



Algorithme de recherche de point-selle de lagrangien non strictement convexe. Application a l'optimisation des investissements pour un réseau electrique.

Jean-François Balducchi

► To cite this version:

Jean-François Balducchi. Algorithme de recherche de point-selle de lagrangien non strictement convexe. Application a l'optimisation des investissements pour un réseau electrique.. Automatique / Robotique. École Nationale Supérieure des Mines de Paris, 1982. Français. NNT: . pastel-00833942

HAL Id: pastel-00833942

<https://pastel.hal.science/pastel-00833942>

Submitted on 13 Jun 2013

HAL is a multi-disciplinary open access archive for the deposit and dissemination of scientific research documents, whether they are published or not. The documents may come from teaching and research institutions in France or abroad, or from public or private research centers.

L'archive ouverte pluridisciplinaire **HAL**, est destinée au dépôt et à la diffusion de documents scientifiques de niveau recherche, publiés ou non, émanant des établissements d'enseignement et de recherche français ou étrangers, des laboratoires publics ou privés.

THESE

présentée à

L'ECOLE NATIONALE SUPERIEURE DES MINES DE PARIS

par

Jean-François BALDUCCHI

en vue de l'obtention

DU TITRE DE DOCTEUR INGENIEUR

EN

MATHEMATIQUES ET AUTOMATIQUE

ALGORITHME DE RECHERCHE DE POINT-SELLE
DE LAGRANGIEN NON STRICTEMENT CONVEXE.
APPLICATION A L'OPTIMISATION DES
INVESTISSEMENTS POUR UN RESEAU ELECTRIQUE.

Soutenue le 10 Décembre 1982 devant le Jury composé de

<i>Messieurs. Guy CHAVENT</i>	<i>Président</i>
<i>Guy COHEN</i>	<i>Rapporteur</i>
<i>Jean-Claude DODU ..</i>	<i>Examineur</i>

REMERCIEMENTS

Je tiens à remercier Monsieur CHAVENT, Professeur à Paris IX-Dauphine, pour avoir bien voulu accepter de présider ce jury de thèse.

Je remercie tout particulièrement Monsieur COHEN, Maître de Recherche à l'Ecole Nationale Supérieure des Mines de Paris, qui a dirigé mes travaux, pour l'aide et les précieux conseils qu'il m'a apportés.

Je remercie également Monsieur DODU, Ingénieur à l'EDF qui a bien voulu participer à ce jury.

J'exprime également toute ma reconnaissance aux Chercheurs du Centre d'Automatique et Informatique, qui d'une manière ou d'une autre m'ont aidé dans ce travail, et en particulier Monsieur CARPENTIER dont l'aide et les conseils amicaux sur le plan informatique m'ont été d'un grand secours.

Je remercie enfin et très chaleureusement Madame LE GALLIC pour la réalisation matérielle de ce document.

Présentation du Problème

1. Énoncé du Problème

Etant donnés, d'une part un ensemble de moyens de production électrique, d'autre part la consommation d'énergie électrique, on se propose d'optimiser les capacités de transport des lignes du réseau électrique national, de façon à minimiser la somme des coûts de production électrique et des coûts d'investissement des lignes de transport électrique.

Cette étude, proposée par EDF dans le cadre d'un contrat passé avec l'INRIA et le CAI de l'Ecole des Mines de Paris, constitue le support et la motivation essentiels de cette thèse.

Ainsi posé, ce problème présente plusieurs particularités importantes :

- En premier lieu, il est de caractère dynamique : chaque année il convient d'estimer la valeur des nouveaux investissements à ajouter à ceux déjà existants.
- Une deuxième caractéristique du problème est qu'il se pose en variables mixtes (entières et réelles). En effet, les capacités des lignes de transport électrique sont dans la pratique choisie parmi une gamme discrète de valeurs.
- Enfin, et surtout, le système que nous étudions est soumis à un certain nombre d'aléas : indisponibilité des moyens de production, des lignes de transport, et caractère aléatoire de la demande en électricité.

Ceci nous amène à formuler un problème d'optimisation stochastique, avec deux types de paramètres :

- Les capacités des lignes de transport électrique, à valeurs entières, et qui sont en boucle ouverte (c.a.d. indépendantes des aléas).
- Les niveaux de production, qui sont en boucle fermée sur les aléas (c.a.d. fonctions des aléas réalisés).

D'un point de vue physique, on se soumet à l'approximation dite du courant continu, ce qui conduit à prendre en compte les deux lois de Kirchhoff régissant le transport de l'énergie électrique sous forme linéaire en termes d'intensités et de tensions.

2. Quelques hypothèses simplificatrices.

Dans notre étude nous ne nous sommes intéressés qu'au seul problème statique, c'est-à-dire correspondant à une année.

De plus, la complexité du problème réel a conduit EDF à attaquer celui-ci par étapes :

La première étape, dans laquelle nous nous situons, consiste à "oublier" le caractère entier des capacités des lignes de transport, et à ne conserver que la première Loi de Kirchhoff (la somme des intensités incidentes en un noeud est égale à zéro) pour un réseau agrégé où un noeud représente une région. En effet la prise en compte de la seconde loi de Kirchhoff (intensités dans les lignes proportionnelles aux différences de tension aux extrémités) aurait peu de sens pour le réseau agrégé. De plus le problème d'optimisation deviendrait alors non convexe. Enfin, on ne tient compte ni des indisponibilités des lignes de transport, ni du caractère aléatoire de la demande en électricité - on considère uniquement le premier palier de la "monotone de charge" correspondant à la puissance maximale appelée dans l'année - et seules les indisponibilités éventuelles des groupes de production sont donc prises en considération.

Ajoutons que le réseau national 400 KV est schématisé pour les besoins de cette étude par un graphe comprenant 48 sommets et 82 arcs (chaque sommet est un noeud de consommation représentant une région). Le parc de production est constitué par 185 centrales.

3. Le Modèle mathématique

3.1. Le graphe associé au réseau de transport.

- $S = \{\sigma_1, \dots, \sigma_n\}$ représente l'ensemble des n sommets du réseau R de transport.
- σ_0 est un sommet fictif qui joue le rôle de source pour la production, et de puits pour la demande.

$$\text{On pose } \bar{S} = S \cup \{\sigma_0\}$$

- $\Gamma : \bar{S} \times \bar{S} \rightarrow N$ est une fonction indiquant le nombre d'arcs orientés
- $\Gamma(\sigma_i, \sigma_j)$ - reliant σ_i et σ_j (ici Γ prend les valeurs 0 ou 1).

On définit alors l'ensemble des arcs orientés de la manière suivante :

$$A = \{(\sigma_i, \sigma_j) \mid \sigma_i \in \bar{S}; \sigma_j \in \bar{S}; \Gamma(\sigma_i, \sigma_j) = 1\} \quad k \leq \Gamma(\sigma_i, \sigma_j)$$

Dans cet ensemble $\mathcal{A}$, on distingue trois types d'arcs :

- les arcs de transport $(\sigma_i, \sigma_j) \in \mathcal{A}$; $\sigma_i \in S$; $\sigma_j \in S$
- les arcs de production $(\sigma_0, \sigma_j) \in \mathcal{A}$; $\sigma_j \in S$
- les arcs de consommation $(\sigma_i, \sigma_0) \in \mathcal{A}$; $\sigma_i \in S$

Les noeuds et les arcs du graphe Q étant ainsi définis, on désigne par A la matrice d'incidence de Q (noeuds/arcs) :

Chaque colonne A^r de A est un vecteur de dimension $n + 1$ qui correspond à un élément de $\mathcal{A}$ (arc orienté (σ_i, σ_j)).

Si R désigne l'ensemble des indexes des éléments de $\mathcal{A}$, on écrit alors :

$$\forall s \in \bar{S}, \forall r \in R \quad A_s^r = +1 \quad \text{si } s = j \quad (\text{noeud d'arrivée})$$

$$A_s^r = -1 \quad \text{si } s = i \quad (\text{noeud de départ})$$

$$A_s^r = 0 \quad \text{sinon.}$$

3.2. Notations.

On désigne par :

- T l'ensemble des indices des arcs de transport associés aux lignes de transport.
 - P l'ensemble des indices des arcs de production associés aux moyens de production.
 - D l'ensemble des indices des arcs de consommation associés à la demande.
- On indice par T, P, D les vecteurs relatifs au transport, à la production, à la demande.

Les variables réelles liées au transport (respectivement à la production, à la demande) sont indicées par $t \in T$ (respectivement $p \in P, d \in D$).

On définit ainsi :

- Pour un arc de transport $t \in T$
 - un coût unitaire d'investissement c_t .
 - une capacité maximale de transport q_t^V .
 - une puissance transistée q_t .

c_T, q_T^V et q_T sont les vecteurs correspondants.

- Pour un arc de production $p \in P$
 - un coût unitaire de production c_p .

- une capacité maximale de production $\overset{V}{q}_p$.
- une puissance produite q_p .

On désigne par q_p , q_p et $\overset{V}{q}_p$ les vecteurs correspondants.

- Pour un arc de consommation $d \in D$
 - . un coût unitaire de défaillance c_d (défaillance = demande non satisfaite).
 - . une demande à satisfaire $\overset{V}{q}_d$.
 - . une énergie réellement consommée q_d .

q_D , $\overset{V}{q}_D$ et q_D représentent les vecteurs correspondants.

Remarques : $\overset{V}{q}_p$ est un vecteur de caractère aléatoire.

q_p , q_d , q_t sont les variables de décision du problème de dispatching optimal. Elles sont déterminées en fonction du vecteur aléatoire $\overset{V}{q}_p$ mais dépendent également de $\overset{V}{q}_T$ qui est un vecteur de dimensionnement à valeur déterministe.

Les vecteurs sans indice désignent la concaténation des vecteurs avec indice P, D, T.

$$(\text{Exemple : } q = \begin{bmatrix} q_p \\ q_D \\ q_T \end{bmatrix}).$$

3.3. Formulation du problème.

Les contraintes s'expriment de la manière suivante :

$$Aq = 0 \quad (1.1)$$

traduit la première loi Kirchhoff.

Par définition, on a aussi :

$$0 < q_p < \overset{V}{q}_p \quad (1.2)$$

$$0 < q_d < \overset{V}{q}_d \quad (1.3)$$

$$|q_t| < \overset{V}{q}_t \quad (1.4)$$

On considère alors :

$$\Phi(q_T^V) = \min_{q_t, q_d, q_p} \{c_F q_P + c_D [q_D^V - q_D]\} \quad (1.5)$$

sous les contraintes (1.1) à (1.4).

où $c_P q_P$ par exemple désigne le produit scalaire $\sum_{p \in P} c_P q_P$.

Φ est une fonction aléatoire (par q_P^V) de la variable vectorielle q_T^V , convexe mais non différentiable en général.

Posons maintenant :

$$\Delta = \{q_T^V | q_t^0 < q_t^V < \bar{q}_t^V \quad \forall t \in T\} \quad (1.6)$$

où q_t^0 représente l'investissement existant et $\bar{q}_t^V$ un investissement limite fixé par E.D.F.

Le problème se formule alors ainsi :

$$\text{Coût minimum : } \min_{q_T^V \in \Delta} \{c_T q_T^V + E[\Phi(q_T^V)]\} - c_T q_T^0 \quad (1.7)$$

où E désigne l'espérance mathématique.

Dans la suite on posera :

$$\varphi(q_T^V) = c_T q_T^V + E[\Phi(q_T^V)] \quad (1.8)$$

CHAPITRE II

Approches primale et duale

1. Méthode primale.

Tel qu'il est formulé en (1.7), le problème peut se résoudre au moyen de méthodes de type primal consistant à essayer de minimiser la fonctionnelle $\varphi[\vec{d}_T]$ pour $\vec{d}_T \in \Delta$.

Pour cela on itère sur les valeurs des capacités des lignes de transport en appliquant une méthode de sous-gradient *. Le sous-gradient de φ peut-être obtenu grâce aux multiplicateurs de Lagrange relatifs aux contraintes (1.4) :

Soit α_T le vecteur des multiplicateurs de Lagrange associés aux contraintes (1.4).

On sait que $-\alpha_T \in \partial\Phi(\vec{d}_T)$ où $\partial\Phi$ représente le sous-différentiel de Φ . Les multiplicateurs optimaux n'étant pas toujours uniques, la fonction Φ n'est pas différentiable en général.

- 6 -

* f étant une fonction convexe dans R^n , on appelle sous-différentiel de f au point x , noté $\partial f(x)$, l'ensemble des vecteurs α de R^n tels que :

$$f(x+y) \geq f(x) + (\alpha, y) \quad \forall y \in R^n.$$

Les éléments de $\partial f(x)$ sont appelés sous-gradients. Pour une fonction convexe en dimension finie, un sous-gradient existe $\forall x \in \text{dom } f$ (intérieur du domaine de f défini par : $\text{dom } f : \{x | f(x) < +\infty\}$) mais peut ne pas être unique. Lorsque f est différentiable, il est unique et coïncide avec le gradient $\nabla f(x)$. Prenons l'exemple de $f(x) = \max_i f_i(x)$ où les f_i sont différentiables. Soit $I(x) = \{i | f_i(x) = f(x)\}$. Toute combinaison convexe $\sum_{i \in I(x)} \alpha_i \nabla f_i(x)$ est un élément du sous-différentiel $\partial f(x)$.

Un algorithme de sous-gradient consisterait alors à faire évoluer q_T^V de la manière suivante :

$$q_T^{k+1} = P(q_T^k - \rho^k \beta_T^k) \quad (2.0)$$

où les ρ^k sont des réels non négatifs tels que $\sum (\rho^k)^2 < +\infty$ et $\sum \rho^k = +\infty$. P est la projection sur Δ et $\beta_T^k = c_T - E(q_T^k)$.

Du fait de la présence de l'espérance mathématique dans (1.7), on peut soit en faire une approximation par une moyenne sur un grand nombre de réalisations comme nous le verrons plus loin (le calcul exact de l'espérance nécessiterait la moyenne sur 2^{185} réalisations), soit calculer le sous-gradient à chaque étape sur une seule réalisation tirée au hasard.

Cette dernière option est la méthode du sous-gradient stochastique mise au point à l'INRIA par J.P. QUADRAT et M. VIOT (J.C. Dodu et al, 1981).

II.2. Une approche duale.

On peut aussi reformuler différemment (1.7) de manière à obtenir un problème de Min-Max.

Reprenons le problème (1.7) et plaçons-nous à l'optimum :

soit q_t^* la valeur optimale de q_t et $q_t(\omega)$ la valeur optimale de q_t dans le problème (1.5), correspondant à q_t^* dans (1.4) et à l'aléa $\omega^{(1)}$ (c'est à dire à $q_p(\omega)$ dans (1.2)). Les coordonnées de c_T étant positives, il est clair que :

$$q_T^* = \text{Max}[\text{Max}_{\omega} |q_T(\omega)|, q_T^0] \quad (2.1)$$

En effet, s'il n'en était pas ainsi, la capacité maximale ne serait utilisée complètement dans aucune des situations aléatoires.

Introduisons un aléa supplémentaire, fictif, ω_0 , de probabilité nulle, auquel on associe un coût de production nul et une solution optimale $q_T(\omega_0) = q_T^0$ (l'investissement existant).

(1) L'indice ω repère toutes les réalisations possibles du vecteur aléatoire q_p . Ces réalisations sont en nombre fixé : $2m$ où m est le nombre de groupes de production.

(2) Les opérations de Max et de valeur absolue sont à interpréter coordonnées

La notation Max_{ω} porte désormais sur ω_0 et sur toutes les réalisations possibles. Le coût minimum (1.7) peut alors se réécrire :

$$c_D q_D - c_T q_T + \text{Min}_{q(\omega)} \{ E_{\omega} [c_P q_P(\omega) - c_D q_D(\omega)] + c_T [M_{\omega}^{\text{Max}} |q_T(\omega)|] \}$$

$$\text{sur les contraintes } Aq(\omega) = 0 \quad (\text{i})$$

$$|q_t(\omega)| \leq \check{q}_t \quad (\text{ii})$$

$$0 < q_d(\omega) \leq \check{q}_d \quad (\text{iii}) \quad (2.3)$$

$$0 < q_p(\omega) \leq \check{q}_p(\omega) \quad (\text{iv})$$

L'introduction de la contrainte (2.3)-(ii) permet de tenir compte de la contrainte correspondante sur $\check{q}_t$ [cf. (1.6)] puisque cette variable a maintenant disparue.

On considère ensuite une approximation de l'espérance mathématique E_{ω} par la moyenne arithmétique sur un nombre fini N de réalisations tirées suivant un processus de génération de nombres pseudo-aléatoires conforme à la loi des pannes des moyens de production.

De plus, on remplace l'opération de maximisation discrète sur les aléas par une maximisation continue en faisant la remarque suivante : le maximum de N nombre réels est égal au maximum de leurs combinaisons convexes. On obtient donc, avec $\omega = 0$ désignant l'aléa fictif introduit plus haut :

$$\text{Max}_{\omega=0, N} |q_t(\omega)| = \text{Max}_{p_t(\omega)} \sum_{\omega=0}^N [|q_t(\omega)| p_t(\omega)] \quad (2.4)$$

Les $N+1$ vecteurs $p(\omega)$ sont de dimension cardT (nombre de lignes de transport) et vérifient pour tout $\omega = 0, \dots, N$:

$$\forall t \in T \quad p_t(\omega) > 0$$

$$\sum_{\omega=0}^N p_t(\omega) = 1. \quad (2.5)$$

Le coût minimal s'écrit donc finalement :

$$c_D^V - c_T^V q_T^0 + \min_{q(\omega)} \max_{p(\omega)} \left\{ \sum_{\omega=1}^N [c_P q_P(\omega) - c_D q_D(\omega)] + \sum_{\omega=0}^N [c_T(\omega) |q_T(\omega)| x_P(\omega)] \right\}^{(3)} \quad (2.6)$$

sous les contraintes (2.3) et (2.5).

On peut parvenir au même résultat en partant de (1.7), en remplaçant l'espérance mathématique par la moyenne arithmétique et en dualisant les contraintes $|q_t(\omega)| \leq q_t^V$ (multiplicateurs $\lambda(\omega) \geq 0$) et $q_0^V \leq q_t^V$ (multiplicateur $\mu \geq 0$). Par ailleurs, l'autre contrainte sur q_t^V dans (1.6) peut-être supprimée à condition d'introduire les contraintes (2.3)-(ii). On obtient alors :

$$\min_{q_T, q} \max_{\substack{\lambda \geq 0 \\ \mu \geq 0}} \left\{ c_P q_P^V + \frac{1}{N} \sum_{\omega=1}^N [c_P q_P(\omega) - c_D q_D(\omega)] - c_T^V q_T^0 + \mu (-q_T^V + q_0^V) + \sum_{\omega=1}^N \lambda(\omega) [-q_T^V + |q_T(\omega)|] \right\} \quad (2.7)$$

sous les contraintes (2.3).

Compte-tenu de la linéarité en q_T^V , et du fait qu'aucune contrainte ne pèse plus sur q_T^V , il est nécessaire que le coefficient de q_T^V soit nul afin que le minimum en q_T^V soit fini. Par conséquent $\sum_{\omega=1}^N \lambda(\omega) + \mu = c_T^V$.

En posant $p_t(\omega) = \frac{1}{c} \lambda(\omega) \geq 0$ pour $\omega = 1, \dots, N$ et $p_t(0) = \frac{1}{c} \mu \geq 0$ on retrouve alors (2.5). Ces conditions étant vérifiées à l'optimum on peut les inclure dans les contraintes sans modifier le problème et il est clair que l'on retrouve ainsi (2.6).

Le critère est additif en ω , et à l'exception de la contrainte de normalisation dans (2.5), toutes les contraintes sont découplées, ω par ω . Il est donc intéressant de pouvoir intervertir le Min et le Max afin de décomposer la minimisation globale en N minimisations indépendantes. La convexité rend possible cette intervention grâce à des théorèmes

- (3) Opération $\times$ est un produit de vecteurs dont le résultat est un vecteur de même dimension et dont la $i^{\text{ème}}$ coordonnée est le produit des $i^{\text{ème}}$ coordonnées.

classiques de point-selle.

On aboutit ainsi à une reformulation duale du problème initialement posé : il s'agit de déterminer un point-selle, les variables duales étant astreintes à appartenir au produit cartésien de sous-ensembles compacts de cônes positifs définis par (2.5).

CHAPITRE III

Résolution des Problèmes de Min-Max par l'approche duale.

Nous abordons dans ce chapitre, avec des notations générales, l'étude des problèmes de Min-Max et leur résolution par méthode duale.

1 - Notations

Nous considérons le problème suivant :

$$\inf_{u \in U} \sup_{i \in I} \theta_i(u_i) \quad (3.1)$$

Sous les hypothèses :

- (i) $I = \{1, \dots, n\}$, $n \in \mathbb{N}$ (Le sup est donc un max dans (3.1))
 U est un sous ensemble convexe fermé d'un espace de Hilbert U .
- (ii) Les fonctions θ_i sont convexes propres, semi-continues inférieurement sur U et telles que $U \subset \bigcap_{i \in I} \text{dom } \theta_i$ (ceci implique donc que les fonctions θ_i sont continues sur U et semi-continues inférieurement pour la topologie faible).
- (iii) Si U n'est pas borné, $\exists i_0 \in I$ tel que
 $\forall \{u^k | k \in \mathbb{N}, u^k \in U\}$ avec $\lim_{k \rightarrow +\infty} \|u^k\| = +\infty$,
on a alors $\lim_{k \rightarrow +\infty} \theta_{i_0}(u^k) = +\infty$. (3.2)

Sous ces hypothèses, la fonction : $u \mapsto \text{Max}_{i \in I} \theta_i(u)$ satisfait également les conditions (ii) et (iii).

Par suite l'infimum (3.1) est fini et il est atteint en un point $u^* \in U$. De plus si les fonctions θ_i sont strictement convexes, ce minimum est unique.

2 - Transformation du problème

La simple remarque triviale que le maximum de n nombres réels est égal au maximum de leurs combinaisons convexes nous permet de dire que le problème (3.1) est équivalent au problème suivant :

- (4) Cette formulation de l'hypothèse (iii) est suffisante pour l'existence d'une solution de (3.1). Du point de vue algorithmique (cf. plus loin), il faudra renforcer (iii) en supposant la propriété vérifiée pour tout i .

$$\min_{u \in U} \max_{p \in \Sigma} \langle p, \theta(u) \rangle \quad (3.3)$$

avec

$$\Sigma = \{p \in \mathbb{R}^n ; p_i \geq 0 ; i = 1, \dots, n ; \sum_{i=1}^n p_i = 1\} \quad (3.3-1)$$

Remarque 1: La formulation (3.3) peut comme au chapitre précédent être obtenue à partir du problème suivant :

$$\min_{u \in U, y \in \mathbb{R}} y \quad (3.4)$$

$$\text{sous les contraintes } \theta_i(u) \leq y, \quad i \in I \quad (3.5)$$

En introduisant des multiplicateurs $p_i \geq 0$ pour les contraintes (3.5) on obtient le Lagrangien

$$L(u, y, p) = y + \langle p, \theta(u) - y \rangle$$

et la condition d'optimalité $L'_y = 0$ donne

$$\langle p, 1 \rangle = 1$$

où 1 représente le vecteur de $\mathbb{R}^n$ de coordonnées égales à 1. Alors le Lagrangien minimisé en y vaut :

$$L(u, p) = \langle p, \theta(u) \rangle.$$

La fonction $(u, p) \rightarrow \langle p, \theta(u) \rangle$ est convexe et continue en u sur U , continue et linéaire en p , donc concave en p .

Par ailleurs, grâce à l'hypothèse (3.2)-(iii) on peut affirmer que si U n'est pas borné, $\exists p^0 = (0, \dots, 0, 1, 0, \dots, 0)$ où 1 est en position i_0 tel que $\forall \{u^k\} \subset U$ avec $\lim_{k \rightarrow +\infty} \|u^k\| = +\infty$, on a alors

$$\lim_{k \rightarrow +\infty} \langle p^0, \theta(u^k) \rangle = +\infty.$$

Σ est borné et on peut donc utiliser les théorèmes classiques d'existence de point-selle concernant les fonctionnelles convexes-concaves qui nous permettent d'affirmer que le problème (3.3) admet un point-selle $(u^*, p^*) \in U \times \Sigma$, c'est à dire réalisant les inégalités suivantes :

$$\forall u \in U, \forall p \in E, \langle p, \theta(u^*) \rangle \leq \langle p^*, \theta(u^*) \rangle \leq \langle p^*, \theta(u) \rangle \quad (3.6)$$

et on a alors la propriété :

$$\min_{u \in U} \max_{p \in E} \langle p, \theta(u) \rangle = \max_{p \in E} \min_{u \in U} \langle p, \theta(u) \rangle \quad (3.7)$$

3 - Algorithme.

La remarque 1 du paragraphe précédent concernant l'équivalence des problèmes (3.3.) et (3.4)-(3.5) nous amène à penser à un algorithme d'Uzawa pour la recherche du point-selle du Lagrangien du problème (3.4)-(3.5).

Malheureusement la linéarité en y fait qu'un tel algorithme ne permet pas de trouver une solution primale optimale en général (cf. chapitre IV). On peut alors essayer de résoudre directement le problème dual de (3.3) (problème de Max-Min) par l'algorithme d'Uzawa, ce qui conduit à une projection sur E à chaque étape de remise à jour de p . L'algorithme ainsi obtenu a été proposé par Cea (1971) et Medanic-Andjelic (1972).

L'inconvénient de cet algorithme réside dans le calcul de la projection sur E qui nécessite un processus itératif assez lourd du point de vue du temps calcul sur ordinateur. Afin d'éviter le calcul de cette projection sur E , tout en conservant les multiplicateurs duaux dans E , l'on procède en deux étapes :

a) première étape :

On considère un algorithme comportant une projection sur C uniquement ($C = \{p \in R^n, p_1 > 0\}$) et résultant de deux lemmes qui suivent, énoncés et démontrés par G. Cohen (1981).

Lemme 1. Soit $\bar{e} \in R^n$ le vecteur dont toutes les coordonnées sont égales à 1.

Soit Φ la fonction de $U \times R^n \rightarrow \bar{R}$ définie ainsi :

$$\Phi \begin{cases} (u, q) \rightarrow \frac{\langle q, \theta(u) \rangle}{\langle q, \bar{1} \rangle} & \text{quand } q \neq 0 \\ (u, 0) \rightarrow -\infty & \forall u. \end{cases}$$

Une condition nécessaire et suffisante pour que $u^* \in U$ soit une solution de (3.1) est que $\exists q^* \in C$ tel que (u^*, q^*) réalise un point-selle de Φ sur $U \times C$.

Remarque 2 : L'introduction de Φ repose sur l'idée suivante :

Puisqu'on s'intéresse à la recherche de point-selle de la fonctionnelle $(u,p) \rightarrow \langle p, \theta(u) \rangle$ sur $U \times E$, on peut toujours modifier celle-ci à l'extérieur de E . L'avantage de Φ est qu'on peut étendre la recherche du point-selle de $U \times E$ à $U \times C$ car Φ est positivement homogène de degré zéro ($\Phi(u, \lambda p) = \Phi(u, p)$, $\forall \lambda \in \mathbb{R}^+$) et quasi-concave en p .

Lemme 2 : Une condition nécessaire et suffisante pour que (u^*, q^*) soit un point-selle de Φ sur $U \times C$ est que :

$$\forall u \in U, \Phi(u^*, q^*) \leq \Phi(u, q^*) \quad (3.8)$$

$$\forall p > 0, q^* = P[q^* + pA(u^*)q^*] \quad (3.9)$$

où $A(u)$ est l'opérateur défini par :

$$q \rightarrow \langle q, \theta(u) \rangle - \langle q, \theta(u) \rangle q$$

et P est la projection sur le cône positif C .

En fait (3.9) est équivalente à l'inégalité de gauche du point-selle de Φ sur $U \times C$.

L'algorithme obtenu est un algorithme de point fixe pour la résolution de l'équation (3.9), en vérifiant (3.8) à chaque étape.

Algorithme 1 :

(i) soit $q^0 \neq 0$ élément de C .

(ii) $\min_{u \in U} \langle q^k, \theta(u) \rangle$

et soit u^k une solution

(iii) $q^{k+1} = P[q^k + pA(u^k)q^k]$

(iv) Retour à (ii) avec $k \leftarrow k+1$ ou arrêt, suivant le résultat du test de convergence.

b) deuxième étape :

Dans un second temps, on fait la remarque suivante :

Remarque 3 : Si à l'étape k , on change q^k en αq^k avec $\alpha > 0$, alors u^k n'est pas modifié et les q^j suivants pour $j > k$ sont changés en αq^j . Ceci signifie que l'algorithme génère une suite de directions q^k et non pas seulement une suite de points.

En conséquence, il est possible de modifier légèrement l'algorithme 1.

Algorithme 2 :

- (i) soit $p^0 \in \Sigma$.
- (ii) $\text{Min}_{u \in U} \sum_{i \in I} p_i^k \theta_i(u)$. Soit u^k une solution.
- (iii) $q_i^{k+1} = \text{Max}\{0, p_i^k + \rho[\theta_i(u^k) - \sum_{j \in I} p_j^k \theta_j(u^k)]\}$ $\forall i \in I$ (3.10)
- et $p_i^{k+1} = [\sum_{j \in I} q_j^{k+1}]^{-1} q_i^{k+1}$ $\forall i \in I$ (3.11)

Cet algorithme dont les opérations de remise à jour de p^k , sont ici détaillées est donc une variante de l'algorithme 1, où q^k est remplacé à chaque étape par le p^k du simplexe, correspondant à cette même direction q^k .

Définition : Une fonctionnelle K définie sur U est "localement" fortement convexe sur U si pour tout sous-ensemble borné B de U il existe un nombre positif α tel que :

$$\forall u, v \in B \text{ et } \forall \alpha \in [0, 1]$$

$$K[\alpha u + (1-\alpha)v] \leq \alpha K(u) + (1-\alpha)K(v) - \alpha(1-\alpha)a\|u-v\|^2. \quad (3.12)$$

Si K est sous-différentiable, ceci équivaut à dire que son sous-différentiel est localement fortement monotone, c'est à dire que :

$$\exists a > 0 : \forall (u, v) \in B, \forall c \in \partial K(u), \forall d \in \partial K(v)$$

$$\langle c-d, u-v \rangle \geq 2a\|u-v\|^2. \quad (3.13)$$

Théorème : Si en plus des hypothèses (3.2) (où (3.2-iii) est renforcée comme indiqué dans la note de bas de page), θ_i est "localement" fortement convexe sur U et Lipschitzienne sur tout borné de U ⁽⁵⁾, l'algorithme 2 génère une suite $\{u^k\}$ unique, pour un p^0 donné, et il existe des valeurs $\rho > 0$

(5) C'est assuré lorsque U est de dimension finie.

tels que cette suite $\{u^k\}$ converge vers l'unique solution u^* du problème (3.1) ; la suite $\{\Phi(u^k, p^k)\}$ est monotone croissante et tout point d'adhérence de la suite $\{p^k\}$ réalise avec u^* un point-selle.

La preuve de ce théorème est donnée par G. Cohen (1981).

Remarque 4 : Lorsque les fonctions θ_i ne sont pas strictement convexes, l'unicité de u^k , solution de l'étape (ii) de minimisation dans l'algorithme, n'est plus assurée. Il en résulte que la fonction ϕ définie par

$$\phi : p \mapsto \min_{u \in U} \langle p, \theta(u) \rangle \quad (3.14)$$

n'est pas différentiable mais sous-différentiable en général.

L'algorithme 2, résolvant $\max_p \phi(p)$, est alors un algorithme de sous-gradient projeté. Il ne converge pas avec un pas ρ fixe mais avec un pas ρ^k décroissant et tendant vers 0, et tel que :

$$\sum_0^\infty (\rho^k)^2 < +\infty \quad \text{et} \quad \sum_0^\infty \rho^k = +\infty.$$

Cependant on n'est pas du tout assuré de la convergence de la suite u^k comme cela est discuté au chapitre suivant.

CHAPITRE IV :

Recherche de point-selle dans le cas d'un Lagrangien non strictement convexe.

1 - Introduction

Dans le problème du chapitre précédent, la non unicité de la solution du problème de minimisation intervenant dans (3.14) soulève en plus de la non différentiabilité de ϕ (définie par (3.14)), une difficulté supplémentaire : même si l'on parvient à trouver un système de multiplicateurs duaux optimaux (maximisant la fonction duale ϕ), on n'a pas pour autant déterminé la solution primale.

En effet, reprenons les notations du chapitre précédent et plaçons nous dans le cas où la suite $\{p^k\}$ converge vers un p^* optimal ; il y a un risque alors (et en fait une forte probabilité) pour que la solution $\hat{d}(p^*) \in \text{Arg Min } \langle p^*, \theta(u) \rangle$ obtenue au premier niveau de l'algorithme ne soit pas une solution u^* du problème (3.3), étant donné que l'ensemble $\text{Arg Min } \langle p^*, \theta(u) \rangle$ contient d'autres points que les solutions de (3.3). $u \in U$

Cela signifie donc que l'algorithme à la Uzawa détaillé dans le chapitre précédent n'a pratiquement aucune chance de fournir une solution du problème initial.

Dans ce chapitre, nous allons considérer de la manière la plus générale possible le problème suivant (le même que celui du chapitre précédent). Trouver un point-selle de

$$L(u, p) = J(u) + \langle p, \theta(u) \rangle \quad (4.1)$$

sur l'ensemble $U \times P$.

Les hypothèses (H) sont les suivantes :

(H-1) J est une fonction convexe, semi-continue inférieurement, définie sur l'espace de Hilbert U et U est un sous-ensemble convexe fermé de U .

Quand U n'est pas borné, nous supposons de plus que

$$\begin{aligned} J(u) &\rightarrow +\infty \\ \text{si } \|u\| &\rightarrow +\infty, \quad u \text{ restant dans } U. \end{aligned}$$

On dira alors que J est coercive sur U .

(H-ii) θ est une fonction de U dans C (espace de Hilbert) et C est un cône convexe fermé de C . θ est C -convexe, ce qui signifie que :

$$\forall u, v \in U, \forall \alpha \in [0, 1], \alpha \theta(u) + (1-\alpha)\theta(v) - \theta[\alpha u + (1-\alpha)v] \in C.$$

On suppose de plus que la fonctionnelle

$$u \mapsto \langle p, \theta(u) \rangle \text{ est semi-continue inférieurement, } \forall p \in C^*$$

C^* étant le cône conjugué de C .

(H-iii) P est un sous-ensemble convexe fermé de C^* .

Quand P n'est pas borné, nous ferons l'hypothèse supplémentaire suivante, appelée condition de Slater, ou de qualification des contraintes :

$$\exists u^0 \in U \text{ tel que } \theta(u^0) \in -\overset{\circ}{C}$$

où $\overset{\circ}{C}$ représente l'intérieur, supposé non vide, de C .

Lorsque $P = C^*$, L est le lagrangien du problème

$$\min_{u \in U} J(u) \text{ sous } \theta(u) \in -C \quad (4.2)$$

et la recherche du point-selle sur $U \times C^*$ fournit la solution de (4.2). Si $P = \Sigma$ (cf (3.3)-1)) et si $J = 0$, on retrouve le problème (3.3) du chapitre précédent. On notera alors que l'hypothèse de coercivité (si U n'est pas borné) est reportée sur θ (cf (3.2-iii)), ce qui n'est pas gênant car $0 \notin \Sigma$. On peut aussi rester dans le contexte du problème (4.2) en utilisant la formulation (3.4)-(3.5) (même observation concernant la coercivité de $J(u, y) = y$).

Dans le cadre de cette formulation, nous pouvons illustrer la remarque faite plus haut concernant l'unicité, par l'exemple suivant :

$$\text{soit } U = C = R, \quad C = R^+, \quad U = [-10, +10]$$

$$J(u) = -u \quad \theta(u) = |u| - 1.$$

Dans ce cas, l'unique solution du problème (4.2) est $u^* = 1$ et le multiplicateur optimal unique est $p^* = 1$.
Cependant, $\text{ArgMin} L(u, 1)$ n'est pas réduit à $\{1\}$ mais est égal au segment $[0, 10]$. Cela signifie que toute valeur de $[0, 10]$ est une solution du problème $\text{Min}_{u \in U} L(u, p^*)$.

Mais tout algorithme fournissant une suite $\{p^k\}$ convergente, approchera p^* soit par la droite, soit par la gauche; étant donné que

$\text{ArgMin} L(u, p) = \{10\}$ pour $p < 1$ et que $\text{ArgMin} L(u, p) = \{0\}$ pour $p > 1$,

il en résulte que l'algorithme considéré fournit une suite $\{u^k\}$ ayant soit 0, soit 10, soit les deux valeurs, pour points d'adhérence.

Par suite, l'unique bonne solution $u^* = 1$ ne peut en aucun cas être obtenue au moyen d'un tel algorithme, et ce bien que la suite $\{p^k\}$ générée converge bien vers $p^* = 1$.

2 - Stabilité.

L'ensemble des points-selle du Lagrangien $L(u, p)$ est toujours de la forme $U^* \times P^*$, où U^* et P^* sont convexes, fermés, bornés.

Il en résulte que si l'on désigne par $\hat{U}(p)$ l'ensemble des $u \in U$ minimisant L pour un p donné, la difficulté que nous venons de mettre en évidence ne provient pas du fait que U^* n'est pas un singleton, mais plus généralement du fait que $\hat{U}(p^*) \neq U^*$ (on a toujours $\hat{U}(p^*) \supset U^*$).

En effet, dans le cas où $\hat{U}(p^*) = U^*$, l'algorithme considéré dans le chapitre précédent fournit comme points d'adhérence de la suite $\{u^k\}$ des éléments de $\hat{U}(p^*)$, et donc des solutions du problème primal. Dans ce cas, favorable, où $\hat{U}(p^*) = U^*$ $\forall p^* \in P^*$, l'ensemble des points-selle est dit stable.

Il s'agit là d'une stabilité en u , pertinente pour les algorithmes du type Uzawa. Pour les algorithmes de type Arrow-Hurwicz où u et p jouent un rôle symétrique il faudrait exiger également la stabilité en p :

$$\hat{P}(u^*) = P^*, \forall u^* \in U^* \text{ avec des notations similaires.}$$

Cependant celle-ci est illusoire pour des fonctionnelles linéaires en p telles que L , sauf si $\forall u^* \in U^*, \theta(u^*) \in -\hat{C}$, c'est à dire $P^* = \{0\}$.

Cette propriété de stabilité en u , moins restrictive que l'unicité de u^* , n'est cependant vérifiée à coup sûr que dans le cas où $\hat{U}(p)$ est un

singleton (quand L est strictement convexe en u). Elle est donc rarement vérifiée en pratique (le cas des lagrangiens augmentés sera discuté plus loin).

3 - Approximations successives de ϕ

Pour maximiser la fonctionnelle duale

$$\phi : p \mapsto \min_u L(u, p) \quad (4.3)$$

On peut utiliser l'approche développée dans Lasdon (1970). Cette approche consiste à résoudre un programme linéaire comportant un nombre croissant de contraintes à chaque itération. La méthode s'appuie sur une approximation externe de ϕ par une enveloppe inférieure concave, constituée d'un nombre fini d'hyperplans support de ϕ . L'approximation est redéfinie à chaque itération par l'addition d'un nouvel hyperplan au point maximum de l'enveloppe précédente. La recherche de ce point maximum à l'étape n de l'algorithme se fait au moyen d'un programme linéaire dont le programme dual fournit une solution primale u^n admissible convergeant vers une solution u^* quand $n \rightarrow +\infty$.

Décrivons l'algorithme dans le cas du problème (4.2) (dans le cas du problème formulé au chapitre précédent, on peut utiliser (3.4)-(3.5)). La recherche du maximum de l'enveloppe inférieure d'hyperplans support de ϕ aux points $p^k (k=1, \dots, n)$ d'équations :

$$z = \phi(p^k) + \langle p - p^k, \theta(v^k) \rangle = J(v^k) + \langle p, \theta(v^k) \rangle$$

$$\text{où } v^k \in \arg \min_{u \in U} L(u, p^k), \text{ donc } \theta(v^k) \in \partial \phi(p^k),$$

se formule

$$\max_{z, p \in C^*} z \quad (4.4)$$

$$\text{sous la contrainte } z \leq J(v^k) + \langle p, \theta(v^k) \rangle, \quad k = 1, \dots, n.$$

Soit p^{n+1} la solution de ce problème.

Le problème dual se formule alors :

$$\text{Min } \sum_{k=1}^n \alpha^k J(v^k)$$

$$\text{sous les contraintes } \sum_{k=1}^n \alpha^k = 1, \quad \alpha^k \geq 0, \quad k = 1, \dots, n \quad (4.5)$$

$$\sum_{k=1}^n \alpha^k \theta(v^k) \in C.$$

Soit $\{\alpha_n^k\}$ $k = 1, \dots, n$ une solution.

Alors, en raison de la C-convexité de θ et grâce à (4.5), il est clair que

$$u^n = \sum_{k=1}^n \alpha_n^k v^k \in U \quad \text{vérifie la contrainte du problème primal (4.2). On}$$

démontre que $u^n \rightarrow u^*$ quand $n \rightarrow +\infty$.

L'étape n de l'algorithme consiste donc, connaissant $\{(v^k, p^k); k=1, \dots, n\}$ à résoudre par la programmation linéaire les problèmes duaux (4.4) et (4.5), ce qui définit p^{n+1} et u^n (et les α_n^k). Puis on doit résoudre :

$$\text{Min } L(u, p^{n+1}) \\ u \in U$$

ce qui fournit v^{n+1} et donc une nouvelle paire (v^{n+1}, p^{n+1}) pour formuler le problème (4.4) à l'étape $n+1$.

Dans le cas où J (et θ) sont additifs ($J(u) = \sum_{i=1}^N J_i(u_i) \dots$), il est

possible d'ajouter N contraintes-supplémentaires à (4.4), à chaque itération. Mais comme on ne peut en contrepartie supprimer au cours de l'algorithme certaines contraintes qui ne sont pas saturées, sous peine de ne plus assurer la convergence en cas de dégénérescence, il en résulte alors une augmentation très sensible de la taille de la mémoire de stockage. Ceci constitue donc pour les problèmes de grande dimension un handicap sérieux. Cette orientation ne semble donc pas devoir être retenue dans notre contexte.

4 - . Solutions à la "Russe"

Polyak (1978) mentionne d'autres méthodes pour résoudre les problèmes de point-selle sans hypothèse de "stabilité". Ce sont des méthodes qui s'inspire davantage de l'algorithme d'Arrow-Hurwicz que de celui d'Uzawa.

Pour forcer l'unicité de la solution à chaque étape du calcul, Bakushinskii et Polyak (1974) proposent par exemple d'ajouter au Lagrangien une fonctionnelle strictement convexe-concave (convexe en u et concave en p),

tendant vers 0 au voisinage de l'optimum. Cette idée, à priori très séduisante, présente l'inconvénient suivant lors de son application au schéma de l'algorithme d'Uzawa.

Dans le cas d'un problème de minimisation linéaire (ou linéaire par morceaux) au premier niveau, pour lequel des méthodes de résolution très performantes existent, l'introduction d'un terme strictement convexe interdit leur utilisation. C'est justement le cas du problème qui a motivé notre étude et pour lequel nous utilisons un programme linéaire très performant mis au point à l'E D F par Maurras (1972).

Toujours en s'inspirant de Arrow-Hurwicz, Nurminskii et Verchenko (1977) proposent pour leur part l'algorithme suivant :

$$u^{k+1} = \Pi^* [u^k - \epsilon^k L_u^* (u^k, p^k)]$$

$$p^{k+1} = \Pi [p^k + \rho^k L_p^* (u^k, p^k)]$$

où Π^* représente la projection sur U et Π la projection sur P .

L'examen des conditions de convergence de ce nouvel algorithme révèle que en sus des conditions sur les "petits pas" ρ^k et ϵ^k , si ρ^k / ϵ^k tend vers 0 (à une certaine vitesse) quand k tend vers $+\infty$, toute limite de la suite $\{p^k\}$ appartient à P^* , mais dans ce cas aucun u^* de U^* n'est généralement atteint. De même, si ϵ^k / ρ^k tend vers 0 (à une certaine vitesse) lorsque k tend vers $+\infty$, c'est la suite $\{u^k\}$ qui converge vers un élément u^* de U^* alors que la suite $\{p^k\}$ ne converge généralement pas dans P^* .

L'obtention d'un point-selle $(u^*, p^*) \in U^* \times P^*$ au moyen d'un tel algorithme n'est donc pas possible (à moins de faire "tourner" deux fois l'algorithme avec des conditions différentes sur le ratio ϵ^k / ρ^k). De plus, l'extension à un algorithme de type Uzawa n'est pas immédiate.

5 - Lagrangien augmenté

Les lagrangiens augmentés permettent de garantir la stabilité en u (cf. Rockafellar (1973)). Par contre, ils ont comme pour l'algorithme de Bakhuisinkii et Polyak, l'inconvénient de transformer le programme primal linéaire en programme non linéaire.

Le lagrangien augmenté peut s'introduire de la manière suivante : partant du problème (4.2), on transforme la contrainte inégalité en contrainte égalité en introduisant en terme $\xi \in C$ tel que l'on se ramène au problème suivant :

$$\min_{\xi, u} J(u) \quad (i)$$

$$\text{sous } \theta(u) + \xi = 0 \quad (ii) \quad (4.6)$$

$$\text{et sous } \xi \in C, u \in U \quad (iii)$$

On peut ensuite se libérer de la contrainte (4.6-ii), soit en ajoutant un terme quadratique de pénalisation, soit en dualisant ces contraintes en introduisant des multiplicateurs $p(\in C^*)$.

L'idée du lagrangien augmenté est de superposer les deux procédés, ce qui aboutit à :

$$\bar{L}_c(u, \xi, p) = J(u) + c \langle p, \theta(u) + \xi \rangle + \frac{c}{2} \|\theta(u) + \xi\|^2$$

avec $c > 0$.

En minimisant par rapport à $\xi \in C$ on obtient :

$$L_c(u, p) = J(u) + \frac{c}{2} [\|\Pi[\theta(u) + \frac{p}{c}]\|^2 - \|\frac{p}{c}\|^2] \quad (4.7)$$

qui représente le lagrangien augmenté du problème (4.2).

On démontre que l'ensemble des points-selle de L_c sur $U \times C^*$ (et non pas seulement sur $U \times C^*$) est le même que pour L sur $U \times C^*$. On le notera donc encore $U^* \times P^*$ et U^* est l'ensemble des solutions de (4.2).

$$\text{Soit } \phi_c(p) = \min_{u \in U} L_c(u, p) \quad (4.8)$$

Il est clair que si J n'est pas strictement convexe, L_c ne l'est pas en général non plus et $\hat{U}_c(p) = \text{Arg Min}_{u \in U} L_c(u, p)$ n'est pas en général un

singleton (penser au cas où U^* lui-même n'est pas un singleton). Cependant, comme démontré par Rockafellar (1973) en dimension finie (cf. Zhu (1982) en dimension infinie), le lagrangien augmenté défini par (4.7) a deux propriétés intéressantes :

a) la fonction ϕ_c définie par (4.8) est différentiable. Ceci s'explique par le fait que ϕ_c est reliée à ψ (définie par (4.3)) par la relation

$$\phi_c(p) = \max_{q \in C^*} \left\{ \psi(q) - \frac{1}{2c} \|q-p\|^2 \right\}. \quad (4.9)$$

Celle-ci se démontre en utilisant le problème dual du problème

$$\min_{\xi \in C} \bar{L}_c(u, \xi, p).$$

définissant L_c . La différentiabilité est une conséquence du fait que le Max est atteint en un point unique dans (4.9).

Du point de vue algorithmique, on pourra donc utiliser un algorithme d'Uzawa à "grand pas" ($\rho > 0$) ce qui est un avantage à la fois par rapport à la méthode de Lasdon (où il faut résoudre un programme linéaire complet pour remettre à jour p) et par rapport aux méthodes "à la Russe" qui utilisent des "petits pas".

b) Le lagrangien augmenté est stable en u . En effet

$$\forall p^* \in P^*, (\phi_c)'(p^*) = 0 \quad (4.10)$$

puisque $P^* = \text{Arg Max}_{p \in C^*} \phi_c(p)$ et que ϕ_c est différentiable.

Mais :

$$(\phi_c)'(p^*) = (L_c)'_p(\hat{d}, p^*) \quad (4.11)$$

pour tout $\hat{d} \in \hat{U}(p^*)$ (cf. (4.18) ci-après et la remarque qui suit (4.18)), c'est à dire vérifiant

$$L_c(\hat{d}, p^*) \leq L_c(u, p^*), \quad \forall u \in U. \quad (4.12)$$

D'après (4.10), (4.11), on déduit que p^* réalise aussi le maximum de $p \rightarrow L_c(\hat{d}, p)$ sur C^* .

On a donc :

$$L_c(\hat{d}, p) \leq L_c(\hat{d}, p^*), \quad \forall p \in C^*. \quad (4.13)$$

Les inéquations variationnelles (4.12) et (4.13) montrent que $(\hat{u}, p^*) \in U^* \times P^*$ et que donc $\hat{u} \in \hat{U}(p^*)$ implique $\hat{u} \in U^*$. Les algorithmes de type Uzawa construisent alors des suites $\{u^k\}$ convergeant vers U^* (cf. Zhu (1982)). Cependant ils ont l'inconvénient de transformer un problème primal éventuellement linéaire ou linéaire par morceaux en un programme non linéaire. Dans notre cas, cela empêchera l'utilisation du programme FLOMAX (très performant) au niveau primal (cf. Maurras (1972)). On se tourne donc vers un nouvel algorithme préservant la nature linéaire du problème primal mais assurant malgré cela la convergence des solutions primales.

6 - Un nouvel algorithme de sous-gradient

Considérons le problème suivant :

$$\begin{array}{l} \text{Max } \phi(p) \\ p \in P \end{array} \quad (4.14)$$

P est un sous ensemble fermé convexe d'un espace de Hilbert $\mathcal{P}$. ϕ est une fonction concave, sous-différentiable : ceci veut dire que $\forall p, \exists \partial \phi(p) \subset P^*$ tel que $\forall r, \exists \partial \phi(p)$ la propriété suivante est vérifiée :

$$\phi(p+p') \leq \phi(p) + \langle r, p' \rangle, \quad \forall p' \in P$$

(propriété transposée de la sous-différentiabilité des fonctions convexes).

- Rappelons ici la propriété classique :

Une condition nécessaire et suffisante pour que p^* soit solution de (4.14) est que :

$$p^* \in P, \quad \exists q^* \in \partial \phi(p^*) : \langle q^*, p - p^* \rangle \leq 0, \quad \forall p \in P. \quad (4.15)$$

Enonçons maintenant le lemme suivant, dont la preuve est fournie en annexe :

Lemme 1 : L'ensemble des paires (p^*, q^*) satisfaisant (4.15) est de la forme $P^* \times Q^*$, où P^* et Q^* sont deux sous-ensembles convexes fermés, et la fonction : $(p, q) \rightarrow \langle p, q \rangle$ est constante sur $P^* \times Q^*$.

Nous proposons un algorithme qui permette d'atteindre l'un des quelconques des éléments (p^*, q^*) de $P^* \times Q^*$.

Considérons l'algorithme suivant :

Algorithme 1

(i) $p^0 \in P$, $q^0 \in \partial \psi(p^0)$, $k = 0$.

(ii) $\forall k \in \mathbb{N}$ $p^{k+1} = \Pi[p^k + \rho^k q^k]$ (4.16)

$$q^{k+1} = (1 - \varepsilon^{k+1}) q^k + \varepsilon^{k+1} r^{k+1} \quad (4.17)$$

avec $r^{k+1} \in \partial \psi(p^{k+1})$

et où Π désigne la projection sur P .

$\{p^k\}$ et $\{\varepsilon^k\}$ sont des suites réelles positives qui satisfont aux conditions de théorème suivant :

Théorème 1 :

Soit ψ une fonction concave, semi-continue supérieurement, sous-différentiable sur P et à sous-gradients bornés. Si le sous-ensemble convexe fermé de P de P n'est pas borné, on suppose que $\psi(p^k) \rightarrow -\infty$ lorsque $\|p^k\| \rightarrow +\infty$, pour toute suite $\{p^k\} \subset P$.

De plus, si les hypothèses ci-dessous sont vérifiées :

(i) $\rho^k = \gamma \varepsilon^k$, $\gamma > 0$, $\varepsilon^k \in]0, 1]$

$$\sum_{k \in \mathbb{N}} \varepsilon^k = +\infty \quad \sum_{k \in \mathbb{N}} (\varepsilon^k)^2 < +\infty$$

(ii) La suite $\left\{ \frac{(\varepsilon^k)^2}{\varepsilon^{k-1}} \right\}$ est décroissante.

(iii) La série de terme général.

$$\left(\frac{\varepsilon^{k-1} - \varepsilon^k}{\varepsilon^k} \right)^2 \quad \text{est convergente.}$$

alors la suite $\{(p^k, q^k)\}$ est bornée et tout point d'adhérence pour la topologie faible est une solution (p^*, q^*) de (4.15). Enfin, si ψ est fortement concave (cf. Auslender (1976)), la suite $\{p^k\}$ converge fortement vers l'unique solution p^* .

La preuve de ce théorème est fournie en annexe.

Il semble possible d'éviter de faire l'hypothèse de sous-gradients bornés (mais celle-ci est vérifiée pour le problème considéré dans cette thèse) et d'envisager pour ρ^k et ε^k du choix moins restrictifs du type

$$\rho^k = \gamma^k \varepsilon^k \text{ avec des conditions sur la suite } \{\gamma^k\} \text{ telles que } 0 < \gamma_1 < \gamma^k < \gamma_2.$$

7 - Application de l'algorithme précédent à la recherche du point-selle

Revenons maintenant au problème de recherche d'un point-selle. du Lagrangien $L(u, p) = J(u) + \langle p, \theta(u) \rangle$ sur $U \times P$ sous les hypothèses (H) définies au début du chapitre et rappelons les résultats suivants qui sont classiques :

Théorème 2 : Sous les hypothèses (H), le Lagrangien L possède un point-selle (u^*, p^*) sur $U \times P$.

Lemme 2 : La fonction ϕ définie par

$$p \mapsto \min_{u \in U} L(u, p)$$

est concave et sous-différentiable, et le sous-différentiel $\partial \phi(p)$ de ϕ en un point p vérifie

$$\partial \phi(p) = \overline{\text{co}} \, \theta[\hat{U}(p)] \quad (4.18)$$

où $\overline{\text{co}}$ représente la fermeture de l'enveloppe convexe.

Remarque : D'après (4.18), on s'aperçoit que ϕ sera différentiable si $\hat{U}(p)$ est un singleton, ou si θ applique $\hat{U}(p)$ sur un point-unique.

A l'aide de (4.18) et des deux inégalités variationnelles caractérisant le point-selle (u^*, p^*) on obtient le résultat suivant :

$$p^* \in P^*, \theta(u^*) \in \partial \phi(p^*) : \langle \theta(u^*), p - p^* \rangle \leq 0 \quad \forall p \in P. \quad (4.19)$$

Si l'on rapproche maintenant ce résultat de (4.15), on constate que la recherche d'un point-selle présente des liens étroits avec le problème traité à la section 6. Reprenons donc l'algorithme 1 en remplaçant dans (4.17) r^{k+1} par $\theta(u^{k+1})$, $u^{k+1} \in \hat{U}(p^{k+1})$.

Avec cette modification on peut alors intégrer (4.16) et (4.17) dans le schéma de l'algorithme d'Uzawa (cf. Algorithme 1 du chapitre III) en lieu et place de (iii) au deuxième niveau, sans toucher à l'étape de minimisation du premier niveau (ii) qui fournit u^k .

Le lemme 2 nous assure d'autre part que tout point d'adhérence de la suite $\{q^k\}$ appartient à $\overline{\Theta}[\Theta(\hat{U}(p^*))]$ et satisfait à l'inéquation variationnelle (4.15), associé à un $p^* \in P^*$.

Enfin, s'il existe $\bar{u} \in \hat{U}(p^*)$ tel que $q^* = \Theta(\bar{u})$, alors $(\bar{u}, p^*)$ est un point-selle car il vérifie (4.19) (inégalité de gauche du point-selle) et évidemment l'inégalité de droite par définition de $\bar{u}$.

Le point non résolu est l'existence d'une représentation de q^* sous la forme $\Theta(\bar{u})$ avec $\bar{u} \in \hat{U}(p^*)$. Pour tourner cette difficulté, on propose une adaptation de l'algorithme 1 aux problèmes de point-selle :

Algorithme 2 :

$$(i) \quad p^0 \in P, \quad v^0 \in \hat{U}(p^0), \quad q^0 = \Theta(v^0), \quad k = 0$$

$$(ii) \quad \forall k \in \mathbb{N} \quad p^{k+1} = \Pi[p^k + p^k q^k]$$

$$q^{k+1} = (1 - \varepsilon^{k+1}) q^k + \varepsilon^{k+1} \Theta(u^{k+1}) \quad (4.20)$$

$$\text{avec } u^{k+1} \in \hat{U}(p^{k+1}).$$

$$v^{k+1} = (1 - \varepsilon^{k+1}) v^k + \varepsilon^{k+1} u^{k+1} \quad (4.21)$$

Remarque : v^k est une combinaison convexe de toutes les solutions obtenues au premier niveau de l'algorithme, au cours des itérations successives. Cette combinaison convexe est récursive. On rapprochera cette observation de ce qui se passe dans l'algorithme décrit à la section 3 où la solution u^* est également obtenue comme limite de combinaisons convexes de solutions du premier niveau, mais où les combinaisons convexes ne sont pas obtenues récursivement (il faut à chaque fois résoudre un programme linéaire de taille croissante pour obtenir les "poids").

Théorème 3 : Sous les hypothèses (H) et sous les hypothèses du théorème 1 concernant p^k et ε^k , la suite $\{(v^k, p^k)\}$ est bornée et tout point d'adhérence pour la topologie faible est un point-selle (u^*, p^*) de L sur $U \times P$.

Une variante plus simple de l'algorithme 2 peut-être également considérée :

Algorithme 3 :

$$(i) \quad p^0 \in P, \quad v^0 \in \hat{U}(p^0), \quad q^0 = \theta(v^0), \quad k = 0$$

$$(ii) \quad \forall k \in \mathbb{N} \quad p^{k+1} = \Pi[p^k + p^k \theta(v^k)] \quad (4.22)$$

$$v^{k+1} = (1 - \epsilon^{k+1})v^k + \epsilon^{k+1}u^{k+1} \quad (4.23)$$

$$\text{avec } u^{k+1} \in \hat{U}(p^{k+1}).$$

Ces deux algorithmes coïncident de manière évidente lorsque la fonction θ est affine. Quand θ est convexe sans être affine, nous émettons la conjecture que le théorème 3 reste valable pour l'algorithme 3, sans toutefois avoir pu démontrer cette conjecture.

CHAPITRE V

Application au problème EDF

Nous présentons dans ce chapitre un algorithme de résolution du problème posé par EDF (cf. chapitre I et II). Cet algorithme est une combinaison de l'algorithme 2 du chapitre III et de l'algorithme 1 proposé au chapitre IV. Il résout le problème de la non-unicité de la solution primale à chaque étape au niveau de la minimisation, tout en conservant l'originalité de l'algorithme 2 du chapitre III qui consiste à éviter le calcul de projection sur autre chose que le cône positif.

Nous reprenons ici les notations spécifiques du problème EDF. Sans décrire à nouveau en détail toutes ces notations (cf. II), rappelons simplement que le coût minimum s'exprime de la manière suivante (cf. (2.6)) :

$$c_D^Y + \min_{q(\omega)} \max_{p(\omega)} \left\{ \sum_{\omega=1}^N \left[\frac{1}{N} (c_P q_P(\omega) - c_D q_D(\omega)) + c_T (|q_T(\omega)| \times p(\omega)) \right] - c_T [q_T^0 \times (1 - p(\omega_0))] \right\} \quad (5.1)$$

sous les contraintes :

$$\begin{aligned} Aq(\omega) &= 0 \\ 0 &\leq q_d(\omega) \leq q_d^Y \\ 0 &\leq q_p(\omega) \leq q_p^Y \\ |q_t(\omega)| &\leq q_t^Y \end{aligned} \quad (5.2)$$

soit k l'indice de l'itération. L'algorithme se déroule ainsi :

1 - Initialisation

Pour $k = 0$, il faut se donner une valeur initiale des paramètres duaux :

$$\{p_t^0(\omega)\} \quad t \in T \quad ; \quad \omega = 0, \dots, N$$

N étant le nombre d'aléas considérés dans le calcul de l'approximation de l'espérance mathématique.

Nous avons choisi pour notre part de prendre 1 pour valeur du poids mis sur l'aléa fictif (ω_0) qui représente l'investissement existant ; ceci reviendrait, dans le cadre d'une approche primale, à initialiser les capacités des lignes du réseau à la capacité existante.

On pose donc :

$$\begin{aligned} p_t^0(0) &= 1 ; & p_t^0(\omega) &= 0 & \forall \omega = 1, \dots, N \\ & & & & \forall t \in T \end{aligned} \quad (5.3)$$

2 - Résolution du problème de minimisation à l'itération k.

Cette phase constitue le premier niveau de l'algorithme. Elle est identique pour tous les algorithmes considérés dans cette étude.

Pour $\omega = 1, \dots, N$ on résoud donc les problèmes de minimisation suivants :

$$\text{Min}_{q(\omega)} \frac{1}{N} \left[\sum_p c_p(\omega) q_p(\omega) - \sum_d c_d q_d(\omega) \right] + \sum_t c_t |q_t(\omega)| p_t^k(\omega) \quad (54)$$

sous les contraintes (5.2).

Cette résolution se fait au moyen d'un programme linéaire de flot à coût minimum mis au point à l'EDF par Maurras, (c'est ce programme particulièrement performant, FLOMAX, qui nous a fait rejeter des solutions détruisant la linéarité de (5.4) - cf. chapitre IV, section 4 et 5).

Remarque 1. Pour utiliser ce programme on pose classiquement :

$|q_t(\omega)| = q_t^+(\omega) + q_t^-(\omega)$ afin de linéariser la fonctionnelle à minimiser. On rappelle que

$$q_t^+(\omega) = \text{Max}(0, q_t(\omega)) \quad \text{et} \quad q_t^-(\omega) = [-q_t(\omega)]^+.$$

Remarque 2. Dans l'approche simplifiée considérée par Dodu, on ne prend en compte que le coût de défaillance ; en effet, après avoir classé les groupes de production disponibles (déterminés par tirage aléatoire) dans l'ordre des coûts croissants de production, on ne retient que les premiers de la liste ainsi ordonnée, de telle sorte que la somme de leurs capacités satisfasse au plus juste la demande totale sur le réseau. Dans ce cas on pose $c_p(\omega) = 0$ dans les formules ci-dessus, et l'on modifie le vecteur $\bigvee_p(\omega)$ en conséquence, en donnant la valeur 0 aux coordonnées correspondant à des groupes qui, bien que disponibles, ont été écartés par la sélection ainsi opérée.

La solution optimale de (5.4)-(5.2), non nécessairement unique, calculée par FLOMAX, est notée $q^k(\omega)$, $\omega = 1, \dots, N$.

3- Calcul du coût dual et du coût primal.

a) Le coût dual (correspondant à la fonction ψ du chapitre IV) qu'il nous faut maximiser, et qui inférieur ou égal à la valeur $\psi(p^*)$ du point-selle, est évalué de la manière suivante :

$$\begin{aligned} \psi(p^k) = & \sum_{\omega=1}^N \left\{ \frac{1}{N} \sum_p c_p q_p(\omega) + \sum_d c_d (q_d^k - q_d^k(\omega)) \right\} + \sum_t [c_t |q_t^k(\omega)| v_t^k(\omega)] \\ & - \sum_t c_t^0 (1 - p_t^k(0)) \end{aligned} \quad (5.5)$$

b) Pour $k > 0$ aller en c)

Pour $k = 0$, on pose $z_t^0(\omega) = |q_t^0(\omega)|$ pour $\omega = 0, \dots, N$ et $t \in T$

$$\text{et } a_t^0 = \sum_{\omega=0}^N p_t^0(\omega) |q_t^0(\omega)| \quad t \in T$$

Puis aller en d)

Remarque 3 : On rappelle que, $\forall k \in N$, $q_t^k(0) = q_t^0$ par convention.

c) Pour $k > 0$ on calcule

$$z_t^k(\omega) = (1 - \varepsilon^k) z_t^{k-1}(\omega) + \varepsilon^k |q_t^k(\omega)| \quad (5.6)$$

$$a_t^k = (1 - \varepsilon^k) a_t^{k-1} + \varepsilon^k \sum_{\omega=0}^N p_t^k(\omega) |q_t^k(\omega)|. \quad (5.7)$$

d) Le coût primal, φ^k , supérieur ou égal à la valeur du point-selle, est évalué à l'itération k de la manière suivante :

$$\varphi^k = \frac{1}{N} \sum_{\omega=1}^N \left[\sum_p c_p q_p^k(\omega) + \sum_d c_d (q_d^k - q_d^k(\omega)) \right] + \sum_t c_t (y_t^k - q_t^k) \quad (5.8)$$

où

$$y_t^k = \max_{\omega=0, N} z_t^k(\omega) \quad \text{pour } t \in T \quad (5.9)$$

4. Remise à jour des paramètres duaux

On calcule :

$$\Pi_t^{k+1} = \text{Max} \{0, p_t^k(\omega) + \rho^k \alpha_t [z_t^k(\omega) - a_t^k]\} \quad (5.10)$$

$$p_t^{k+1}(\omega) = \frac{\Pi_t^{k+1}(\omega)}{\sum_{\omega=0}^N \Pi_t^{k+1}(\omega)} \quad \text{avec } \omega = 0, \dots, N \text{ et } t \in T \quad (5.11)$$

5. Test d'arrêt

a) Ce test porte normalement sur l'écart entre le coût dual $\phi(p^k)$ et le coût primal φ^k , c'est à dire sur la quantité suivante :

$$\Delta^k = \sum_t [y_t^k - \sum_{\omega=0}^N p_t^k(\omega) |q_t^k(\omega)|] \quad (5.12)$$

qui tend vers 0 lorsque l'on se rapproche de l'optimum.

b) Si le test d'arrêt n'est pas satisfait, faire $k \leftarrow k+1$ et retourner en 2.

6. Remarques et Commentaires.

a) Le schéma (5.10)-(5.11) est la transposition immédiate du schéma (iii) (3.10)-(3.11) de l'algorithme 2 du chapitre III.

b) Cependant, l'élément du sous-gradient de la fonctionnelle quasi-concave

$$\phi(p) / \langle p, 1 \rangle \text{ au point } (|q_t^k(\omega) - \sum_{\omega=0}^N p_t^k(\omega) |q_t^k(\omega)|) \quad t \in T, \omega=0, \dots, N$$

est remplacé par un terme résultant du "moyennage" des sous-gradients (cf. (5.6), (5.7) et (5.10)), comme pour le passage de l'algorithme de sous-gradient simple à l'algorithme de sous gradient moyenné ((4.16), (4.17)). L'algorithme proposé consiste donc bien en une combinaison des algorithmes 2 de III et 1 de IV.

c) Pour appliquer l'algorithme 2 de IV, il faudrait en toute rigueur faire les moyennes des $q^k(\omega)$ (cf.(4.21)) dont on se servirait d'une part pour évaluer le coût primal φ^k (puisque $q^k(\omega)$ tend vers q^* d'après le théorème 3), et d'autre part pour évaluer les capacités optimales des lignes $t \in T$ de transport électrique : $\bar{q}_t^k(\omega)$ représentant la moyenne (au sens de (4.21)) des $q_t^l(\omega)$ pour $1 \leq k$, le calcul de la limite quand k tend vers $+\infty$ de $\text{Max}_{\omega} |\bar{q}_t^k(\omega)|$ donnerait les valeurs de ces capacités optimales.

La mise en oeuvre de cette procédure nécessiterait une taille mémoire considérable sur ordinateur. Mais :

d) D'une part, $q_t^k(\omega)$ ne change plus de signe après quelques itérations, et il revient donc au même d'utiliser dans l'algorithme $z_t^k(\omega)$ (cf. (5.6)) au lieu de $|\bar{q}_t^k(\omega)|$.

e) D'autre part, il semble que la non-unicité porte sur les variables $q_t(\omega)$ (relatives à l'acheminement de la puissance électrique) et non sur les variables $q_p(\omega)$ et $q_d(\omega)$ (représentant la façon de produire la puissance électrique et de satisfaire la demande) ; on a donc confondu $\bar{q}_p^k(\omega)$ (respectivement $\bar{q}_d^k(\omega)$) avec la dernière valeur obtenue au cours des itérations, $q_p^k(\omega)$ (respectivement $q_d^k(\omega)$) dans le calcul du coût primal φ^k . A ces heuristiques près, l'algorithme utilisé est alors l'algorithme 2 du chapitre IV (combiné à l'algorithme 2 du chapitre III).

f) L'algorithme 3 consisterait à faire les moyennes des $q^k(\omega)$ comme en (4.21) pour obtenir $\bar{q}_t^k(\omega)$, puis à utiliser l'expression :

$$|\bar{q}_t^k(\omega)| - \sum_{\omega=0}^N p_t^k(\omega) |\bar{q}_t^k(\omega)| \text{ à la place de } z_t^k(\omega) - a_t^k \text{ dans la formule (5.10).}$$

Si, comme observé en d), les termes $z_t^k(\omega)$ et $|\bar{q}_t^k(\omega)|$ vont pratiquement coïncider (on est donc pratiquement dans le cas θ linéaire) par contre a_t^k donné par (5.7) et $\sum_{\omega=0}^N p_t^k(\omega) |\bar{q}_t^k(\omega)|$ ne coïncident pas. Ceci est dû au fait que la dérivée partielle de Φ en p (ou en q) fait encore intervenir p explicitement (cf. Lemme 1 du Chapitre III), alors que ce n'est pas le cas pour le lagrangien (4.1) linéaire en p considéré au chapitre IV.

g) Si l'on n'accepte pas les considérations heuristiques des alinéas d) et e), l'algorithme décrit ci-dessus est l'algorithme 1 du chapitre IV (combiné avec l'algorithme 2 du chapitre III). On a vu que cet algorithme appliqué au problème du point-selle fournit un couple (p^*, q^*) vérifiant (4.15).

Cependant, si q^* n'est pas de la forme $\theta(\bar{u})$ pour $\bar{u} \in \theta[\bar{U}(p^*)]$ on ne peut pas tirer de conclusion sur l'appartenance de $\bar{u}$ à U^* . Ici, on souhaite pouvoir calculer des expressions du type

$$y^* = \max_i \theta_i(u^*) = \sum_i p_i^* \theta_i(u^*) \text{ pour } u^* \in U^* \text{ et } p^* \in P^*,$$

pour évaluer les investissements optimaux. Il suffit donc de pouvoir affirmer le corollaire suivant du théorème 1 du chapitre IV.

Corollaire

Tout point d'adhérence (p^*, q^*) de la suite construite par l'algorithme 1 du chapitre IV est tel que :

$$\langle p^*, q^* \rangle = \langle p^*, \theta(u^*) \rangle \quad \text{pour } u^* \in U^*$$

Preuve : On a vu que $(p^*, q^*) \in P^* \times Q^*$ (solution de (4.15)) et que la fonctionnelle $(p, q) \rightarrow \langle p, q \rangle$ est constante sur cet ensemble (Lemme 1, chapitre IV). Comme, $\forall u^*, \in U^*$, on a aussi $(p^*, \theta(u^*)) \in P^* \times Q^*$, le résultat est immédiat ■

S'agissant du cas C de dimension finie et $P = \Sigma$ (cf(3.3-1)) cela prouve (avec (4.15)) que

$$\max_i \theta_i(u^*) = \sum_i p_i^* \theta_i(u^*) = \sum_i p_i^* q_i^* = \max_i q_i^*$$

On pourrait de plus prouver que

$$\forall i, \theta_i(u^*) \leq q_i^*$$

et que $p_i^* > 0$ entraîne que $\theta_i(u^*) = q_i^*$.

Ces remarques justifient de toute façon l'évaluation des investissements optimaux dans notre problème par la limite des expressions (5.9).

CHAPITRE VI

Résultats Numériques

1. Exemple 1.

Nous avons appliqué l'algorithme 2 du chapitre IV à un exemple test tiré de l'ouvrage de Lemaréchal et Mifflin (1978) .

La fonctionnelle à minimiser est définie ainsi :

$$k = \{1, \dots, 5\}$$

$$i, j \in \{1, \dots, 10\}$$

$$a_k(i, j) = e^{\frac{i}{j}} \cos(i \cdot j) \sin(k) \quad i \neq j$$

$$a_k(j, i) = a_k(i, j)$$

$$a_k(i, i) = \sin(k) \times \frac{i}{10} + \sum_{\substack{j=1 \\ j \neq i}}^{10} a_k(i, j)$$

$$b_k(i) = e^{\frac{i}{k}} \times \sin(i \cdot k)$$

On cherche à minimiser $\text{Max}_{k=1,5} \sum_{i=1,10} \{ [\sum_{j=1,10} a_k(i, j) x_j] x_i - b_k(i) x_i \}$

Le meilleur résultat sur cet exemple a été obtenu avec

$$\epsilon^n = \frac{1}{[(1+0,25(n-1))]}; \quad \rho^n = 0,80 \cdot \epsilon^n \times \left(\sum_{l=1}^n \epsilon^l \right)^{1/2} / \sqrt{1+5/|g^n - g^*|}$$

$0,80 \times \left(\sum_{l=1}^n \epsilon^l \right)^{1/2}$ est un terme qui permet à ρ^n de décroître beaucoup moins vite que ϵ^n dans les premières itérations.

Le facteur $1/\sqrt{1+5/|g^n - g^*|}$ permet d'accélérer la décroissance de ρ^n sa valeur g^n parvient au voisinage de g^* .

g^n est la valeur courante de la fonctionnelle et g^* sa valeur optimale (supposée connue.)

En plus de la courbe 1 on peut également donner le tableau suivant, où q désigne le vecteur des sous-gradients moyennés courant (voir (4.20)), x désigne le vecteur courant, et x^* le vecteur solution.

Iterations	ϵ	$\ x-x^*\ $	$\ q\ $
1	533,7	3,163	1
50	-0,62	0,052	0,0051
100	-0,8274	0,017	0,06
200	-0,8362	0,0078	0,03
300	-0,8408	0,0054	0,013
400	-0,8411	0,0042	0,0055
500	-0,8412	0,0031	0,013

2. Exemple 2

On a fait tourner l'algorithme du chapitre V sur un mini réseau (3 arcs, 3 noeuds) avec 10 groupes de production. Toutes les données ont été tirées de celles du réseau EDF : les 3 premières valeurs pour tout ce qui concerne les arcs et les noeuds, le 10 premières valeurs pour les groupes de production. Nous avons également considéré 500 aléas.

Si P désigne le coût primal et D le coût dual, le rapport $\frac{P-D}{P}$ est inférieur à 10 % en 90 itérations et inférieur à 4 % en 150 itérations. Valeur des pas dans le cas présenté ici :

$$\rho^n = \frac{0,15 \times 10^{-6}}{1+0,08(n-1)} \quad \epsilon^n = \frac{1}{1+0,5(n-1)}$$

Valeur maximale du coût dual $0,227 \cdot 10^6$ obtenue après 150 itérations.

Valeur minimale du coût primal correspondant $0,240 \cdot 10^6$.

3. Exemple 3

On a appliqué le même algorithme à ce mini réseau avec $\epsilon^n = 1 \forall n$ et $\rho^n = \frac{0,7|D-P|}{(c^n)^2}$.

D^* est la valeur estimée du coût optimal (valeur du point-selle).

g^n est la norme du vecteur du sous-gradient intervenant dans l'algorithme.

Il s'agit alors d'un algorithme de sous-gradient avec le choix de ρ^n proposé par Polyak (1978). La valeur 0,7 s'est avérée la meilleure valeur parmi celles essayées pour k qui doit être comprise entre 0 et 2. Comme ϵ^n reste constant le coût primal reste divergent. Le coût dual ne converge pas plus vite vers l'optimum que précédemment :

$0,215 \times 10^6$ en 150 itérations

$0,224 \times 10^6$ en 300 itérations.

4. Exemple 4.

La taille du réseau EDF ne nous a pas permis d'obtenir plus de 50 à 60 itérations pour l'algorithme du chapitre V.

Dans le cas présenté ici, $\rho^n = 0,15 \cdot 10^{-6} \cdot \forall n$.

$$\text{et } \epsilon^n = \frac{15}{1+20n} \text{ pour } n \geq 1.$$

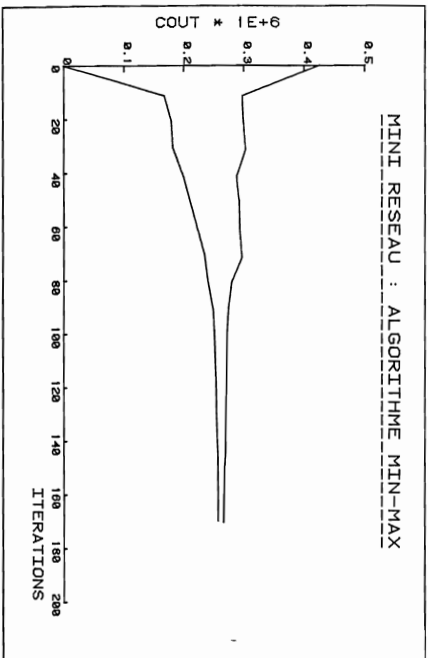
Valeur maximale obtenue pour le coût dual $0,12 \cdot 10^7$

Valeur minimale obtenue pour le coût primal $0,28 \cdot 10^7$.

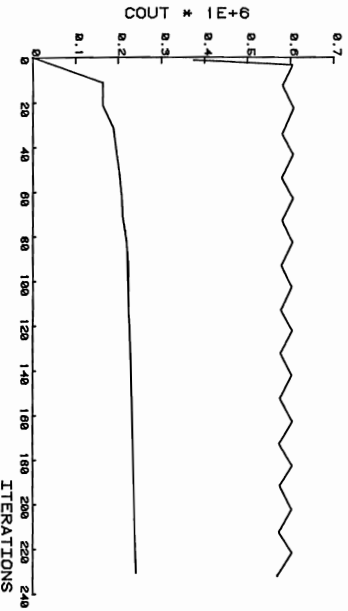
Une estimation du coût optimal calculé pour la solution fournie par EDF sur les mêmes situations aléatoires donnant une valeur de l'ordre de $0,2 \cdot 10^7$, il apparaît que la convergence de cet algorithme pour ce problème est au moins aussi lente que pour le petit réseau test de l'exemple 2 : 150 à 200 itérations, chaque itération comportant 500 calculs de FLOMAX.

5. Conclusions.

Il serait souhaitable que d'autres expériences numériques permettent de mieux évaluer la qualité des algorithmes présentés au chapitre IV. Sur les exemples 2 et 3 notamment, on a vu que leur comportement se compare assez favorablement à celui d'autres algorithmes plus connus (Polyak) mais qui ne donnent pas le résultat escompté dans un problème de point-selle. L'échec obtenu sur le problème EDF en vraie grandeur semble plus dû au mauvais conditionnement de l'approche duale pour ce problème qu'à la qualité intrinsèque des algorithmes utilisés.

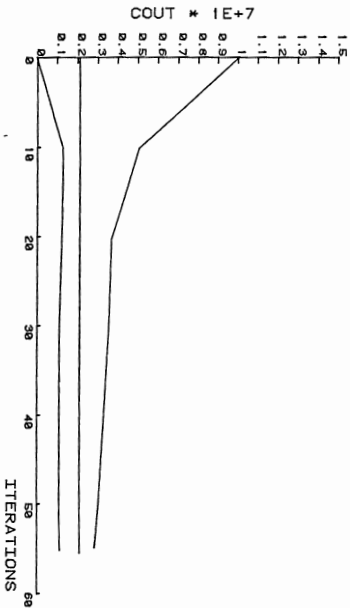


MINI RESEAU : METHODE DE POLYAK



Exemple n° 4

RESEAU EDF ----- exemple no 4



PREUVES des LEMMES et THEOREMES du chapitre IV

1. PREUVE du LEMME 1

Soit (p_1^*, q_1^*) et (p_2^*, q_2^*) satisfaisant (4.15)

$$\left. \begin{array}{l} \text{Alors } q_1^* \in \partial\psi(p_1^*) \\ q_2^* \in \partial\psi(p_2^*) \\ \psi \text{ concave} \end{array} \right\} \Rightarrow \langle q_1^* - q_2^*, p_1^* - p_2^* \rangle \leq 0. \quad (1)$$

De plus, d'après (4.15)

$$\langle q_1^*, p_2^* - p_1^* \rangle \leq 0 \quad (2)$$

$$\langle q_2^*, p_1^* - p_2^* \rangle \leq 0. \quad (3)$$

Par (1), (2), (3)

$$0 \leq \langle q_1^*, p_1^* - p_2^* \rangle \leq \langle q_2^*, p_1^* - p_2^* \rangle \leq 0.$$

$$\text{D'où } \langle q_1^*, p_1^* \rangle = \langle q_2^*, p_2^* \rangle = \langle q_1^*, p_2^* \rangle = \langle q_2^*, p_1^* \rangle. \quad (4)$$

Alors, d'après (4.15)

$$\langle q_1^*, p - p_1^* \rangle \leq 0, \quad \forall p \in P$$

d'où, d'après (4) ci-dessus

$$\langle q_1^*, p - p_2^* \rangle \leq 0, \quad \forall p \in P. \quad (5)$$

De plus, puisque $q_1^* \in \partial\psi(p_1^*)$,

$$\psi(p) - \psi(p_1^*) \leq \langle q_1^*, p - p_1^* \rangle = \langle q_1^*, p - p_2^* \rangle, \quad \forall p \quad (6)$$

la dernière égalité d'après (4) ci-dessus.

Par ailleurs, on a par hypothèse

$$\psi(p_1^*) = \psi(p_2^*) = \max_{p \in P} \psi(p). \quad (7)$$

De (6) et (7) on déduit que $q_1^* \in \partial\psi(p_2^*)$, ce qui montre avec (5) que la paire (p_2^*, q_1^*) satisfait (4.15). Il en va évidemment de même pour la paire (p_1^*, q_2^*) .

On a donc montré que si (p_1^*, q_1^*) et (p_2^*, q_2^*) sont solutions de (4.15), il en est de même pour (p_2^*, q_1^*) et (p_1^*, q_2^*) , c'est-à-dire que l'ensemble des solutions est de la forme $P^* \times Q^*$. De plus (4) montre que la fonctionnelle $(p, q) \mapsto \langle p, q \rangle$ est constante sur ce sous-ensemble.

Montrons maintenant que $P^* \times Q^*$ est un convexe fermé.

Soit $\bar{p} = \alpha p_1^* + (1-\alpha)p_2^*$, $\alpha \in [0, 1]$.

Alors $\phi(\bar{p}) = \max_{p \in P} \phi(p)$ d'après (7) et la concavité de ϕ . Donc $\bar{p} \in P^*$, et d'après ce qui précède $(\bar{p}, q_1^*)$ et $(\bar{p}, q_2^*)$ appartiennent aussi à $P^* \times Q^*$ donc $q_1^* \in \partial\phi(\bar{p})$ et $q_2^* \in \partial\phi(\bar{p})$. Par convexité du sous-différentiel on a que

$$\bar{q} = \alpha q_1^* + (1-\alpha)q_2^* \in \partial\phi(\bar{p}).$$

Par ailleurs

$$\left. \begin{aligned} \forall p \in P, \quad & \langle q_1^*, p - \bar{p} \rangle \leq 0 \\ & \langle q_2^*, p - \bar{p} \rangle \leq 0 \end{aligned} \right\} \Rightarrow \langle \bar{q}, p - \bar{p} \rangle \leq 0.$$

Ceci montre $(\bar{p}, \bar{q}) \in P^* \times Q^*$, c'est à dire que ce sous-ensemble est convexe.

Montrons qu'il est fermé.

Comme P^* est exactement égal à l'ensemble $\{p^* | \phi(p^*) = \max_{p \in P} \phi(p)\}$, cet ensemble est fermé en supposant ϕ s.c.s.. Soit $\{q_k\}$ une suite dans Q convergente vers $\bar{q}$. Alors, $\forall p^* \in P^*$

$$\langle q_k, p - p^* \rangle \leq 0, \quad \forall p \in P$$

et par passage à la limite on a la même inéquation variationnelle pour $\bar{q}$. D'autre part, $q^k \in \partial\phi(p^*)$, et comme le sous-différentiel est fermé, $\bar{q} \in \partial\phi(p^*)$. Donc $\bar{q} \in Q^*$, c.q.f.d.

2. PREUVE du THEOREME 1.

2.1. Supposons d'abord que l'Algorithme 1 (4.16)-(4.17) devienne stationnaire $(p^{k+1} = p^k, q^{k+1} = q^k)$ au bout d'un nombre fini de pas et montrons que le point stationnaire (notons le $(\bar{p}, \bar{q})$) est un point de $P^* \times Q^*$. En effet, d'après (4.16)-(4.17), on a

$$\forall p \in P, \quad \langle p - p_{k+1}, \gamma \varepsilon_k q_k + p_k - p_{k+1} \rangle \leq 0. \quad (8)$$

$$r_{k+1} = (q_{k+1} - q_k) / \varepsilon_{k+1} + q_k \in \partial \psi(p_{k+1}).$$

et donc

$$\forall p \in P, \quad \langle \bar{q}, p - \bar{p} \rangle \leq 0$$

$$\text{avec } \bar{q} \in \partial \psi(\bar{p}),$$

et donc $(\bar{p}, \bar{q}) \in P^* \times Q^*$.

2.2. On considère l'expression suivante

$$\Phi_k = \psi(p_k) - \frac{1}{2\gamma} \|p_{k+1} - p^*\|^2 - \frac{1}{2\gamma} \|(p_{k+1} - p_k) / \varepsilon_k\|^2 - \langle q_k, p_{k+1} - p^* \rangle. \quad (9)$$

D'après (8) ci-dessus

$$- \langle q_k, p_{k+1} - p^* \rangle \leq - \frac{1}{\gamma \varepsilon_k} \langle p_{k+1} - p^*, p_{k+1} - p_k \rangle$$

D'où

$$\begin{aligned} \Phi_k &\leq \psi(p_k) - \frac{1}{2\gamma} \|p_{k+1} - p^*\|^2 + (p_{k+1} - p_k) / \varepsilon_k \|^2 \\ &\leq \psi(p_k) \leq \psi(p^*). \end{aligned} \quad (10)$$

2.3. Etudions la variation $\Phi_{k+1} - \Phi_k$. On va étudier séparément la variation des quatre termes de (9) ci-dessus.

$$a) \quad \psi(p_k) - \psi(p_{k+1}) \leq \langle r_{k+1}, p_k - p_{k+1} \rangle \quad (11)$$

d'après la définition du sous-différentiel d'une fonction concave.

$$b) \quad \frac{1}{2\gamma} \|p_{k+2} - p^*\|^2 \leq \frac{1}{2\gamma} \|p_{k+1} - p^*\|^2 + \gamma \varepsilon_{k+1} q_{k+1} \|^2$$

d'après (4.16) et le fait que la projection est contractante.

$$\frac{1}{2\gamma} \|p_{k+2} - p^*\|^2 \leq \frac{1}{2\gamma} \|p_{k+1} - p^*\|^2 + \frac{\gamma \varepsilon_{k+1}^2}{2} M^2 + \varepsilon_{k+1} \langle q_{k+1}, p_{k+1} - p^* \rangle \quad (12)$$

où on a utilisé le fait que $\|q_{k+1}\| \leq M$ qui se déduit aisément de l'hypothèse $\|r_k\| \leq M, \forall k$.

- c) $\frac{1}{2\gamma} \| (p_{k+2} - p_{k+1}) / \varepsilon_{k+1} \|^2 \leq \frac{1}{2\gamma \varepsilon_{k+1}^2} \| p_{k+1} - p_k + \gamma (\varepsilon_{k+1} q_{k+1} - \varepsilon_k q_k) \|^2$
 (d'après (4.16) et le fait que la projection est contractante).
 En utilisant (4.17), on a

$$\begin{aligned} & \frac{1}{2\gamma} \| (p_{k+2} - p_{k+1}) / \varepsilon_{k+1} \|^2 \leq \frac{1}{2\gamma \varepsilon_{k+1}^2} \| p_{k+1} - p_k + \gamma \varepsilon_{k+1}^2 r_{k+1} + \gamma (\varepsilon_{k+1} (1 - \varepsilon_{k+1}) - \varepsilon_k) q_k \|^2 \\ = & \frac{1}{2\gamma} \left(\frac{1}{\varepsilon_{k+1}^2} - \frac{1}{\varepsilon_k^2} \right) \| p_{k+1} - p_k \|^2 + \frac{1}{2\gamma} \| (p_{k+1} - p_k) / \varepsilon_k \|^2 + \\ & + \frac{\gamma}{2\varepsilon_{k+1}^2} \| \varepsilon_{k+1}^2 r_{k+1} + (\varepsilon_{k+1} (1 - \varepsilon_{k+1}) - \varepsilon_k) q_k \|^2 + \\ & \cdot \langle r_{k+1}, p_{k+1} - p_k \rangle + \frac{(\varepsilon_{k+1} (1 - \varepsilon_{k+1}) - \varepsilon_k)}{\varepsilon_{k+1}^2} \langle q_k, p_{k+1} - p_k \rangle. \end{aligned} \quad (13)$$

D'après (8) ci-dessus on a

$$\langle q_k, p_{k+1} - p_k \rangle \geq \frac{1}{\gamma \varepsilon_k} \| p_{k+1} - p_k \|^2. \quad (14)$$

Par ailleurs

$$\frac{\varepsilon_{k+1} (1 - \varepsilon_{k+1}) - \varepsilon_k}{\varepsilon_{k+1}^2} = \frac{\varepsilon_{k+1} - \varepsilon_k}{\varepsilon_{k+1}^2} - 1$$

et $\varepsilon_{k+1} - \varepsilon_k \leq 0$ si la suite $\{\varepsilon_k\}$ est non croissante. Alors pour le dernier terme de (13) on a

$$\frac{(\varepsilon_{k+1} (1 - \varepsilon_{k+1}) - \varepsilon_k)}{\varepsilon_{k+1}^2} \langle q_k, p_{k+1} - p_k \rangle \leq - \langle q_k, p_{k+1} - p_k \rangle + \frac{\varepsilon_{k+1} - \varepsilon_k}{\gamma \varepsilon_k^2 \varepsilon_{k+1}} \| p_{k+1} - p_k \|^2.$$

En considérant maintenant la somme du premier et du dernier terme dans le second membre de (13), elle est majorée par

$$- \langle q_k, p_{k+1} - p_k \rangle - \frac{1}{2\gamma} \left(\frac{\varepsilon_k - \varepsilon_{k+1}}{\varepsilon_k \varepsilon_{k+1}} \right)^2 \| p_{k+1} - p_k \|^2 \leq - \langle q_k, p_{k+1} - p_k \rangle.$$

Finalement, de (13) et des considérations ci-dessus, on déduit

$$\frac{1}{2\gamma} \| (p_{k+2} - p_{k+1}) / \varepsilon_{k+1} \|^2 \leq \frac{1}{2\gamma} \left(\frac{\varepsilon_k - \varepsilon_{k+1}}{\varepsilon_k \varepsilon_{k+1}} \right)^2 \| p_{k+1} - p_k \|^2 + \frac{\varepsilon_{k+1} - \varepsilon_k}{2\gamma \varepsilon_k^2 \varepsilon_{k+1}} \| p_{k+1} - p_k \|^2.$$

$$+ \langle r_{k+1}, p_{k+1} - p_k \rangle - \langle q_k, p_{k+1} - p_k \rangle. \quad (15)$$

$$\begin{aligned} d) \quad & \langle q_k, p_{k+1} - p^* \rangle - \langle q_{k+1}, p_{k+2} - p^* \rangle = \\ & \langle q_{k+1}, p_{k+1} - p_{k+2} \rangle + \langle q_k - q_{k+1}, p_{k+1} - p^* \rangle = \\ & \langle q_{k+1}, p_{k+1} - p_{k+2} \rangle + \epsilon_{k+1} \langle q_k - r_{k+1}, p_{k+1} - p^* \rangle. \end{aligned}$$

(en utilisant (4.17)).

De plus

$$\phi(p^*) - \phi(p_{k+1}) \leq \langle r_{k+1}, p^* - p_{k+1} \rangle, \quad (r_{k+1} \in \partial \phi(p_{k+1}))$$

d'où

$$\begin{aligned} & \langle q_k, p_{k+1} - p^* \rangle - \langle q_{k+1}, p_{k+2} - p^* \rangle \geq \\ & \langle q_{k+1}, p_{k+1} - p_{k+2} \rangle + \epsilon_{k+1} \langle q_k, p_{k+1} - p^* \rangle + \epsilon_{k+1} (\phi(p^*) - \phi(p_{k+1})). \end{aligned} \quad (16)$$

e) En réunissant les résultats de (11), (12), (15), (16), on obtient

$$\begin{aligned} \phi_{k+1} - \phi_k & \geq -\frac{\gamma}{2} M^2 \left[\left(\frac{\epsilon_k - \epsilon_{k+1} + 2\epsilon_{k+1}^2}{\epsilon_{k+1}} \right)^2 + \epsilon_{k+1}^2 \right] \\ & + \epsilon_{k+1} \langle q_k - q_{k+1}, p_{k+1} - p^* \rangle + \langle q_{k+1}, p_{k+1} - p_{k+2} \rangle \\ & - \langle q_k, p_k - p_{k+1} \rangle + \epsilon_{k+1} (\phi(p^*) - \phi(p_{k+1})). \end{aligned} \quad (17)$$

f) En utilisant (4.17), le deuxième terme dans le second membre de (17) est égal à

$$\epsilon_{k+1}^2 \langle q_k - r_{k+1}, p_{k+1} - p^* \rangle \geq \epsilon_{k+1}^2 \langle q_k, p_{k+1} - p^* \rangle.$$

car

$$\langle r_{k+1}, p^* - p_{k+1} \rangle \geq \phi(p^*) - \phi(p_{k+1}) \geq 0.$$

De plus, de (B) ci-dessus, on déduit

$$\langle q_k, p_{k+1} - p^* \rangle \geq \frac{1}{\gamma \varepsilon_k} \langle p_{k+1} - p^*, p_{k+1} - p_k \rangle \geq \frac{1}{2\gamma \varepsilon_k} (\|p_{k+1} - p^*\|^2 - \|p_k - p^*\|^2).$$

2.4. En utilisant les résultats de l'alinéa f) dans (17), on obtient

$$\begin{aligned} \Phi_{k+1} - \Phi_k &\geq -\frac{\gamma}{2} N^2 (5\varepsilon_{k+1}^2 + 4(\varepsilon_k - \varepsilon_{k+1}) + (\frac{\varepsilon_k - \varepsilon_{k+1}}{\varepsilon_{k+1}})^2) \\ &\quad + \frac{\varepsilon_{k+1}^2}{2\gamma \varepsilon_k} (\|p_{k+1} - p^*\|^2 - \|p_k - p^*\|^2) + \\ &\quad \langle q_{k+1}, p_{k+1} - p_{k+2} \rangle - \langle q_k, p_k - p_{k+1} \rangle + \varepsilon_{k+1} (\psi(p^*) - \psi(p_{k+1})). \end{aligned}$$

Soit, en sommant de $k = 0$ à $N-1$,

$$\begin{aligned} \Phi_N - \Phi_0 &\geq -\frac{\gamma}{2} N^2 \sum_{k=1}^N (5\varepsilon_k^2 + (\frac{\varepsilon_{k-1} - \varepsilon_k}{\varepsilon_k})^2) - 2\gamma N^2 (\varepsilon_0 - \varepsilon_N) \\ &\quad + \frac{1}{2\gamma} [\frac{\varepsilon_N^2}{\varepsilon_{N-1}} \|p_N - p^*\|^2 - \frac{\varepsilon_1^2}{\varepsilon_0} \|p_0 - p^*\|^2 + \sum_{k=1}^{N-1} (\frac{\varepsilon_k^2}{\varepsilon_{k-1}} - \frac{\varepsilon_{k+1}^2}{\varepsilon_k}) \|p_k - p^*\|^2] \\ &\quad + \langle q_0, p_1 - p_0 \rangle - \langle q_N, p_{N+1} - p_N \rangle \\ &\quad + \sum_{k=1}^N \varepsilon_k (\psi(p^*) - \psi(p_k)). \end{aligned} \quad (18)$$

On a

$$-\langle q_N, p_{N+1} - p_N \rangle \geq -\|q_N\| \|p_{N+1} - p_N\| \geq -\gamma \varepsilon_N \|q_N\|^2 \geq -\gamma M^2 \varepsilon_N$$

(où on a utilisé $\|p_{N+1} - p_N\| \leq \gamma \varepsilon_N \|q_N\|$ comme conséquence de (4.16)).

On rappelle les hypothèses suivantes (H1). La suite $\{\frac{\varepsilon_k^2}{\varepsilon_{k-1}}\}$ est décroissante.

(H2). La série de terme général $(\frac{\varepsilon_{k-1} - \varepsilon_k}{\varepsilon_k})^2$ est convergente.

Remarque. Ces deux hypothèses sont toujours vérifiées pour des suites $\{\varepsilon_k\}$ tel que $\varepsilon_k \sim k^{-\alpha}$ avec $1/2 < \alpha \leq 1$ (suites vérifiant les hypothèses du Théorème 1) mais nous n'avons pu vérifier que (H1) et (H2) sont toujours une conséquence des hypothèses du Théorème 1.

Avec (H1), (H2) et les hypothèses du Théorème 1, on déduit de (17) que

$$\begin{aligned} \Phi_N - \Phi_0 &\geq -\Lambda_N - 2\gamma M^2 \epsilon_0 - \frac{\epsilon_1^2}{2\gamma \epsilon_0} \|p_0 - p^*\|^2 + \langle q_0, p_1 - p_0 \rangle \\ &\quad - \gamma M^2 \epsilon_N + \sum_{k=1}^N \epsilon_k (\phi(p^*) - \phi(p_k)) \end{aligned} \quad (19)$$

où Λ_N est la somme jusqu'à N d'une série convergente.

2.5. De (19) on tire deux conclusions

a) En notant que le dernier terme de (19) est non-négatif, on conclue que le second membre de (19) est supérieur à une certaine constante a , d'où

$$a \leq \Phi_N - \Phi_0 \leq \phi(p_N) - \Phi_0$$

d'après (10) ci-dessus. Si P n'est pas borné, et si l'on suppose que la suite $\{p_k\}$ n'est pas bornée, on obtient une contradiction car on aurait, d'après une hypothèse de théorème 1, que

$$\liminf_{N \rightarrow +\infty} \phi(p_N) = -\infty > a.$$

b) De (19) on conclue aussi que

$$\sum_{k=1}^{+\infty} \epsilon_k (\phi(p^*) - \phi(p_k)) < +\infty. \quad (20)$$

Grâce à (20) on montre que

$$\lim_{k \rightarrow +\infty} \phi(p_k) = \phi(p^*). \quad (21)$$

2.6. On va montrer que

$$\forall \alpha > 0, \exists K(\alpha) : \forall k \geq K(\alpha), 0 \leq \phi(p^*) - \phi(p_k) \leq \alpha \quad (22)$$

ce qui est équivalent à (21).

Soit $N_1(\alpha) = \{k \in \mathbb{N} \mid \phi(p^*) - \phi(p_k) > \alpha/2\}$.

Alors

$$\sum_{k \in N_1(\alpha)} \epsilon_k < +\infty \quad (23)$$

En effet

$$\frac{\alpha}{2} \sum_{k \in N_1(\alpha)} \epsilon_k < \sum_{k \in N_1(\alpha)} \epsilon_k (\phi(p^*) - \phi(p_k)) < \sum_{k \in N} \epsilon_k (\phi(p^*) - \phi(p_k)) < +\infty$$

d'après (20).

Par ailleurs, comme la suite $\{p_k\}$ est bornée et comme ϕ est concave s.c.s, ϕ est localement lipschitzienne et on désigne par L la constante de Lipschitz sur un ensemble contenant la suite $\{p_k\}$. D'après (23)

$$\exists K(\alpha) : \forall k > K(\alpha), \sum_{1 \geq k; 1 \in N_1(\alpha)} \epsilon_1 < \frac{\alpha}{2\gamma LM} \quad (24)$$

Alors $\forall k > K(\alpha)$, ou bien $k \notin N_1(\alpha)$ et (22) est assuré, ou bien $k \in N_1(\alpha)$ et l'on définit

$$k_0 = \min\{1 | 1 \notin N_1(\alpha) \text{ et } 1 > k\}$$

Remarquons que k_0 existe toujours car le complément de $N_1(\alpha)$ dans N est nécessairement infini d'après (23) (la série de terme général ϵ_k est divergente).

Alors

$$\begin{aligned} \phi(p^*) - \phi(p_k) &\leq \phi(p^*) - \phi(p_{k_0}) + |\phi(p_{k_0}) - \phi(p_k)| \\ &\leq \frac{\alpha}{2} + \sum_{l=k}^{k_0-1} L \|p_{l+1} - p_l\| \quad (\phi \text{ Lipschitz et } k_0 \notin N_1(\alpha)) \\ &\leq \frac{\alpha}{2} + L \sum_{l=k}^{k_0-1} \gamma \epsilon_1 \|q_1\| \quad (\text{d'après (4.16)}) \\ &\leq \frac{\alpha}{2} + L \gamma M \frac{\alpha}{2\gamma LM} = \alpha \quad (\text{d'après (24) et } \|q_1\| \leq M) \end{aligned}$$

ce qui démontre (22) (donc (21)).

2.7. a) La suite $\{p_k\}$ étant bornée est compacte dans la topologie faible. Pour tout point d'accumulation $\bar{p}$ et toute sous-suite $\{p_{k_i}\}$ convergente vers $\bar{p}$, on a avec la semi-continuité supérieure de ψ et (21) ci-dessus

$$\psi(p^*) = \limsup \psi(p_{k_i}) \leq \psi(\bar{p})$$

ce qui montre que $\bar{p} \in P^*$.

b) Si ψ est fortement concave

$$\exists a > 0, \forall p : \psi(p^*) - \psi(p) \geq a \|p - p^*\|^2$$

et $p^* \in P^*$ est unique. On en déduit immédiatement que $p_k \rightarrow p^*$ dans la topologie forte.

2.8. Il reste à montrer un résultat de convergence sur la suite $\{q_k\}$. Pour cela on étudie d'abord la quantité.

$$\Delta_k = \psi(p_k) - \frac{1}{2\gamma} \|(p_{k+1} - p_k) / \epsilon_k\|^2 \leq \psi(p^*). \quad (25)$$

$$\text{et on montre que } \lim_{k \rightarrow +\infty} \|(p_{k+1} - p_k) / \epsilon_k\| = 0. \quad (26)$$

a) En considérant les inégalités (11) et (15) ci-dessus, on établit que

$$\begin{aligned} \Delta_{k+1} - \Delta_k &\geq -\frac{M}{2\epsilon^2} (\epsilon_k - \epsilon_{k+1} + 2\epsilon_{k+1}^2)^2 + \langle q_k, p_{k+1} - p_k \rangle \\ &\geq -\frac{M^2}{2} \left(\left(\frac{\epsilon_k - \epsilon_{k+1}}{\epsilon_{k+1}} \right)^2 + 4\epsilon_{k+1}^2 + 4(\epsilon_k - \epsilon_{k+1}) \right) \\ &\quad + \frac{\epsilon_k}{\gamma} \|(p_{k+1} - p_k) / \epsilon_k\|^2 \end{aligned} \quad (27)$$

en utilisant l'inégalité (14).

En sommant ces inégalités de $k = 0$ à $N-1$, on obtient

$$\sum_{k=0}^{N-1} \varepsilon_k \|(p_{k+1} - p_k) / \varepsilon_k\|^2 \leq \gamma(\Delta_N - \Delta_0) + B_N \quad (28)$$

où B_N est la somme jusqu'à N d'une série convergente (d'après les hypothèses (H1) et (H2) et celles du théorème 1). Compte-tenu de (25) ($\Delta_N \leq \psi(p^*)$), (28) montre qu'au moins une sous-suite de la suite $\{\|(p_{k+1} - p_k) / \varepsilon_k\|^2\}$ tend vers zéro. En fait on va montrer que toute la suite tend vers zéro par un raisonnement analogue à celui du §2.6.

b) En effet, en reconsidérant les inégalités (27) et en sommant de $k = 1$ à $m-1$, on obtient

$$\Delta_m - \Delta_1 \geq C_{1,m}$$

où $C_{1,m}$ est la somme des termes 1 à $m-1$ d'une série convergente. L'inégalité ci-dessus se réécrit

$$\frac{1}{2\gamma} \|(p_m - p_{m-1}) / \varepsilon_{m-1}\|^2 \leq \frac{1}{2\gamma} \|(p_{1+1} - p_1) / \varepsilon_1\|^2 + \psi(p_{m-1}) - \psi(p_1) - C_{1,m} \quad (29)$$

Mais, d'après (21),

$$\forall \alpha > 0, \exists K_1(\alpha) : \forall l, m \geq K_1(\alpha), |\psi(p_{m-1}) - \psi(p_l)| < \frac{\alpha}{3} \quad (30)$$

$$\exists K_2(\alpha) : \forall l, m \geq K_2(\alpha), |C_{1,m}| < \frac{\alpha}{3} \quad (31)$$

Enfin, du fait que d'après (28), au moins une sous-suite $\{\|(p_{k_1+1} - p_{k_1}) / \varepsilon_{k_1}\|^2\}$ tend vers zéro,

$$\exists K_3(\alpha) : \forall k_1 \geq K_3(\alpha), \frac{1}{2\gamma} \|(p_{k_1+1} - p_{k_1}) / \varepsilon_{k_1}\|^2 < \frac{\alpha}{3} \quad (32)$$

(en prenant la précaution de choisir $K_3(\alpha)$ comme indice de la sous-suite $\{k_1\}$). Alors en réunissant (29) à (32), et en posant

$K(\alpha) = \max(K_1(\alpha), K_2(\alpha), K_3(\alpha))$ on voit que

$$\forall \alpha, \exists K(\alpha) : \forall n \geq K(\alpha), \frac{1}{2\gamma} \|(p_n - p_{n-1}) / \varepsilon_{n-1}\|^2 < \alpha$$

ce qui démontre (26).

2.9. On montre maintenant que, pour tout $p^* \in P^*$,

$$\lim_{k \rightarrow +\infty} \langle q_k, p_{k+1} - p^* \rangle = 0. \quad (33)$$

a) De (8), on tire

$$\langle q_k, p_{k+1} - p^* \rangle \geq \langle p_k - p_{k+1} \rangle / \varepsilon_k \langle p^* - p_{k+1} \rangle.$$

Comme la suite $\{p_k\}$ est bornée et avec (26), on a que

$$\liminf \langle q_k, p_{k+1} - p^* \rangle \geq 0. \quad (34)$$

b) Par ailleurs, de (16) on déduit que

$$t_{k+1} \leq (1 - \varepsilon_{k+1}) t_k + \varepsilon_{k+1} \langle q_{k+1}, (p_{k+2} - p_{k+1}) / \varepsilon_{k+1} \rangle$$

en ayant posé $t_k = \langle q_k, p_{k+1} - p^* \rangle$. Remarquant que $\lim \langle q_{k+1}, (p_{k+2} - p_{k+1}) / \varepsilon_{k+1} \rangle = 0$ et utilisant le lemme ci-dessous (§2.13) on en déduit que $\limsup t_k \leq 0$ ce qui avec (34) établit (33).

2.10. Comme la suite $\{q_k\}$ est bornée, soit $\bar{q}$ un point d'adhérence dans la topologie faible et $\{q_{k_i}\}$ une sous-suite convergente vers $\bar{q}$. D'après (8)

$$\forall p \in P, \langle q_{k_i}, p - p_{k_i+1} \rangle \leq \| (p_{k_i+1} - p_{k_i}) / \varepsilon_{k_i} \| \| p - p_{k_i+1} \|.$$

Le second membre tend vers zéro, et pour le premier, on a

$$\lim \langle q_{k_i}, p \rangle = \langle \bar{q}, p \rangle, \quad \forall p \quad (\text{convergence faible}),$$

$$\lim \langle q_{k_i}, p_{k_i+1} \rangle = \lim \langle q_{k_i}, p^* \rangle = \langle \bar{q}, p^* \rangle \quad (\text{d'après (33) et la convergence faible}).$$

d'où finalement par passage à la limite ($k_i \rightarrow +\infty$) dans l'inégalité ci-dessus, avec (26),

$$\forall \bar{q} \in \bar{Q}, \quad \langle \bar{q}, p^* \rangle \leq 0. \quad (35)$$

2.11. Un calcul analogue à celui du §2.3.d) (mais pour p quelconque au lieu de p^*) conduit à

$$t_{k+1} \leq (1 - \epsilon_{k+1}) t_k + \epsilon_{k+1} (\psi(p_{k+1}) - \psi(p) + \langle q_{k+1}, (p_{k+2} - p_{k+1}) / \epsilon_{k+1} \rangle)$$

où l'on a posé maintenant $t_k = \langle q_k, p_{k+1} - p \rangle$. Alors par le lemme ci-dessous (§2.13) on a $\limsup t_k \leq \psi(p^*) - \psi(p)$ et pour les indices de la sous-suite $\{k_i\}$ comme précédemment on obtient, $\lim t_{k_i} = \langle \bar{q}, p^* - p \rangle$ et donc

$$\forall p, \quad \langle \bar{q}, p - p^* \rangle \geq \psi(p) - \psi(p^*)$$

ce qui prouve que

$$\bar{q} \in \partial \psi(p^*). \quad (36)$$

La réunion de (35) et (36) montre que $\bar{q} \in Q^*$ ce qui termine la preuve du Théorème 1.

2.12. Remarque : $\lim_{k \rightarrow +\infty} \langle q_k, p_k \rangle = \langle q^*, p^* \rangle$ (37)

où la valeur du second membre de (37) est indépendante du couple (p^*, q^*) choisi dans $P^* \times Q^*$ d'après le Lemme 1. En effet, on a vu que pour toute sous-suite d'indices k_i tel que $\{q_{k_i}\}$ converge faiblement vers $\bar{q}$,

$$\bar{q} \in Q^* \quad \text{et} \quad \lim \langle q_{k_i}, p_{k_i+1} \rangle = \lim \langle q_{k_i}, p^* \rangle = \langle \bar{q}, p^* \rangle$$

d'après (33) et la convergence faible. Mais comme la valeur $\langle \bar{q}, p^* \rangle$ est indépendante des points $(p^*, \bar{q})$ dans $P^* \times Q^*$, c'est toute la suite $\langle q_k, p_{k+1} \rangle$ qui converge vers cette valeur.

D'autre part $\lim \langle q_k, p_{k+1} \rangle = \lim \langle q_k, p_k \rangle$ d'après (26).

2.13. Au cours de la démonstration on a utilisé à deux reprises le lemme suivant.

Lemme. Soit trois suites réelles $\{t_k\}$, $\{s_k\}$ et $\{\epsilon_k\}$ telles que

$$0 < \varepsilon_k < 1, \quad \sum_{k=1}^{+\infty} \varepsilon_k = +\infty \quad (38)$$

$$\lim_{k \rightarrow +\infty} s_k = s^* \quad (39)$$

$$t_{k+1} \leq (1 - \varepsilon_{k+1}) t_k + \varepsilon_{k+1} s_k, \quad l = 0, \dots, +\infty \quad (40)$$

$$\text{Alors} \quad \limsup t_k \leq s^*. \quad (41)$$

Preuve. Il n'y a pas de perte de généralité à supposer $s^* = 0$ (dans le cas général, on considère les quantités $t_k - s^*$ et $s_k - s^*$ pour lesquelles (40) est encore valable).

Notons que

$$\begin{aligned} \log(1 - \varepsilon_k) &< -\varepsilon_k \\ \text{donc} \quad \log \prod_{k=1}^{+\infty} (1 - \varepsilon_k) &< -\sum_{k=1}^{+\infty} \varepsilon_k = -\infty \\ \text{d'où} \quad \lim_{k \rightarrow +\infty} \prod_{k=1}^{+\infty} (1 - \varepsilon_k) &= 0 \end{aligned} \quad (42)$$

On a

$$\forall \alpha > 0 : \exists K_1(\alpha) : \forall k > K_1(\alpha) \quad , \quad \prod_{l=0}^{k-1} (1 - \varepsilon_{l+1}) < \frac{\alpha}{3|s_0|} \quad (43)$$

d'après (42). D'après (39)

$$\exists K_2(\alpha) : \forall k > K_2(\alpha) \quad , \quad |s_k| < \frac{\alpha}{3} \quad (\text{rappel : } s^* = 0). \quad (44)$$

Posons $K_3(\alpha) = \max(K_1(\alpha), K_2(\alpha))$ et

$$M(\alpha) = \sum_{l=0}^{K_3(\alpha)} \varepsilon_{l+1} |s_l| \prod_{m=l+1}^{K_3(\alpha)} (1 - \varepsilon_{m+1}) \quad (45)$$

(avec la convention : $\forall K, \prod_{m=K+1}^K (1 - \varepsilon_{m+1}) = 1$).

Alors d'après (42)

$$\exists K(\alpha) > K_3(\alpha) : \forall k > K(\alpha) \quad , \quad \prod_{m=K_3(\alpha)}^{k-1} (1 - \varepsilon_{m+1}) < \frac{\alpha}{3M(\alpha)}. \quad (46)$$

En utilisant (40) récursivement, on obtient, $\forall k > K(\alpha)$:

$$\begin{aligned}
t_k &\leq \sum_{l=0}^{k-1} \varepsilon_{l+1} s_l \prod_{m=l+1}^{k-1} (1-\varepsilon_{m+1}) + s_0 \prod_{l=0}^{k-1} (1-\varepsilon_{l+1}) \\
&= \left[\sum_{l=0}^{K_3(\alpha)} \varepsilon_{l+1} s_l \prod_{m=l+1}^{K_3(\alpha)} (1-\varepsilon_{m+1}) \right] \times \prod_{m=K_3(\alpha)}^{k-1} (1-\varepsilon_{m+1}) \\
&\quad + \sum_{l=K_3(\alpha)+1}^{k-1} \varepsilon_{l+1} s_l \prod_{m=l+1}^{k-1} (1-\varepsilon_{m+1}) + s_0 \prod_{l=0}^{k-1} (1-\varepsilon_{l+1}) \\
&\leq M(\alpha) \frac{\alpha}{3M(\alpha)} + \frac{\alpha}{2} \sum_{l=K_3(\alpha)+1}^{k-1} \varepsilon_{l+1} \prod_{m=l+1}^{k-1} (1-\varepsilon_{m+1}) + |s_0| \frac{\alpha}{3|s_0|}
\end{aligned} \tag{47}$$

en passant aux valeurs absolues et en utilisant (43) à (46). Enfin en remarquant que

$$\sum_{l=K_3(\alpha)+1}^{k-1} \varepsilon_{l+1} \prod_{m=l+1}^{k-1} (1-\varepsilon_{m+1}) \leq \sum_{l=0}^{k-1} \varepsilon_{l+1} \prod_{m=l+1}^{k-1} (1-\varepsilon_{m+1}) + \prod_{l=0}^{k-1} (1-\varepsilon_{l+1}) = 1$$

(considérer $t_k = s_k = 1$, $\forall k$, et l'égalité dans (40)), on obtient par (47)

$$\forall \alpha > 0, \quad \exists K(\alpha), \quad \forall k > K(\alpha), \quad t_k \leq \alpha,$$

ce qui démontre le lemme.

Corollaire : Si on a l'égalité dans (40), alors $\lim t_k = s^*$.

3. PREUVE du THEOREME 3

$$3.1. \text{ Soit } t_{k+1} = (1-\varepsilon_{k+1}) t_k + \varepsilon_{k+1} \langle p_{k+1}, \theta(u_{k+1}) \rangle, \quad k = 0, 1, \dots \tag{48}$$

partant de $t_0 = \langle p_0, \theta(u_0) \rangle$. Alors

$$\lim_{k \rightarrow +\infty} t_k = \langle p^*, \theta(u^*) \rangle \tag{49}$$

où (u^*, p^*) est un point-selle quelconque (rappelons que la valeur du second membre de (49) est indépendante du point-selle choisi).

On a en effet

$$t_{k+1} - \langle q_{k+1}, p_{k+1} \rangle = (1-\varepsilon_{k+1}) (t_k - \langle q_k, p_k \rangle) + (1-\varepsilon_{k+1}) \varepsilon_{k+1} \langle q_k, (p_k - p_{k+1}) \rangle / \varepsilon_{k+1}$$

en se servant de (4.20). Comme, d'après (26),

$$\lim (1-\varepsilon_{k+1}) \langle q_k, (p_k - p_{k+1}) / \varepsilon_{k+1} \rangle = 0$$

on a d'après le corollaire ci-dessus que

$$\lim (t_k - \langle q_k, p_k \rangle) = 0, \text{ ce qui établit (49) grâce à (37).}$$

3.2. Avec {4.23} et la convexité de J

$$J(v_{k+1}) \leq (1-\varepsilon_{k+1})J(v_k) + \varepsilon_{k+1}J(u_{k+1}).$$

En utilisant (48), il vient

$$J(v_{k+1}) + t_{k+1} \leq (1-\varepsilon_{k+1})(J(v_k) + t_k) + \varepsilon_{k+1}L(u_{k+1}, p_{k+1}). \quad (50)$$

Par ailleurs puisque $u_{k+1} \in \hat{U}(p_{k+1})$ et par l'une des inégalités du point-selle

$$L(u_{k+1}, p_{k+1}) \leq L(u^*, p_{k+1}) \leq L(u^*, p^*)$$

d'où avec (50)

$$J(v_{k+1}) + t_{k+1} \leq (1-\varepsilon_{k+1})(J(v_k) + t_k) + \varepsilon_{k+1}L(u^*, p^*)$$

d'où, par le lemme du §2.13

$$\limsup (J(v_k) + t_k) \leq L(u^*, p^*),$$

et, avec (49),

$$\limsup J(v_k) \leq J(u^*). \quad (51)$$

Par la coercivité de J, ceci montre que la suite $\{v_k\}$ est bornée. Soit $\bar{v}$ un point d'adhérence faible et $\{v_{k_1}\}$ une sous-suite convergente vers $\bar{v}$.

Alors, puisque J est s.c. i (même dans la topologie faible par convexité) et par (51)

$$J(\bar{v}) \leq \liminf J(v_{k_i}) \leq \limsup J(v_k) \leq J(u^*). \quad (52)$$

3.3. On a également par convexité de l'application $u \mapsto \langle p^*, \theta(u) \rangle$

$$\langle p^*, \theta(v_{k+1}) \rangle \leq (1 - \varepsilon_{k+1}) \langle p^*, \theta(v_k) \rangle + \varepsilon_{k+1} \langle p^*, \theta(u_{k+1}) \rangle$$

mais avec (4.20).

$$\langle p^*, \theta(v_{k+1}) - q_{k+1} \rangle \leq (1 - \varepsilon_{k+1}) \langle p^*, \theta(v_k) - q_k \rangle$$

$$\text{et donc } \limsup \langle p^*, \theta(v_k) - q_k \rangle \leq 0. \quad (52 \text{ bis})$$

Comme les points d'adhérence faible de $\{q_k\}$ sont des $q^* \in Q^*$ et comme $\langle p^*, q^* \rangle$ est constant sur $P^* \times Q^*$, comme par ailleurs il est évident que pour tout point-selle (u^*, p^*) , on a $\theta(u^*) \in Q^*$, (52 bis) implique

$$\limsup \langle p^*, \theta(v_k) \rangle \leq \langle p^*, \theta(u^*) \rangle$$

et en considérant la sous-suite $\{v_{k_i}\}$ du §3.2 ci-dessus et avec la semi-continuité inférieure de l'application $u \mapsto \langle p^*, \theta(u) \rangle$

$$\begin{aligned} \langle p^*, \theta(\bar{v}) \rangle &\leq \liminf \langle p^*, \theta(v_{k_i}) \rangle \leq \limsup \langle p^*, \theta(v_k) \rangle \\ &\leq \langle p^*, \theta(u^*) \rangle. \end{aligned} \quad (53)$$

3.4. Alors en additionnant (52) et (53)

$$L(\bar{v}, p^*) \leq L(u^*, p^*) \quad (54)$$

mais seule l'égalité est possible car $u^* \in \hat{U}(p^*)$.

Mais ceci implique que l'on a simultanément les égalités

$$J(\bar{v}) = J(u^*) \quad \text{et} \quad \langle p^*, \theta(\bar{v}) \rangle = \langle p^*, \theta(u^*) \rangle \quad (55)$$

car si l'on suppose par exemple l'inégalité stricte dans (52), cela implique à cause de l'égalité dans (54) une inégalité stricte en sens inverse dans (53) ce qui est une contradiction.

Avec (55), il est trivial de vérifier les deux inégalités du point-selle pour le couple $(\bar{v}, p^*)$ ce qui achève la démonstration.

3.5. Remarque : De façon analogue à (48)-(49), on peut montrer facilement en utilisant les résultats précédents que

$$\lim t_k = J(u^*)$$

pour t_k défini récursivement par

$$t_0 = J(u_0), \quad t_{k+1} = (1 - \epsilon_{k+1})t_k + \epsilon_{k+1} J(u_{k+1})$$

Note sur l'utilisation du programme

I- Lexique des variables du programme

1-Variables de transit et de capacité de transport.

IX(260) : transit sur chaque arc dans Flomax.

IBP (260) : Capacité maximale de chaque arc.

IEC(82) : Capacité existante sur chaque arc du réseau ($q_t^{V_0}$)

IXX(82x501) : Tableau des transits circulant alébroiquement sur chaque arc pour tous les aléas.

IX2(48x501) : Tableau des capacités des arcs de production pour chaque aléa.

IX3(48x501) : Tableau des capacités des arcs de défaillance pour chaque aléa.

IQX : Valeur absolue du transit sur un arc pour 1 aléa $|q_t^k(\omega)|$

AA(82) : $\sum_{\omega} p_t^k(\omega) |q_t^k(\omega)|$ pour chaque arc.

BX(82x501) : Tableau des $z_t^k(\omega)$ pour tous les arcs t et les aléas ω .

AB(82) : $\sum_{\omega} p_t^k(\omega) z_t^k(\omega)$ pour chaque arc.

INV(82) : tableau suivant à stocker $y_t^k(\omega) = \max_{\omega} z_t^k(\omega)$ ainsi que $y_t^k - q_t^{V_0}$ pour tous les arcs.

JP(48) : Production fixe au noeuds MIP(185) : Puissance maximale de chaque centrale.

JD(48) : Consommation aux noeuds.

MIPTO : Production totale dissipée sur le réseau pour une situation aléatoire.

ICSOM : Consommation totale réelle du réseau.

JBK (48x501) : Tableau servant à conserver les capacités maximales des arcs de production IBP, pour tous les aléas - JBK est issu de ALEAT.

2. Variables de Coût.

ICT (82) Coût unitaire d'investissement pour tous les arcs (c_t).

IC (260) Coût pour chaque arc et pour un aléa intervenant dans Flomax.

IC comprend les coûts de production $c_p(0)$ et les coûts de défaillance (24 000 ou 2400) c_d .

CF : Coût de défaillance.

AG : Coût dual.
COT : Coût primal.

3. Autres variables.

IPA (185) : Taux de panne des centrales.

ING1 (182) : sommets de départ des arcs de transport.

ING2 (182) : sommets d'arrivée des arcs de transport.

ING1(165 - 349) : Sommets reliés aux groupes de production.

B(82x501) : Tableau contenant les paramètres duaux ($\pi_t^k(\omega)$ et $p_t^k(\omega)$) pour tous les arcs et tous les aléas .

V(82) : $\sum_{\omega} \pi_t^k(\omega)$ pour tous les arcs.

ROU(82) : tableau servant à stocker le pas pour chaque arc t.

R , P , T , : variables relatives au pas ρ^k de l'algorithme.

PI : Pas ϵ^k de l'algorithme.

RI : Terme de série divergente pouvant servir au calcul de ϵ_t^k ou ρ_t^k

NAL : Nb d'aléas.

NAP : Nb d'aléas y compris l'aléa fictif = NAL + 1 .

NL : nb d'arcs (82).

NN : nb de noeuds (48).

NG : nb de groupes de production (185).

IPAS : N^0 de l'itération.

IPAMAX : Numéro de la dernière itération désirée.

IA, IB, Z : variables d'Aléat.

M, N, NS, NC, NG1, NG2, IBAS, NUM, IU, ISOL, JC00L . Variables de FLOMAX.

MUM, NNC : Variables relatives à Flomax. (cf. II).

II - Quelques précisions sur l'utilisation du programme.

1 - Lecture des données.

Les valeurs continues dans ICT sont multipliées par 1000.

JD - JP représente la demande réelle en chaque noeud.

2 - Initialisation des tableaux

* B(82x501) est initialisé à 0 sur tous les arcs pour les 500 aléas et à 1 pour tous les arcs dans la 501^e situation aléatoire, fictive.

* NG1(260) et NG2 (260) sont des tableaux entrant dans FLOMAX et relatifs au graphe :

- pour un arc I_t de transport, $NG1(I_t)$ représente le sommet de départ, $NG2(I_t)$ représente le sommet d'arrivée.
 Dans FLOMAX, les transits doivent être positifs.

Chaque arc est donc dédoublé en deux arcs de sens contraire, I_t et

$$I_t + NL : \begin{cases} NG1(I_t) = NG2(I_t + NL) = ING1(I_t) \\ NG2(I_t) = NG1(I_t + NL) = ING2(I_t) \end{cases}$$

NB $0 \leq I_t \leq 164$

- Pour un arc I_p de production ($165 \leq I_p \leq 212$) qui n'est pas un arc vrai, $NG1(I_p) = I_p - 164$ représente le sommet, l'arrivée de cet arc qui est dit arc rentrant. $NG2(I_p) = 0$

- Pour un arc de défaillance I_d ($213 \leq I_d \leq 260$) on a aussi $NG2(I_d) = 0$ car I_d n'est pas un arc vrai.

$NG1(I_d) = I_d - 212$ si l'arc est rentrant (défaillance réelle).

$NG1(I_d) = 212 - I_d$ si l'arc est sortant (défaillance négative).

NB Le tableau $ING1(165 - 349)$ contenant les numéros des sommets reliés aux groupes de production n'apparaît plus dans la suite du programme car on compile tous les groupes sur les arcs de production rentrant sur les sommets du graphe, et à jusqu'à satisfaction de la demande.

* Le tableau $IBP(260)$ contient les bornes supérieures des transits pour tous les arcs. (tableau nécessaire à FLOMAX).

- pour un arc de transport ou de défaillance il n'y a pas réellement de limitation. On prend donc un nombre très grand : 150 000.

- pour un arc de production on calcule la somme des puissances maximales des centrales empilées sur l'arc considéré.

3. Première partie de l'itération.

Elle est constituée par la résolution des 500 problèmes de Flot (call FLOMAX) correspondant aux 500 situations aléatoires (Boucle 14).

* Le coût IC est obtenu par la répartition du coût unitaire par arc, ICT sur tous les aléas (Boucle 20).

La suite du programme, jusqu'à l'étiquette 722, initialise les variables de FLOMAX à la première itération.

* Boucle 17 : IA étant initialisé, chaque appel de ALEAT fournit une valeur uniforme entre 0 et 1, Z, à condition de reprendre IA = IB après chaque appel.

Chaque situation aléatoire correspondant à une valeur de L (cad de ω), est donc constitué par 185 valeurs Z aléatoirement tirées par ALEAT. Chaque Z ainsi obtenu est comparé au taux de panne IPA de la centrale correspondante. Si Z est inférieur au taux de panne, la centrale est en panne.

Chaque fois qu'un groupe est déclaré en service à l'issue de ce test, sa puissance maximale de production, MIP, s'ajoute à IBP sur l'arc de production correspondant. Quand, au cours de cette opération, MIPTO atteint on dépasse IC50M, on sort de la boucle.

Les 500 situations aléatoires étant déterminées une fois pour toutes, on conserve pour chaque L, les 48 valeurs de IBP ainsi obtenues dans le tableau JBK (48x500), ce qui évite de refaire le tirage à chaque itération.

* Vient ensuite l'initialisation des variables de transit IX. Au départ, les arcs de transport et de production sont initialisés à 0 ($IX(1 - 212)$) IBAS = -1 car IX est en butée inférieure sur les contraintes. Pour les arcs de base (IBAS=0), IX (213-260) est initialisé au signe près à la valeur du second membres des contraintes : $|JD-JP|$.

En effet FLOMAX impose IX positif ou nul.

On a donc : si $JD-JP > 0$ IX = JD - JP.

et $NG(I) = I-212$ Arc de défaillance rentrant.

si $JD - JP < 0$ IX = JP-JD

et $NG(I) = 212 - I$ Arc de défaillance sortant.

NUM sert à numérotter les arcs de base. Au départ $NUM(I) = I-212$ puisque ici $213 \leq I \leq 260$.

* La partie suivante du programme jusqu'à l'appel de FLOMAX est constituée par la récupération des paramètres issus de FLOMAX à l'itération précédente : BBP, IX, NUM, IBAS, NC.

Cette récupération permet d'initialiser les variables de FLOMAX après la première itération de manière très performante et entraîne ainsi un gain de temps appréciable.

La partie suivant l'appel de FLOMAX est constituée bien sûr par le stockage de ces mêmes variables à l'aide des tableaux IXX, IX2, IX3, MUM,

Deux remarques :

- IBAS ne pouvant contenir que les valeurs suivantes :

- 1 (Butée inférieure), 0 (base) et +1 (butée supérieure),

il n'a pas été nécessaire de conserver IBAS sous la forme d'un tableau (260 x 500). On ajoute à chaque IX issu de FLOMAX, la quantité entière $10\,000 \times (1 + \text{IBAS})$. Deux tests simples permettant ensuite de récupérer la valeur exacte de IX ainsi que IBAS.

- Il est aisé de constater que les arcs dédoublés avant de passer dans FLOMAX, n'ont jamais tous les deux un transit non nul. Pour des raisons de coût minimum, le transit passe soit dans 1 sens soit dans l'autre. Ceci nous permet de ne garder en mémoire que la différence $\text{IXX}(I,L) = \text{IX}(I) - \text{IX}(I+2)$ pour chaque arc.

Le signe de IXX permet ensuite de récupérer le sens dans lequel le transit est non nul.

4. La seconde partie de l'itération

Une fois les transits calculés pour les 500 situations aléatoires, on aborde l'algorithme de MINMAX proprement dit.

Les équations (2) et (3) ont indiquées ainsi que les phases du calcul des coûts dual (AA,AG) et primal (BX,INV et COT).

Enfin les équations (6) et (7) constituent le deuxième niveau au cours duquel sont remis à jour les paramètres duaux B (82x500).

```

      REAL*4 IC,IX,IU,IBP,IOX,IXX,IX2,IX3,INV
      COMMON /FLO/PI,N,NS,NC,NG1(510),NG2(510),IC(510),IX(510),IBAS(510),
      MUH(100),IIP(510),IU(510),ISOL(510),JCOL(510)
      DIMENSION AA(43),AB(43),ROI(43)
      COMMON IXX(43,502),IX2(49,502),IX3(49,502),JOK(49,502)
      COMMON RX(43,502)
      COMMON B(43,502),MUM(49,502)
      DIMENSION MIP(200),IEC(200),INV(100)
      DIMENSION ING1(500),ING2(500),ICT(100),V(100)
      INTEGER BLANC,CROI,ETOIL
      DIMENSION HNC(501)
      DIMENSION JD(100),IPA(200),JP(100)
      DATA BLANC/1H /,CROI/1H*/ETOIL/1H*/
      LEC=105
      IMP=108
      IGR=100000
      IPAS=1
      AG=0.
3307 F-IRMAT(1X,"TILT A L'ITERATION",IX,I3)
      L=500
      F-AL=FLOAT(NAL)
      P=185
      =48
      =82
      N.=NL+1
      NS=2*NL
      I-AMAX=68
      TPI=1.
      PI=1.
      P=5.E-4
C
C
C      LECTURES DES DONNEES
C
C      *****
1 F-IRMAT(26I3)
6 F-IRMAT(13I6)
9 F-IRMAT(16I5)
10 F-IRMAT(20I4)
11 F-IRMAT(26I3)
      READ(LEC,1)(ING1(I),I=1,NL)
      READ(LEC,1)(ING2(I),I=1,NL)
      READ(LEC,6)(ICT(I),I=1,NL)
      READ(LEC,1)(ING1(I),I=165,349)
      READ(LEC,9)(IEC(I),I=1,NL)
      READ(LEC,10)(MIP(I),I=1,NG)
      READ(LEC,11)(IPA(I),I=1,NG)
      READ(LEC,9)(JP(I),I=1,NN)
      READ(LEC,9)(JD(I),I=1,NN)
      WRITE(IMP,9)(JD(I),I=1,NN)
332 F-IRMAT(1X,14I5)
      WRITE(IMP,3307) IPAS
98A8 F-IRMAT(1X,"PI=",E16.8)
C
C
C      INITIALISATION DES TAPLFAHX
C
C      *****
      IG1=NS+1
      IG2=NS+NN
      ID=IG2+1
      IT=NS+2*NN
1A10 F-IRMAT(1X,I2,F16.8,2I0)

```

```

NAF=NAL+1
FNAF=FLOAT(NAF)
R(1,1)=0.
DO 701 L=1,NAL
DO 701 I=1,NL
R(I,L)=B(1,1)
701 CONTINUE
DO 700 I=1,NL
AR(I)=0.
BX(I,NAF)=FLOAT(TEC(I))
IXX(I,MAF)=FLOAT(IEC(I))
B(I,NAF)=1.
700 CONTINUE
DO 8 I=IG1,IT
NG2(I)=0
8 CONTINUE
DO 18 I=1,NL
NG1(I)=ING1(I)
NG2(I)=ING2(I)
JK=I+NL
NG1(JK)=ING2(I)
NG2(JK)=ING1(I)
IBP(I)=150000
IBP(JK)=150000
18 CONTINUE
ICSOM=0
DO 118 I=1,NN
ICSOM=ICSOM+JD(I)-JP(I)
K=I+NS
IC(K)=0
NG1(K)=I
KK=I+IG2
IBP(KK)=150000
118 CONTINUE

```

INITIALISATION D'ALEAT

IA=99

DEBUT DE L'ITERATION

```

CONTINUE
FORMAT(1X,'R0=',E16.8)
CF=0.
DO 14 L=1,NAL
M=NN
N=II

```

CALCUL DES COUTS D'INVESTISSEMENT SUR CHAQUE ARC DE TRANSPORT

```

DO 20 I=1,NL
LI=I+NL
TC(I)=FLOAT(ICT(I))*B(I,L)
IC(LI)=IC(I)
905 CONTINUE
20 IF(IPAS-1) 21,21,722
21 CONTINUE

```



```

      NG1(I)=IG2-I
      GO 10 721
422  NG1(I)=I-IG2
721  CONTINUE
      NC=0
      GO 10 7274
722  CONTINUE
C
C
C      *****
C      REINITIALISATION DE FLOMAX A L'ITERATION K(K SUP A 2)
C      A PARTIR DE L'ITERATION K=1 POUR CHAQUE ALEA
C      *****
C
      NC=NC(L)
      DO 1723 I=1,M
      NI=I+NS
      IBP(NI)=JBK(I,L)
      IX(NI)=IX2(I,L)
      KI=I+IG2
      IX(KI)=IX3(I,L)
1723  CONTINUE
      DO 723 I=1,NL
      LI=I+NL
      IF (IXX(I,L)) 7221,7222,7222
7222  IX(I)=IXX(I,L)
      IX(LI)=0.
      GO 10 723
7221  IX(LI)=-IXX(I,L)
      IX(I)=0.
      723  CONTINUE
      DO 888 K=1,IT
      IF (IFIX(IX(K)+1.)-IGR) 7201,7202,7202
7201  IBAS(K)=-1
      GO 10 888
7202  IF (IFIX(IX(K)+1.)-2*IGR) 7203,7204,7204
7203  IBAS(K)=0
      IX(K)=IX(K)-IGR
      GO 10 888
7204  IBAS(K)=1
      IX(K)=IX(K)-2*IGR
      888  CONTINUE
      DO 47 K=1,M
      NIJM(K)=MIJM(K,L)
      47  CONTINUE
7274  CONTINUE
C
C      *****
C      **** APPEL DE FLOMAX ****
C      *****
C
      CALL FLOMAX
C
C      COUT DE PRODUCTION ET DE DEFAILLANCE PAR ALEA
C
      DO 29 I=10,IT
      CF=CF+IX(I)*IC(I)
      29  CONTINUE
C

```



```

C      *****
C
C      EQUATION (2)
C
C      *****
C
C      BX(I,LZ)=(1.-PI)*BX(I,LZ)+PI*IQX
C
C      AA(I)=AA(I)+B(I,LZ)*IQX
C      INV(I)=AMAX1(INV(I),BX(I,LZ))
331  CONTINUE
C      DO 333 I=1,NL
C      AA(I)=AA(I)+B(I,NAF)*FLOAT(IEC(I))
C      AG=AG+(AA(I)-FLOAT(IEC(I)))*FLOAT(ICT(I))*1.E-3
C
C      *****
C
C      EQUATION (3)
C
C      *****
C
C      AB(I)=(1.-PI)*AB(I)+PI*AA(I)
C
C      333  CONTINUE
C      338  FORMAT(1X,'TRANSITS')
C      335  FORMAT(1X,8F14.2)
C
C      COUT MOYEN DE PRODUCTION ET DE DEFAILLANCE
C
C      CF=CF*1.E-3
C      1982  WRITE(IMP,1982) CF
C      FORMAT(1X,'CF=',E16.8)
C      AG=AG+CF
C      1979  FORMAT(1X,'AG=',E16.8)
C      WRITE(IMP,1979) AG
C      COT=0.
C      DO 810 I=1,NL
C      V(I)=0.
C      INV(I)=AMAX1(INV(I)-IEC(I),0.)
C      COT=COT+INV(I)*FLOAT(ICT(I))*1.E-3
C      INV(I)=INV(I)+IEC(I)
C      810  CONTINUE
C      COT=COT+CF
C      9980  FORMAT(1X,'COT=',E16.8)
C      WRITE(IMP,9980) COT
C      WRITE(IMP,338)
C      WRITE(IMP,335) (AA(I),I=1,NL)
C      WRITE(IMP,335) (INV(I),I=1,NL)
C
C      *****
C
C      CALCUL DE LA NORME DE LA DIRECTION DE MONTEE
C
C      *****
C
C      BR=0.
C      DO 8001 LZ=1,NAF
C      DO 8001 I=1,NL
C      BR=BR+(FLOAT(ICT(I))*(IX(I,LZ)-AB(I)))*2

```

```

8001 CONTINUE
RB=SQRT(RR)*1.E-3

C
C *****
C
C AJUSTEMENT DES PAS DE L'ALGORITHME
C
C *****
C
PPS=ABS(AG-0.28E+07)
RI=1./(1.+FLOAT(IPAS))
TPI=TPI+RI
PI=0.05*RI*TPI**2
TO=PO*0.001*PPS/BH
T0=10*RI
RO=10*1.E+3
TO=10/BB
DO 8881 I=1,NL
ROU(I)=TO*FLOAT(ICT(I))
8881 CONTINUE

C
C *****
C
C REMISE A JOUR DES PARAMETRES DUAUX
C
C *****
C
DO 1330 LZ=1,NAF
DO 1330 I=1,NL
C *****
C
C EQUATION (6)
C
C *****
C
B(I,LZ)=B(I,LZ)+ROU(I)*(BX(I,LZ)-AB(I))
IF(B(I,LZ))137,136,136
137 B(I,LZ)=0.
136 CONTINUE
C
V(I)=V(I)+B(I,LZ)
1330 CONTINUE
DO 1369 L=1,NAF
DO 1369 I=1,NL
C *****
C
C EQUATION (7)
C
C *****
C
B(I,L)=B(I,L)/V(I)
1369 CONTINUE
WRITE(IMP,3307) IPAS
IF(IPAS-IPAMAX)801,800,800
801 CONTINUE
IPAS=IPAS+1
WRITE(IMP,45) RO
C *****
C
C FIN DE L'ITERATION
C
C *****

```

800 GO TO 60
CONTINUE
STOP
END

```

C
C
C
C      SOUS-PROGRAMME DE GENERATION
C      D'UNE SUITE DE NOMBRES PSEUDO-ALEATOIRES
C
C      SUBROUTINE  ALEAT(IA,IB,Z)
C
C      IA:QUAND ON APPELLE ALEAT POUR LA PREMIERE FOIS,IA DOIT
C      AVOIR UNE VALEUR ENTIERE,IMPAIRE ET INFERIEURE A 9
C      CHIFFRES; DANS LES APPELS SUIVANTS,ON REINITIALISE LA
C      VALEUR DE IA EN LUI DONNANT LA VALEUR IB GENEREE PAR
C      LE PROGRAMME
C      IB:VALEUR ENGENDREE PAR LE PROGRAMME
C      Z:VALEUR UNIFORME SUR (0,1) ENGENDREE PAR ALEAT
C
C      2**16+3=65539
C      IB=IA*65539
C      IF (IB,GE,0) GOTO 1
C      IB=IB+2147483647+1
C      2**31=2147483647+1
C      1  Z=IB
C      Z=Z*0.4656613E-9
C      2**(-31)=0.4656613E-9
C      RETURN
C      END

```

REFERENCES

- A. AUSLENDER : Optimisation, Méthodes Numériques (Masson, Paris, 1976).
- A. BAKUSHINSKII, B.T. POLYAK : On the solution of variational Inequalities. Soviet Math. Doklady, Vol.15, pp 1705-1710, 1974.
- J. CEA : Optimisation : Théorie et Algorithmes. (Dunod, Paris 1971).
- G. COHEN : An algorithm for convex constrained minimax optimization based on duality. (Applied Mathematics and Optimization, 1981).
- J.C. DODU, M. COURSAT, A. HERTZ, J.P. QUADRAT, M. VIOT : Méthodes de gradient stochastique pour l'optimisation des investissements dans un réseau électrique. EDF publication. Etude et Recherche. Série C, Mathématique, informatique n°2, (1981).
- L.S. LASDON : Optimization theory for Large System, The Macmillan Company New York, 1970.
- C. LEMARECHAL et R. MIFFLIN : Nonsmooth optimization. (proceedings of a IIASA Workshop, 1977).
- J.F. MAURRAS : Flot optimal sur un graphe avec multiplicateurs, EDF publication, Etudes et Recherches, Clamart (France). Série S, n°1, 1972.
- J. MEDANIC, M. ANDJELIC : Minimax Solution of the Multiplier Target Problem, IEEE Transactions on Automatic Control, 1972.
- E.A. NURMINSKI et P.I. VERCHENKO : Convergence of Algorithms for finding Saddle points, Cybernetics, Vol 13, n°3, 1977.
- B.T. POLYAK : Subgradient Methods. A survey of Soviet Research. in Nonsmooth Optimization par Lemaréchal et Mifflin.
- R.T. ROCKAFELLAR "A dual approach to Solving Non-linear programming. Problems by Unconstrained optimization". Mathematical Programming 5, pp.354-373 (1973).
- ZHU DAO LI : Optimisation sous-différentiable et méthodes de décomposition.