

CO optimisation numérique 2023-24 : corrigé

Ex1 (11 pts)

1) (2 pts) La fonction H0phm renvoie le résultat approché (après N itérations de l'algorithme d'Uzawa) du problème :

$$\begin{aligned} \text{Min } x &\mapsto J(x) \\ \text{sous la contrainte} \\ C_E(x) &= 0 \end{aligned}$$

2) (4 pts) l'exemple étudié
ici est :

$$\text{Min } J(x) = \frac{1}{2} \langle Ax, x \rangle - \langle b, x \rangle$$

$$\text{avec } A = \begin{pmatrix} 2 & -1 & 0 & 0 \\ -1 & 2 & -1 & 0 \\ 0 & -1 & 2 & -1 \\ 0 & 0 & -1 & 2 \end{pmatrix}$$

$$b = \begin{pmatrix} 1 \\ 1 \\ 1 \\ 1 \end{pmatrix}$$

sous la contrainte pour $x \in \mathbb{R}^4$:

$$\sum_{i=1}^4 x_i = 1$$

Sur cet exemple, la fonction J est coercive et strictement convexe

car la matrice A est définie positive.
 $\langle x | Ax \rangle = \sum_{i=1}^{n+1} (x_i - x_{i-1})^2$, avec
 la convention $x_{n+1} = x_0 = 0$, et
 donc $\langle x | Ax \rangle > 0$ si $x \neq 0$. (*)

L'ensemble $\Omega = \{x \in \mathbb{R}^4, \sum x_i = 1\}$
 est fermé et convexe. le problème
 admet donc une unique solution.

Pour la trouver, on résout KKT:

$$\begin{cases} Ax - b + \lambda \begin{pmatrix} 1 \\ 1 \\ 1 \\ 1 \end{pmatrix} = 0 \\ \sum_{i=1}^4 x_i = 1 \end{cases}$$

(*) on peut aussi utiliser Sylvester

soit:

$$\begin{cases} 2x_1 - x_2 + \lambda - 1 = 0 \\ -x_1 + 2x_2 - x_3 + \lambda - 1 = 0 \\ -x_2 + 2x_3 - x_4 + \lambda - 1 = 0 \\ -x_3 + 2x_4 + \lambda - 1 = 0 \\ x_1 + x_2 + x_3 + x_4 = 1 \end{cases}$$

ceci implique (avec $\lambda' = \lambda - 1$)
 $x_1 + x_4 + 4\lambda' = 0 \quad (L_1 + L_2 + L_3 + L_4)$

mais aussi :

$$\underbrace{x_2 + x_3}_{1 - (x_1 + x_4)} = (x_1 + x_4) + 2\lambda' = 0 \quad (L_2 + L_3)$$

soit

$$8\lambda' + 1 + 2\lambda' = 0$$

$$\text{et } \lambda' = -\frac{1}{10} \text{ puis } x_1 + x_4 = \frac{4}{10}$$

$$\text{et } x_2 + x_3 = \frac{6}{10}$$

$2L_1 + L_2$ donne :

$$3x_2 - 2x_3 = \frac{3}{10}, \text{ soit}$$

$$x_2 = x_3 = \frac{3}{10} \text{ puis}$$

$$x_1 = x_2 = \frac{2}{10} \text{ (et } \lambda = \frac{9}{10})$$

Il s'agit de l'unique solution du
KKT et donc forcément du problème
de minimisation //

3) (2pts) La ligne

$$X = X - \alpha \nabla J(X) +$$

$$\lambda \nabla CE(X) \text{ consiste}$$

à effectuer un pas de gradient

(pas constant égal à $\alpha = 1$)

pour minimiser la fonction

$$x \mapsto \mathcal{L}(x, \lambda) \text{ (Lagrangien du problème)}$$

La ligne

$$\lambda = \lambda + \alpha \nabla CE(X)$$

consiste à effectuer un pas de

gradient (pas égal à $\alpha = 2$)

pour maximiser la fonction
 $\lambda \mapsto \mathcal{L}(x, \lambda)$, ceci afin de
 rechercher un point selle de $\mathcal{L}$.

4) (3 pts) On effectue une
 itération de l'algorithme :

$$X_0 = \begin{pmatrix} 0 \\ 0 \\ 0 \\ 0 \end{pmatrix}, \lambda = 3$$

↙
 ↘

$$X = 0 - 0.1 * \left(A * 0 - \begin{pmatrix} 1 \\ 1 \\ 1 \\ 1 \end{pmatrix} \right) + 3 \begin{pmatrix} 1 \\ 1 \\ 1 \\ 1 \end{pmatrix}$$

soit

$$X = - \begin{pmatrix} 0.2 \\ 0.2 \\ 0.2 \\ 0.2 \end{pmatrix}$$

et

$$\lambda = 3 + 0.1 * (4 * (-0.2) - 1)$$

$$= 3 - 0.18 = 2.82$$

Ex 2 (16 pts)

3

1) (2 pts) L'ensemble

$S = \{x \in \mathbb{R}^n / Dx = d\}$ est un
 ensemble convexe et fermé. La
 fonction J étant continue, coercive et
 strictement convexe, le problème
 admet bien une unique solution.

2) (1 pt) Attention : il

manque une hypothèse du type :

$$J(y) \geq J(x) + \langle \nabla J(x), y - x \rangle + \underbrace{\mu}_{>0} \|y - x\|^2$$

(J fortement convexe) par conclure

Dans ce cas, comme

$$\begin{cases} H_x \mathcal{L} = H_x J \text{ et} \\ \nabla_x \mathcal{L} = \nabla_x J + \lambda D \end{cases}$$

on garde les mêmes hypothèses ($\mathcal{L}(\cdot, \lambda)$ fortement convexe et coercive) et donc les mêmes conclusions (existence et unicité de x_λ^*).

3) (3 pts) la fonction $\lambda \mapsto R(x, \lambda)$ est affine et donc concave. Or si f, g sont 2 fonctions concaves, alors

4

$h = \min(f, g)$ est également concave (en effet)

$$f(\alpha u + (1-\alpha)v) \geq \alpha f(u) + (1-\alpha)f(v)$$

$$\geq \alpha h(u) + (1-\alpha)h(v)$$

De même,

$$g(\alpha u + (1-\alpha)v) \geq \alpha h(u) + (1-\alpha)h(v)$$

D'où $h(\alpha u + (1-\alpha)v) \geq \alpha h(u) + (1-\alpha)h(v)$ et h concave). Il en va de même par $\lambda \mapsto \min_x (\mathcal{L}(x, \lambda))$

4) (7 pts)

On a dans ce cas

$$L(x, \lambda) = \frac{1}{2} \langle Ax, x \rangle - \langle b, x \rangle + \langle \lambda^t D x - d \rangle$$

et on a alors

$$A x_\lambda^* = b - b^t D \lambda$$

soit

$$x_\lambda^* = A^{-1} (b - b^t D \lambda) \quad (2 \text{ pts})$$

On a alors

$$H(\lambda) = \frac{1}{2} \langle A x_\lambda^*, x_\lambda^* \rangle - \langle b, x_\lambda^* \rangle + \langle \lambda, D x_\lambda^* - d \rangle = -D A^{-1} b^t D \lambda + D A^{-1} b - d$$

puis

$$H(\lambda) = \frac{1}{2} \langle \bar{A} (b - b^t D \lambda), b - b^t D \lambda \rangle - \langle b, \bar{A} (b - b^t D \lambda) \rangle + \langle \lambda, D \bar{A}^{-1} (b - b^t D \lambda) - d \rangle$$

En prenant le gradient en λ , on trouve:

$$\nabla H(\lambda) = D A^{-1} b^t D \lambda$$

$$\begin{aligned} & - \frac{1}{2} D A^{-1} b \\ & - \frac{1}{2} b^t (A^{-1} D) b \\ & + b^t (A^{-1} D) b \\ & + D A^{-1} b - d \\ & - 2 D A^{-1} b^t D \lambda \end{aligned}$$

Gr

$$D x_\lambda^* - d = D A^{-1} b - D A^{-1} b^t D \lambda - d$$

soit: $\nabla H(\lambda) = D x_\lambda^* - d$ (3 pts)

H est de la forme

$$H(\lambda) = \frac{1}{2} \langle \lambda, B\lambda \rangle - \langle f, \lambda \rangle$$

avec $B = -DA^{\top}D$ qui est définie négative. Son unique maximum est donc atteint en λ_0 tq $-DA^{\top}D\lambda_0 = DA^{\top}b - d$

Gn a alors :

$$\begin{cases} Dx_{\lambda_0}^* = d \text{ (voir précédemment) et} \\ Ax_{\lambda_0}^* = b - D^{\top}\lambda_0 \text{ (idem)} \end{cases}$$

Cela correspond exactement au système KKT vérifié par (x^*, λ^*)
(2 pts)

5) (3 pts) Gn a de même / 6

$$H(\lambda) = J(x_{\lambda}^*) + \langle \lambda, Dx_{\lambda}^* - d \rangle$$

En supposant $\lambda \mapsto x_{\lambda}^*$ différentiable la règle de la chaîne donne :

$$\begin{aligned} \underbrace{\nabla H(\lambda)}_{\in \mathbb{R}^m} &= \underbrace{\frac{\partial x_{\lambda}^*}{\partial \lambda}}_{\in \mathcal{L}(\mathbb{R}^n, \mathbb{R}^n)} \cdot \underbrace{\nabla J(x_{\lambda}^*)}_{\in \mathbb{R}^n} + (Dx_{\lambda}^* - d) \\ &+ \left(\frac{\partial x_{\lambda}^*}{\partial \lambda} \right)^{\top} D \lambda \in \mathbb{R}^m \\ &\in \mathcal{L}(\mathbb{R}^n, \mathbb{R}^m) \end{aligned}$$

$$\begin{aligned} \text{soit } \nabla H(\lambda) &= \frac{\partial x_{\lambda}^*}{\partial \lambda} \left[\nabla J(x_{\lambda}^*) + D^{\top} \lambda \right] \\ &+ Dx_{\lambda}^* - d \\ &= 0 + Dx_{\lambda}^* - d \end{aligned}$$

car x_λ^* point critique de

$$x \mapsto J(x) + {}^t\lambda D\psi$$

$$\text{et donc } \nabla J(x_\lambda) + {}^t D\lambda = 0$$

7