

CURS 8: METODE DE OPTIMIZARE PARAMETRICĂ

Problemele de optimizare vizează extremizarea (maximizarea sau minimizarea) unui **criteriu de performanță**. Acesta din urmă poate fi o *funcție* – caz în care este vorba despre *optimizarea parametrică* – sau o *funcțională* (o funcție de funcții), când este vorba despre *optimizarea dinamică*. Acest paragraf este dedicat ilustrării comparative a principalilor algoritmi de optimizare parametrică.

1.1 Formularea problemei. Definiții

Criteriul de performanță care trebuie extremizat se mai numește **funcție scop**, **funcție obiectiv** sau **funcție criteriu** fără impuneri de legături sau/și restricții. Aceasta este o funcție de n argumente reale, $f(x_1, x_2, \dots, x_n)$, a cărei extremizare înseamnă găsirea unui vector $\underline{x}^* = \begin{bmatrix} x_1^* & x_2^* & \dots & x_n^* \end{bmatrix}^T \in \mathbb{R}^n$, corespunzător unui punct din spațiul n -dimensional, pentru care funcția este *optimă* (maximă sau minimă). Maximizarea și minimizarea unei funcții nu se deosebesc formal, întrucât maximumul lui f se obține pentru minimumul lui $-f$.

Pentru $n \leq 2$, funcția obiectiv admite o reprezentare geometrică, și anume:

- bidimensională: $y = f(x_1)$ este ecuația unei curbe în planul (x_1, y) ;
- tridimensională: $y = f(x_1, x_2)$ este ecuația unei suprafețe în sistemul de coordonate (x_1, x_2, y) (*suprafața de răspuns*).

Generalizarea acestor reprezentări pentru o funcție obiectiv de n variabile independente necesită un sistem de $n+1$ axe independente, perpendiculare două câte două într-un sistem cartezian. Pentru $n \geq 3$ aceasta este evident imposibil de realizat în spațiul tridimensional și de aceea dezvoltările teoretice – cu introducerea noțiunilor de „hiperplan”, „hipersuprafață”, „hiperspațiu” – sunt uneori mai dificil de urmărit.

Este utilă reluarea a două definiții cunoscute din matematică.

Definiția 1: Fie funcția scalară *continuă* și *derivabilă* de variabilă vectorială $f(\underline{x})$, $\underline{x} \in \mathbb{R}^n$. Se numește **gradientul** funcției f funcția *vectorială* de variabilă vectorială notată cu $\nabla f(\underline{x})$, care conține derivatele parțiale de ordinul întâi ale lui f în raport cu fiecare dintre componentele lui $\underline{x}$:

$$(1) \quad f'(\underline{x}) = \nabla f(\underline{x}) \triangleq \begin{bmatrix} \frac{\partial f(\underline{x})}{\partial x_1} & \frac{\partial f(\underline{x})}{\partial x_2} & \dots & \frac{\partial f(\underline{x})}{\partial x_n} \end{bmatrix}^T$$

Funcția $-\nabla f(\underline{x})$ se numește **antigradientul** funcției f . Punctele de *extrem*, $\underline{x}^*$, ale unei funcții de variabilă vectorială sunt caracterizate de *anularea gradientului*: $\nabla f(\underline{x}^*) = 0$. Direcția vectorului gradient indică spre maximumul unei funcții concave, iar a antigradientului indică spre minimumul unei funcții convexe.

Definiția 2: Se numește (**matricea**) **hessian** a funcției de variabilă vectorială $f(\underline{x})$ funcția **matricială** de variabilă vectorială notată cu $H(\underline{x})$, obținută printr-o nouă derivare a gradientului în raport cu componentele vectorului $\underline{x}$:

$$(2) \quad f''(\underline{x}) = H(\underline{x}) \triangleq \begin{bmatrix} \frac{\partial^2 f(\underline{x})}{\partial x_1^2} & \frac{\partial^2 f(\underline{x})}{\partial x_1 \partial x_2} & \cdots & \frac{\partial^2 f(\underline{x})}{\partial x_1 \partial x_n} \\ \frac{\partial^2 f(\underline{x})}{\partial x_2 \partial x_1} & \frac{\partial^2 f(\underline{x})}{\partial x_2^2} & \cdots & \frac{\partial^2 f(\underline{x})}{\partial x_2 \partial x_n} \\ \cdots & \cdots & \cdots & \cdots \\ \frac{\partial^2 f(\underline{x})}{\partial x_n \partial x_1} & \frac{\partial^2 f(\underline{x})}{\partial x_n \partial x_2} & \cdots & \frac{\partial^2 f(\underline{x})}{\partial x_n^2} \end{bmatrix}$$

1.2 Clasificarea metodelor de optimizare parametrică

Principalele metode de optimizare parametrică sunt sintetizate în tabelul 2.9; ele sunt detaliate în paragrafele următoare. Cele două mari clase de metode – *indirecte* și *directe* – se deosebesc atât principial, cât și din punctul de vedere al performanțelor (viteza de convergență, timpul de calcul și memoria necesară algoritmilor rezultați). Astfel, metodele directe sunt în general mai *lent convergente* decât metodele indirecte, dar necesită un volum de memorie mai mic. Nivelul admisibil al compromisului acceptat în ce privește performanțele este un factor în alegerea uneia sau alteia dintre metode.

A. indirecte (folosesc valorile funcției $f(\underline{x})$ și pe cele ale derivatelor ei)	A.1 metode de gradient (Cauchy)	
	A.2 metode Newton (de gradient de ordinul al II-lea)	
	A.3 metode ale direcțiilor conjugate	
B. directe (folosesc numai valorile funcției $f(\underline{x})$, fără a le utiliza pe cele ale derivatelor ei)	B.1 metoda explorării exhaustive	
	B.2 metode de eliminare	<i>unidimensională</i>
		<i>multidimensională</i>
	B.3 metode de căutare pe bază de hiperpoliedre (simplexuri) exploratoare	
	B.4 metode de căutare aleatoare (Monte Carlo)	
	B.5 metode de căutare unidirecțională (unidimensională) $\equiv$ metode de relaxare (Gauss)	

Tabel 1 Clasificarea metodelor de optimizare parametrică

1.3 Metode indirecte

Metodele indirecte se mai numesc și *metode de urcare (coborâre)*; pentru simplificare, pînă-un lejer abuz de limbaj, termenul de „metode de gradient” denotă în mod generic întreaga clasă a metodelor indirecte. Esența lor constă în **găsirea punctului de anulare a gradientului** (acesta fiind punctul de extrem căutat), pornind dintr-un punct inițial dat. În calculul aproximativ, termenul de „anulare” înseamnă, de fapt, situarea valorii absolute sub o anume limită (toleranță) dată, considerată ca suficientă.

Metodele de gradient (Cauchy) și cele de gradient de ordinul al II-lea (Newton) se

bazează pe aproximări de ordinul I, respectiv de ordinul al II-lea, ale dezvoltării în serie Taylor a funcției obiectiv $f(\underline{x})$ în jurul punctului de optim, $\underline{x}^*$:

$$(3) \quad f(\underline{x}) \cong f(\underline{x}^*) + f'(\underline{x}^*)^T \cdot (\underline{x} - \underline{x}^*)$$

$$(4) \quad f(\underline{x}) \cong f(\underline{x}^*) + f'(\underline{x}^*)^T \cdot (\underline{x} - \underline{x}^*) + \frac{1}{2} \cdot (\underline{x} - \underline{x}^*)^T \cdot f''(\underline{x}^*) \cdot (\underline{x} - \underline{x}^*)$$

Metodele Newton au o convergență mai bună decât cele de gradient simplu, dar prezintă inconvenientul unui timp de calcul și al unui volum de memorie mai mari, pentru că necesită calculul matricei hessian la fiecare iterație de căutare.

În formula (23) gradientul arată direcția ratei maxime de creștere a funcției f . Pe această formulă se bazează **metoda celei mai mari pante** (dacă optimul este un **maxim**) sau **metoda celei mai abrupte coborâri** (dacă optimul este un **minim**). Mai jos se prezintă descrierea în meta-limbaj a celui mai simplu algoritm de gradient (notația $\|\cdot\|$ semnifică norma vectorială).

Date de intrare: ε (toleranța suficientă), $\underline{x}^{(0)}$ (punctul inițial)

$i \leftarrow 0$

Repetă

#1. Se determină direcția de căutare la pasul i (dată de versorul asociat gradientului):

$$(5) \quad \underline{d}^{(i)} = \pm \frac{\nabla f(\underline{x}^{(i)})}{\|\nabla f(\underline{x}^{(i)})\|}, \text{ unde „+/-” corespund maximizării/minimizării}$$

#3. Se alege arbitrar $p^{(i)}$, pasul de deplasare pe direcția $\underline{d}^{(i)}$.

#2. Se calculează noul punct de evaluare a gradientului:

$$(6) \quad \underline{x}^{(i+1)} = \underline{x}^{(i)} + p^{(i)} \cdot \underline{d}^{(i)}$$

$i \leftarrow i+1$

Până când $\|\nabla f(\underline{x}^{(i)})\| \leq \varepsilon$.

Se observă că algoritmul de mai sus necesită cunoașterea a priori a tipului de extrem căutat (maxim sau minim). Implementarea în Matlab se poate face sub forma unei funcții care primește punctul de start al căutării, x_0 (la cazul general acesta este un vector) și toleranța dorită, ε , și returnează valoarea de extrem a unei funcții (la cazul general, vectoriale) cunoscute, x , și numărul de iterații în care s-a obținut aceasta, nr_it . S-a utilizat o variabilă locală, $itermax$, pentru a forța ieșirea din ciclare dacă extremul nu poate fi găsit.

```
function [x,nr_it]=opt_grad(x0,epsilon)
    nr_it=1;x_curent=x0;itermax=1500;%număr maxim de iterații
    h=0.001;%variație utilizată în calculul gradientului
    p=0.001;%se alege un pas constant de deplasare
    for k=1:n,
        v=zeros(1,n);v(k)=1;
        grad_curent(k)=(f1(x0+h/2*v)-f1(x0-h/2*v))/h;
    end;
```

```

while (norm(grad_curent)>=epsilon)&(nr_it<=itermax),
    %implementarea formulei (2.29) folosind (2.28);
    %pentru o funcție de maximizat „-” se înlocuiește cu „+”
    x_viitor=x_curent-p*grad_curent/norm(grad_curent);
    for k=1:n,
        v=zeros(1,n);v(k)=1;
        grad_viitor(k)=(f1(x_viitor+h/2*v)-f1(x_viitor-h/2*v))/h;
    end;
    x_curent=x_viitor;
    grad_curent=grad_viitor;
    nr_it=nr_it+1;
end;
x=x_curent;

```

Funcția de mai sus face uz de apelul funcției Matlab `norm`, care calculează norma vectorială sau matricială; ea realizează *minimizarea* unei funcții care trebuie să se afle în fișierul `f1.m`; în particular ea poate fi o funcție scalară:

```

function y=f1(x)
y=5*x^2+2*x+13;

```

Valoarea în care se atinge minimul acestei funcții este -0.2 , care se obține după un număr de iterații mai mic sau mai mare, depinzând de condiția inițială și de pasul de deplasare ales. Într-adevăr, apelul:

```
[x,nr_it]=opt_grad(1,1e-3)
```

produce rezultatele:

```

x =
    -0.2000

nr_it =
    1201

```

iar apelul:

```
[x,nr_it]=opt_grad(-1,1e-3)
```

are drept rezultat:

```

x =
    -0.2000

nr_it =
    801

```

Numărul mare de iterații se datorează folosirii unui pas foarte mic ($p=0.001$); folosirea unui astfel de pas se justifică în cazul extremelor abrupte („creste” ascuțite sau „văi” abrupte), unde componentele gradientului au variații mari.

Algoritmul Fletcher-Reeves (*metoda gradientilor conjugați*) este o procedură puternică de determinare a *minimului* local al unei funcții generale, $f(\underline{x})$. De această dată, la fiecare iterație i se definește o nouă direcție de căutare, $\underline{p}^{(i)}$, ca o combinație liniară între vectorul gradient la iterația curentă, $\nabla f(\underline{x}^{(i)})$, și direcțiile de la iterațiile anterioare, $\left\{ \underline{p}_j^{(i)} \right\}_{j=0, i-1}$.

Etaple algoritmului sunt listate mai jos.

1. Se alege un punct inițial caracterizat de vectorul $\underline{x}^{(0)}$.
2. Se stabilește direcția inițială drept direcția negativă a gradientului în $\underline{x}^{(0)}$:

$$\underline{p}^{(0)} = -\nabla \left(f \left(\underline{x}^{(0)} \right) \right)$$

3. La fiecare iterație i se determină minimul în raport cu parametrul α al funcției obiectiv în direcția $\underline{p}^{(i)}$, conform relației:

$$(7) \quad f \left(\underline{x}^{(i+1)} \right) = f \left(\underline{x}^{(i)} + \alpha \cdot \underline{p}^{(i)} \right)$$

Se determină astfel punctul $\underline{x}^{(i+1)}$.

4. Se determină noua direcție de căutare, $\underline{p}^{(i+1)}$, din punctul $\underline{x}^{(i+1)}$:

$$(8) \quad \underline{p}^{(i+1)} = -\nabla f \left(\underline{x}^{(i+1)} \right) + \frac{\left\| \nabla f \left(\underline{x}^{(i+1)} \right) \right\|^2}{\left\| \nabla f \left(\underline{x}^{(i)} \right) \right\|^2} \cdot \underline{p}^{(i)}$$

5. Dacă $f \left(\underline{x}^{(i+1)} \right) \geq f \left(\underline{x}^{(i)} \right)$, atunci punctul de minim a fost găsit: $\underline{x}^* \leftarrow \underline{x}^{(i)}$, STOP;

altfel se reia de la etapa 3.

În cele ce urmează este dat programul Matlab (`fletreev.m`) care implementează algoritmul Fletcher-Reeves; acesta va fi testat pentru aceeași funcție scalară de gradul al II-lea conținută în fișierul `f1.m`, cu un minim la -0.2 .

```
%inițializări
i=1;x(i,:)=x0;n=length(x0);itermax=1000;%număr maxim de iterații
h=0.0001;%variație utilizată în calculul gradientului
for k=1:n,
    v=zeros(1,n);v(k)=1;
    %la fiecare pas gradientul este un vector de dimensiune n
    grad(i,k)=(f1(x0+h/2*v)-f1(x0-h/2*v))/h;
end;
p(i,:)=-grad(i,:);
stop=0;%variabilă booleană care arată găsirea minimului local căutat
while (~stop)&(i<=itermax),
    alfa=0:0.01:10;
    y_min=f1(x(i,:));
    %se determină minimul funcției când alfa variază în intervalul (ales arbitrar) [0;10]
    %cu pasul 0.01
    stop=1;
    for j=1:length(alfa),
        y=f1(x(i,:)+alfa(j)*p(i));%formula (2.30)
        if y<y_min,
            y_min=y;
            x(i+1,:)=x(i,:)+alfa(j)*p(i);
            stop=0;
        end;
    end;
```

```

end;%în acest moment noul punct de căutare este x(i+1)
if (~stop)
    for k=1:n,
        v=zeros(1,n);v(k)=1;
        grad(i+1,k)=(f1(x(i+1,:)+h/2*v)-f1(x(i+1,:)-h/2*v))/h;
    end;
    %noua direcție de căutare se determină cu formula (2.31)
    p(i+1,:)=-grad(i+1,:)+...
        norm(grad(i+1,:))^2/norm(grad(i,:))^2*p(i,:);
end;
i=i+1;
end;

[nr_it n]=size(x);nr_it
minim=x(nr_it,:);

```

Execuția programului din linie de comandă necesită inițializarea variabilei x_0 , punctul de start al căutării. De exemplu, pentru:

```
x0=1;fletreev
```

rezultatul:

```

nr_it =
      2

minim =
    -0.2000

```

arată o convergență mai rapidă decât cea a metodei de gradient anterioare, implementată prin funcția utilizator `opt_grad`.

Observație:

Pentru o funcție de n variabile care *nu* este pătratică, după fiecare n iterații se reinițializează direcția de căutare la direcția antigradientului, ca la primul pas. Scopul este eliminarea erorilor datorate faptului că funcția se poate aproxima bine cu o funcție pătratică *numai* în stricta vecinătate a optimului.

1.4 Metode directe. Metode de relaxare

Din clasa metodelor directe de căutare a optimului, cea mai simplă, dar și cea mai costisitoare ca timp, este cea a **explorării exhaustive**. Metodele **de eliminare** se folosesc când funcția obiectiv are *un singur optim (funcție unimodală)*; ele se bazează pe eliminarea unei regiuni a domeniului de variație a variabilelor independente care nu conține optimul.

Căutarea aleatoare (metoda Monte Carlo) constă în evaluarea funcției obiectiv într-un set de puncte generate pseudoaleator la fiecare iterație a algoritmului. Domeniul de explorare din jurul optimului aflat la fiecare iterație se restrânge, până când el devine mai mic decât cel impus; astfel, optimul de la ultima iterație se declară drept optim global.

Spre deosebire de metodele de gradient – care efectuează modificări *simultane* ce produc deplasări în spațiul n -dimensional – metodele **de căutare unidimensională (Gauss)** se fondează pe modificarea *succesivă* a componentelor vectorului $\underline{x}$. Aceste metode se mai numesc și metode **de relaxare** sau de *optimizare ciclică* de-a lungul axelor de coordonate (engl. *cyclic coordinate search*). Căutarea multidimensională este astfel transformată într-o

succesiune de căutări unidimensionale, fără a prospecta direcția înaintării, ci doar prin relaxarea rând pe rând a tuturor direcțiilor axelor de coordonate. Există mai mulți algoritmi bazați pe metoda relaxării, care diferă după modul în care se face varierea pasului de căutare.

Metodele de relaxare sunt în general mai lent convergente decât metodele obișnuite de gradient, fiindcă pot conține mai multe iterații de căutare unidimensională, în funcție de alegerea punctului de start și a pasului. Ele pot deveni ineficiente sau cel puțin foarte lent convergente dacă funcția obiectiv prezintă o vale (sau o creastă) care nu are direcția paralelă cu axele de coordonate.

Se prezintă mai jos etapele unui algoritm de relaxare [CEAN 84], în care pasul este menținut constant pe durata unui ciclu de explorare a tuturor direcțiilor, după care pasul este micșorat în progresie geometrică de rație r , până când devine mai mic decât un prag dat, ϵ .

1. Se alege un punct inițial caracterizat de vectorul $\underline{x}^{(0)} = \begin{bmatrix} x_1^{(0)} & x_2^{(0)} & \dots & x_n^{(0)} \end{bmatrix}$,

valoarea inițială a pasului de căutare, p , și pragul ϵ .

2. Se inițializează indexul coordonatei de relaxat: $k \leftarrow 1$.
3. Se inițializează contorul de iterații din optimizarea în raport cu coordonata k : $i \leftarrow 0$.
4. La o iterație i se relaxează coordonata de index k , obținându-se vectorul:

$$\underline{x}^{(i+1)} = \begin{bmatrix} x_1^{(i)} & \dots & \underbrace{x_k^{(i)} + p}_{x_k^{(i+1)}} & \dots & x_n^{(i)} \end{bmatrix}$$

Dacă $f(\underline{x}^{(i+1)}) \leq f(\underline{x}^{(i)})$, atunci $i \leftarrow i+1$ și se reia etapa 4;

altfel dacă $i=1$ (prima iterație), atunci se face relaxarea în sens opus:

$$\underline{x}^{(i+1)} = \begin{bmatrix} x_1^{(i)} & \dots & x_k^{(i)} - 2p & \dots & x_n^{(i)} \end{bmatrix}$$

$$p \leftarrow -p$$

$$i \leftarrow i+1$$

și se reia etapa 4;

altfel

dacă $f(\underline{x}^{(i)}) \leq f(\underline{x}^{(i-1)})$, atunci valoarea de optim a coordonatei k este:

$$x_{k_{opt}} = x_k^{(i-1)}$$

$$k \leftarrow k+1$$

și se reia etapa 3 (se relaxează următoarea coordonată, $k+1$).

5. Se modifică pasul p , înmulțindu-l cu rația $r < 1$: $p \leftarrow p \cdot r$.

Dacă $p \leq \epsilon$, atunci s-a obținut extremul funcției, STOP; altfel se reia de la etapa 2.

Funcția Matlab de mai jos implementează algoritmul anterior pentru o funcție scop de o variabilă vectorială de o dimensiune oarecare (numele fișierului ce conține funcția scop este transmis ca parametru de intrare, iar valorile funcției se evaluează cu feval).

```
function [x_opt,nr_it]=opt_rlx(fct_scop,x0,pas_init,r,tol)
%optimizare parametrică prin METODA RELAXĂRII

%fct_scop - șirul de caractere ce desemnează funcția de optimizat
```

`%x0` - punctul de start (vector de dimensiunea variabilei funcției de optimizat, n)
`%pas_init` - pasul inițial de căutare
`%r` - rația subunitară de modificare a pasului după fiecare ciclu de relaxare a tuturor
`%` celor n coordonate
`%tol` - valoarea de prag a pasului, când căutarea se oprește

`%inițializări`

```

i=1;x(i,:)=x0;p=pas_init;
n=length(x0);

while (p>tol),
    p0=p*r;p=p0;
    k=1;d=zeros(1,n);
    stop_global=0;inapoi=0;
    while (~stop_global),
        d(k)=1;j=0;
        stop=0;
        while (~stop)
            if inapoi,
                p=p/2;
                inapoi=0;
            end;
            x(i+1,:)=x(i,:)+p*d;
            j=j+1;
            if feval(fct_scop,x(i+1,:))>feval(fct_scop,x(i,:)),
                if j==1,
                    p=-2*p;
                    i=i+1;
                    x(i+1,:)=x(i,:)+p*d;
                    inapoi=1;
                else stop=1;
                    x(i+1,:)=[];
                    i=i-1;
                    d(k)=0;
                    k=k+1;
                    p=p0;
                end;
            end;
            i=i+1;
        end;
        if k<=n,
            stop_global=0;
        else stop_global=1;
        end;
    end;

```

```
end;

[nr_it n]=size(x);
x_opt=x(nr_it,:);
```

Pentru exemplificare, s-a folosit o funcție de două argumente reale:

```
function y=f(x)
y=4+6*x(1)-4*x(2)+x(1)^2+2*x(2)^2-...
    6*x(1)*x(2)+x(1)^4+2*x(1)^2*x(2);
```

care descrie o suprafață (figura 2.9), și care are trei minime locale, calculabile analitic, situate aproximativ în punctele de coordonate (0;1), (0.3117;1.419) și (-4.8117;-17.794).

Cele două apeluri de mai jos arată că, din același punct de start, (-1;3), dar cu alt pas inițial ($p=1$, respectiv $p=0.5$), metoda relaxării poate furniza rezultate semnificativ diferite.

```
[x_opt,nr_it]=opt_rlx('f', [-1 3], 1, 0.5, 1e-6)
x_opt =
    0.0000    1.0000
nr_it =
    84
[x_opt,nr_it]=opt_rlx('f', [-1 3], 0.5, 0.5, 1e-6)
x_opt =
    0.3118    1.4191
nr_it =
   407
```

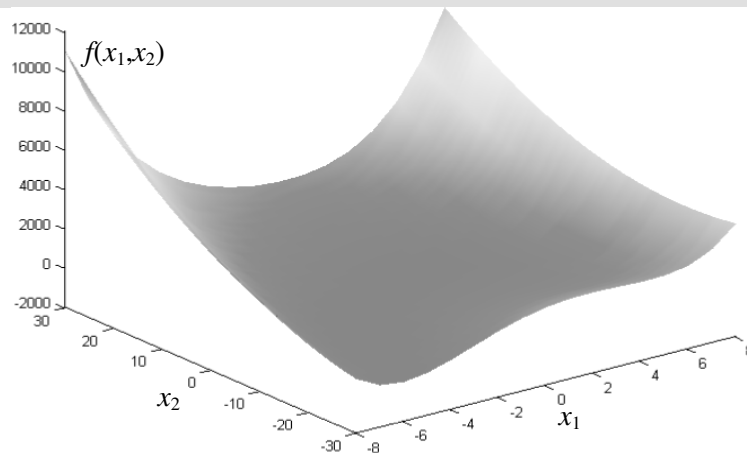


Fig. 1 O suprafață cu trei minime locale:

$$f(x_1, x_2) = 4 + 6x_1 - 4x_2 + x_1^2 + 2x_2^2 - 6x_1x_2 + x_1^4 + 2x_1^2x_2$$

Dintr-un alt punct de start se poate obține un alt punct de minim local:

```
[x_opt,nr_it]=opt_rlx('f', [-3 -7], 2, 0.5, 1e-5)
x_opt =
   -4.8158  -17.8200
nr_it =
   3577
```