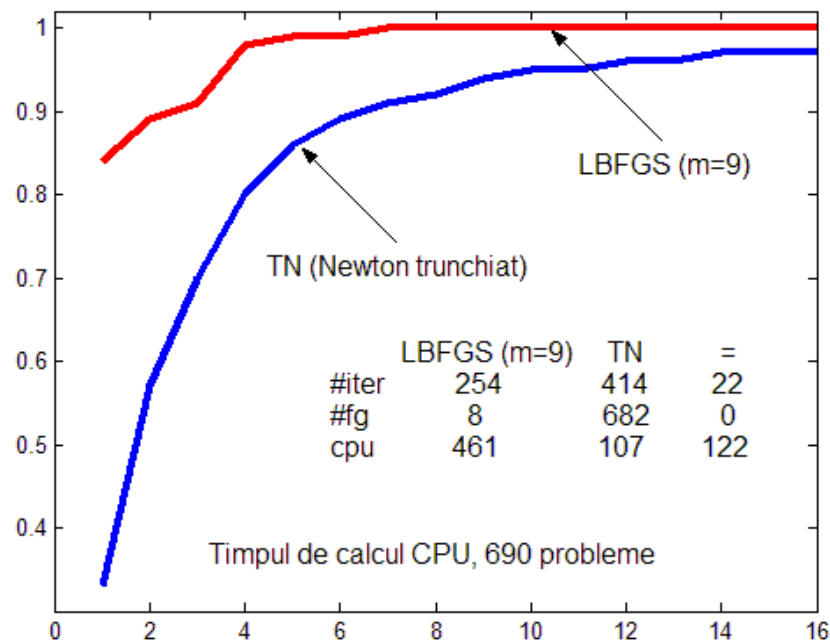


Fig. 1.3.4. Profilul de performanță L-BFGS ($m = 9$) versus TN (Newton Trunchiat).

Algoritmul TN a fost mai bun în doar 105 probleme, ambii algoritmi au realizat același timp de calcul în 107 probleme.

Din punctul de vedere al iterațiilor și al numărului de evaluări ale funcției de optimizat și a gradientului ei, vedem că metoda TN este mai bună. Din figura 1.3.1 notăm că referitor la numărul de iterații algoritmul TN a fost mai bun în 418 probleme, adică pentru 418 probleme din cele 682 algoritmul TN le-a rezolvat într-un număr de iterații mai mic decât LBFGS ($m = 3$). Metoda TN cere mai puține iterații, dar timpul de calcul pe iterație este mult mai mare.

Partea din stânga din figurile de mai sus arată procentajul din problemele considerate în acest studiu pentru care un algoritm a fost mai bun. Partea din dreapta arată procentajul problemelor care au fost rezolvate cu succes de fiecare algoritm în parte. Curba superioară corespunde algoritmului care a rezolvat cele mai multe probleme într-un timp de calcul care este un multiplu al celui mai bun timp de calcul. Vedem că, cel puțin pentru acest set de probleme de test, în privința timpului de calcul, algoritmul LBFGS pentru toate cele 4 valori ale lui m este superior algoritmului TN. Diferențele între acești doi algoritmi sunt semnificative.

Este important de observat faptul că valori mari ale lui m nu conduc la o creștere a performanței metodei quasi-Newton cu memorie limitată LBFGS. Creșterea numărului de perechi (s_i, y_i) implică un timp de calcul mai mare în implementarea lui (1.2.15). Aceasta arată că metoda LBFGS are un anumit caracter conservativ în sensul că în produsul $-H_{k+1}g_{k+1}$ care definește direcția d_{k+1} nu este nevoie de toată actualizarea inversei matricei Hessien, câteva perechi de vectori (s_i, y_i) sunt suficiente.

În capitolele următoare ale acestei lucrări vom utiliza în mod constant acești algoritmi în sensul de a prezenta profilurile de performanță Dolan-Moré ale algoritmilor de gradient conjugat versus LBFGS și TN.