

# Boîte à outils “Gradient stochastique”

Pierre Carpentier (ENSTA — UMA)  
Guy Cohen (ENPC — CERMICS)  
Jennie Lioris (DEA MMME)

7 mai 2004

## Résumé

Le but de cette note est de décrire la mise en œuvre du gradient stochastique sur un exemple quadratique gaussien et de comparer les comportements :

- du gradient stochastique normal,
- du gradient stochastique moyenné,
- de l'estimateur optimal de la solution.

Cette comparaison se fait, d'une part sur les statistiques portant sur plusieurs exécutions des algorithmes, et d'autre part sur l'analyse des matrices de covariance des algorithmes et de la borne de Cramer-Rao.

## 1 Description du problème traité

On se limite ici à l'étude d'une fonction coût aléatoire quadratique, avec un bruit gaussien agissant de façon linéaire sur les variables de décision. Le critère à optimiser est l'espérance de la fonction coût :

$$\min_{x \in \mathbb{R}^p} \mathbb{E} \left\{ \frac{1}{2} x^\top B x + \xi(\omega)^\top x \right\}, \quad (1)$$

où :

- $B$  est une matrice symétrique définie positive d'ordre  $p$ ,
- $\xi$  est une variable aléatoire définie sur  $(\Omega, \mathcal{F}, \mathbb{P})$  à valeur dans  $\mathbb{R}^p$  qui suit une loi normale de moyenne  $\mu$  et de matrice de covariance  $Q$ .

On montre très facilement que ce système admet une solution unique que l'on note  $x^\sharp$  :

$$x^\sharp = -B^{-1}\mu. \quad (2)$$

Étant donné un  $N$ -échantillon  $(\xi^{(1)}, \dots, \xi^{(N)})$  de la variable aléatoire  $\xi$ , et  $\hat{x}_N$  un estimateur sans biais de  $x^\sharp$  (considéré comme une fonction du paramètre  $\mu$ ) défini sur cet échantillon, on sait que la matrice de covariance de cet estimateur est plus grande (au sens des matrices définies positives) que la borne de **Cramer-Rao** associée :

$$\mathbb{V} \{ \hat{x}_N \} \geq \frac{1}{N} B^{-1} Q B^{-1}. \quad (3)$$

De plus, on sait que l'estimateur de type Monte-Carlo de la solution :

$$\tilde{x}_N = -\frac{1}{N} \sum_{k=1}^N B^{-1} \xi^{(k)} , \quad (4)$$

est un estimateur sans biais *efficace*<sup>(1)</sup> de  $x^\sharp$ .

## 1.1 Algorithme de gradient stochastique normal

Les itérées de l'algorithme de gradient stochastique standard sont fournies par la formule récurrente :

$$x^{(k+1)} = x^{(k)} - \epsilon^{(k)} (Bx^{(k)} + \xi^{(k+1)}) , \quad (5)$$

la suite des pas de l'algorithme étant de la forme :

$$\epsilon^{(k)} = \frac{\alpha}{k^\lambda + \theta} . \quad (6)$$

Il est de plus possible de calculer itérativement la matrice de covariance de cet algorithme :

$$\begin{aligned} \mathbb{V} \{x^{(k+1)}\} &= \mathbb{V} \{x^{(k)} - \epsilon^{(k)} (Bx^{(k)} + \xi^{(k+1)})\} \\ &= \mathbb{V} \{(I - \epsilon^{(k)} B)x^{(k)} - \epsilon^{(k)} \xi^{(k+1)}\} \\ &= \mathbb{V} \{(I - \epsilon^{(k)} B)x^{(k)}\} + \mathbb{V} \{\epsilon^{(k)} \xi^{(k+1)}\} \\ &= (I - \epsilon^{(k)} B) \mathbb{V} \{x^{(k)}\} (I - \epsilon^{(k)} B) + \epsilon^{(k)^2} Q . \end{aligned}$$

En posant  $S^{(k)} = \mathbb{V} \{x^{(k)}\}$  on obtient :

$$S^{(k+1)} = (I - \epsilon^{(k)} B) S^{(k)} (I - \epsilon^{(k)} B) + \epsilon^{(k)^2} Q . \quad (7)$$

## 1.2 Algorithme de gradient stochastique moyenné

L'algorithme du gradient stochastique moyenné est le suivant :

$$x^{(k+1)} = x^{(k)} - \epsilon^{(k)} (Bx^{(k)} + \xi^{(k+1)}) , \quad (8)$$

$$\bar{x}^{(k+1)} = \frac{k}{k+1} \bar{x}^{(k)} + \frac{1}{k+1} x^{(k+1)} , \quad (9)$$

la suite des pas de l'algorithme étant de la forme :

$$\epsilon^{(k)} = \frac{\alpha}{k^\lambda + \theta} . \quad (10)$$

Cet algorithme s'écrit de façon matricielle :

$$\begin{pmatrix} x^{(k+1)} \\ \bar{x}^{(k+1)} \end{pmatrix} = \begin{pmatrix} I - \epsilon^{(k)} B & 0 \\ \frac{1}{k+1} (I - \epsilon^{(k)} B) & \frac{k}{k+1} I \end{pmatrix} \begin{pmatrix} x^{(k)} \\ \bar{x}^{(k)} \end{pmatrix} - \epsilon^{(k)} \begin{pmatrix} \xi^{(k+1)} \\ \frac{1}{k+1} \xi^{(k+1)} \end{pmatrix} , \quad (11)$$

que l'on met sous la forme compacte suivante :

$$z^{(k+1)} = H^{(k)} z^{(k)} - \epsilon^{(k)} W^{(k+1)} . \quad (12)$$

On en déduit l'équation d'évolution de la matrice de covariance :

$$\mathbb{V} \{z^{(k+1)}\} = H^{(k)} \mathbb{V} \{z^{(k)}\} H^{(k)\top} + \epsilon^{(k)^2} \mathbb{V} \{W^{(k+1)}\} . \quad (13)$$

La matrice de covariance de  $W^{(k)}$  se déduit aisément de celle de  $\omega^{(k)}$  :

$$\mathbb{V} \{W^{(k)}\} = \begin{pmatrix} Q & \frac{1}{k} Q \\ \frac{1}{k} Q & \frac{1}{k^2} Q \end{pmatrix} . \quad (14)$$

Enfin, la matrice de covariance de  $\bar{x}^{(k)}$  s'obtient à partir de celle de  $z^{(k)}$  par :

$$\mathbb{V} \{\bar{x}^{(k)}\} = (O, I) \mathbb{V} \{z^{(k)}\} (O, I)^\top . \quad (15)$$

---

<sup>1</sup>dont la variance atteint la borne de Cramer-Rao

### 1.3 Buts de l'étude

On souhaite tout d'abord comparer le comportement des algorithmes de gradient stochastique normal et moyenné, tant durant la phase transitoire que du point de vue du comportement asymptotique. Pour cela, on obtient par tirages aléatoires une trajectoire des bruits  $(\xi^{(1)}, \dots, \xi^{(N)})$ , on fait se dérouler en parallèle les deux algorithmes de gradient stochastique sur ces tirages aléatoires et on compare les trajectoires obtenues. On calcule aussi, avec ces mêmes valeurs des bruits, la trajectoire de l'estimateur (4) de Monte-Carlo. Comme une telle simulation est relativement sensible aux valeurs prises par les bruits, on s'autorise à enchaîner plusieurs simulations des algorithmes et comparer la moyenne arithmétique des trajectoires.

On désire aussi comparer l'efficacité des deux algorithmes en tant qu'estimateurs, et donc comparer leurs matrices de covariance asymptotique à la borne de Cramer-Rao. Pour cela, on peut :

- ou bien enchaîner plusieurs simulations des algorithmes et faire ensuite des statistiques sur l'échantillon de trajectoires ainsi obtenu,
- ou utiliser les formules itératives (déterministes) donnant l'évolution des matrices de covariance, dont les expressions ont été obtenues dans les deux paragraphes précédents, et comparer les valeurs propres de ces matrices à celles de la borne de Cramer-Rao.

Les principales conclusions (à vérifier par l'utilisateur) sont les suivantes :

1. le gradient stochastique normal est plus difficile à régler (coefficients définissant la série des pas  $\epsilon^{(k)}$ ) que le gradient stochastique moyenné ; les conclusions qui suivent supposent que l'algorithme du gradient stochastique normal a été bien réglé ;
2. le gradient stochastique normal est beaucoup plus rapide que le gradient stochastique moyenné durant la phase transitoire de l'algorithme ;
3. le comportement asymptotique du gradient stochastique moyenné est meilleur que celui du gradient stochastique normal ;
4. la plus grande valeur propre de la matrice de covariance asymptotique du gradient stochastique normal est égale à celle et de la borne de Cramer-Rao ;
5. les valeurs propres de la matrice de covariance asymptotique du gradient stochastique moyenné et de la borne de Cramer-Rao sont identiques.

## 2 Description de la boîte à outils

La boîte à outils, entièrement écrite en **Scilab**, est constituée des modules suivants.

### 2.1 Moniteur d'enchaînement

**moniteur.sce** : module permettant d'exécuter toutes les fonctionnalités de la boîte à outils ; pour l'utiliser, lancer la commande : `exec('moniteur.sce')` ; dans la fenêtre Scilab et se laisser guider.

### 2.2 Données du problème

**donn-critere.dat** : fichier contenant la matrice  $B$  définissant la partie quadratique du critère à minimiser (de taille (10,10)).

**donn-bruit.dat** : fichier contenant la matrice  $q$  permettant de calculer la variance de covariance  $Q$  du bruit :  $Q = q * q^\top$  (de taille (10,10)).

Ces deux fichiers, qui ont été générés par Scilab, sont lus à partir du module **param-probleme.sce** (voir paragraphe suivant).

## 2.3 Mise en forme du problème

**param-probleme.sce** : module lisant les données fixes du problème ( $B, q, \dots$ ) et calculant les variables qui en dépendent (dont la borne de Cramer-Rao).

**param-simulation.sce** : module de dialogue d'entrée permettant de définir le nombre de simulations à effectuer.

**param-gradient\_n.sce** : module de dialogue d'entrée des données relatives à l'algorithme de gradient stochastique normal (coefficients définissant le pas de gradient, nombre d'itérations de l'algorithme). Noter que, dans la formule du pas de gradient, l'exposant  $\lambda$  de l'indice d'itération  $k$  est toujours égal à 1 :

$$\epsilon^{(k)} = \frac{\alpha}{k + \theta} ,$$

de telle sorte qu'il est inutile de le lire.

**param-gradient\_m.sce** : module de dialogue d'entrée des données relatives à l'algorithme de gradient stochastique moyenné (coefficients définissant la série des pas de gradient, dont l'exposant  $\lambda$  (sa valeur est alors de l'ordre de 0.6), et nombre d'itérations de l'algorithme).

**param-gradient\_2.sce** : module de dialogue d'entrée des données relatives aux deux algorithmes de gradient stochastique précédents.

## 2.4 Fonctions utilitaires

**Mmvp.sci** : fonction fournissant les valeurs propres extremes du spectre d'une matrice.

**affiche.sci** : fonction pour l'affichage des trajectoires générées par les algorithmes de gradient stochastique. Comme on constate souvent des déviations de grande amplitude lors des premières itérations, la fonction d'affichage permet de n'afficher (sur demande) que les dernières itérations.

## 2.5 Fonctionnalités de la boîte à outils

Il y a trois fonctions possibles.

1. *Simulation unique* (un seul "run" de l'algorithme de gradient stochastique) ; cette simulation est effectuée par les modules **p12.sce**, **p13.sce** ou **p14.sce** suivant que l'on met en oeuvre le gradient stochastique normal, le gradient stochastique moyenné ou les deux simultanément ; dans tous les cas, on fait en parallèle le calcul (4) de l'estimateur optimal de l'arg-min par la méthode de Monte-Carlo, ce qui permet de comparer avec le gradient stochastique. Les résultats représentés sont les trajectoires au cours des itérations de la norme des itérées de l'algorithme. Pour le gradient stochastique moyenné, on affiche les résultats concernant les itérées engendrées par l'algorithme avant et après moyennisation.

2. *Simulation multiple* (plusieurs runs de l'algorithme de gradient stochastique et calcul de statistique sur ces runs : moyenne et écart-type) ; elle est effectuée par les modules **p22.sce**, **p23.sce** ou **p24.sce** suivant que l'on utilise le gradient stochastique normal, le gradient stochastique moyenné ou les deux simultanément ; dans tous les cas, on effectue en parallèle le calcul de l'estimateur optimal de l'arg-min par la méthode de Monte-Carlo. Dans le cas du gradient stochastique moyenné, on n'affiche que les résultats concernant les itérées engendrées par l'algorithme après moyennisation. Les statistiques consistent à calculer à chaque itération la moyenne et l'écart-type empiriques des normes des itérées sur l'ensemble des runs, et à afficher les trajectoires de ces moyennes plus ou moins leur écart-type.
3. *Evaluation des matrices de covariance* (à l'aide de leurs équations récurrentes déterministes) : elle est faite par les programmes **p32.sce**, **p33.sce**, **p34.sce** suivant le type de l'algorithme de gradient stochastique. On affiche les valeurs propres extrêmes des matrices de covariance à la fin des itérations, que l'on compare aux valeurs propres extrêmes de la borne de Cramer-Rao.

Comme on l'a indiqué, chacune de ces fonctions peut donc être utilisée :

- avec le gradient normal seul (programmes **p12.sce**, **p22.sce**, **p32.sce**),
- avec le gradient moyenné seul (programmes **p13.sce**, **p23.sce**, **p33.sce**),
- avec les deux algorithmes (programmes **p14.sce**, **p24.sce**, **p34.sce**).

Il y a donc 9 combinaisons possibles (et donc 9 programmes) qu'il faut choisir après le lancement de **moniteur.sce** en sélectionnant des cases dans les dialogues proposés (la case **Non** signifiant que l'on ne désire pas exécuter la fonction correspondante).

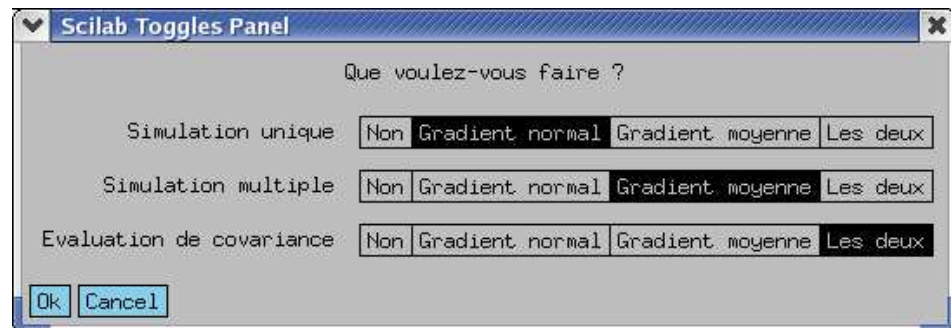


Figure 1: dialogue de choix des opérations à effectuer (sous Linux).